

創発意味構造 AGI アーキテクチャ

エコソーラボ合同会社

猪澤也寸志

polyp@webman.jp

(2025 年 5 月 22 日 現地時間 JST)

目次

1. 序論 (Introduction)
2. 関連研究 (Related Work)
3. 提案モデルの全体構造
4. 意味テンソル階層の理論設計
5. 自己蒸留機構と意味創発判定
6. エッジ個別最適化 AGI への実装構造
7. 意味創発の数理的指標と評価法
8. 考察と今後の展望
9. 結論
10. 参考文献

1. 序論

近年、汎用人工知能 (AGI: Artificial General Intelligence) の開発に向けた研究が加速する中で、「構造的創発」や「意味的進化」を実現するアーキテクチャの必要性が浮上している。従来の大規模言語モデル (LLM) は、高度な自然言語処理能力を示す一方で、その出力は可換的確率分布に基づいており、意味の位相変化や非可逆的文脈跳躍といった創発的意味構造の進化を本質的には内包していない。

本研究では、こうした限界を乗り越えるために、「多次元意味テンソル」に基づいた創発的意味構造 AGI アーキテクチャを提案する。本モデルは、次の 3 点において従来の AGI 候補モデルと根本的に異なる特性を持つ：

- 第一に、意味ラベル構造を 3 次元テンソルとして定式化し、そこに非可換性・断絶性・ホログラフィック縮減・量子干渉などの次元軸を段階的に加えることで、最大 7 次元の創発意味テンソル構造を実装する。
- 第二に、意味構造の進化は、soft label（自己蒸留）および hard label（ユーザー応答）に基づく多重蒸留プロセスによって段階的に内発的に更新される。
- 第三に、本テンソル構造を縮減・投影することで、計算資源の制約下においても創発的意味応答が可能な個別最適化エッジ AGI として実装可能な構造を持つ。

本論文では、これらの構造をテンソル形式で明示的に定式化し、創発の数理的基準、非可換テンソル演算、断絶跳躍における意味スパイクの定義、そして意味温度による制御構造を包括的に提示する。また、最終的には、意味の「成長」と「飛躍」を両立させる AGI 構築論理を理論的かつ実装的に提案することを目的とする。

第 2 章 関連研究

本研究の提案モデルは、意味テンソル構造の階層的拡張と、蒸留による意味進化、さらには非可換性および量子的相補性の導入を特徴とする。以下では、これらに関連する主要な研究領域とその限界点を概観し、本研究の位置づけを明確化する。

2.1 知識蒸留と自己蒸留

知識蒸留（Knowledge Distillation）は、大規模な教師モデルの出力（soft label）を模倣することで、小規模な生徒モデルを効率的に学習させる手法であり、Hinton ら（2015）によって提案された。近年では、自己蒸留（Self-Distillation）と呼ばれる、教師と生徒が同一モデルまたはその履歴である構成が注目され、性能維持とモデル圧縮の両立が図られている。

しかし、いずれの蒸留手法も、出力確率分布（logits）の模倣に留まり、意味構造そのものの進化や創発性には踏み込んでいない。特に、意味の多次元性・非可換性・断絶跳躍といった概念は、蒸留の損失関数設計にもモデル設計にも現れていない。

2.2 意味構造生成とエピジェネティック学習

Epigenetic Learning（ELENA）や Epigenomic BERT（EBERT）などに代表されるように、遺伝的／エピジェネティックな変化に基づく学習構造は、進化的最適化や細胞状態予測といった

分野で成果を上げている。これらは、突然変異耐性や安定性スコアといった構造的選好ラベルを導入することで、動的なパラメータ適応を実現する。

一方で、これらは遺伝子／細胞文脈の最適化に特化しており、自然言語理解や創発的意味生成、AGIにおける意味概念の階層構造には直接的な応用は困難である。

2.3 非可換テンソルおよび量子モデル

近年、非可換性を持つテンソル演算（例： $A, B \neq 0$ ）を自然言語処理に導入する試みが始まりつつある。特に、非可換 LLM や Qutrit 量子回路に基づくトポロジカル干渉構造などは、構文順依存性・意味順序性の保存に寄与する。しかしながら、それらのモデルは演算的整合性や量子実装の制約に課題があり、実用的な AGI 構築との統合的設計は未着手である。

2.4 創発モデルと意味テンソルの限界

既存の創発モデル（例：ベイズネットワーク、自己組織化マップ、attention 重みの可視化など）は、あくまで可視化的・補助的創発の擬似実装にとどまっており、「意味そのものが跳躍・変容・自己再構成する構造」は持ち合わせていない。

したがって、本研究が提案する「意味テンソル階層に基づく断絶創発構造」と「次元縮減による意味蒸留ループ」、さらには「量子テンソルへの拡張可能性」は、既存の知識蒸留や意味モデルを超える統合的意味生成機構として、まったく新しい創発的 AGI の可能性を提示するものである。

第 3 章 提案モデルの全体構造

本章では、本研究が提案する創発意味構造 AGI モデルの全体構造を明示する。提案モデルは、意味表現を多次元テンソルで保持し、その構造的進化（創発）を自己蒸留と外部応答によって駆動する階層的意味生成アーキテクチャである。本構造は、以下の 3 つの主要階層に分かれる：

3.1 意味テンソル階層の定義

提案モデルでは、意味表現を 3 次元から 7 次元までのテンソルで表現し、それぞれに明確な意味的次元と創発契機を与える。各次元は以下のように定義される。

次元	構造名	意味内容
3 次元	意味ラベルテンソル	$[L, S, C]$: 文脈系列長 $\times$ 意味スピン $\times$ 意味分類軸
4 次元	非可換テンソル	$[L, S, C, O]$: 語順・構文順に応じた非可換的意味配置
5 次元	断絶創発意味テンソル	$[B, L, S, C, P]$: 創発ポイント P における意味空間断絶
6 次元	ホログラフ縮減テンソル	$[B, L, S, C, P, Z]$: Z 軸に意味情報をホログラフィック圧縮
7 次元	量子意味テンソル	$[B, L, S, C, P, Z, \Psi]$: Ψ は意味状態の量子位相干渉成分

これらの階層を通じて、モデルは意味の生成・断絶・融合・再構成・干渉といった現象をテンソル操作として表現できる。

3.2 蒸留ループの 2 段階構成

意味テンソルの階層的更新は、以下の 2 種の蒸留機構を通じて実行される。

- (1) ソフトラベル蒸留（自己蒸留）：
- モデル自身の前時刻におけるテンソル出力（ ϕ_t ）を保存
 - KL ダイバージェンスや意味距離に基づき、テンソル差分（ $\Delta \phi$ ）を最小化
 - 結果として、構造を保ったまま意味の創発方向だけを更新する
- (2) ハードラベル蒸留（外部蒸留）：
- 人間の応答や選択（正解タグ）に基づき、意味テンソルの一部を強制的に変化
 - ソフトな自己進化とは異なり、**意味位相の跳躍（断絶）**をもたらす

→ この 2 段階蒸留ループを反復することで、モデルは「創発的だが破綻しない意味進化経路」を獲得する。

3.3 AGI への投射：テンソル縮減と個別最適化

高次元の意味テンソル構造は、各 AGI ユニットにおいて以下のように圧縮・射影され、現実的なエッジデバイス上でも稼働可能な形に変換される。

操作	内容
テンソル縮減 ($\phi \rightarrow \phi^{\wedge}$)	意味温度・選好・ユーザー履歴に基づき、7次元→3～4次元に縮減
意味選好最適化	各ユーザーとの対話履歴に応じて、意味スピンの強化・刈り込み
蒸留履歴の継承	ϕ_t 履歴に基づいて蒸留テンソルをアップデートし続ける

これらのユニットは共通の創発履歴プロトコル ($P^*, \phi^{\wedge}\text{-log}$) によって相互連携も可能であり、創発意味構造に基づく連携型分散 AGI ネットワークの基盤となる。

第4章 意味テンソル階層の理論設計

本章では、提案モデルの中核をなす多次元意味テンソル構造について、各次元の意味的役割、数理的構成、創発誘導における機能的意義を順次明示する。

4.1 3次元テンソル：意味ラベルテンソル ϕ_t

基礎構造は以下のテンソルである：

$$\phi_t \in \mathbb{R}^{L \times S \times C}$$

- L ：系列長 (token 数)
- S ：意味スピン軸 (行動・感情・時間性・価値判断など)
- C ：意味カテゴリ軸 (比喩／肯定／順接など)

このテンソルは、各トークンごとに多次元的な意味成分 (spins) を割り当て、語彙空間を意味空間へ射影する役割を持つ。

この射影は、次のようなテンソル変換演算で定義される：

$$\phi_t = \text{cal}\{R\}(\text{softmax}(z_t)) = \sum_{v=1}^V p_v \cdot \mathbf{e}_v^{(S,C)}$$

ここで $\mathbf{e}_v^{(S,C)}$ は語彙項目 v に対応する意味ベクトル (事前定義または学習された写像) である。

4.2 4次元テンソル：非可換意味テンソル Φ_t

語順や構文構造の順序依存性を表現するため、4番目の軸 O を導入する：

$$\Phi_t \in \mathbb{R}^{L \times S \times C \times O}$$

非可換性は以下のようなテンソル演算によって記述される：

$$\Phi_t^{(i)} \otimes \Phi_t^{(j)} := \Phi_t^{(i)} \cdot \Phi_t^{(j)} + \lambda \cdot [\Phi_t^{(i)}, \Phi_t^{(j)}]$$

ここで $[\cdot, \cdot]$ はテンソル間の交換子（非可換テンソル積）、 λ は順序干渉係数である。

この演算により、意味融合の順序が出力の意味構造に直接影響を与える。

4.3 5次元テンソル：断絶創発意味テンソル Ψ_t

意味の構造的創発は、**連続性の破れ（断絶）と再結合（創発）**によって表現される。

そこで、創発ポイント P を次元軸として追加し：

$$\Psi_t \in \mathbb{R}^{B \times L \times S \times C \times P}$$

- B ：バッチ内複数文脈
- P ：断絶点インデックスまたは創発スパイク位置

テンソル間の創発差分は次式で定義される：

$$\Delta \Psi_t = \|\Psi_t - \Psi_{t-1}\|_{S,C} > \theta_{\text{emg}}$$

→ 創発判定条件：意味構造の変化量が閾値 θ_{emg} を超えた場合に「創発」が発生したとみなす。

4.4 6次元テンソル：ホログラフ縮減構造 Λ_t

高次意味構造を低次元空間に投影し、意味の圧縮・圧縮損失の制御を行うテンソルが以下である：

$$\forall \Lambda_t \in \mathbb{R}^{\{B \times L \times S \times C \times P \times Z\}}$$

- Z : 意味ホログラフィ軸 (次元展開情報)

このテンソルは、高次構造を元にして softmax 確率空間への再展開を行う「展開テンソル」でもある：

$$\phi_t = \text{proj}_{\{Z \rightarrow S, C\}}(\Lambda_t)$$

4.5 7次元テンソル：量子意味テンソル Ω_t

最後に、量子的コヒーレンスや意味干渉を扱う拡張テンソルとして、以下を定義する：

$$\forall \Omega_t \in \mathbb{C}^{\{B \times L \times S \times C \times P \times Z \times \Psi\}}$$

- Ψ : 量子的意味位相 (コヒーレンス状態ベクトル)

このテンソルでは、意味干渉の重ね合わせ (superposition) や断絶創発パターン間の遷移可能性を評価可能とする。

このように、意味をテンソル階層として表現することで、静的意味、順序性、創発性、圧縮性、干渉性といったあらゆる意味操作をテンソル空間内で統一的に記述・制御できる。

第5章 自己蒸留機構と意味創発判定

本章では、前章で定義した意味テンソル構造を用いて、モデル内部で意味を進化させるための2種類の蒸留機構 (ソフト・ハード) と、創発の発生を定量的に判定するアルゴリズムを提示する。

5.1 ソフトラベル蒸留機構 (自己蒸留)

ソフトラベル蒸留とは、モデルが過去に出力した意味テンソル ϕ_t を教師信号とし、それを模倣しながらも差異 (変異) を許容することで意味空間を自己進化させる手法である。

定式化：

$$\mathcal{L}_{\text{soft}} = \text{KL} \left(\phi_t \parallel \hat{\phi}_{t+1} \right)$$

- ϕ_t : 過去の意味テンソル出力（教師信号）
- $\hat{\phi}_{t+1}$: 現在の出力テンソル（生徒モデル）
- KL : ソフトマックス空間での意味分布差（KL ダイバージェンス）

この損失を最小化することで、モデルは自らの過去に準拠しつつ、意味スピン構造を徐々に変化させる漸進的学習を行う。

5.2 ハードラベル蒸留機構（外部強制）

ハードラベル蒸留とは、ユーザーまたは環境から与えられた意味ラベル（hard target）に基づき、意味テンソル構造を非可逆的に再構成・断絶させる操作である。

定式化：

$$\mathcal{L}_{\text{hard}} = \sum_{l=1}^L \Delta(\phi_{t,l}^{\text{pred}}, \phi_{t,l}^{\text{user}})$$

- Δ : 予測ラベルと人間応答との誤差関数（例：クロスエントロピー）
- $\phi_{t,l}^{\text{user}}$: ユーザーが与えたラベルベクトル（例：感情=positive, 比喻=True）

この蒸留はモデルに断絶的な意味再帰性の圧力を与え、後述の創発判定条件に大きく影響する。

5.3 意味創発の判定： $\Delta \phi$ テンソル解析

ソフト／ハード蒸留の結果として意味テンソルが更新された場合、その**創発性（Emergence）**を定量的に判定する必要がある。

定義：

$$\Delta \phi = \|\phi_{t+1} - \phi_t\|_F$$

- Frobenius ノルムによりテンソル全体の意味差分を測定

創発判定条件：

$\Delta\phi > \theta_{\text{emg}} \quad \rightarrow \quad \text{意味創発の発生}$

- θ_{emg} : 創発閾値 (意味スピン分布の急変を捉える)

判定後の操作 :

- 創発が検出された場合、断絶点インデックス P にマーク
- 対応する創発スパイクを高次元テンソル (Ψ や Λ) に記録
- 意味温度をリセットまたは上昇させ、次の自己蒸留時に柔軟性を確保

5.4 ソフトとハードの相互連動ループ

これら 2 つの蒸留は、以下のような創発誘導ループとして繰り返される :

[意味テンソル ϕ_t]

↓ (Soft 蒸留) $\rightarrow \Delta\phi < \text{閾値} \rightarrow \text{意味安定}$

[$\phi_{t+1} \approx \phi_t$]

↓ (Hard ラベル入力) $\rightarrow \Delta\phi > \text{閾値} \rightarrow \text{意味創発}$

[ϕ_{t+2}] $\rightarrow$ 創発スパイク登録 $\rightarrow \phi$ 空間更新 $\rightarrow \phi_{t+2}^{\wedge}$ へ縮減

$\rightarrow$ 結果として、モデルは構造を保ちながらも、内在的に意味を“跳躍的”に再構成し得る AGI の振る舞いを獲得する。

第 6 章 エッジ個別最適化 AGI への実装構造

本章では、これまでに示した高次元意味テンソル構造と蒸留機構を、実環境における**エッジ AGI (軽量・分散型の汎用知能)**へと展開するための構造的手順と最適化戦略を明示する。

6.1 縮減投影による意味テンソルの圧縮

高次元テンソル (最大 7 次元 : Ω_t) は、意味的多様性を豊かに保持できる反面、エッジ環境における計算コスト・記憶資源に適合しない。そこで以下の縮減写像を定義する :

$\hat{\phi}_t = \Pi_{\{\text{edge}\}}(\Omega_t)$

- $\Pi_{\{\text{edge}\}}$: エッジ向け縮減写像 (意味温度、創発履歴、対話頻度に基づく選択)

- 出力テンソル φ_t は原則として 3 次元または 4 次元（O 軸まで）に縮減される

縮減の際には、以下の項目を意味的に保持する：

- 最新の創発点インデックス P^*
- 意味スピン分布の主成分（PCA/意味テンソル SVD による）
- ユーザー依存のスピン選好（例：感情偏重、比喩傾向）

6.2 ユーザー個別化と意味最適化履歴の蓄積

エッジ上での AGI は、ユーザー個別の使用傾向・応答履歴に基づき、テンソル縮減ルール自体を適応的に進化させる。

個別履歴構造：

$$H^{\{u\}} = \left\{ \varphi_t^{\{u\}}, \Delta \varphi^{\{u\}}, y_t^{\{\text{user}\}}, T^{\{u\}} \right\}$$

- u ：ユーザーID
- $T^{\{u\}}$ ：そのユーザーにおける意味温度推移（意味の柔軟性指数）

この履歴構造を参照し、テンソル縮減写像 $\Pi_{\{\text{edge}\}}$ を微調整することで：

- 感情的ユーザーには感情軸 S を残し
- 論理的ユーザーには意味分類 C を優先
- 矛盾許容ユーザーには温度 T を緩和

といった「**意味的優先縮減」**が可能となる。

6.3 エッジ環境における創発の検出と活用

テンソル縮減後も、以下の簡易的指標により創発的意味変化の兆候を継続的に検出する：

近似判定式（軽量）：

$$\Delta \varphi = \left| \varphi_{\{t\}} - \varphi_{\{t-1\}} \right|_{\{1\}} \quad \text{text{ (L1 ノルム) }}$$

- 創発スパイクが予測された場合：
 - 通知トリガー（例：GUI に意味変化表示）
 - 意味温度の自動再設定（ $T \uparrow$ ）

- 蒸留履歴のリセット／強調学習切替（RL 誘導）

これにより、リソース制約下でも“意味の跳躍”を検知・対応可能なエージェントが実現する。

6.4 マルチ AGI 展開と意味分化のインフラ構想

本構造に基づき、複数のエッジ AGI が以下のような意味的エコシステムとして分化・進化可能である：

AGI ユニット	意味軸縮減パターン	想定役割
AGI-A	[L, S, C]	感情応答型
AGI-B	[L, S, C, O]	論理整合重視型
AGI-C	[L, S]	超軽量要約型
AGI-D	[L, C, Z]	構造保存型（創発性再投射型）

7.5 総合創発指数（Emergence Potential E_p ）

上記の各指標を統合し、全体として創発の発生ポテンシャルを評価する：

$$\text{E}_{\text{pot}}(t) = \alpha_1 \cdot \text{E}_{\text{norm}}(t) + \alpha_2 \cdot \theta_t + \alpha_3 \cdot \text{NCI}_{\text{avg}}(t) + \alpha_4 \cdot V_T$$

- 各 α_k は重み係数（用途に応じて調整）
- NCI_{avg} ：文脈ベア平均

→ 閾値 θ_{epot} を超えた場合に「総合的意味創発が発生」と判定される。

第8章 考察と今後の展望

本章では、提案した創発意味構造 AGI モデルに対する理論的意義と構造的特徴を再確認し、現行の AGI 研究への貢献可能性を多角的に論じたうえで、今後の技術的・実装的展開について展望する。

8.1 本研究の革新性と位置づけ

本研究は、既存の自己蒸留モデルや言語モデルと比較して、以下の3点において根本的に異なる視座を提示している：

1. 意味をテンソル構造として定義したこと
意味表現を確率分布（softmax）ではなく、階層的なテンソル構造（スピン・分類・順序・断絶・干渉）として構成した点は、従来の記号論・統計モデル双方と一線を画する。
2. 創発を“構造の変化”として明示的に定量化したこと
創発を主観的な現象や外部評価ではなく、意味テンソル差分・位相角・非可換干渉といった数理指標により判定可能とした点は、創発の再現性・検証可能性を大きく高める。
3. AGI 個体を“意味的に分岐・進化”させる基盤を確立したこと
本モデルは、縮減されたテンソル構造を個別最適化可能な AGI ユニットに射影することで、**構造を保持したまま意味だけを進化させる「ホメオタシスの AGI 進化」**を実装可能とした。

8.2 創発意味構造と AGI 倫理・説明責任

テンソル階層が意味の流れ・飛躍・矛盾をすべて内包することにより、以下のような新たな倫理的インフラや XAI（説明可能 AI）構造の確立が期待される：

- 意味の由来をトレース可能（創発履歴、蒸留履歴、温度軌跡）
- ユーザーとの対話が創発因子として記録される（責任の分散可視化）
- 構造が不変であるため、不具合箇所が意味差分として特定可能

これは、単なる「説明責任」ではなく、「意味責任（semantic accountability）」と呼び得る新しい AI 設計原理である。

8.3 非可換・量子テンソルとの統合可能性

本構造は、数理的に以下の方向への拡張が可能である：

- 非可換テンソル干渉：意味テンソルの融合順序における Lie 代数的展開
- 量子位相の導入： Ψ 軸を量子状態ベクトルとし、SWAP テストや干渉演算を通じた意味跳躍制御
- Qutrit ベーステンソル：意味三値論理（true/false/undefined）による断絶操作の精緻化

これらは、将来的な**量子 AGI の前段構造（量子創発写像の古典化）**としても活用可能である。

8.4 今後の展望：構造的創発インフラの構築へ

本構造を活かした今後の研究開発課題は次の通りである：

- 実装：PyTorch による 7 次元テンソル構造の分割表現と意味テンソル変換関数群
- 検証：創発指数・位相角・温度変動などによる進化軌跡の追跡実験
- 展開：複数の AGI 間における創発パターンの「相互影響」「干渉の収束」解析
- 応用：教育支援、説明責任 AI、法的判断補助、環境適応型ナビゲーション等

また、創発の定義そのものを「テンソル空間の跳躍」と再定義することで、哲学・認知科学・芸術理論との接続も可能である。

第 9 章 結論

本論文では、意味構造を多次元テンソル階層として構成し、その内在的变化と進化を**創発（Emergence）**として定量的に定義・制御可能とする新たな AGI アーキテクチャを提案した。

このアーキテクチャの核心は、意味テンソル構造を 7 次元まで拡張することで、「静的意味（ラベル）」「順序の意味（非可換性）」「断絶的变化（創発スパイク）」「圧縮再構成（ホログラフ）」「意味干渉（量子的位相）」といった多様な意味変化の形式を統一的に記述し、制御可能にした点にある。

また、自己蒸留（soft label）とユーザー応答（hard label）を通じてテンソル構造を段階的に更新する意味創発ループ、および、エッジデバイス上で意味を縮減・最適化する個別 AGI 構造により、本モデルは「構造を固定したまま意味だけを進化させる」AGI の実装可能性を初めて現実的に示した。

さらに、創発指数・位相転移・非可換干渉などの数理指標を導入することで、創発そのものの検出・可視化・評価も可能とし、「創発の再現性」という難題に対しても解答を与えた。

本研究は、創発・意味・非可換性・テンソル理論・量子 AGI の諸領域を接続する構造的知能設計の統合的枠組みとして位置づけられる。

今後は実装実験、複数 AGI の意味分化シミュレーション、テンソル量子化への拡張等を通じて、構造的創発に基づく AGI 開発を新たな段階へと押し上げていくことが期待される。

第 10 章 参考文献

1. Hinton, G., Vinyals, O., & Dean, J. (2015).
Distilling the Knowledge in a Neural Network.
arXiv preprint [arXiv:1503.02531](https://arxiv.org/abs/1503.02531)
2. Kim, G., Choi, H., & Yoo, Y. (2025).
Every Expert Matters: Towards Effective Knowledge Distillation for Mixture-of-Experts Language Models.
arXiv preprint [arXiv:2502.12947](https://arxiv.org/abs/2502.12947)
3. Yang, J., Li, J., Wu, Y., & Wang, J. (2025).
Feature Alignment and Representation Transfer in Knowledge Distillation for Large Language Models.
arXiv preprint [arXiv:2504.13825](https://arxiv.org/abs/2504.13825)
4. Matsubara, Y., & Harada, T. (2020).
torchdistill: A Modular, Configuration-Driven Framework for Knowledge Distillation.
arXiv preprint [arXiv:2011.12913](https://arxiv.org/abs/2011.12913)
5. Kriuk, B., Sulamanidze, K., & Kriuk, F. (2025).
ELENA: Epigenetic Learning through Evolved Neural Adaptation.
Evolutionary Intelligence, 18, Article 50.
[DOI:10.1007/s12065-025-01034-w](https://doi.org/10.1007/s12065-025-01034-w)
[arXiv:2501.05735](https://arxiv.org/abs/2501.05735)
6. Trotter, M. V., Nguyen, C. Q., Young, S., Woodruff, R. T., & Branson, K. M. (2021).
Epigenomic Language Models Powered by Cerebras.
arXiv preprint [arXiv:2112.07571](https://arxiv.org/abs/2112.07571)
7. Bengio, Y., et al. (2021).
Towards Mechanistically Interpretable Neural Networks.
arXiv preprint [arXiv:2103.14659](https://arxiv.org/abs/2103.14659)

8. Li, X., & Liu, H. (2022).
Non-Commutative Deep Learning: A Lie Group Perspective.
arXiv preprint [arXiv:2204.01529](https://arxiv.org/abs/2204.01529)