

機械学習を用いたキャップレート予測モデルの活用可能性

2025 年 7 月 30 日

株式会社三井住友トラスト基礎研究所

私募投資顧問部 上席主任研究員 米倉勝弘

- ◆ キャップレート(Cap Rate)は、不動産投資における主要な指標の一つであり、資産価格の評価にとどまらず、取引市況の動向把握や類似性の高い物件間におけるリスクとリターンのバランス判断にも資する指標である。
- ◆ 本稿では、「重回帰分析」、「ランダムフォレスト」、「XGBoost」の3手法を用いて予測モデルを構築し、それぞれの予測精度を比較検討することで、不動産取引市場におけるキャップレート予測に適した分析手法を模索した。
- ◆ 各モデルの予測精度については、テストデータにおける RMSE (Root Mean Squared Error) を用いて評価した。その結果、最も予測誤差が小さかったモデルは「XGBoost」となったものの、3手法間に顕著な差は確認されず、各モデルとも概ね同程度の予測精度を示した。
- ◆ ただし、今後、サンプル数の増加や属性情報(特徴量)の拡充が進めば、これら機械学習モデルの予測精度はさらに向上する可能性がある。
- ◆ 「ランダムフォレスト」および「XGBoost」は、欠損値の存在を前提とした内部処理機構を備えており、欠損値を含むデータに対しても比較的安定した予測性能を維持できる。一方で、属性情報の拡充については、Web スクレイピング等の技術を用いた自動収集に加え、人手による補足的なデータ収集も依然として必要とされる。不動産業界においては、情報の非公開性やデータの分散性が高く、これらが機械学習の導入を困難にしている要因の一つである。こうした課題を克服し、より網羅的かつ高品質なデータセットの整備が進めば、機械学習モデルの有効性は一層高まることが期待される。
- ◆ もっとも、実務の現場では、モデルの可読性や解釈性が依然として重視される傾向にある。「ランダムフォレスト」および「XGBoost」も特徴量における影響の強さ(重要度)は把握可能であるが、たとえば、「駅からの距離が遠いほど Cap Rate が上昇する」といった変数間の統計的関係を直感的に捉えることが可能な点において、「重回帰分析」は依然として高い支持を得ている。このような背景を踏まえると、機械学習モデルの導入に際しては、まず「重回帰分析」を主軸に据えたうえで、その予測精度および変数の寄与度を検証・補完する目的で機械学習モデルを併用する、いわゆるシャドーモードでの運用が、実務への円滑な導入を図るうえで有効なアプローチであると考えられる。

1.はじめに

キャップレート(Cap Rate)は、不動産投資における投資判断の主要な指標の一つであり、対象不動産の純収益(NOI または NCF)をもとに資産価格を算出する際に広く用いられている。加えて、資産価格の評価にとどまらず、取引市況の動向把握や類似性の高い物件間におけるリスクとリターンのバランス判断にも資する指標である。

キャップレートは、取引時点や立地条件、建物特性、さらには用途制限・法令上の制限等の行政的要因を含む複数の価格形成要因により水準が変動する。そのため、不動産投資分析においては、不動産取引市場で観測された実際の取引キャップレートをもとに価格形成要因との関係を考慮した予測モデルを構築し、当該モデルにより

算出された品質補正済みのキャップレートを活用することが一般的である。

キャップレートの予測モデルには、線形回帰、とりわけ「重回帰分析」が広く用いられているが、キャップレートと価格形成要因との関係には一部で非線形性が認められ、価格形成要因間に複雑な相互作用が含まれていることから、これらを捉えることができる決定木ベースのアンサンブル学習（複数の機械学習モデルを組み合わせ、単一のモデルよりも高い予測精度を目指す手法）が有効となる可能性がある。特に「ランダムフォレスト」および「XGBoost」は、非線形関係や特徴量（説明変数）間の相互作用を自動的に学習可能なモデルとして注目されている。

そこで本稿では、「重回帰分析」、「ランダムフォレスト」、「XGBoost」の3手法を用いて予測モデルを構築し、それぞれの予測精度を比較検討することで、不動産取引市場のキャップレート予測に適した分析手法を模索した。

2. 分析データおよび分析モデル

本稿では、2001年以降に全国で観測されたオフィスビルの取引キャップレート約1,700事例をもとに、3 σ 法等の統計的手法により外れ値を除去したデータセットを分析対象とした。目的変数として取引キャップレートを用い、特徴量には、「時点要因」「市況要因」「行政的要因」「立地要因」「建物要因」に関する各種データを採用した。

分析手法としては、前述のとおり、従来型の線形モデルである「重回帰分析」に加え、非線形性および特徴量間の相互作用を捉える能力に優れるアンサンブル学習モデルとして「ランダムフォレスト」および「XGBoost」を選定した（各手法の内容および特徴については、＜Appendix＞参照）。モデルの性能を正しく評価するためには、モデルが学習に使用したデータとは異なるデータに対してどれだけ正確に予測できるか、すなわち汎化能力を測定する必要がある。このため、データセットを「訓練データ（学習用）」と「テストデータ（評価用）」に7:3で分割して分析を行った。

なお、「ランダムフォレスト」および「XGBoost」は、決定木に基づくアンサンブル学習モデルであり、外れ値に対して一定のロバスト性（モデルが外れ値、ノイズ、未知のデータ、分布の変化等に対し、安定して良い性能を保つ性質）を有することから、厳密には外れ値の除去を必須としない手法である。しかし、本稿ではモデル間の比較においてサンプルデータの一貫性を担保するため、3手法のモデルにおいて共通のデータセット（外れ値除去後）を用いることとした。

2-1 重回帰分析 (Multiple Linear Regression)

「重回帰分析」における予測結果の概要は図表1に示したとおりである。モデルの適合度を示す自由度調整済み決定係数 (Adjusted R-squared) は0.653であり、特徴量によって目的変数を一定程度説明できていることから本モデルは概ね妥当なものと評価できる。

また、回帰係数の符号はすべて事前に設定した理論的仮説と整合しており、各特徴量に対するt検定を実施した結果、5%有意水準において統計的に有意ではないと判断された特徴量は1つのみであった。

図表 1_モデルの予測結果

特徴量	取得時 キャップレート
定数項	(+) ***
時点要因	(+) ***
市況要因	(-) ***
立地要因(大分類)	(-) ***
行政的要因	(-) ***
建物要因Ⅰ	(+) ***
建物要因Ⅱ	(-) ***
建物要因Ⅲ	(+) **
立地要因(小分類)Ⅰ	(+) ***
立地要因(小分類)Ⅱ	(+)
立地要因(小分類)Ⅲ	(-) *
Adjusted R-squared	0.653

注:カッコ内は+および-は各係数の符号を表す。

＋の場合、特徴量が増加すると、目的変数も増加する。

－の場合、特徴量が増加すると、目的変数は減少する。

***は0.1%有意水準、**は1%有意水準、*は5%有意水準を表す。

Adjusted R-squaredは自由度調整済み決定係数を表す。

出所)三井住友トラスト基礎研究所

2-2 ランダムフォレスト(Random Forest)

「ランダムフォレスト」では決定木の最大深さを4に設定し、決定木(弱学習器)の数を1から300まで順次増加させることで、モデルの性能を評価した。図表2に示したように、決定木の数が150を超えるあたりから、テストデータに対するRMSE(Root Mean Squared Error: テストデータにおける実績値と予測値の平方平均二乗誤差)が安定してきており、モデルの汎化性能が向上していることが確認できる。

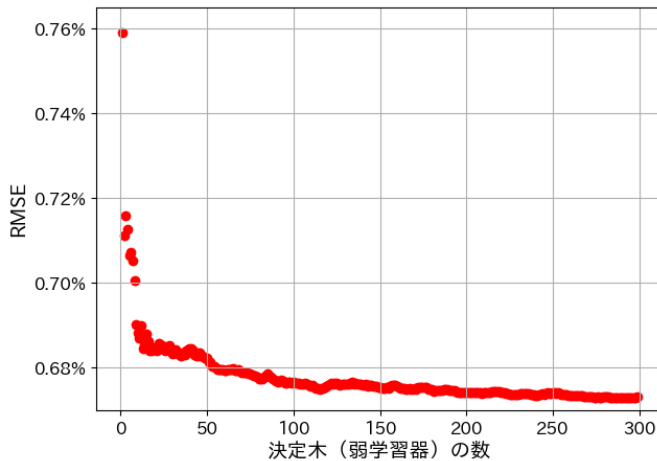
図表3は、本モデルにおける特徴量重要度(Feature Importance)ⁱⁱを示している。本稿では、各特徴量が予測にどれだけ貢献しているかを定量的に示す指標として、「不純度の減少(Impurity Decrease)」をもとに計算している。本モデルにおいては、「時点要因」と「市況要因」が最も重要な特徴量であることが分かる。一方で、「立地要因(小分類)」の重要度が低くなっている。定性的な解釈とはなるが、「立地要因(大分類)」が「立地要因(小分類)」を内包している可能性が考えられる。また、「行政的要因」は建物の規模などに影響を与えるため、「行政的要因」が「建物要因」の一部を代理する役割を果たしている可能性も考えられる。

i 木構造における根ノードから最も遠い葉ノードまでの分岐の数。

ii 各決定木において、ある特徴量がノード分割に使われた回数を記録し、その分割によって得られた不純度の減少量(平均二乗誤差の減少量)を加算。すべての木における貢献度を合計し、全特徴量について正規化する。

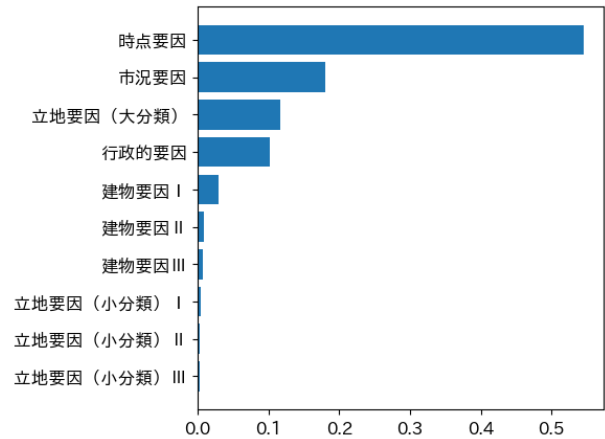
iii 決定木がノードを分割する際に用いるノード内のデータの「混ざり具合」や「ばらつき」の指標。ノードを分割するときに、どこで分割するとより純粋なグループ(似たデータのまとまり)になるかを評価する際に指標として用いる。

図表 2_決定木増加による RMSE の減少効果



出所)三井住友トラスト基礎研究所

図表 3_特徴量重要度 (不純度減少)



出所)三井住友トラスト基礎研究所

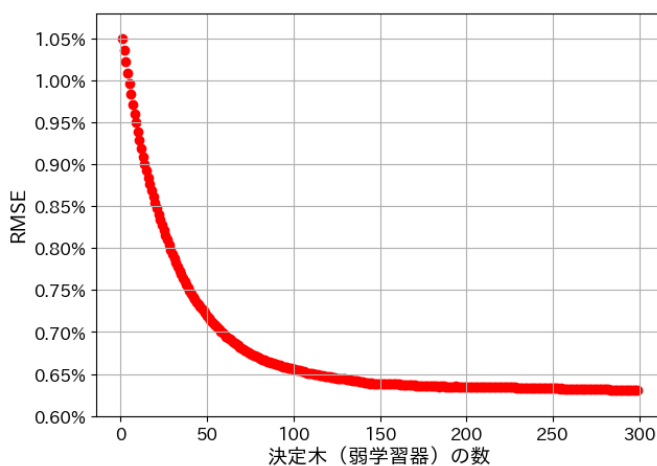
2-3 XGBoost(eXtreme Gradient Boosting)

「XGBoost」においてもランダムフォレストモデルと同様に決定木の最大深さを 4 に設定し、決定木 (弱学習器) の数を 1 から 300 まで順次増加させることで、モデルの性能を評価した。図表 4 に示したように、決定木の数が 150 を超えるあたりから、テストデータに対する RMSE が安定してきており、モデルの汎化性能が向上していることが確認できる。

図表 5 は、本モデルにおける特徴量重要度を示している。「XGBoost」では主に「分割頻度 (Weight)」、「情報利得 (Gain)」、「カバー率 (Cover)」などの指標に基づいて特徴量の寄与度を計算するが、本稿では各特徴量が予測にどれだけ貢献しているかを定量的に示す指標として、「分割頻度」を指標に採用している。これは各特徴量が決定木の分割に使用された回数に基づくシンプルで直感的な重要度指標である。

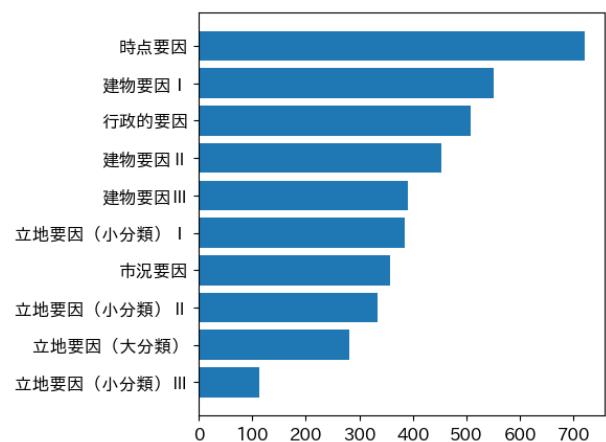
本モデルでは、「時点要因」、「建物要因」、「行政的要因」、「立地要因」の特徴量がバランスよく予測に貢献していることが確認できた。

図表 4_決定木増加による RMSE の減少効果



出所)三井住友トラスト基礎研究所

図表 5_特徴量重要度 (頻度)



出所)三井住友トラスト基礎研究所

3.分析結果

本稿では、「重回帰分析」、「ランダムフォレスト」、および「XGBoost」の3手法を用いて予測モデルを構築し、それぞれの予測精度をテストデータに対するRMSEをもとに評価した。なお、キャップレート(%)を目的変数として使用したため、RMSEも同様に%単位で表記している。評価結果としてRMSE(テストデータにおける実績値と予測値の誤差を示す指標であることから値が小さいほど予測精度が高い)は、「重回帰分析」で0.643%、「ランダムフォレスト」で0.659%、「XGBoost」で0.642%となった。「XGBoost」の予測精度が最も高かったものの、3つの予測モデル間に顕著な差は確認されず、各モデルとも概ね同程度の予測精度を示す結果となった。

精度差に顕著な違いがみられなかった理由としては、取引キャップレートと各特徴量の関係がほぼ線形であったことが挙げられる。今回採用した特徴量は非線形性が認められる特徴量であっても、対数変換等により線形近似が可能な特徴量であったことから、非線形性や特徴量間の複雑な相互作用を捉えることが得意な「ランダムフォレスト」や「XGBoost」の強みが十分に発揮されなかったものと考えられる。

また、目的変数との間に非線形性を持つ特徴量や複雑な相互作用を持つ特徴量が欠落変数(本来考慮すべきなのにデータとして含まれていない変数)となっている可能性があり、そのため「重回帰分析」との予測制度の差違が小さくなっていることも考えられる。なお、「ランダムフォレスト」や「XGBoost」は、一般的にサンプル数が増えるほど過学習リスクを減らし、モデルの汎化性能が向上するため、当該モデルの精度向上にはサンプルの増加も重要な要素である。

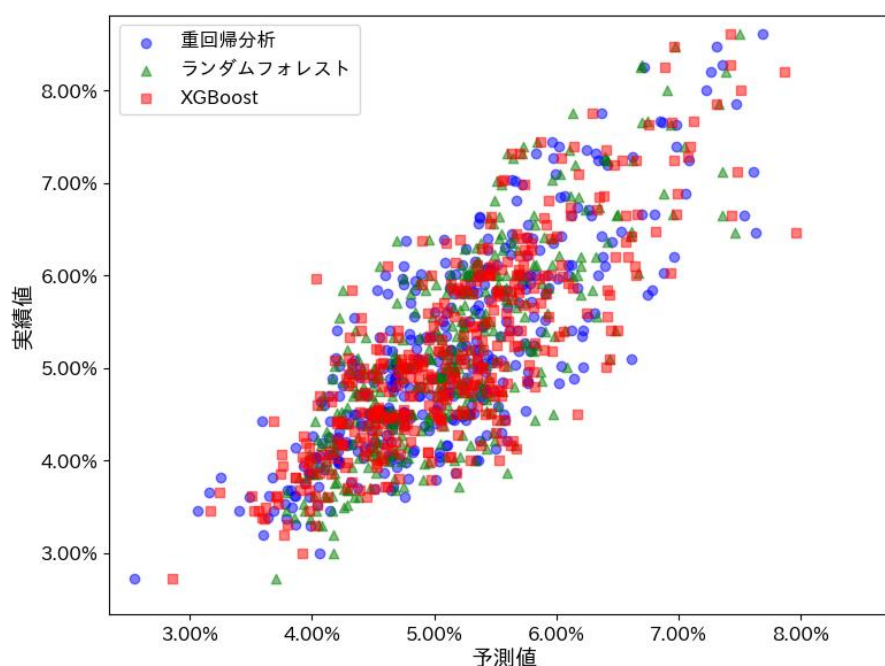
図表 6_各モデルにおける RMSE の結果

	重回帰分析	ランダムフォレスト	XGBoost
RMSE(テストデータ)	0.643%	0.659%	0.642%

出所)三井住友トラスト基礎研究所作成

注)RMSE はモデル作成には使用していないテストデータにおける平均平方二乗誤差

図表 7_実績値と予測値の比較



出所)三井住友トラスト基礎研究所作成

4.機械学習を用いたキャップレート予測モデルの活用可能性

本稿では、「重回帰分析」、「ランダムフォレスト」、「XGBoost」の3手法を用いて予測モデルの構築し、各モデルの予測精度を比較検討した。その結果、いずれのモデルにおいても予測精度に顕著な差異は確認されなかった。しかしながら、この結果は、本稿で用いたサンプルデータの質および量においても、決定木およびアンサンブル学習に基づく機械学習モデルが一定の予測性能を発揮し得ることを示唆している。

今後、サンプル数の増加や属性情報(特微量)の拡充が進めば、これら機械学習モデルの予測精度はさらに向上する可能性が高い。とりわけ、サンプル数の増加という観点においては、従来であれば分析対象から除外せざるを得なかった欠損値を含むデータの有効活用が鍵となる。

不動産関連の実データにおいては、情報取得コストの高さや契約内容の非開示といった実務上の制約により、欠損値を含むケースが少なくない。「重回帰分析」では、こうした欠損値の存在に対して脆弱であり、欠損値補完手法の選択がモデル精度に大きく影響するため、場合によってはデータを一括して除外せざるを得ないこともある。これに対し、「ランダムフォレスト」および「XGBoost」は、欠損値の存在を前提とした内部処理機構を備えており、欠損値を含むデータに対しても比較的安定した予測性能を維持できる。また、「重回帰分析」と比較して、外れ値やノイズの影響を受けにくい構造(高いロバスト性)を有している点も特筆すべきである。

加えて、機械学習モデルの利点として、各特微量が予測に与える影響度(重要度)を数値的に可視化できる点が挙げられる。たとえば、「ランダムフォレスト」や「XGBoost」では、各変数の寄与度を「特微量重要度」として出力でき、これにより実務上も関心の高い特微量における影響の強さ(重要度)の把握が可能となる。また、SHAP (SHapley Additive exPlanations)などの補完的な解釈手法を用いることで、モデルの透明性や説明性を高めることもできる。これは、「重回帰分析」と同様、どの要因が結果に影響を及ぼしているのかを分析したいという実務上の要求に十分に応え得るものであり、機械学習モデルの現場適用を後押しする要素の一つといえる。

一方で、属性情報の拡充については、Webスクレイピング等の技術を用いた自動収集に加え、人手による補足的なデータ収集も依然として必要とされる。不動産業界においては、情報の非公開性やデータの分散性が高く、これらが機械学習の導入を困難にしている要因の一つである。こうした課題を克服し、より網羅的かつ高品質なデータセットの整備が進めば、機械学習モデルの有効性は一層高まることが期待される。

もっとも、実務の現場では、モデルの可読性や解釈性が依然として重視される傾向にある。たとえば、「駅からの距離が遠いほどCap Rateが上昇する」といった変数間の統計的関係を直感的に捉えることが可能な点において、「重回帰分析」は依然として高い支持を得ている。このような背景を踏まえると、機械学習モデルの導入に際しては、まず「重回帰分析」を主軸に据えたうえで、その予測精度および変数の寄与度を検証・補完する目的で機械学習モデルを併用する、いわゆるシャドーモードでの運用が、実務への円滑な導入を図るうえで有効なアプローチであると考えられる。

<Appendix>

重回帰分析(Multiple Linear Regression)

「重回帰分析」は、式(1)のように複数の特徴量(説明変数)を用いて、1つの目的変数を予測・説明するための基本的な統計手法である。単回帰分析の拡張形にあたり、各特徴量が目的変数に対して線形な関係性を持つと仮定する点が特徴である。

$$y = Xw + \varepsilon \quad (1)$$

y : キャップレートの実績値

w : 重みベクトル(回帰係数)

X : サンプル数×キャップレート形成要因数(定数項を含む特徴量)の行列

ε : 誤差項

「重回帰分析」は、外れ値の存在や特徴量間に高い相関関係(多重共線性)が認められる場合、推定結果の信頼性や回帰係数の解釈性が低下する傾向にある。一方で、適切な前処理や変数選択を行うことにより、回帰係数を通じて各特徴量の目的変数に対する影響の方向性および大きさを解釈可能であり、変数間の統計的関係の分析に適した手法とされる。実務上、機械学習を含むデータ分析の結果に接する機会が必ずしも多くない不動産実務者等にとっても感覚的に理解されやすいという特徴がある。

本稿における「重回帰分析」は、Python のライブラリである statsmodels (<https://www.statsmodels.org/>) を用いて実施した。なお、分析にあたっては、特徴量の標準化^{iv}やモデルへの正則化項^v(L1 または L2 など)は導入していない。

ランダムフォレスト(Random Forest)

「ランダムフォレスト」は、式(2)のように複数の決定木を構築し、それらの予測結果を統合するアンサンブル学習モデルである。個々の決定木は、単独では過学習しやすく分散が大きい弱学習器であるが、複数の木を組み合わせることで、予測の安定性と汎化性能を高めることが可能となる。

本手法では、学習データからブートストラップサンプリング^{vi}によって複数のサブセットを作成し、それぞれに対して独立した決定木を学習させる。また、各ノード^{vii}の分割においては、全特徴量の中からランダムに抽出された一部の特徴量に限定して最適な分割を行う。このようなランダム性の導入により、個々の木の相関を低減し、アンサンブル全体の性能向上につなげている。

「ランダムフォレスト」は明示的な正則化項(L1 または L2 など)を導入しないが、式(4)のように木の深さや葉ノードの最小サンプル数などのハイパーパラメータ^{viii}を設定を通じて、構造的な正則化を行っている。この結果、過学習のリスクを抑えつつ、柔軟かつ安定的な予測モデルを構築することが可能である。

モデルは以下のような形で表現される。

iv 各特徴量の値を平均 0、標準偏差 1 に変換する前処理の手法。重回帰分析においては、これにより回帰係数の大きさによって変数の重要性を比較できる利点がある。

v モデルの損失関数にペナルティを加える手法。過学習を抑えるために有効な手法で、未知のデータに対しての予測精度を上げることに用いられる。

vi ブートストラップ法は、観測されたサンプルデータから復元抽出(抽出したデータを元に戻して、再び同じデータが選ばれる可能性を許す抽出方法)によって多数の擬似データセット(リサンプル)を作成し、統計量の分布を推定する再標本化手法。

vii 幹から枝分かれする分岐点や葉に分かれていく分岐点・終着点(結節点)を指す。

viii モデルの学習前に設定する調整可能なパラメータ。

$$\hat{y}_i = \frac{1}{M} \sum_{j=1}^M T_j(x_i) \quad (i = 1, 2, 3, \dots, n) \quad (2)$$

$$\mathcal{L}(y_i, \hat{y}_i) = (y_i - \hat{y}_i)^2 \quad (3)$$

$$\Omega(T_j) = \lambda \cdot \text{Depth}(T_j) \quad (4)$$

$$L(\phi) = \sum_{i=1}^n \mathcal{L}(y_i, \hat{y}_i) + \sum_{j=1}^M \Omega(T_j) \quad (5)$$

y : キャブプレートの実績値

$\hat{y}$: キャブプレートの予測値

x : 入力値

M : 決定木の数

$T_j(x_i)$: j 番目の決定木がサンプル x_i に対して出す予測値

$\mathcal{L}$: 損失関数

Ω : 正則関数(「木の深さ: 数式(4)」や「特徴量選択」など、モデルの複雑さを制御する決定木の正則化項)

λ : 重みの正則化パラメータ

L : 目的関数(損失関数および正則化項の組み合わせ)

n : サンプル数

「ランダムフォレスト」は、各特徴量が予測に与える寄与度を測定することで特徴量重要度を定量的に評価することが可能である。しかしながら、アンサンブル学習モデルとしての性質により、「ランダムフォレスト」はブラックボックスな側面を持ち、個々の予測結果に対する定量的な解釈が困難である場合がある。この点において、実務上、機械学習を含むデータ分析の結果に接する機会が必ずしも多くない不動産実務者等に対する予測結果の可視化や解釈の難易度が課題となることがある。

一方で、「ランダムフォレスト」は外れ値や特徴量間の相関に対して高い頑健性を持ち、ロバスト性に優れた分析手法である。データに外れ値が含まれている場合でも、モデルの予測精度が大きく損なわれることは少なく、ノイズや外れ値が学習に与える影響を抑えることができる。また、「ランダムフォレスト」は欠損値に対してもある程度頑健であり、欠損値の推定および再推定を行うことで欠損値を含むデータに対して適切に対応できる。

さらに、「ランダムフォレスト」は非線形性や特徴量間の複雑な相互作用を自動的に捉える能力があり、「重回帰分析」と比較しても高い予測性能を示す場合が多い。特に、予測対象と特徴量との間に非線形な関係や相互作用が存在する場合、「ランダムフォレスト」はこれらを効果的に学習するため、より精度の高い予測が可能となる。

本稿では、Python のライブラリである scikit-learn (<https://scikit-learn.org/>) を使用してモデルを構築した。

XGBoost(eXtreme Gradient Boosting)

「XGBoost」は、勾配ブースティング (Gradient Boosting)^{ix} アルゴリズムの一種であり、従来のブースティング手法を高速化し、さらに高精度な予測を可能にする機械学習モデルである。「XGBoost」の特徴的な点は、逐次的に構築される決定木(弱学習器)を用いて、誤差を段階的に補正しながら学習を進めることである。この手法は、各新しい決定木が前回の決定木の誤差を修正する形でモデルを改善していく。

また、「XGBoost」では、他の決定木系手法とは異なり、損失関数に L1 正則化 (Lasso) および L2 正則化

ix 勾配ブースティングは、式(6)のように複数の弱学習器(浅い決定木)を逐次的に学習・加算することで、予測精度の高いアンサンブルモデルを構築する手法である。

(Ridge)の項を明示的に組み込んでいる(式(8)第2項では「L2正則化:ユークリッド距離」のみ反映)。L1正則化(Lasso)は特徴量選択を促進し、L2正則化(Ridge)はモデルの重みの大きさを抑制することによって、より簡潔で解釈可能なモデルの構築に寄与する。この正則化手法の導入により、「XGBoost」はモデルの複雑性を適切にコントロールし、過剰適合(オーバーフィッティング)のリスクを低減させる。

さらに、「XGBoost」は、式(8)第1項のように木の複雑度(ノード数や葉の重み等)を罰則項として損失関数に組み込み、より単純でかつ効率的なモデルを構築する。このアプローチにより、高次元かつ大量のデータを扱う場合でも計算資源を効率的に利用しながら高精度の予測を実現することが可能となっている。

$$\hat{y}_i = \sum_{k=1}^K f_k(x_i) \quad (i = 1, 2, 3, \dots, n) \quad (6)$$

$$\mathcal{L}(y_i, \hat{y}_i) = (y_i - \hat{y}_i)^2 \quad (7)$$

$$\Omega(f) = \gamma T + \frac{1}{2} \lambda \|w\|^2 \quad (8)$$

$$L(\phi) = \sum_{i=1}^n \mathcal{L}(y_i, \hat{y}_i) + \sum_{k=1}^K \Omega(f_k) \quad (9)$$

y : キャブプレートの実績値

$\hat{y}$: キャブプレートの予測値

x : 入力値

K : 決定木数

f_k : 決定木 K に対する重み

$\mathcal{L}$: 損失関数

Ω : 正則関数

γ : 木の複雑さに対するペナルティのパラメータ

T : 葉ノードの数

λ : 重みの正則化パラメータ

w : 重み

L : 目的関数(損失関数および正則化項の組み合わせ)

n : サンプル数

「XGBoost」は、アンサンブル学習に基づくため、モデル全体がブラックボックスとなりがちで、個々の予測に対する定量的な説明が難しいことがある。一方で、外れ値に対して比較的ロバスト性が高く、ノイズに強い分析モデルとして知られており、情報利得(Gain)や分割頻度(Frequency)をもとに特徴量の寄与度を可視化することも可能である。ただし、「ランダムフォレスト」に比べると外れ値への耐性がやや劣るため、ロバスト損失関数の使用が必要となる場合がある。また、特徴量間に高い相関が存在する場合、冗長性の問題が生じやすく、主成分分析(PCA)のような手法で冗長性を削減する工夫が求められる場合がある。

さらに、「XGBoost」は決定木の分岐において欠損値をどの分岐先に送るかを自動的に学習するため、欠損値を含むデータに対して適切に対応できる特徴があり、他の手法と比較して優位性が認められる。

「XGBoost」は、「ランダムフォレスト」と同様に、非線形性や特徴量間の複雑な相互作用を自動的に捉えることができ、「重回帰分析」と比較しても高い予測精度を示す場合が多い。

本稿では、Python のライブラリである scikit-learn (<https://scikit-learn.org/>) を使用してモデルを構築した。

【本件のお問い合わせ先】

私募投資顧問部 担当: 米倉

TEL: 080-7207-5117(直)

<https://fofa.jp/smtri/a.p/116/>**株式会社三井住友トラスト基礎研究所**

〒105-8574 東京都港区芝 3-33-1 三井住友信託銀行芝ビル 11 階

<https://www.smtri.jp/>

1. この書類を含め、当社が提供する資料類は、情報の提供を唯一の目的としたものであり、不動産および金融商品を含む商品、サービスまたは権利の販売その他の取引の申込み、勧誘、あっ旋、媒介等を目的としたものではありません。銘柄等の選択、投資判断の最終決定、またはこの書類のご利用に際しては、お客さまご自身でご判断くださいますようお願いいたします。
2. この書類を含め、当社が提供する資料類は、信頼できると考えられる情報に基づいて作成していますが、当社はその正確性および完全性に関して責任を負うものではありません。また、本資料は作成時点または調査時点において入手可能な情報等に基づいて作成されたものであり、ここに示したすべての内容は、作成日における判断を示したものです。また、今後の見通し、予測、推計等は将来を保証するものではありません。本資料の内容は、予告なく変更される場合があります。当社は、本資料の論旨と一致しない他の資料を公表している、あるいは今後公表する場合があります。
3. この資料の権利は当社に帰属しております。当社の事前の了承なく、その目的や方法の如何を問わず、本資料の全部または一部を複製・転載・改変等してご使用されないようお願いいたします。
4. 当社は不動産鑑定業者ではなく、不動産等について鑑定評価書を作成、交付することはありません。当社は不動産投資顧問業者または金融商品取引業者として、投資対象商品の価値または価値の分析に基づく投資判断に関する助言業務を行います。当社は助言業務を遂行する過程で、不動産等について資産価値を算出する場合があります。しかし、この資産価値の算出は、当社の助言業務遂行上の必要に応じて行うものであり、ひとつの金額表示は行わず、複数、幅、分布等により表示いたします。