

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

宮内 和義 Kazuyoshi Miyauchi 独立研究者 / Independent Researcher miyauchikazuyoshi@gmail.com



論文 PDF



Landing



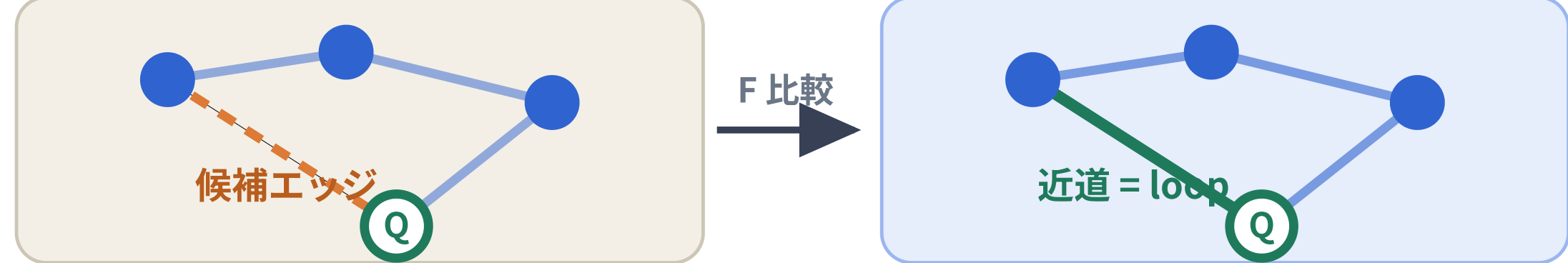
GitHub

THE “WHEN” PROBLEM

動的知識グラフが検索すべき「何を (What)」は GraphRAG・Self-RAG が最適化してきた。だが新情報を「いつ (When)」構造として統合し、いつ探索を続けるか——その原理的な基準は無かった。

geDIG は自由エネルギー原理 (FEP) と最小記述長 (MDL) を操作的に橋渡しし、このトレードオフを単一スカラー F で評価する。

どちらが良い? — 候補を仮に足し F (更新前) と F (更新後) を比較



「いつ統合するか」= 候補エッジを仮に足し F (更新前) と F (更新後) を比較。ループ短絡で近道を生む ($\Delta F < 0$) 場合のみコミット、希薄なら棄却。

GEDIG | GRAPH EDIT DISTANCE AND INFORMATION GAIN

$$F = \Delta EPC_{\text{norm}} - \lambda (\Delta H_{\text{norm}} + \gamma \cdot \Delta SP_{\text{rel}})$$

λ, γ = 情報温度 (構造の複雑さと圧縮率のバランス)

ΔEPC_{norm}

構造変更コスト

GED の局所近似。構造を変えるほど大きい。

FEP 複雑さ・MDL L(M) 増・抑制

ΔH_{norm}

伸展利得 (optionality)

候補分布のエントロピー差。探索圧 = 温度項。

FEP 探索圧・温度項・拡散

ΔSP_{rel}

経路短縮 (近道)

平均最短経路長の相対短縮。効率化を測る。

MDL L(D|M) 圧縮・効率

核心仮説 — 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。

01 CONTROL

AG/DG 二段ゲート制御

エージェントは F を用いて 2 つのゲートを事象駆動的に開閉し、探索モードと統合モードを自律的に切り替える。

AG Attention Gate | 0-hop

$$g_0 = F|_{\text{local}} > \theta_{\text{AG}}$$

局所的な「違和感・未知」を検知 → 探索モードを起動。

DG Decision Gate | multi-hop

$$g_{\text{min}} = \min_h F^{(h)} < \theta_{\text{DG}}$$

大域的な「近道・洞察」を発見 → 統合モードでエッジを確定。

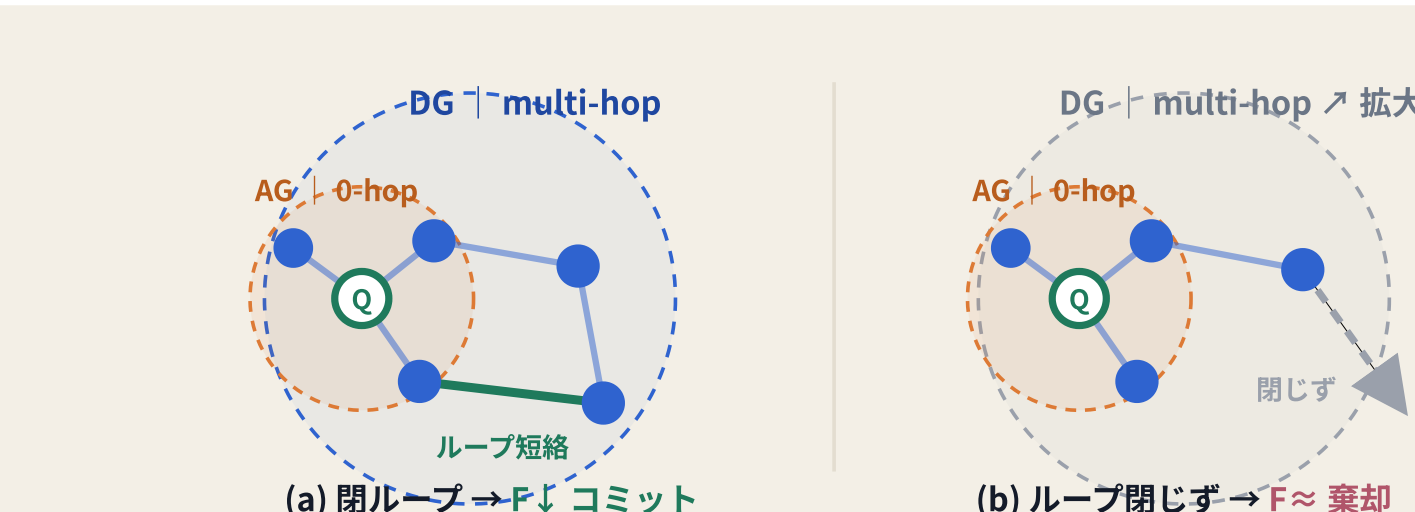


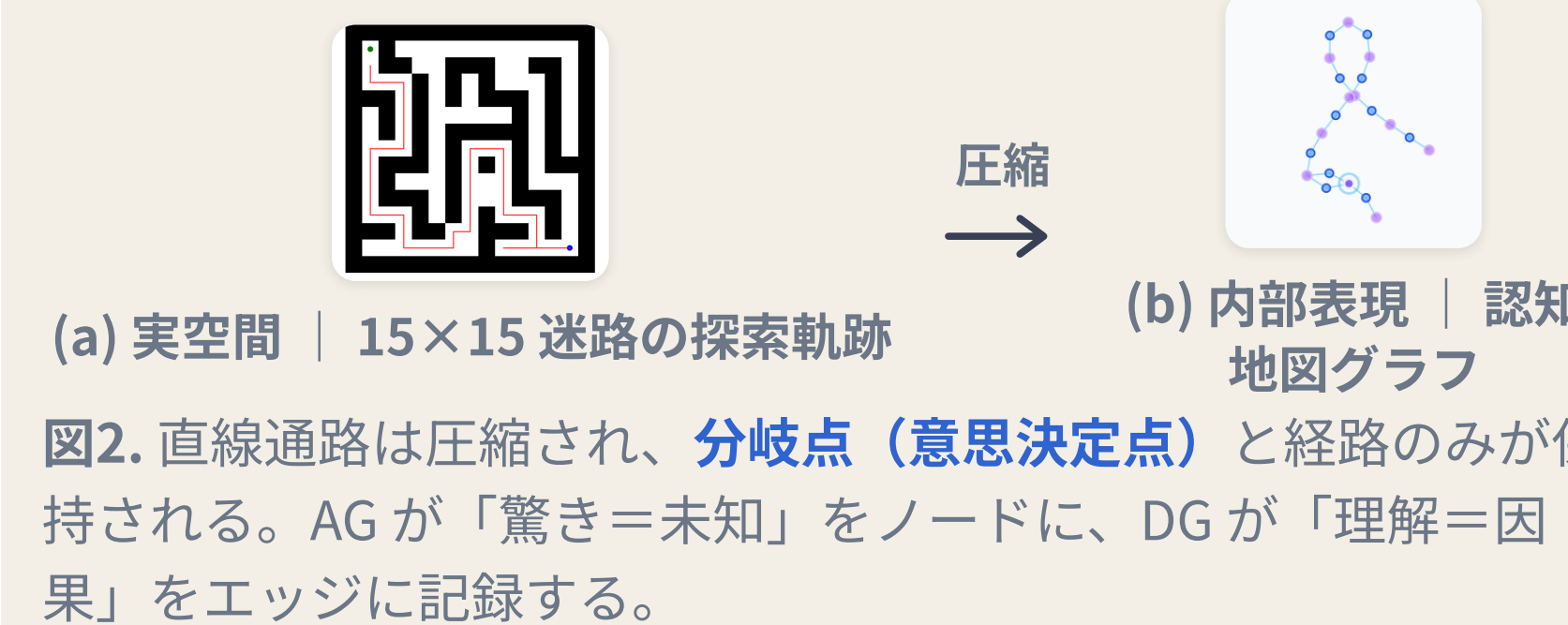
図1. AG は Q と直接つながる 0-hop 近傍 (1次接続) で違和感を検知、DG は multi-hop 範囲で候補がループを短絡=近道になるかを判定。なる場合のみコミットする。

「わからないときだけ調べ (AG)、
わかったときだけ覚える (DG)」

02 VALIDATION I | SPATIAL KG

検証 I: 迷路探索

部分観測迷路でエージェントが築く認知地図は動的知識グラフの一形態。周囲1マスしか見えず、「全通路を記憶」か「分岐点のみ記憶」かの選択を迫られる。



移動エピソードベクトル 8D

位置・相対移動・通行可否・訪問回数・success/goal を符号化 (weights: wall 6 / visit 4)。クエリとの重み付き距離で候補分布を生成する。ゴール座標・正解経路は与えない——移動履歴のみから記憶接続を選ぶ。

98%

15×15 迷路の成功率
(Greedy DFS と同等)

98%

候補エッジの棄却率
= 記憶圧縮率

表2 | 迷路ナビゲーション結果 (部分観測)

手法	成功率	平均ステップ	圧縮率
15×15 (n=100, max=250)			
Random Walk	0.75 ± .08	105.8	—
Greedy DFS	0.99 ± .03	82.1	—
geDIG	0.98 ± .03	89.6	98%
25×25 (n=60, max=500)			
Random Walk	0.40 ± .12	242.2	—
Greedy DFS	0.95 ± .06	242.2	—
geDIG	0.75 ± .11	254.1	99%

geDIG は「必要なトポロジー骨格」のみを選択的に記憶する。vs Random Walk は McNemar 検定で有意 (15×15: $p < 10^{-5}$)。

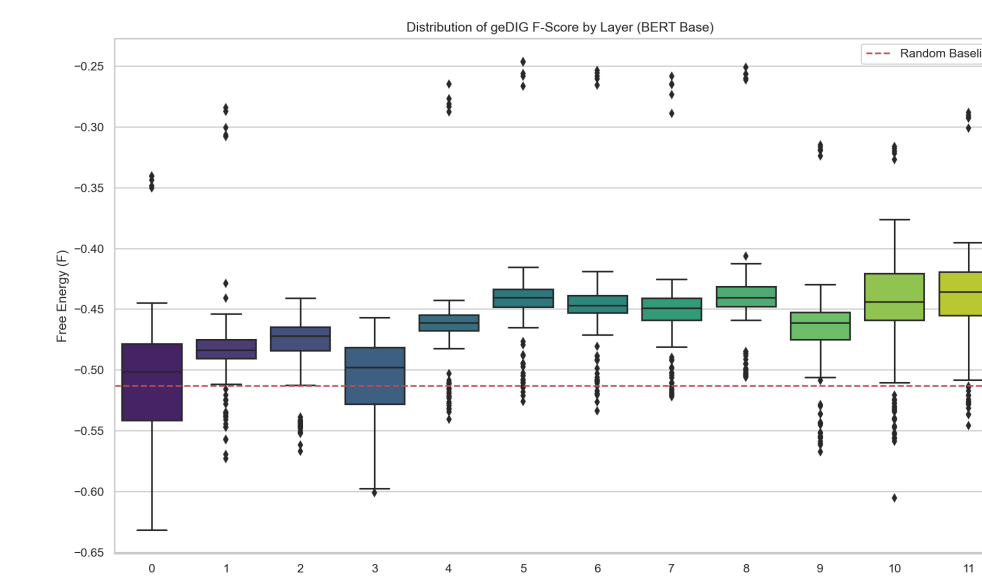
03 VALIDATION II | SEMANTIC KG

検証 II: Transformer

Attention 行列 A_{ij} を関係グラフと見なすと意味的 KG が得られる (BERT 12層・SST-2・n=200)。同じ F の3項を Attention 上で定義:

グラフ化 ノード = トークン、エッジ = Attention A_{ij} (上位 10% を保持)	
ΔEPC	Attention 構成の変化 = 編集コスト
ΔH	Attention 分布の拡散 = エントロピー
ΔSP	Attention 上の最短経路の短縮

図4 | BERT 12層における geDIG F 値の分布 (BERT-base, SST-2, n=200)



- 浅層 0-1 | 探索相 ($F \approx -0.50$)
- 中層 4-5 | 構造化が進行 ($F \approx -0.45$)
- 深層 10-11 | 構造相へ相転移 ($F \approx -0.44$)

Layer 0-1 vs 10-11 は有意 (paired $t(199)=32.6$, $p < 10^{-80}$, Cohen's $d=2.31$)。ヘッドは「探索担当/統合担当」に機能分化。

表3 | F 正則化の因果的介入 ($L = L_{\text{CE}} + \alpha \cdot F$)

α	Accuracy (%)	Δ (pp)
0 (基準)	86.00 ± .53	—
0.001	86.33 ± .46	+0.33*
0.01	86.27 ± .70	+0.27
0.1	85.93 ± 1.1	-0.07
1.0	78.93 ± 1.6	-7.07

*paired t -test $p=0.038$ 。逆U字=微弱な正則化は改善、過剰は悪化。適切な「構造化圧」が汎化を引き出す因果的証拠。学習済み Transformer は既に F 最小化に近い。

THERMODYNAMIC INTERPRETATION

熱力学的自由エネルギーとの同型性

$$\begin{aligned} \text{geDIG } F &= \Delta EPC - \lambda(\Delta H + \gamma \Delta SP) \\ \text{展開} &= (\Delta EPC - \lambda \gamma \Delta SP) - \lambda \Delta H \\ \text{同型} &\cong E - T \cdot S \end{aligned}$$

構造項 ($\Delta EPC - \lambda \gamma \Delta SP$) $\leftrightarrow$ エネルギー E 、 $\lambda \Delta H \leftrightarrow$ 温度 $\times$ エントロピー $T \cdot S$ ($T = \lambda$: 情報温度)。層別 F 推移は「高温 $\rightarrow$ 低温への冷却 = アニリング過程」と読める。

高温・高エントロピー | 探索

低温・低エントロピー | 構造

結論と展望

- 2種の動的 KG (迷路・Transformer) で、 F が共通の構造化原理として機能することを実証。
- いずれも探索相 $\rightarrow$ 構造相への相転移を示し、F 正則化により F と性能の因果関係を確認。

FUTURE WORK

- $\Delta SP \rightarrow \beta_1$ への位相拡張。近道 (ΔSP) を第一ベッチ数 β_1 に一般化し、ループ・短絡を位相不変量として扱う。
- GraphRAG / Self-RAG への応用。AG/DG ゲートで現行 RAG の「いつ探索し、いつ構造化するか」を原理的に自律制御する。