

機械学習を用いた指定変形を有する格子構造のトポロジー設計

Topology Design of Lattice Structures with Specified Deformations using Machine Learning

○松岡 知哉^{*1}, 大崎 純^{*2}, 林 和希^{*3}

Tomoya MATSUOKA^{*1}, Makoto OHSAKI^{*2} and Kazuki HAYASHI^{*3}

^{*1} 京都大学工学研究科 大学院生

Graduate Student, Graduate School of Eng., Kyoto Univ.

^{*2} 京都大学工学研究科 教授・博士(工学)

Prof., Graduate School of Eng., Kyoto Univ., Dr. Eng.

^{*3} 京都大学工学研究科 助教・博士(工学)

Asst. Prof., Graduate School of Eng., Kyoto Univ., Dr. Eng.

キーワード：格子構造; ポアソン比; トポロジー設計; 機械学習; 特徴選択; 転移学習

Keywords: Lattice structure; Poisson's ratio; topology design; machine learning; feature selection; transfer learning.

1. 序

トポロジー設計は、構造要素の有無を変数とする設計問題であり、建築分野ではトラスの部材配置や骨組のブレース配置が代表的である¹⁾。本稿で扱う格子構造は、軽量な壁構造あるいは構造要素として利用可能であり、変形特性を最適化することで、構造としてのポアソン比などの様々な機械的特性を実現できる。例えば、Reis et al.²⁾は遺伝的アルゴリズムを用いて2次元の直交異方性のメタマテリアルのパラメータを最適化した。

格子構造のトポロジー最適化問題は、組合せ最適化問題であり、近似最適解を得るために多くの計算量を必要とする。したがって、機械学習を用いて計算負荷を低減する手法が近年多く提案されている³⁾。さらに、データのラベルあるいは評価値を精度よく表現できる特徴(変数)の部分集合を選ぶ特徴選択を入力データに行うことで、学習・予測コストを低減できる。また、転移学習は、学習済みモデルの一部を異なるモデルに転用する手法であり、学習コストを大幅に低減できる。

一方、機械学習に用いられるニューラルネットワーク(以下 NN)などの、入力と出力の複雑な依存関係を近似する関数モデルは、その予測結果を解釈するのは難しいという側面がある。しかし、設計の実務に機械学習を応用する場合、その予測の信頼性は重要である。そのため、予測の妥当性を検討するため、複数の予測結果を解析することで予測に対する特徴量の寄与を分析する手法が開発されている^{4,5)}。

本稿では、特定の節点変位を目的関数として、指定変形を有する格子構造のトポロジー設計を扱う。さらに、規模の大きい格子構造に対し既習の NN を転用する手法と予測の解釈性を得る手法について述べる。

2. 設計問題の定式化

2.1. 格子構造モデル

図 1 のように、正方形ユニットと斜材で構成された $m \times m$ グリッドからなる格子構造を考える。下辺はピン支持され、上辺に鉛直方向下向きの分布荷重が作用する。図 1 の赤色の節点をターゲット節点とし、その水平変位が格子構造のポアソン比を代表するものとする。

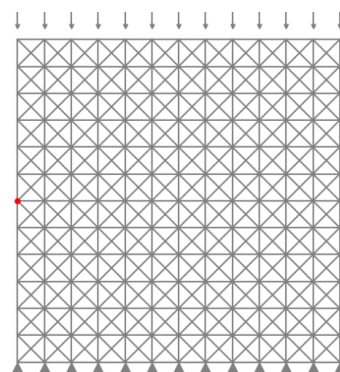


図 1. 基本とする格子構造 ($m = 12$)

構造解析のモデルとパラメータは以下のように定める。

- 部材は Euler-Bernoulli 梁要素でモデル化する。
- グリッドの 1 辺の長さを 0.10 m とする。部材の材料は GFRP を想定し、ヤング係数 $E = 4.1 \times 10^4$ [N/mm²]、断面積 $A = 1.0 \times 10^2$ [mm²]、断面 2 次モーメント $I = 1.0 \times 10^4$ [mm⁴] とする。
- 節点荷重は、分布荷重を想定し上辺の端点に 0.5 kN、上辺の他の節点にそれぞれ 1.0 kN とする。
- 斜材はグリッドの中央で接合されていない。

2.2. 設計変数の設定

以下のような部材除去によって最適位相を求める。

- 図 1 に示す初期位相から部材を除去し、設計変数は部材の有無を表す 0-1 変数とする。
- 1 つのグリッドについて、隣接するグリッドと重複がないように数えると、縦、横、右斜め、左斜めの 4 部材を選ぶことができ、それぞれのグリッドで 1 つの部材のみを除去する。
- 格子構造は中央の鉛直軸に関して対称とする。
- 対称軸上の部材は全て削除し上辺の部材は全て残す。

3. ニューラルネットワークの構築

3.1 フィルターを用いたデータ変換

グリッドパターンはそれぞれのユニットが独立な特性を持つのではなく、その隣接関係が重要な情報であることが想定できる。そのため、 2×2 のグリッドで構成される部分領域における部材配置パターンを表すフィルターを設定する。このとき、部分領域ごとに部材配置パターンは 4^4 通り、領域数は設計の対称性から $(m-1)(m/2-1)$ であり、入力変数の次元は $4^4 \times (m-1)(m/2-1)$ となる。

3.2 特徴選択

2.2 節で示したルールに従う 100000 通りのランダムな部材接続関係を生成して構造解析を行い、ターゲット節点の水平変位を求める。値が大きい 10% の解をラベル 1、値が小さい 10% の解をラベル -1 とし、20000 個のカテゴリーデータを用意する。特徴とラベルの独立性を表すカイ 2 乗統計量が大きい順に特徴を選択し、 $m = 12$ の格子構造を対象にデータの次元を $14080 \rightarrow 5000$ へと削減する。

3.3 ニューラルネットワークの学習

上記と同様の手順で 100000 個の学習データと 10000 個のテストデータを作成し、以下のような NN の構成および学習条件で訓練を行う。

- ネットワーク構成は、括弧内をノード数として、入力層(5000)、中間層 1(1000)、中間層 2(100)、出力層(1) とする。
- 活性化関数は、中間層 1, 2 に ReLU を用い、出力層には用いない。
- それぞれの層は全結合層である。
- 最適化アルゴリズムに Adam[®] を用いる。
- ミニバッチサイズ 256、エポック数 50 とする。

テストデータに対する平均二乗誤差 (MSE) を計算し、表 1 に、 $m = 12$ の構造の場合の結果を示す。特徴選択を行うことで学習・予測時間を大きく削減しつつ、十分な予測精度を実現していることが分かる⁷⁾。

この学習済み NN を焼きなまし法 (SA) を用いた最適化に組み込むことで、効率的に解を探索できる。図 2 に、ポアソン比最大化を目的とした最適化における最良解を示

す。荷重の作用方向と垂直な方向に広がる変形が顕著になるような、特徴的な部材配置が見て取れる。

表 1. $m = 12$ の構造の NN の学習⁷⁾

	特徴選択しない場合 (特徴数 14080)	特徴選択する場合 (特徴数 5000)
MSE[mm ²]	7.201×10^{-5} [mm ²]	9.128×10^{-5} [mm ²]
学習時間[s]	2011	779
予測時間[s]	0.547	0.213

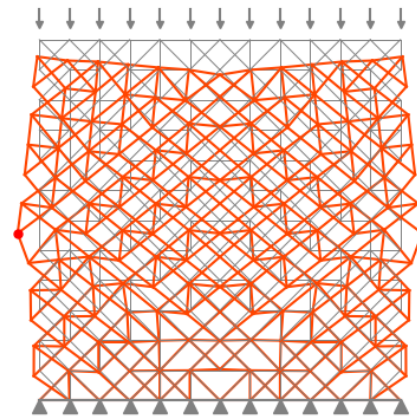


図 2. ポアソン比最大問題の最良解⁷⁾

4. 既習モデルの大規模な構造への転用

学習済み NN を、異なるグリッド数を持つ構造の予測にそのまま用いることはできない。そこで、隣接する部材配置の特徴を縮約表現する畳み込み層を追加することで、特徴量を抽出し、元の NN の入力層と次元を揃えて大規模な構造へ NN を転用する。例えば、 $m \times m$ の構造で学習済みの NN を $2m \times 2m$ の構造に転用する場合、図 3 に示すように、 2×2 のグリッドパターンが 1 つのグリッドパターンに変換される。

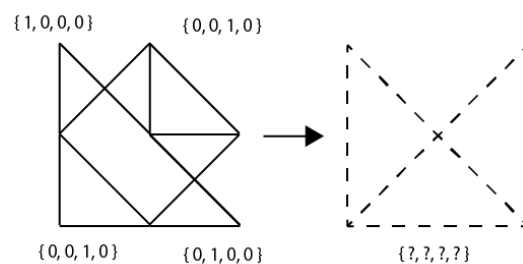


図 3. グリッドパターン抽出のイメージ

畳み込み層において、畳み込み処理を行うフィルターの種類をチャンネル、大きさをカーネルサイズ、フィルターをずらす幅をストライドと呼ぶ。今回は、4 つのパターンを各チャンネルに割り当て、入力データを (チャンネル数, $2m$, m) の 3 次元配列とする。図 4 に示すように、畳み込み層をストライド数 2、チャンネル数 $4 \rightarrow 4$ 、カーネルサ

イズ 2×2 と設定することで、データは $(4, 2m, m) \rightarrow (4, m, m/2)$ と変換され、1次元化することで既習モデルに転用できる。このとき、既習モデルのパラメータを固定し、畳み込み層のパラメータだけを学習で更新するため、大幅な学習コスト削減が可能である。なお、フィルターによるデータ変換は行わない。

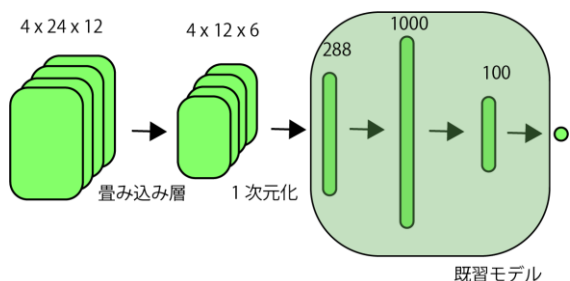


図4. $m = 12$ の構造の既習モデルを $m = 24$ の構造に転用する場合のネットワーク構造

$m = 12$ の構造で学習したモデルを $m = 24$ の構造に転用した場合の結果を表2に、 $m = 24$ の構造で学習したモデルを $m = 48$ の構造に転用した場合の結果を表3に示す。どちらの場合も50000個のテストデータで学習し、10000個のテストデータに対してMSEを計算する。

表2. $m = 24$ の構造のNNの学習

	直接学習	$m = 12$ の転移学習
MSE [mm ²]	2.592×10^{-4}	6.108×10^{-4}
学習時間 [s]	98.5	26.6
予測時間 [s]	5.98×10^{-2}	3.59×10^{-2}

表3 $m = 48$ の構造のNNの学習

	直接学習	$m = 24$ の転移学習
MSE [mm ²]	3.959×10^{-4}	6.839×10^{-4}
学習時間 [s]	376.1	80.2
予測時間 [s]	1.97×10^{-1}	8.48×10^{-2}

転移学習により学習時間は、 $m = 24$ の場合約73%、 $m = 24$ の場合約79%低減された。MSEは、 $m = 24$ の場合約2.36倍、 $m = 48$ の場合約1.73倍となった。予測精度の悪化は大きくなく、畳み込み層によるグリッドパターンの抽出は正しく機能していると考えられる。

畳み込み層を用いたNNによってテストデータの中で最もポアソン比が大きいと予測された構造について表4にまとめ、 $m = 24$ の場合について図5に、 $m = 48$ の場合について図6に示す。

表4. テストデータ中ポアソン比最大と予測された構造

m	予測値 [mm]	実際の値 [mm]	平均値 [mm]	標準偏差 [mm]
24	0.2053	0.1803	0.1615	0.0237
48	0.3673	0.3716	0.3287	0.0261

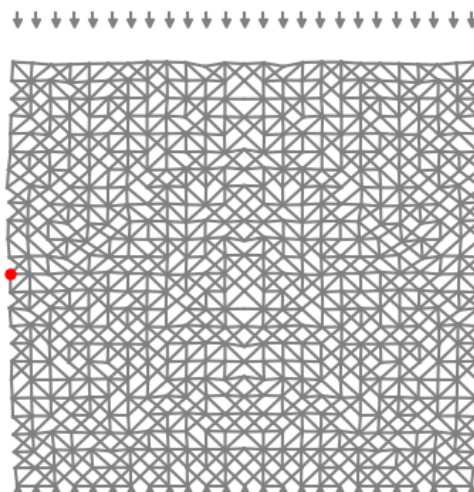


図5. テストデータ中ポアソン比最大と予測された構造 ($m = 24$)

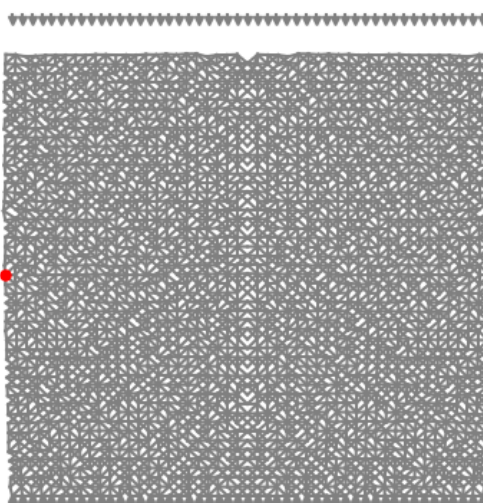


図6. テストデータ中ポアソン比最大と予測された構造 ($m = 48$)

5. 説明可能性の検討

5.1 Lime

説明可能性の検討に、Lime (Local Interpretable Model-Agnostic Explanations)⁴⁾を用いる。Limeは、ある解の予測に対する特徴量寄与度を得るために、その解の近傍解をサンプリングし局所的に忠実な説明可能モデルで学習する手法であり、任意の機械学習モデルに適用できる。

$f(\mathbf{x})$ を解釈したい予測関数として、説明モデル $\xi(\mathbf{x})$ は損失関数 L と、非零な係数の数などの関数の複雑さを表す関数 Ω の和を最小とするような説明可能な関数 g として式(1)

で決定される。ここで誤差関数 L は、 $\mathbf{x}$ 近傍でランダムにサンプリングした $\mathbf{z}, \mathbf{z}'$ における f と g の差に、 $\mathbf{x}$ に対するサンプリング解の類似度関数 π で重みづけして式(2)で定義され、 π は距離関数 D とカーネル指数 σ^2 を用いて式(3)で定義される。上記によって決定される説明可能モデル $\xi(\mathbf{x})$ によって予測に対する特徴の寄与が得られ、これは対象とする機械学習に局所的な線形性がある限り有効な手法である。

$$\xi(\mathbf{x}) = \operatorname{argmin}_{g \in G} L(f, g, \pi_x) + \Omega(g) \quad (1)$$

$$L(f, g, \pi_x) = \sum_{\mathbf{z}, \mathbf{z}'} \pi_x(\mathbf{z}) (f(\mathbf{z}) - g(\mathbf{z}'))^2 \quad (2)$$

$$\pi_x(\mathbf{z}) = \exp\left(-\frac{D(\mathbf{x}, \mathbf{z})^2}{\sigma^2}\right) \quad (3)$$

5.2 NN の予測に対する Lime の適用

本稿では、 $m = 12$ の構造のテストデータ中でポアソン比最大と予測された、図7に示す解について、Limeでその特徴を解釈する。この解の正確な目的関数値は0.2540 mm、NNの予測値は0.2648 mm、Limeの近似モデルの予測値は0.2337 mmであった。

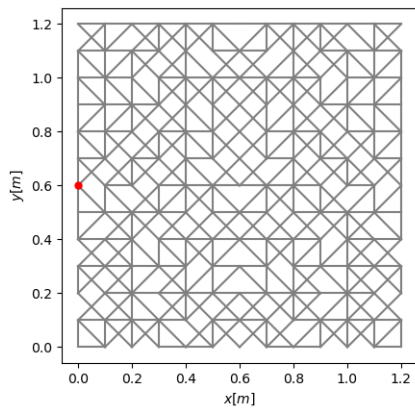


図7. Limeで説明する解

表5にLimeによって算出された寄与度の大きい特徴をまとめる。位置は、 2×2 の部分領域の中心座標を示す。グリッドパターンを、下辺削除:1, 右辺削除:2, 右斜めの辺削除:3, 左斜めの辺削除:4とし、パターンにその組み合わせを示す。また、寄与の向きとは実際の値が予測値の正負どちらに寄与しているかを示す。さらに、寄与度の大きい上位1000個の特徴を、特徴が示す部分領域ごとに分類し、特徴数が多い順に数え上げた結果を表6に示す。

表5から、位置(0.1,0.7)において横部材が存在しないグリッドが斜めに連続したパターンが最も予測に寄与していることが分かる。また、表6から、ターゲット節点に近い位置(0.1,0.6), (0.1,0.7)の部分領域のパターンの予測への寄与が大きいことが分かる。これは、極めて妥当な結果であり、NNが正しく特徴を解釈していることが確認

できる。

表5. ポアソン比最大の予測解に対する特徴量寄与の分析

位置 ([m], [m])	パターン	値	寄与度	寄与の向き
(0.1,0.7)	(1,2,3,1)	1	0.03336	正
(0.1,0.7)	(1,2,1,1)	0	0.02787	負
(0.1,0.7)	(1,2,2,1)	0	0.02599	負
(0.1,0.7)	(1,2,4,1)	0	0.02201	負
(0.1,0.6)	(4,1,1,2)	0	0.02188	負

表6. 部分領域ごとに集計した寄与の大きい特徴数

位置([m],[m])	特徴数
(0.1,0.6)	108
(0.1,0.7)	99
(0.5,0.6)	61
(0.5,0.7)	51
(0.3,0.6)	41

6. 結

本稿では、NNを用いた格子構造のトポロジー設計において、特徴選択や転移学習を用いたNNが、目的関数の大きい解を予測するために十分な精度を持つことを示した。また、NNの予測に対するLimeによる特徴量寄与の分析により、妥当性のある結果が得られた。

本稿では、ポアソン比をターゲット節点の水平変位により評価したが、周期的なユニットごとに設計を行ったうえで、より厳密にポアソン比を評価する方法も考えられる。しかし、その場合、機械学習による特徴量の認識が難しくなるため、より発展したアプローチが求められる。

[参考文献]

- 1) M. Ohsaki, Optimization of Finite Dimensional Structures, CRC Press, 2010.
- 2) F. D. Reis and N. Karathanasopoulos, Inverse metamaterial design combining genetic algorithms with asymptotic homogenization schemes, Int. J. Solids and Structures, Vol. 250, Paper No. 111702, 2022.
- 3) T. Tamura, M. Ohsaki and J. Takagi, Machine learning for combinatorial optimization of brace placement of steel frames, Japan Architectural Review, Vol. 1(4), pp. 419-430, 2018.
- 4) M. T. Ribeiro, S. Singh and C. Guestrin, Why Should I Trust You?: Explaining the Predictions of Any Classifier, Association for Computing Machinery, New York, 2016.
- 5) S. Lundberg and S.-I. Lee, A unified approach to interpreting model predictions, Advances in Neural Information Processing Systems, Vol. 30, 2017.
- 6) Diederik P. Kingma, Jimmy Ba, Adam: A Method for Stochastic Optimization, arXiv preprint arXiv:1412.6980, 2014.
- 7) 松岡知哉, 大崎純, 林和希, 特徴選択を用いた機械学習モデルによる格子構造の最適化, 日本建築学会大会学術講演梗概集(近畿), 20166, 2023