

金融情報学研究会(第26回)

日時 2021 年 3 月 6 日(土)

会場 オンライン

SIG-FiN
JSAI Special Interest Group on
Financial Informatics

人工知能学会
金融情報学研究会

第26回 人工知能学会 金融情報学研究会 (SIG-FIN)

2021 年 3 月 6 日 (土) オンライン

01. LDA で求めた新聞記事のトピック割合と株価指数の関係性	1
竹澤譲, 伊藤真利子, 大西立顕(立教大学)	
02. BERT を用いた有価証券報告書からの ESG 関連文抽出	9
土橋諒太, 中田和秀(東京工業大学)	
03. キーワード検索数とツイートの情報を用いたビットコイン価格の騰落予測	16
松田周也, 原尚幸(同志社大学)	
04. FinMegatron: Large Financial Domain Language Models	22
Xianchao Wu (NVIDIA)	
05. 株価の残差リターンに注目した深層学習ポートフォリオ最適化	26
今城健太郎, 南賢太郎, 伊藤克哉(Preferred Networks), 中川慧(野村アセットマネジメント)	
06. 状態空間モデルを用いたバランス型投資信託のリバランス推定	32
今井崇公(野村アセットマネジメント)	
07. シュレーディンガー・リスクパリティポートフォリオ	40
内山祐介(MAZIN), 中川慧(野村アセットマネジメント)	
08. 人工市場を用いたリポートが取引シェアに与える影響調査	46
星野真広(神奈川工科大学), 水田孝信(スパークスアセットマネジメント), 八木勲(神奈川工科大学)	
09. 金融データにおける VC 関連の応用	53
落合友四郎(大妻女子大学), ホセ ナチエル(東邦大学)	
10. 非同期時系列の Lead-lag 効果推定のための新しい推定量	56
伊藤克哉(Preferred Networks), 中川慧(野村アセットマネジメント)	
11. 予約価格を用いた指値・成行注文選択による約定寿命の冪分布の再現	64
吉村勇志, 陳昱(東京大学)	
12. 高頻度注文情報を用いたデータの層化と多段階事前学習による株価動向予測	69
松原冬樹, 和泉潔, 坂地泰紀(東京大学)	
13. Restricted Genetic Network Programming によるリスク回避型外国為替取引戦略の構築	77
内田純平, 穴田一(東京都市大学)	
14. テンソルを用いたマルチモーダル学習による株式価格の変動予測	82
笹尾知広, 中田和秀(東京工業大学)	

15. 効率的な Deep Hedging のためのニューラルネットワーク構造	88
今木翔太(東京大学), 今城健太郎, 伊藤克哉, 南賢太郎(Preferred Networks), 中川慧(野村アセットマネジメント)	

LDA で求めた新聞記事のトピック割合と株価指数の関係性

Correlation between topics in newspaper articles calculated by LDA and stock index

竹澤 譲¹ 伊藤 真利子¹ 大西 立顕¹

Jo Takezawa¹, Mariko I. Ito¹, and Takaaki Ohnishi¹

¹立教大学 大学院人工知能科学研究科

¹Graduate School of Artificial Intelligence and Science, Rikkyo University

概要: 社会の諸事象を伝える新聞記事のテキストから経済の変動の兆候を捕捉できないか探るため、記事のトピックの割合と株価指数の相関を分析した。まず、国内一紙の2003年から2012年までの10年間の77,814記事の単語を潜在的ディリクレ配分法(LDA)によって30のトピックに分類して、日ごとの各トピックが占める割合を求めた。次に、そのトピック割合と34種の東証株価指数の収益率及びボラティリティの時系列データとの間の順位相関を計算した。結果として、トピックの割合と株価指数の収益率との間に相関は見られなかったが、一部のトピックと株価指数のボラティリティとの間には時期に依存した有意な相関があることを発見した。

Abstract: We analyzed the correlation between topics in newspaper articles and stock indices to see if we could capture signs of economic changes from the text of the articles that report on various social events. First, we classified the words of 77,814 articles from a single Japanese newspaper over a 10-year period from 2003 to 2012 into 30 topics, and calculated the distribution of each topic on each day by Latent Dirichlet Allocation (LDA). Next, we calculated the rank correlation between the topics and the profitability and volatility of the 34 Tokyo Stock Exchange indices. As a result, we found that although there was no correlation between the topics and the profitability of stock indices, there was a significant time-dependent correlation between some topics and the volatility of the stock indices.

1. はじめに

政治、経済、その他社会の諸事象の報道や論評を速やかに伝えるための媒体である新聞記事は株価の予測に役立つ。過去及び過去を踏まえた未来の収益に関する予測情報は即時に現在の株価に織り込まれる、と考える効率的市場仮説[1]を支持するならば否定される。新聞記事の情報は半日から一日前のもので、新規性がないからである。

もっとも、情報が即時に株価に織り込まれるという点については議論の余地がある。未来の株価暴落につながる情報が顕在化していても、株価に織り込まれていないとは言い難い暴落や市場の混乱が生じた場面がしばしばあったからである。リーマンブラザーズ社の倒産をきっかけとして、金融機関に対する信用不安と金融収縮の連想が生じて、世界的に株安が連鎖したリーマンショックがその一場面である。

図1は、2007年から2年間で、リーマンブラザーズ社が破綻する原因となった”subprime”の文字が読売新聞英字版(以下、読売新聞という)で出現した

回数を表したものである。リーマンブラザーズ社が倒産した2008年9月15日の前月までに”subprime”の文字は135回出現した。記事の数が多かっただけでなく、内容としても、2007年8月19日の時点で、世界経済の方向性を決める中心的な役割を担うFRBが”downside risks to growth have increased appreciably.”と表明したことを報じるなど、先立つ危機を予期するものが多かった。

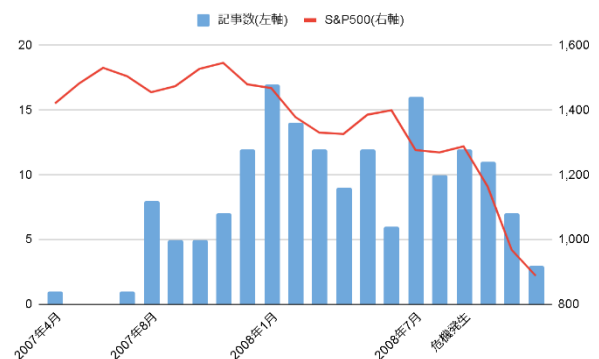


図1: 読売新聞「subprime」出現頻度とS&P 価格

しかし、その危機感が株価に織り込まれていたとは評価し難い。2008 年 9 月に最高値 1,274 ドルであった S&P500 は、半年後の 2009 年 3 月に最安値 667 ドルまで下落し、FRB 議長をして「100 年に 1 度の津波」と言わしめるほど大きな混乱を来した。これは情報が即時に株価に反映されるわけではないことを表す一事例である。

そして、社会の情報が市場に伝わってから時間をかけて株価に反映されていくのだとすれば、過去の情報を掲載した新聞記事も株価の予測に役立つ可能性がある。本稿が分析の対象とした読売新聞は一般紙に当たる。一般紙は、経済のみならず政治、国際問題、文化、スポーツなど多様なトピックを掲載する社会のミニモデルとも言える媒体で、社会の動向に合わせて紙面の構成が変わる。新聞記事を占めるトピックの割合の遷移過程を追うことで、経済変動の兆候を捕捉できる可能性がある。

そのような視点に立って、本稿では新聞記事のテキストを潜在的ディリクレ配分法 (Latent Dirichlet Allocation, 以下 LDA という) により分類し、日ごとの各トピックが紙面を占める割合を求め、株価指数の収益率やボラティリティとの相関を計算した。

2. 関連研究

2-1. テキストの極性から予測を行った研究

金融ニュースのテキストから株価を予測する研究としては、テキストに含まれる単語の極性を分析して、その出現頻度から株価予測を行う研究が多い。具体的には、金融ニュースの極性をネガティブ、ポジティブ、中立に分類してサポートベクターマシンで株価騰落の予測モデルを作った研究、金融ニュースのトーンに基づいてモデルを作った研究、心理学辞書と金融感情辞書を組み合わせてモデルを作った研究がある[2-4]。近年ではニュース記事の極性分析に深層学習のアルゴリズムを組み込んで精度を上げる研究も増えている[5]。

2-2. 新聞記事と株価の関係に着目した研究

新聞記事のテキストマイニングに関する研究としては、日本銀行の金融経済月報のテキストを共起解析、主成分分析、回帰分析を行って株価を予測する CPR 法[6]を応用してモデルを作り、株価指数の予測精度を高めた研究のほか、株価に影響を与える新聞記事と与えない新聞記事を識別した研究がある[7-8]。

2-3. トピックと株価の関係に着目した研究

トピックと株価の関係に着目した研究としては、米国の公開会社が証券取引委員会に届け出ることを

義務付けられている 8-K filings (日本の決算短信に該当する文書) を LDA でトピック分類して、特定のトピックが株価に影響を及ぼすことを発見した研究のほか、ロイター配信のニュース記事を LDA でトピック分類し、トピックと S&P500 を構成する主要銘柄の取引量との関係性を発見した研究がある[9-10]。

2-4. 新聞を LDA でトピック分類した研究

新聞記事のテキストを LDA で分類して、トピック割合の時系列データをグラフで可視化した上で、地域ごとの話題や関心事を見つけて、地域研究に活用することを提案した研究がある[11]。

2-5. 本研究の特徴

本研究の特徴の一つは、トピックモデルによってテキストを分類したことである。教師なし学習で効率的に分ける点と、単語をどのようなカテゴリーに分けているか解釈しやすい点で、極性分析と似ている。極性分析は、ネガティブな単語が多ければ株価下落、ポジティブな単語が多ければ株価上昇というもので解釈性が高いが、単語の持つ意味を大きく捨象する側面がある。一方で、トピックモデルは極性分析より細かく単語の分類を行うため、より豊かな表現でテキストを分類できる点で利点がある。

本研究のもう一つの特徴は、一般紙の新聞記事のテキストを分析対象としたことである。前述の先行研究が用いた 8-K filings やロイター記事のように、会社の収益や金融情報に限定して言及する媒体の方が株価予測の点で有用とも考えられる。一方で、株価が複合的な社会の状況に呼応して変動する点を踏まえると、一般紙の新聞記事を分析することで金融や経済を専門とする媒体の研究からは見えない知見が得られる可能性があり、その点で利点がある。

3. 分析手法

3-1. 新聞記事のテキスト処理方法

まず、2003 年から 2012 年までの 3,653 日のうち、読売新聞が発刊された 3,554 日の 77,814 記事の英字テキストを抽出し、大文字が出現したら小文字に変換した上で、冠詞、前置詞、代名詞、英数字や記号など、固有の意味を持たない単語、いわゆるストップワードを除去した。

次に、時制が異なる動詞や原形が同じ形容詞・副詞を別単語として識別しないよう、語幹を取り出して原形にするレンマティゼーションを施した。

そして、これらの処理を終えた後に、出現した 23,488 単語に番号を振り、記事の出現箇所を記録し、トピック分類する際に遡れるようにした。

3-2. LDA によるトピック割合の導出方法

LDA は、テキストが潜在的に存在する複数のトピックに分類できると仮定し、そのトピックについて、確率分布を用いてモデリングするものである[12]。

まず、文書の総数を D 、文書 d の総単語数を N_d 、トピックの総数を K 、語彙数を V とする。次に、文書 d でトピック k が出現する確率を $\theta_{d,k}$ 、文書 d のトピック分布を $\theta_d = (\theta_{d,1}, \theta_{d,2}, \dots, \theta_{d,K})$ とする。そして、トピック k から単語 v が作られる確率を $\phi_{k,v}$ 、単語出現分布を $\phi_k = (\phi_{k,1}, \phi_{k,2}, \dots, \phi_{k,V})$ とする。最後に、文書 d の n 番目の単語に割り当てられたトピックを $z_{d,n} \in \{1, 2, \dots, K\}$ として文書 d の n 番目の単語を $w_{d,n} \in \{1, 2, \dots, V\}$ とする。その上で、 K 次元ベクトルのハイパーパラメータを α 、 V 次元ベクトルのハイパーパラメータを β として、

$$\theta_d \sim \text{Dir}(\alpha) \quad (d = 1, 2, \dots, D),$$

$$\phi_k \sim \text{Dir}(\beta) \quad (k = 1, 2, \dots, K)$$

というようにディリクレ分布から θ や ϕ を定めて、各文書 d に対して、

$$z_{d,n} \sim \text{Multi}(\theta_d) \quad (n = 1, 2, \dots, N_d),$$

$$w_{d,n} \sim \text{Multi}(\phi_{z_{d,n}})$$

として文書生成する確率モデルが LDA である。

トピックモデルの評価指標としては、モデルの予測性能を表す Perplexity とトピックの品質を表す Coherence がある。本稿では、トピックの総数 K を決める上で、Coherence の最大化を図った。Coherence の値 C_{UCI} の求め方は以下のとおりである[13]。

$$C_{UCI} = \frac{2}{N(N-1)} \sum_{i=1}^{N-1} \sum_{j=i+1}^N PMI(w_i, w_j),$$

$$PMI(w_i, w_j) = \log \frac{P(w_i, w_j) + \epsilon}{P(w_i)P(w_j)}.$$

$P(w_i, w_j)$ は単語 i と j が同時に、 $P(w_i)P(w_j)$ は単語 i

と j が独立に出現する確率で、前者が大きいほどトピックに含まれる単語が同時に出現しやすいということ、言い換えれば生成したモデルのトピックのまとまりが大きいということで、Coherence が大きくなる。本稿では、新聞記事を分けるのに適切なトピック数を 30 と判断したが、その理由は次章で述べる。

実処理としては Python の gensim モジュールを用いて 77,814 記事の単語のトピック分類を行い、記事ごとのトピック割合を算出した。一日に 25 記事あれば、25 記事のトピック割合を平均して、日ごとのトピック割合を求めた。そして、各トピックに関わる単語含有率が高い上位 10 記事抽出して内容を読み

取り、共通する観念を抽象化して表 1 のとおり名称を付けた。

3-3. 東証株価指数の時系列

分析の対象期間で証券市場が開いた 2,455 日の東証株価指数と東証 33 業種別株価指数あわせて 34 の株価指数の対数収益率 r を求めた。

$$r = \log \frac{x}{y}$$

前営業日後に編集を終えた新聞記事に対する市場の反応を分析することが目的なので、 x は当日始値、 y は前日終値とした。株価の水準が異なっても同じ尺度で変化を評価するために対数をとったが、本稿では、 r を単に収益率という。

また、トピック割合の変化に対する株価の変化の大きさである収益率の絶対値 $|r|$ を本稿ではボラティリティ（以下、図表などではボラと略す）という。

3-4. 順位相関の計算を行うための前処理

トピック割合の時系列データは 3,554 日分あり、株価指数の時系列データは 2,455 日分しかないので、トピック割合のデータ数を株価指数のデータ数に合わせた。具体的には、市場が休みの土日に発刊された新聞のトピック割合は市場が開く月曜日の新聞のトピック割合と平均して、月曜日 1 日のデータにする処理を行った。

No.	トピックの内容	略称	主トピック 該当記事数	占有率
t01	英語学習	t英語	986	1.9%
t02	国際政治	t国際	3,886	4.1%
t03	ゴルフ	tゴルフ	3,476	3.6%
t04	意見・質問・相談	t意見	9,853	10.5%
t05	財政・マクロ経済	t財政	2,516	2.7%
t06	地震・災害	t地震	3,142	3.3%
t07	電車の障害	t電車	730	2.1%
t08	白書系の数値発表	t白書	1,791	2.5%
t09	教育・学校	t教育	2,798	2.5%
t10	政策・省庁の活動	t政策	3,638	9.6%
t11	T V ・メディア	t T V	289	1.7%
t12	裁判	t裁判	1,496	2.0%
t13	スポーツ	tスポ	423	1.2%
t14	音楽・外国文化	t音楽	49	1.8%
t15	ワイン・ファッション	t酒服	108	0.8%
t16	電力・発電	t電力	1,362	1.6%
t17	美術館・博物館	t美術	1,610	1.8%
t18	I T ・ P C	t I T	1,227	2.4%
t19	医療・病院	t医療	2,085	1.6%
t20	イベント	tイベ	2,881	5.0%
t21	宇宙・航空	t宇宙	1,090	1.7%
t22	選挙・政党	t選挙	5,457	4.5%
t23	出資・経営統合	t出資	4,269	4.0%
t24	日本の外交	t外交	2,923	6.0%
t25	地方自治	t地方	2,335	4.0%
t26	企業動向・ミクロ経済	t企業	2,862	2.7%
t27	プロ野球	t野球	5,705	5.2%
t28	逮捕・刑事事件	t逮捕	5,484	4.8%
t29	劇場・公演	t劇場	1,705	2.2%
t30	料理・食文化	t料理	1,638	2.1%
総計			77,814	100.0%

表 1：30 トピックの略称、記事数、占有率

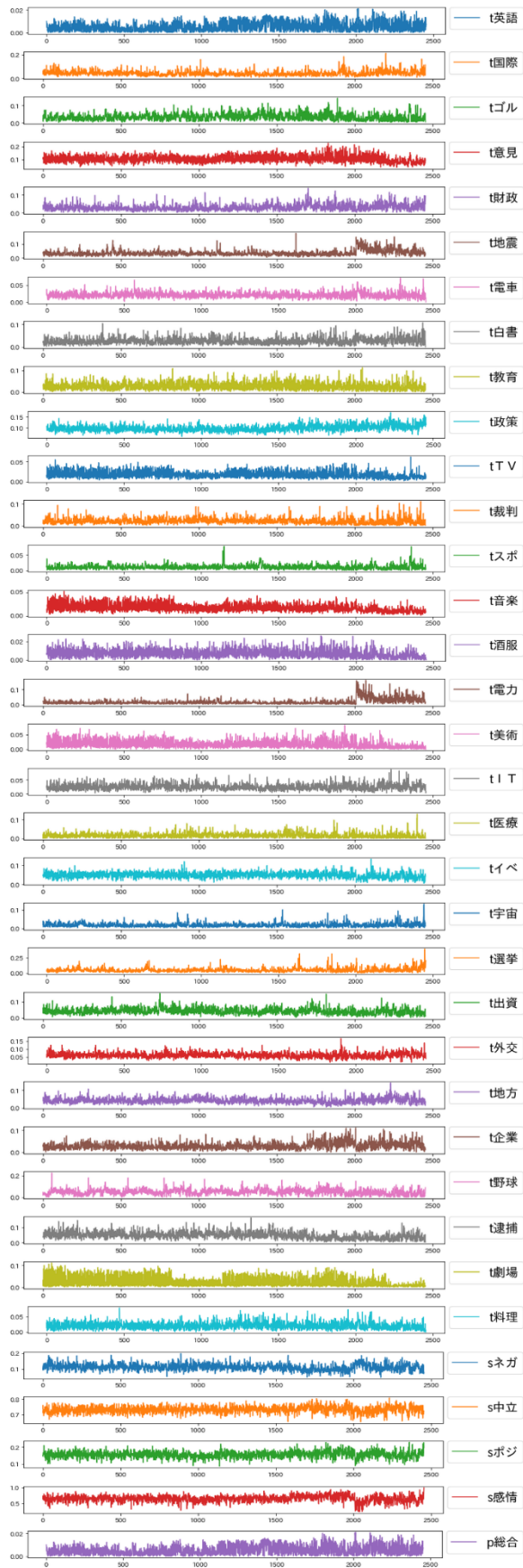


図 2 : 34 指標と TOPIX ボラの時系列

また、Python の nltk.sentiment.vader モジュールを用いてネガティブ、ポジティブ、中立、感情複合値の 4 指標も算出し、トピック割合と同じ処理で 2,455 日分に揃えた。

以下、34 の株価指数を「34 株価」、全業種を総合した東証株価指数を単に「TOPIX」または「総合」、30 トピックと 4 つの感情指標をあわせて「34 指標」という。また、株価、トピック、感情指標いずれか識別するため、それぞれの頭に「p」、「t」、「s」をつけた。株価の収益率は正規分布より裾野の長い分布であるため、スピアマンの順位相関（以下、順位相関という）を用いて相関を調べた。なお、積率相関で計算しても定性的には同様の結果が得られた。

4. 分析結果と考察

4-1. 社会事象とトピック割合の関係を俯瞰

図 2 は 34 指標と TOPIX のボラティリティの 10 年間の遷移をまとめたグラフである。2011 年 3 月（x 軸の補助目盛り 2000 の直後）の東日本大震災の時期に地震や電力トピックの割合が急増したことがわかる。時系列を注視すると、10 年間で起きた社会の諸事象を映し出してトピックの割合が遷移していることが見て取れる。

なお、劇場トピックが 2006～2007 年の間に激減して、2008 年から再び 2003～2005 年の水準に戻ったことがわかる。1 年で見ると定常的である点を踏まえると、社会の事象ではなくて新聞の編集方針が変化したというべきで、注意が必要である。

図 3 は、図 2 と同じ 35 項目の相関行列である。エンタテインメント分野で共通する劇場トピックと TV、音楽、美術、イベントの各トピックの相関が強いことがわかる（相関係数 0.5～0.6）。

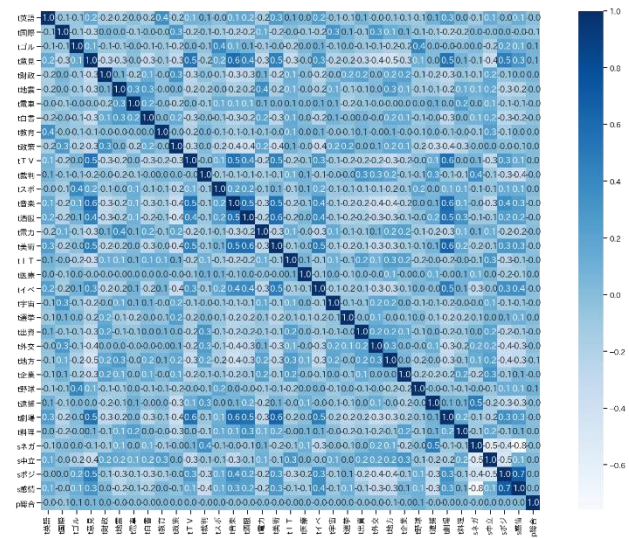


図 3 : 34 指標と TOPIX ボラの相関行列

4-2. トピック割合と TOPIX 収益率の関係

図 4 は 34 指標と 34 株価の収益率との間の 1,156 の組み合わせの相関行列である。t 検定により有意水準 p 値 5%以上となって無相関と推定できる値は非表示にした。 p 値 5%以下で相関があると推定できる値は 1,156 組の約 2%の 23 組しか存せず、トピックと収益率との関係は、ほとんど見られなかった。

そのため、収益率を分析の対象から外し、ボラティリティを分析することとする。

4-3. トピック割合と TOPIX ボラの関係

図 5 は 34 指標と 34 株価のボラティリティの相関行列であるが、図 4 と異なり多数の組み合わせで有意な相関が出たことがわかる。ただし、10 年間の時系列において有意な相関が出ても、特定の期間に強い相関が出ただけで、特定のトピックと株価指数との間で全期間に通じる関係性があるとは言えない可能性がある。そこで、1 年ごとの 34 指標と TOPIX のボラティリティの順位相関を計算した。その相関係数を表示したのが図 6 である。

以下、特定の社会事象と関連して相関係数が相対的に高くなったと考えられる組み合わせ、相関係数が相対的に高く前年と翌年で正と負が反転した組み合わせ等について考察を行う。

4-3-1. 2006 年のライブドアショック

ライブドアショックは、2006 年 1 月 16 日、M&A を繰り返して事業拡大を図り続けていたライブドア社の代表取締役である堀江貴文氏の自宅に強制捜査が入り、23 日に堀江氏が逮捕され、24 日に堀江氏の代表退任が発表されたという一連の事件で、新興市場ひいては日本の株式市場全体が一時的に混乱に陥った。

資金拠出や経営統合に関する出資トピックの新聞紙面に占める割合は、強制捜査の翌日と翌々日に 8.3%、11.3%に達し、堀江氏退任発表の翌日には最大 15.2%まで上昇した。これらは 10 年間の同トピックの平均 4.0%を大きく上回る水準だった。同期間において、ライブドアが関わる東証の情報・通信業のボラティリティが 10 年間の平均 0.5%を大きく上回る 1.4%まで上昇した。

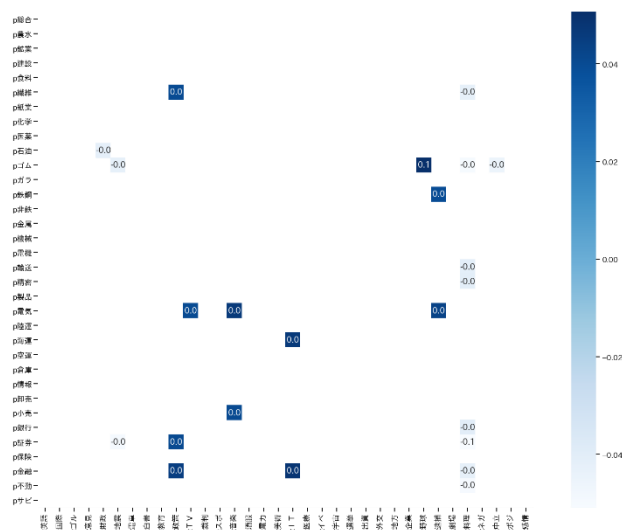


図 4 : 34 指標と 34 株価の収益率の相関行列

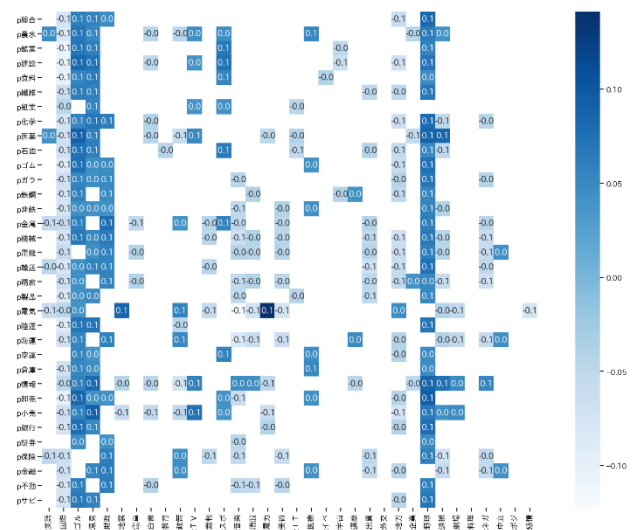


図 5 : 34 指標と 34 株価ボラの相関行列

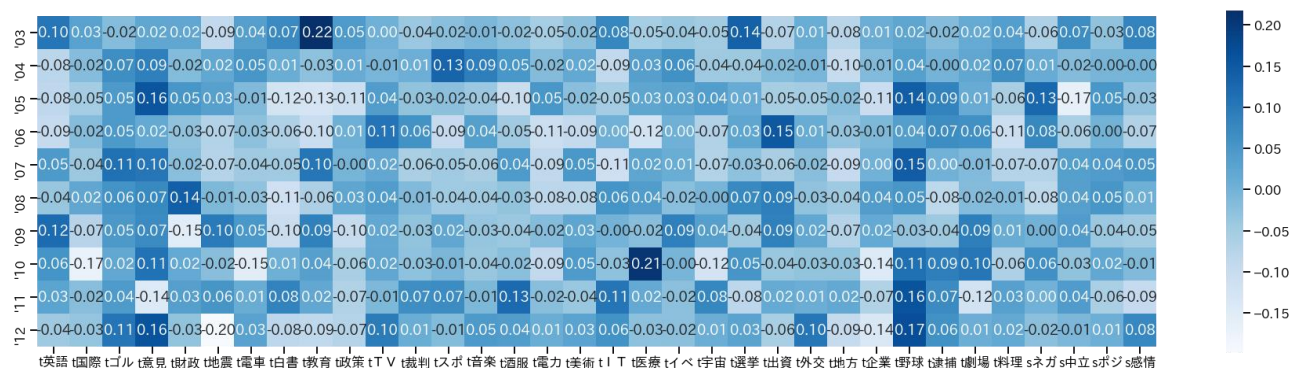


図 6 : 34 指標と TOPIX のボラティリティの 1 年ごとの相関係数

2006 年の出資トピックは、TOPIX ボラティリティとの間で強い相関（相関係数 0.152、 p 値 0.017）があったが、その他、16 業種のボラティリティとの間で有意水準 5%以下となる相関関係があった。図 7 は、同年の出資トピックと情報・通信業ボラティリティの時系列を表したもので（相関係数 0.154、 p 値 0.017）、両者の関係性を視覚的に確認しやすい。

4-3-2. 2008 年のリーマンショック

リーマンショックは、2008 年 9 月 15 日にリーマンブラザーズ社が倒産し、金融機関の信用不安と金融収縮の連想が働き、半年以上に渡って世界の金融経済が混乱に陥った事象である。リーマンブラザーズ社は、米国の信用力が低い層に向けた住宅債権・サブプライムローンを中心に組成したデリバティブを大量に保有していたことで知られ、デリバティブの価格の下落が始まった 2007 年から混乱が収まる 2009 年までの 2~3 年間を指すこともある。

日本の予算や財政政策に関する財政トピックの新聞紙面に占める割合は、10 年間の平均が 3.0%で、例年 11~12 月にかけて高くなる（月平均 3.3~3.9%）。また、経済危機が起こると、特別な出来事から一定期間経過後に大きく高まる傾向があり、例えばリーマンブラザーズ社倒産前の予兆の一つと言われるパニッシュの 1 週後の 2007 年 8 月 14 日に 8.7%、リーマンブラザーズ社倒産から 2 週後の 2008 年 10 月 2 日に 5.8%まで上がった。もっとも例示した 2 日の TOPIX のボラティリティは 0.001 と 0.005 で 10 年間の平均 0.005 と比べて特別に高いわけではない。

しかし、2008 年 1 年を通じると正の相関（相関係数 0.141、 p 値 0.027）、2009 年 1 年を通じると負の相関（相関係数 -0.152、 p 値 0.018）があった。この変動の要因を考察するため、2007 年から 2009 年までの間を 4 か月ごと 12 期間に区切って、同期間内の財政トピックの平均 (A)、TOPIX ボラティリティの平均 (B)、期間内 (4 か月) の相関係数 (C) をまとめて表示したものが表 2 である。

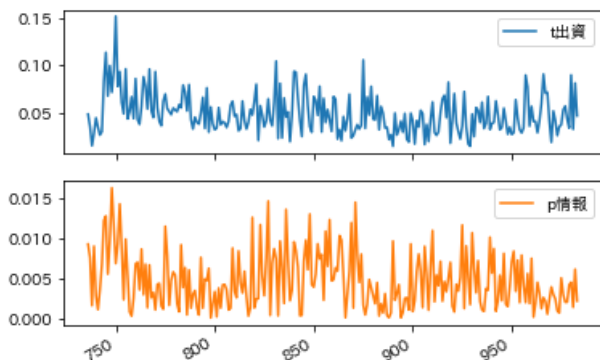


図 7：06 年の出資トピックと情報業ボラの時系列

2009 年に関しては、年始から年末に向けて金融経済が安定していった TOPIX ボラティリティが下がっていったのに対して、実体経済の促進を図るべく、年末に向けて財政政策の議論が多くなされ、新聞を占める財政トピックの割合が増えて、両者が負の相関となったと考えられる。

2008 年に関しては 9~12 月の財政トピックが増えた時期に、TOPIX のボラティリティも高まった（表 2 赤色部分の右）が、その 4 か月間で相関関係は見られなかった。一方で、その直前の 5~8 月は強い相関関係があった（相関係数 0.219、 p 値 0.046、表 2 緑色部分の右）。この時期は予算編成の時期ではないので財政トピックの割合も 2.6%と高くはなかった。また、2008 年の 1~4 月期に TOPIX ボラティリティが期間内平均 0.8%と高まった時期（表 2 赤色部分の左）は財務トピックとの相関は弱かったが、その直前期の 2007 年 9~12 月期は一定の相関があった（相関係数 0.205、 p 値 0.068、表 2 緑色部分の左）。

以上の結果を踏まえると、財政トピックの割合が増えたと TOPIX ボラティリティが高くなる、とは言えない。しかし、リーマンショックの直前期に強い相関があったことは確認できた。今後、データを整理して分析を深めることで、例えば「財務トピックと TOPIX ボラティリティの相関が強くなった後（表 2 緑色部分）はボラティリティ上昇（表 2 赤色部分）の傾向がある」といった確かめられる可能性はある。本稿冒頭で触れた "downside risks to growth have increased appreciably." という FRB 声明を報じた記事も、財政トピックの割合が一番大きい記事だった。市場参加者の不安を高め、ボラティリティを上げる方向に作用したと考えられる。

4-3-3. 2011 年の東日本大震災

東日本大震災は 2011 年 3 月 11 日に起きた日本の観測史上最高の M9.0、震度 7 の地震で、多数の犠牲者を生み、経済にも混乱を生じさせた災害である。

地震の情報や地震からの復興に関する地震トピックは 11 日連続で 9.8%を超えて（10 年間の平均は 3.0%）、TOPIX ボラティリティも 6 日連続で 1.0%以上だったが、2011 年の年間の両者に相関は見られなかった（相関係数 0.055、 p 値 0.391）。

	2007年			2008年			2009年		
	1-4月	5-8月	9-12月	1-4月	5-8月	9-12月	1-4月	5-8月	9-12月
A	2.3%	2.6%	2.3%	2.7%	2.6%	3.6%	2.9%	2.8%	4.3%
B	0.004	0.005	0.006	0.008	0.006	0.009	0.008	0.006	0.005
C	-0.142	-0.139	0.205	0.058	0.219	-0.033	0.046	-0.161	-0.128
	-0.022			0.141			-0.152		

A：財政トピック割合の期間内平均（10年間平均は3.0%）
B：TOPIXボラティリティの期間内平均（10年間平均は0.005）
C：財政トピックとTOPIXボラティリティの相関係数

表 2：07~09 年の財政トピックと TOPIX ボラ

その一方で、翌年 2012 年は強い負の相関関係があった（相関係数-0.198、 p 値 0.002）。地震トピックの割合は期間内平均 4.8%と依然高水準だったが、TOPIX のボラティリティは収まる傾向にあり、負の相関になったと考えられる。

なお、図 9 は電力・ガス業のボラティリティと 34 指標の 1 年ごとの相関係数をまとめたものであるが、TOPIX 総合とは異なり、2011 年に多数のトピックとの間で正の相関が出た（電力トピックでは相関係数 0.351、 p 値<0.001、地震トピックでは相関係数 0.264、 p 値<0.001）。また、感情指標「ネガティブ」に分類される単語割合との間でも正の相関が出た（相関係数 0.283、 p 値<0.001）。2003~2010 年まで相関がなかった中で、この年だけネガティブな記事でボラティリティが上がる、という想定される動きをしたことが特筆に値する。

4-3-4. 10 年を通して見たときの傾向

野球、ゴルフ、意見の 3 トピックは TOPIX ボラティリティとの関係で、2003 年から 2012 年までの 10 年間に於いて有意な正の相関があり、1 年ごとでも 10 年間で 9 年において正の相関があり、これらトピックが増えるとボラティリティが高くなる傾向が認められた（図 6）。意見トピックは、人生相談、質問回答、インタビューなど個人の意見や主観に関するトピックで、30 トピックの中で一番紙面を占める割合が大きなトピックである（表 1）。

これら 3 トピックの割合が大きくなるときは、相対的に重要な政治や経済に関する新情報が少ないときで、そのこととボラティリティ上昇に何らかの関係があると推定した。しかし、大震災の 2011 年は悲観的な意見の記事が増えて意見トピックと TOPIX ボラティリティは正の相関になると考えたが、同年は 10 年間で唯一負の相関が出た（図 6）。現時点で、これら 3 トピックとボラティリティの関係性について、整合的に説明することはできていない。

4-4. 最適なトピック数

トピック数を増やしながら分析対象の記事の分類を行って Coherence が高い数を選べば、最適なトピック数がわかるが、10 年分 78,814 記事をトピック数 1 から順に求めていくのは計算機にかかる負担が極めて大きい。そこで、2008 年 1 年分 8,863 記事でトピック数 1~50、2008 年 4 月 1 か月分 764 記事でトピック数 1~100 の Coherence を求めて、一般紙である読売新聞における最適なトピック数を推定することとした。前者では 28、後者では 30 で Coherence が最大になった。

以上の結果に加えて、(1)図 8 のとおり、1 年分、1 か月分どちらでも概ねトピック数 30 をピークに下落が始まったこと、(2)10 年分 77,814 記事をトピック数 30 で分類したときの Coherence が 0.534 で 1 年分をトピック数 28 で分類したときの Coherence 0.547 と大差がなかったこと、(3)各トピックに名称を付したとき、一貫性のないトピックが存在せず、定性的にもトピック数 30 は適切と評価できたこと、以上 3 点を踏まえ、トピック数 30 が最適な値だと推定した。

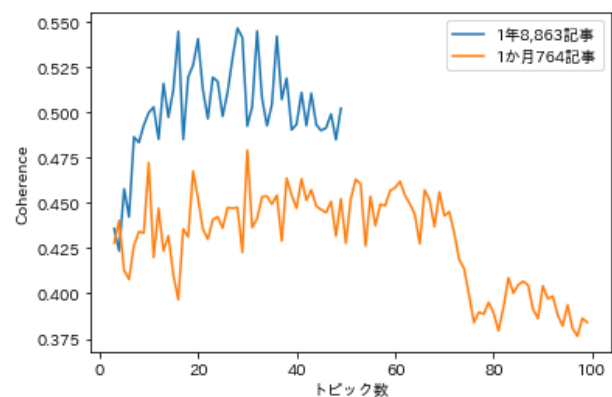


図 8：トピック数と Coherence

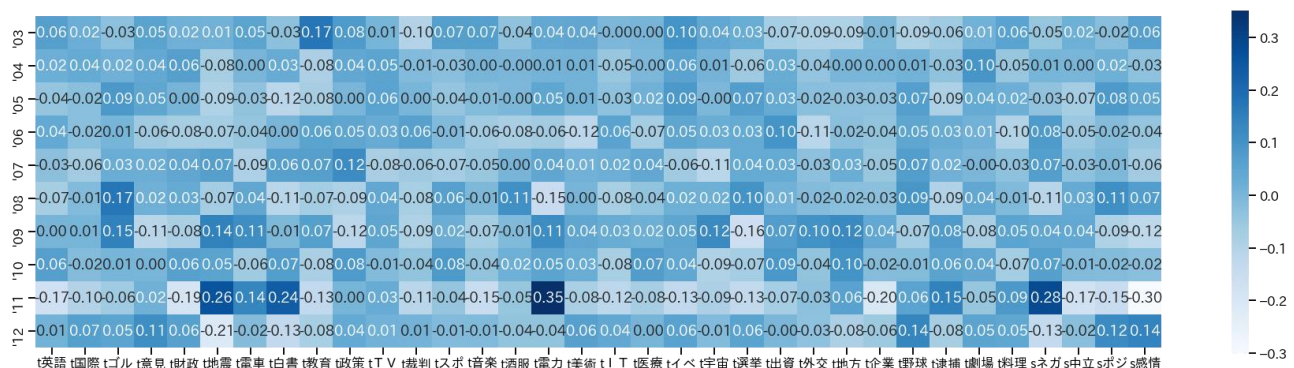


図 9：34 指標と電気・ガス業のボラティリティの 1 年ごとの相関係数

5. おわりに

社会の諸事象を伝える新聞記事のテキストから経済変動の兆候を捕捉できないか探るため、読売新聞の2003年から2012年までの10年間の77,814記事をLDAでトピック分類して、日ごとのトピック割合を求めて株価指数の時系列データとの相関を分析した。その結果、トピック割合とTOPIXの収益率との間に相関は見られなかったが、TOPIXのボラティリティとの間には時間に依存した有意な相関があることを発見した。

TOPIXのボラティリティと相関があった新聞記事のトピックの具体例としては、ライブドアショックのあった2006年の出資トピック、リーマンショックのあった2008年と翌2009年の財政トピックが挙げられる。東日本大震災があった2011年は、地震や電力トピック、ネガティブ割合と電気・ガス業のボラティリティとの間に特に強い正の相関があった。一方で大震災翌年の2012年、それらは負の相関関係へと変わった。野球、ゴルフ、意見の3トピックに関しては、10年を通じた傾向として、それらの割合が増えるとTOPIXのボラティリティが上がる正の相関があった。

現時点では、事後的にトピック割合と株価指数のボラティリティに相関関係があった組み合わせや期間を特定できたに過ぎない。今後、データを整理して分析を深め、関係があった理由を特定し、トピック割合のデータを経済危機の予兆の捕捉に役立てられるようにしていきたい。その他、今後の課題として以下5点を挙げる。(1)新聞記事が株価指数に与える影響を実質的に把握するためには、異時刻の相関分析が不可欠である。(2)株価指数の収益率を当日始値と前日終値に基づいて算定したが、新聞の掲載情報は前日の市場が開いているときに公開されている情報を含む点や当日の市場が開いたときに新聞には未掲載の情報がある点を踏まえると収益率を前日終値と前日始値など基準時を変えて分析を行うことも必要である。(3)一般紙である読売新聞の記事を対象としてトピック分類を行ったが、Webニュース、TwitterほかSNS、TVなど他の媒体のテキストに関しても分類と分析を行って比較することが必要である。(4)LDAによるトピック分類のみならずUniversal sentence encoderなど別の手法で分類と分析を行って比較をすることが必要である。(5)トピック割合や極性分析で得た感情指標の値を活用して株価指数や株価指数ボラティリティ(VIXや日経VI)の予測を試みることも経済危機の予兆の捕捉に役立つと考えられるので、必要である。

謝辞

本研究はJSPS 科研費JP19H01114の助成を受けたものです。

参考文献

- [1] Malkiel, B. G., and Fama, E. F.: Efficient capital markets: A review of theory and empirical work. *The Journal of Finance*, Vol. 25, No. 2, pp. 383-417, (1970)
- [2] Kalyani, J., Bharathi, H. N., and Rao, J.: Stock trend prediction using news sentiment analysis, *arXiv preprint*, arXiv: 1607.01958, (2016)
- [3] Schumaker, R. P., Zhang, Y., Huang, C., and Chen, H.: Evaluating sentiment in financial news articles, *Decision Support Systems*, Vol. 53, No. 3, pp. 458-464, (2012)
- [4] Li, X., Xie, H., Chen, L., Wang, J., and Deng, X.: News impact on stock price return via sentiment analysis, *Knowledge-Based Systems*, Vol. 69, pp. 14-23, (2014)
- [5] Mohan, S., Mullapudi, S., Sammeta, S., Vijayvergia, P., and Anastasiu, D. C.: Stock price prediction using news sentiment analysis, *IEEE 5th International Conference on Big Data Computing Service and Applications*, (2019)
- [6] 和泉潔, 後藤卓, 松井藤五郎: 経済テキスト情報を用いた長期的な市場動向推定, *情報処理学会論文誌*, Vol. 52, No. 12, pp. 3309-3315, (2011)
- [7] 藏本貴久, 和泉潔, 吉村忍, 石田智也, 中嶋啓浩, 松井藤五郎, 吉田稔, 中川裕志: 新聞記事のテキストマイニングによる長期市場動向の分析, *人工知能学会論文誌*, Vol. 28, No. 3, pp. 291-296, (2013)
- [8] 中山大, 坂地泰紀, 勝田研一郎, 酒井浩之: 株価に影響を与える重要な出来事が記載された記事の自動抽出, *成蹊大学理工学研究報告*, Vol. 51, No. 2, pp. 53-60, (2014)
- [9] Feuerriegel, S., and Pröllochs, N.: Investor reaction to financial disclosures across topics: An application of latent Dirichlet allocation, *Decision Sciences*, pp. 1-21, (2018)
- [10] Hisano, R., Sornette, D., Mizuno, T., Ohnishi, T., and Watanabe, T.: High quality topic extraction from business news explains abnormal financial market volatility, *PLoS ONE*, Vol. 8, No. 6, e64846, (2013)
- [11] 山田太造: 新聞記事に対するトピックモデルの適用とトピックの時系列変化に関する考察, *研究報告, 人文科学とコンピュータ (CH)*, Vol. 2017-CH-115, pp. 1-5, (2017)
- [12] Blei, D. M., Ng A. Y., and Jordan M. I.: Latent Dirichlet Allocation, *Journal of Machine Learning Research* 3, pp. 993-1022, (2003)
- [13] Röder, M., Both, A., and Hinneburg, A.: Exploring the space of topic coherence measures, *Proceedings of the Eighth ACM International Conference on Web Search and Data Mining*, pp. 399-408, (2015)

BERT を用いた有価証券報告書からの ESG 関連文抽出

ESG-related sentences extraction from securities reports using BERT

土橋 諒太^{1*} 中田 和秀¹
Ryota Dobashi¹ Kazuhide Nakata¹

¹ 東京工業大学工学院経営工学系

¹ School of Engineering, Department of Industrial Engineering and Economics,
Tokyo Institute of Technology

Abstract: 資産運用分野において、従来から考慮されてきた財務情報に加えて ESG 課題を考慮して投資を行う「ESG 投資」が世界的な潮流となっている。ESG 投資において、企業による ESG 情報開示は重要な情報源となるが、日本において多くの投資家が企業の ESG 情報開示は十分ではないと考えているという調査結果がある。本研究では、上場企業に法的に義務付けられた開示書類である有価証券報告書を使用して、ESG 関連文を抽出するモデルを作成する。具体的には、有価証券報告書の経営方針項目及び事業等のリスク項目の文に対しアノテーションを行うことで ESG 関連文データセットを作成し、その ESG 関連文データセットを利用して BERT のファインチューニングを行って、ESG に関連する情報を抽出する。更に、ファインチューニングした BERT を使用して、ESG 情報開示の動向を可視化した。

1 はじめに

資産運用分野において、従来から考慮されてきた財務情報に加えて ESG 課題を考慮して投資を行う「ESG 投資」が世界的な潮流となっている。ESG は、Environment（環境）・Social（社会）・Governance（ガバナンス）のイニシャルを合わせた言葉である。

ESG 投資に対する関心は世界全体で高まっており、Global Sustainable Investment Alliance の調査 [1] によると、世界全体でのサステナブル投資運用資産残高は 2016 年時点で 22.8 兆米ドルだったが、2018 年には 30.6 兆米ドルを超えた。日本においても、サステナブル投資運用資産残高は、2016 年には 0.5 兆米ドルから 2018 年の約 2.2 兆米ドルに増加している。

ESG 投資が資産運用分野で拡大している一方で、ESG 情報開示には課題が存在している。企業は温暖化防止対策やダイバーシティに関する取り組みなど自社の ESG に関する様々な情報を開示して、ステークホルダーに対して説明責任を果たすことが求められている。しかしながら 2019 年の一般社団法人生命保険協会の調査 [2] により、企業と投資家の ESG 情報開示の現状に関する認識にはギャップが存在し、投資家は現在の企業の開示を十分ではないと考えていることが分かった。調査 [2] によると、「ESG への取組に関する情報開示は十分と

考えるか」という質問に対して、「十分開示している」と開示した企業は 28 %であった一方で、「十分開示している」と回答した投資家は 1 %に留まっている。

企業が ESG 情報を開示する手段として、統合報告書、CSR・サステナビリティレポート、有価証券報告書などが用いられているが、これらの開示手段にはそれぞれ課題が存在する。CSR・サステナビリティレポートと統合報告書は ESG に関する開示手段として豊富な ESG 開示情報を含むが、課題が 3 つ存在する。1 つ目に開示に関する統一的な基準が存在しないため開示内容が企業によって異なっていることがある。開示形式についても、有価証券報告書のような共通した開示形式は存在せず、独自の図表を多用したデザインや pdf フォーマットの使用などから、統計分析や機械学習を行うことが困難である。2 つ目に CSR・サステナビリティレポートと統合報告書は企業の自主的な開示書類であるため発行企業数が全上場企業ではないことがある。KPMG の調査 [3] によれば、統合報告書を発行する企業数は上場・非上場を併せて 2019 年で 512 社に留まり、上場企業における発行企業の割合は少ない。3 つ目に共通の開示プラットフォームが存在しないため、各企業のホームページなどから個別に書類を取得する必要がある。以上 3 つの理由から、投資家が CSR・サステナビリティレポートと統合報告書に対して統計分析や機械学習を行って ESG 情報を活用することは困難であると言える。

*E-mail: dobashi.r.ab@m.titech.ac.jp

一方で、有価証券報告書は、統一的な開示内容・形式が存在し、上場企業に法的に義務付けられた開示書類であり、EDINET という共通の開示プラットフォームが存在するため、統計分析や機械学習に用いるデータとして適している。しかしながら、有価証券報告書は財務情報の記述がメインであり ESG に関する記述は一部であるため、ESG に関する情報を効率的に入手するためには ESG 情報を抽出する必要がある。従って、本研究では、投資家・企業間の ESG に関するコミュニケーションギャップの減少を目的として、有価証券報告書に対して言語モデルを適用することで、有価証券報告書から ESG に関する記述を抽出すると共に ESG 情報開示動向の可視化を行う。

2 関連研究

有価証券報告書、決算短信、経済ニュースから投資判断に有用な情報を抽出する手法の研究が行われている。因果関係に注目した研究として、Sakaji et al.[4]、坂地他 [5]、佐藤他 [6] があり、手がかり表現と機械学習により、企業の業績などに対する因果関係文を新聞の記事・決算短信・有価証券報告書から抽出する手法を提案している。因果関係文とは、結果と原因のペアから構成される文を指し、手がかり表現は、因果関係文を判定する上で重要な手がかりとなる文を指す。佐藤他 [6] では因果関係文の例文として、「猛暑日が連続したため、飲料水の売り上げが伸びた」という文が挙げられており、「ため、」が手がかり表現となる。これらの研究では、手がかり表現を利用して素性を作成し、サポートベクターマシンなどの機械学習により、因果関係文の判別を行った。

業績に注目した研究として、酒井他 [7] や Kitamori et al.[8] がある。酒井他 [7] は決算短信から手がかり表現を用いて業績要因を含む文を抽出し、Kitamori et al.[8] は決算短信から手がかり表現を用いて業績予測を含む文を抽出した。

リスクに注目した研究として、Masson&Montariol[9]、藤井他 [10] がある。Masson&Montariol[9] は、有価証券報告書において重要なリスク要因を企業が省略する場合を検出するため、BERT を使用して有価証券報告書からリスク文を抽出した上で、抽出したリスク要因のクラスタリングを行った。藤井他 [10] は、有価証券報告書からサポートベクターマシン、ロジスティック回帰、ランダムフォレスト、双方向 LSTM、BERT など様々な機械学習手法によるリスク文の抽出を行い、その結果を比較した。リスク文の抽出において、双方向 LSTM、BERT によるリスク文抽出の精度が良いという結果を報告した。

本研究では、ESG に注目し、有価証券報告書から機械

学習を用いて ESG に関連した文を抽出する。手法として、ESG 関連文の抽出においては、因果関係文などの抽出で用いられる手がかり表現の使用は有効ではないと考えられる。本研究では、有価証券報告書の、経営方針、経営環境及び対処すべき課題等及び事業等のリスクの文に対しアノテーションを行うことで ESG 関連文データセットを作成し、その ESG 関連文データセットを利用して BERT のファインチューニングを行い、ESG に関連する情報を抽出した。更にファインチューニングした BERT を使用して ESG 情報開示の動向を可視化を行った。

3 ESG 関連文抽出タスク

3.1 ESG 関連文データセット

有価証券報告書から ESG 関連文を抽出するモデルの作成・検証を行うため、有価証券報告書からサンプルした文のアノテーションを行い、ESG 関連文データセットを作成した。データセットを作成する際に、モデルが年度・企業の差異に頑強であるかを検証するため、異なる企業・年度の有価証券報告書から表 1 のように文のサンプリングを行った。企業群 A は 60 企業、企業群 B は 10 企業を上場企業からランダムに選択した。

表 1: ESG 関連文データセット

	企業群 A	企業群 B
2017 年	1200 文	300 文
2018 年	233 文	205 文

抽出した有価証券報告書の経営方針経営環境及び対処すべき課題等、事業等のリスク項目に含まれる文に対してアノテーションを行い、E・S・G・ESG 非関連の 4 項目にラベリングした。ラベリングは世界で広く採用されているサステナビリティ報告の枠組みである GRI スタンダード [11] に従った。

3.2 GRI スタンダードに基づくラベリング

サンプルした有価証券報告書の文をラベリングするにあたって、E・S・G・ESG 非関連を定義する必要がある。E・S・G・ESG 非関連を定義は、世界で広く利用されているサステナビリティ報告書のガイドラインである GRI スタンダード [11] に従った。GRI が定めた開示項目について E・S・G の分類を行い、E・S・G のいずれかの項目にあてはまるものをそれぞれ E・S・G、あてはまらないものを ESG 非関連と定義してラベリングを行った。

3.3 言語モデルによる ESG 関連文抽出

ESG 関連文データセットに含まれる文章は E・S・G・ESG 非関連の 4 カテゴリであるため、4 クラスの文章分類である ESG 関連文抽出タスクを設定する。ESG 関連文抽出タスクの目的は、ESG 関連文抽出タスクで精度良く ESG 関連文を抽出できるモデルを開発し、そのモデルを任意の有価証券報告書に適用して ESG 関連文を精度良く抽出することである。

ESG 関連文抽出タスクに対して、事前学習済みの言語モデルを利用する。言語モデルには単語レベルの埋め込みを行うモデルと文レベルの埋め込みを行うモデルが存在する。単語レベル・文レベルでの言語モデルの代表的なものとして Word2vec[12] と BERT[13] がある。

Word2vec[12] は単語レベルの言語モデルであり、各単語に対して単一の埋め込みベクトルを生成する。BERT[13] は文レベルの言語モデルであり、入力文に依存した埋め込み表現を生成する。BERT のアーキテクチャは Transformer のエンコーダーを積層したものであり、Masked LM という事前学習のタスクを新たに導入することで、事前学習において双方向で学習することが可能となった。その結果、GLUE など自然言語処理のタスクの精度が大幅に向上した。

本研究では文レベルの言語モデルである BERT が ESG 関連文抽出タスクにおいても有効であると考え、ESG 関連文抽出タスクにおける BERT のファインチューニングの有効性を検証する。ESG 関連文抽出タスクで検証する以下の 3 つのモデルを表したものが図 1 である。

1. Word2vec
2. ファインチューニング無し BERT
3. ファインチューニング有り BERT

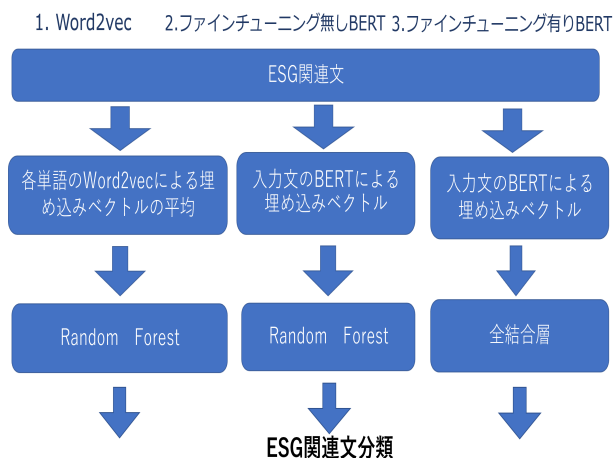


図 1: ESG 関連文抽出手法

Word2vec を用いたモデルについて、入力文の各単語

の Word2Vec による埋め込みベクトルを取得し、全単語の Word2Vec による埋め込みベクトルの平均ベクトルを算出して、Random Forest の入力とした。Random Forest を使用した理由として、ロジスティック回帰やサポートベクターマシンなどと比較して ESG 関連文抽出タスクにおける精度が高かったためである。

ファインチューニング無しの BERT を用いたモデルについて、入力文の BERT による埋め込みベクトルである BERT 最終層の出力を取得して、Random Forest の入力とした。

ファインチューニング有りの BERT を用いたモデルについて、BERT の追加的なアーキテクチャとして、BERT の最終層に 1 層の全結合層を連結することで、文章分類に対応した。全結合層を通して誤差逆伝播により、BERT のパラメーターをファインチューニングすることが出来る。

4 数値実験

4.1 実験設定

ESG 関連文抽出タスクを用いて、有価証券報告書から ESG 関連文を抽出するモデルの作成・検証を行う。年度・企業が異なるデータに対するモデルの汎化性能を検証するため、以下のように、年度・企業が訓練・検証データで同じデータセットと、年度・企業が訓練・検証データで異なるデータセットを使用した。但し、訓練・検証データで年度・企業が重複していたとしても、訓練・検証データに含まれる文は重複していないため、リークは発生していない。

実験 1 2017 年訓練データ, 2017 年検証データ
訓練・検証データの企業重複がある

実験 2 2017 年訓練データ, 2017 年検証データ
訓練・検証データの企業重複がない

実験 3 2017 年訓練データ, 2018 年検証データ
訓練・検証データの企業重複がある

実験 4 2017 年訓練データ, 2018 年検証データ
訓練・検証データの企業重複がない

予測ラベルは E・S・G・ESG 非関連の 4 クラスである。評価指標は、ESG 非関連のラベルが多いという不均衡性を考慮して、F 値を使用する。

比較手法は以下の 3 通りである。

1. Word2vec

入力文を Mecab で分かち書きし、Word2vec による単語ベクトルを文章に含まれる単語数で標準化したものを特徴量として、Random Forest で予測を行った。

2. ファインチューニング無し BERT

日本語 Wikipedia で事前学習済み Bert を訓練データで学習せずに検証データでの文章埋め込みを行い、Random Forest で予測を行った。

3. ファインチューニング有り BERT

日本語 Wikipedia で事前学習済み BERT を訓練データで訓練し、検証データでの予測を行った。

実装について、Word2vec は東北大学乾・鈴木研究室が公開した日本語 Wikipedia エンティティベクトル [14] 及び Gensim を利用した。BERT は東北大学乾・鈴木研究室が公開した日本語 Wikipedia から学習したモデルを利用した [15]。ファインチューニング有り BERT では、訓練データにおける誤差の減少から判断して、9 エポックの学習を行った。

4.2 実験結果

以下の表の値は全て F 値である。Weighted Avg は各クラスの F 値をクラスのデータ数で重みをつけた加重平均である。表の値は小数点 3 桁目を四捨五入した値とした。表 2 は実験 1 の結果をまとめた表であるが、実験 2~4 については、実験 1 と同様の結果であったため、実験 2~4 の結果については省略する。

表 2 において、Word2vec と比較してファインチューニング有り BERT の F 値が全てのクラスで高いことから、文レベルでの言語モデルを使用することが、ESG 関連文の分類に有効であると判断できる。ファインチューニング無し BERT と比較してファインチューニング有り BERT の F 値が全てのクラスで高いことから、ESG 文分類タスクで良い精度を出すためには BERT のファインチューニングが必要であると判断できる。

以上の結果は実験 2~4 についても同様であり、ファインチューニング有り BERT の F 値が全てのクラスで高いという結果であった。従って、ファインチューニングした BERT は企業と年度の差異の両方に頑強であり、

年度・企業が異なるデータに対しても汎化し、精度良く ESG 関連文を抽出できることが分かった。

4.3 考察

実験 1 の検証データにおいて、抽出が適切に行われなかった文について考察する。表 3 は、ファインチューニングした BERT による予測が不正解であった文である。予測が不正解であった文の特徴としては、赤字の部分が判断の根拠となるが、赤字の部分が 1 つの単語など根拠箇所が少なく、予測が難しい文であったと言える。訓練したモデルを用いて有価証券報告書から ESG 関連文を抽出する場合、ESG に関する言及が少ない文を抽出する必要があるかは、投資家のニーズに依存している。このような ESG に関する言及が少ない文まで抽出する必要がある場合は、アノテーションにより、訓練データにおいて ESG に関する言及が少ない文を増やすことなどが必要であると考えられる。

4.4 BERT による ESG 情報開示の可視化

ESG 関連文抽出データセットによりファインチューニングした BERT を用いて、有価証券報告書の各文の予測 ESG ラベルを算出した。経営方針、経営環境及び対処すべき課題と、事業等のリスクのそれぞれの項目において、E・S・G と予測された文の数をカウントし、売上高別にその平均値を算出した。使用した有価証券報告書は、2014 年は 3213 企業、2015 年は 3340 企業、2016 年は 2880 企業、2017 年は 3350 企業、2018 年は 3672 企業である。

図 2~7 の縦軸は予測された ESG 関連文数の平均値を示し、横軸は売上高を示している。横軸において、~1000 億円は売上高 1000 億円未満、1000 億円~5000 億円は売上高 1000 億円以上 5000 億円未満、5000~は売上高 5000 億円以上を表している。

表 2: 実験 1, 2017 年訓練・2017 年検証データ、訓練・検証データの企業重複がある

	Word2vec	ファイン チューニング 無し BERT	ファイン チューニング 有り BERT
E	0.32	0.28	0.84
S	0.41	0.29	0.88
G	0.72	0.64	0.9
ESG 非関連	0.82	0.74	0.94
Weighted Avg	0.71	0.64	0.91

表 3: BERT 不正解例

入力文	正解ラベル	ファインチューニング 有り BERT 予測ラベル	Word2vec 予測ラベル
同プライスを経済性評価に織り込むことで、 CO2 排出量を投資判断の材料としています	E	非関連	S
また、2018 年 3 月には、当社が女性活躍推進 に優れた上場企業として経済産業省と 株式会社東京証券取引所が共同で選定する 「なでしこ銘柄」に初めて選定されました	S	非関連	G
■ポートフォリオ経営の推進■資本効率を重視 した経営指標に基づく事業運営と現場への展開 ■コーポレートガバナンスの変革なお、成長戦略 を織り込んだ新中期経営計画は、構造改革終了後 の 2019 年 4 月のスタートを目指して、改めて発表 する予定です	G	非関連	非関連

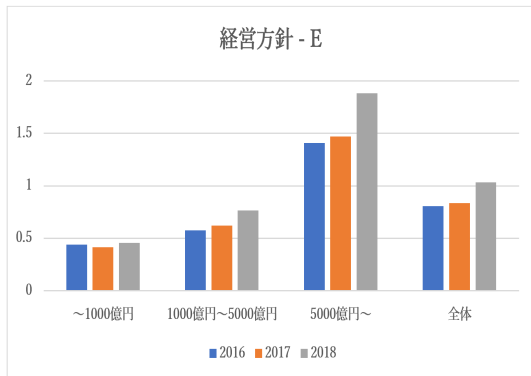


図 2: 経営方針, 経営環境及び対処すべき課題における, 売上高別の予測 E 関連文数の平均値

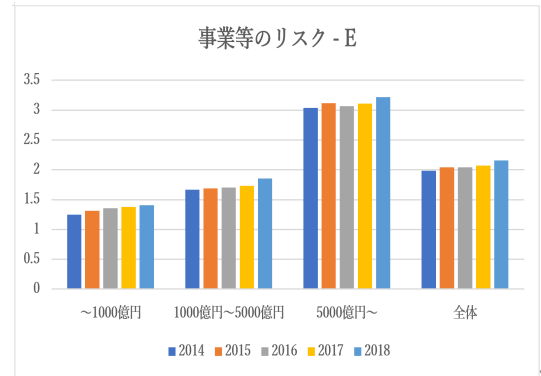


図 3: 事業等のリスクにおける, 売上高別の予測 E 関連文数の平均値

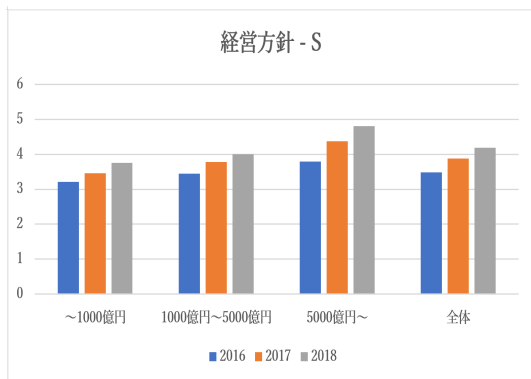


図 4: 経営方針, 経営環境及び対処すべき課題における, 売上高別の予測 S 関連文数の平均値

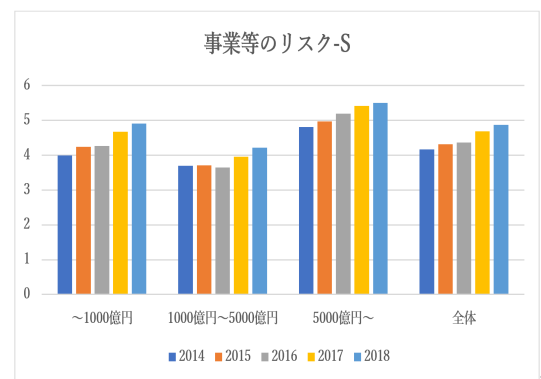


図 5: 事業等のリスクにおける, 売上高別の予測 S 関連文数の平均値

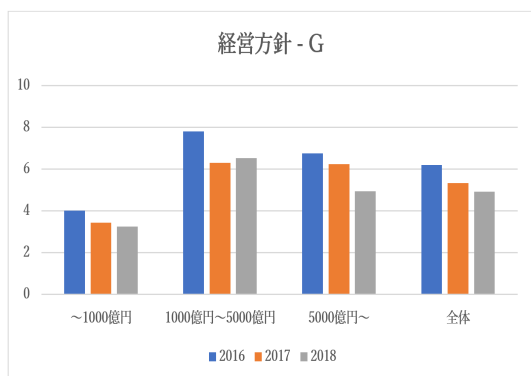


図 6: 経営方針, 経営環境及び対処すべき課題における, 売上高別の予測 G 関連文数の平均値

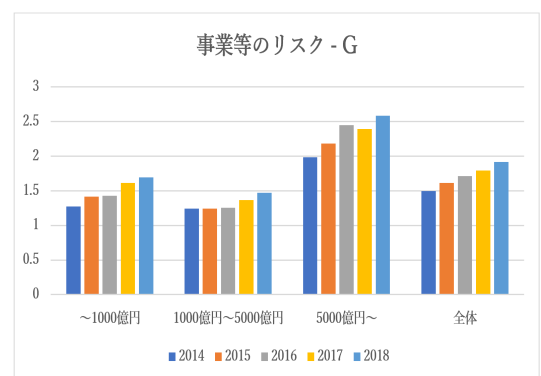


図 7: 事業等のリスクにおける, 売上高別の予測 G 関連文数の平均値

図2・図3は、それぞれ経営方針、経営環境及び対処すべき課題項目・事業等のリスク項目における、予測E関連文数の平均値を表すグラフである。図4・図5は、それぞれ経営方針、経営環境及び対処すべき課題項目・事業等のリスク項目における、予測S関連文数の平均値を表すグラフである。図6・図7は、それぞれ経営方針、経営環境及び対処すべき課題項目・事業等のリスク項目における、予測G関連文数の平均値を表すグラフである。

図2～7により、2つの知見が得られる。1つ目に年度の経過につれて予測ESG関連文数の平均値は増加しているが、経営方針、経営環境及び対処すべき課題項目においては予測G関連文数の平均値が減少傾向にあることが分かる。その要因として、経営方針、経営環境及び対処すべき課題項目におけるG関連文には大量買付・買収防衛策に関する記述が多いが、買収防衛策を廃止・非継続する企業が増加しており、大量買付・買収防衛策に関する記述が減少したことが推測できる。大和総研のレポート[16]によれば、上場企業の内買収防衛策を導入している企業数は2016年度末の450企業から2018年度末の386企業に減少している。2つ目に予測E・G関連文数は売上高5000億円以上の企業で多い傾向にあるが、予測S関連文数は売上高による大きな差は見られない。その要因として、環境報告ガイドラインやコーポレートガバナンスコードなどにより近年積極的な情報開示が求められている環境情報とガバナンス情報について、売上高5000億円以上の大企業がより積極的な開示を行い、投資家と社会のニーズに対応していると考えられる。

5 まとめと今後の課題

本研究の研究課題は、有価証券報告書からESGに関する情報を抽出するモデルを作成することであった。研究課題に対する提案手法として、有価証券報告書の経営方針項目及び事業等のリスク項目の文に対しアノテーションを行うことでESG関連文データセットを作成し、そのESG関連文データセットを利用してBERTのファインチューニングを行い、ESGに関連する情報を抽出した。

ESG関連文抽出の数値実験では、年度・企業が異なる検証データを使って検証を行った結果、BERTをESG関連文抽出タスクでファインチューニングすることで、年度・企業が異なるデータに対しても精度良く汎化し、ESG関連文を抽出できることが分かった。BERTによる予測が失敗したサンプルを分析した結果、予測が失敗した例は、ESGに関する記述量が少なかったことから、本研究で開発したBERTは、有価証券報告書からESGに関する記述が多い文を自動的に収集し、分析し

たいという投資家のニーズに答えることが出来ると言える。本研究では、個別企業レベルでESG関連文を抽出すると共に、複数企業のESG関連文に関して集計することで、ESG情報開示の動向を可視化した。以上をまとめると、本研究の貢献はESG情報の抽出のための精度の良いモデルを開発することで、個別企業レベルでのESG関連文の抽出を可能とすると共にESG情報開示の動向を可視化したことである。

本研究に対する今後の課題として、次の事項が挙げられる。本研究では、機械学習で扱いやすいデータとして有価証券報告書を扱ったが、ESG情報の豊富さという点で、CSR・サステナビリティ報告書、統合報告書などからのESG情報抽出に対する投資家のニーズは高いと考えられる。本研究で得られたBERTモデルをCSR・サステナビリティ報告書、統合報告書、ニュースなどのデータに適用できるかを検証し、投資家にとってより有効なモデルを開発していく必要がある。

謝辞

本研究はみずほ第一フィナンシャルテクノロジー株式会社様のご協力の下に行われました。惜しみないご協力とアドバイスをして下さったみずほ第一フィナンシャルテクノロジー株式会社の皆様に心より感謝を申し上げます。

参考文献

- [1] Global Sustainable Investment Alliance, “Global Sustainable Investment Review,” 2018, http://www.gsi-alliance.org/wp-content/uploads/2019/06/GSIR_Review2018F.pdf, (参照 2021-01-15).
- [2] 一般社団法人生命保険協会, 生命保険会社の資産運用を通じた「株式市場の活性化」と「持続可能な社会の実現」に向けた取組について, 2019, <https://www.seiho.or.jp/info/news/2020/pdf/20200417.4-all.pdf>, (参照 2021-01-15).
- [3] KPMG ジャパン統合報告センター・オブエクセレンス, 日本企業の統合報告に関する調査 2019, 2020, <https://assets.kpmg/content/dam/kpmg/jp/pdf/2020/jp-integrated-reporting.pdf>, (参照 2021-01-15).
- [4] Hiroki Sakaji, Satoshi Sekine, Shigeru Masuyama, “Extracting Causal Knowledge Using

- Clue Phrases and Syntactic Patterns,” 7th International Conference on Practical Aspects of Knowledge Management, pp.111-122, 2008.
- [5] 坂地泰紀, 酒井浩之, 増山繁, 決算短信 PDF からの原因・結果表現の抽出, 電子情報通信学会論文誌 D, J98-D(5), pp.811-822, 2015.
- [6] 佐藤史仁, 佐久間洋明, 小寺俊哉, 田中良典, 坂地泰紀, 和泉潔, 有価証券報告書からの因果関係文の抽出, 第 32 回人工知能学会全国大会, pp.1518-1521, 2018.
- [7] 酒井 浩之, 西沢 裕子, 松並 祥吾, 坂地 泰紀, 企業の決算短信 PDF からの業績要因の抽出, 人工知能学会論文誌, 30(1), pp.172-182, 2015.
- [8] Shiori Kitamori, Hiroyuki Sakai, Hiroki Sakaji, “Extraction of Sentences Concerning Business Performance Forecast and Economic Forecast from Summaries of Financial Statements by Deep Learning,” IEEE Symposium on Computational Intelligence for Financial Engineering & Economics, pp.67-73, 2017.
- [9] Corentin Masson, Syrielle Montariol “Detecting Omissions of Risk Factors in Company Annual Reports,” The 2nd Workshop on Financial Technology and Natural Language Processing at the 29th International Joint Conference on Artificial Intelligence and the 17th Pacific Rim International Conference on Artificial Intelligence, pp.15-21, 2020.
- [10] 藤井元雅, 坂地泰紀, 佐々木一, 増山繁, 有価証券報告書からのリスク文抽出の試み, 25 回 人工知能学会 金融情報学研究会, pp.44-48, 2020
- [11] Global Reporting Initiative, GRI サステナビリティ・レポート・スタンダード 2016(完本版), 2016, <https://www.globalreporting.org/how-to-use-the-gri-standards/gri-standards-japanese-translations/>, (参照 2021-01-15).
- [12] Tomas Mikolov, Ilya Sutskever, Kai Chen, Greg Corrado, Jeffrey Dean, “Distributed Representations of Words and Phrases and their Compositionality,” In Advances in Neural Information Processing Systems, pp.3111-3119, 2013.
- [13] Jacob Devlin, Ming-Wei Chang, Kenton Lee, Kristina Toutanova, “BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding,” In Proceedings of Annual Conference of the North American Chapter of the Association for Computational Linguistics, pp.4171-4186, 2019.
- [14] 鈴木正敏, 松田耕史, 関根聡, 岡崎直観, 乾健太郎, Wikipedia 記事に対する拡張固有表現ラベルの多重付与, 言語処理学会第 22 回年次大会, 2016.
- [15] 東北大学 乾・鈴木研究室, Pretrained Japanese BERT models, <https://github.com/cl-tohoku/bert-japanese>, (参照 2021-01-15).
- [16] 吉川英徳, 2019 年 6 月株主総会シーズンの総括と示唆, 2019, https://www.dir.co.jp/report/consulting/governance/20190821_020969.pdf, (参照 2021-01-15).

キーワード検索数とツイートの情報を用いた ビットコイン価格の騰落予測

Bitcoin price prediction using keyword search volumes and tweet information

松田周也¹ 原尚幸¹
Shuya Matsuda¹ Hisayuki Hara¹

¹ 同志社大学 文化情報学部

¹ Faculty of Culture and Information Science, Doshisha University

Abstract: In this paper, we propose an LSTM for predicting Bitcoin price using Google Trends data and sentiment scores of news and tweets. We used sentiment scores of news weighted by the number of news as inputs. We also used sentiment scores of tweet sentiment scores weighted by tweet information, such as the number of likes, the number of retweets and the number of followers and used them as inputs. The results show better performance of the proposed method than using the non-weighted sentiment scores of news and tweets as inputs.

1 はじめに

近年、テキストデータの取得技術、解析技術双方の進歩に伴い、テキストマイニングは様々な分野に応用されるようになった。テキスト情報と市場動向の関係を考察する金融テキストマイニング技術も、新たな展開を見せている。

和泉ら [1] は、株価変動の要因分析におけるニュース記事のテキストデータ有用性を示している。和泉ら [2] は、テキストデータ内の語句の出現頻度と価格変動の関連を分析するための CPR 法という手法を提案し、株価の月次変動の分析におけるニュース記事のテキストデータの有用性を示している。

また、テキストデータから感情分析によって算出された感情スコア¹を市場動向の分析に用いる研究も多く見られるようになった。投資判断の際にセンチメント指標があるように、感情スコアは金融商品価格の変動を分析する際に有効であると考えられる。深層学習の登場により感情スコアの推定精度が向上したことから、感情スコアの金融分野への応用は国内外で注目されている。

五島ら [3] は、深層学習を用いて算出したニュース記事の感情スコアを用いて、株価変動との関連性について分析を行い、ニュース記事が株価に影響を与えていることを示している。Batra et al.[4] は、2010 年から 2017 年の Apple 社の株価と Twitter のツイートの感情スコアを説明変数に用いて機械学習によって予測モデルを

構築したところ、価格データのみを説明変数とするモデルよりも価格騰落の予測精度が向上することを示している。Kalyani et al.[5] は、2014 年 5 月 1 日から 2015 年 4 月 30 日までのベトナムの株価に関するニュース記事の感情スコアを用いて、機械学習で予測モデルを構築し、翌日の価格騰落の予測に関しては、価格データのみを用いたモデルよりも高い予測精度が得られることを示している。Nayak et al.[6] は、2014 年から 2015 年までのインドの企業の株価に関するニュースの感情スコアとツイートの感情スコアを同時に用いて機械学習でモデル構築することにより、価格データのみを用いた予測モデルよりも翌日の騰落予測の精度を高められることを示している。

このような、テキスト情報を用いて株価を予測する手法は、近年、ビットコイン価格の予測にも適用されるようになってきている。Lamon et al.[7] は、2017 年 9 月 24 日から 2017 年 11 月 30 日のニュースの見出しの感情スコアとツイートの感情スコアを用いて機械学習によってビットコイン価格の予測モデルの構築を行い、価格データのみを用いたモデルよりビットコインの翌日の価格騰落の予測精度を高められることを示している。Karalevicius et al.[8] では、2014 年 1 月から 2017 年 9 月のビットコインに関連したニュース記事の感情スコアとビットコインの価格に関連があり、ビットコインの日間騰落予測モデルの精度向上に利用できることが示されている。Galeshchuk et al.[9] は、2016 年 2 月 1 日から 2016 年 2 月 29 日のビットコインに関連したツイートの感情スコアを用いて深層学習を用いて予測

¹感情分析とは、人工知能を用いてテキストを自然言語処理し、感情を定量的に評価する手法である。感情スコアとはその出力である。

モデル構築することで価格データのみでの騰落予測モデルよりビットコインの翌日の騰落予測の精度が向上することを示している。Kinderis et al.[10] は、2013 年から 2017 年のツイートの感情スコアとニュース記事の感情スコアを用いて深層学習によって構築した予測モデルが、価格データのみを用いるよりも翌日の価格騰落の予測精度が高くなることを示している。

本稿でも、Kinderis et al.[10] にならい、ニュース記事の感情スコアとツイートの感情スコアを説明変数として使用し、深層学習を用いて翌日のビットコイン価格の騰落を予測することを考えていくが、本稿で考えるモデルと上述の先行研究のモデルには以下の 4 つの相違点がある。

Kinderis et al.[10] は、ニュース記事として仮想通貨ニュースサイトである Coindesk の記事を用いている。しかし、Coindesk のみからビットコインに関するニュース記事を取得すると、ビットコインに対してポジティブなニュースが多くなるという偏りが生じることが知られているため、Bloomberg のニュース記事を用いることで精度の向上が期待できる可能性を指摘している。そこで、本稿では全てのニュースサイトのニュース記事を取得できる Bloomberg のニュース記事を用いる。

また、Baig et al.[11] は、Google トレンドで取得したビットコインの検索数のデータ（以下、Google トレンドのデータとする）を用いることでビットコインの価格変動をパターン分けできることを示している。このことは、Google トレンドデータがビットコインの価格変動に対して説明力をもつこと示唆している。そこで、本稿では、先行研究で説明変数として用いられたツイートやニュースの感情スコアに加え、Google トレンドのデータも説明変数に用いることによってさらなる精度の向上を目指す。

近年、ソーシャルネットワーキングサービスにおいては、インフルエンサーの影響力が非常に大きいことから、被リツイート（以下、RT とする）数・いいねの数・フォロワー数なども騰落予測に影響を与えると考えるのはひとつの可能性である。そこで、本稿では、ビットコインに関連したツイートの感情スコアを算出する際に、「RT 数」「いいねの数」「フォロワー数」を用いて重み付けを行い、予測モデルに用いる。また、ニュース記事の感情スコアも、その日のニュースの数で重み付けを行ったものを予測モデルに用いる。

さらに、先行研究では英語のニュースやツイートのテキストデータを用いて分析が行われていたが、日本の仮想通貨保有率は世界の中で最も高く、加えて、ビットコイン取引の約半数が日本円で行われているという現状を踏まえ（図 1・図 2 参照）、ビットコインに関連した日本語のニュース記事とツイートのデータを用いることとする。

本稿では、日本語のニュース記事の重み付き感情ス

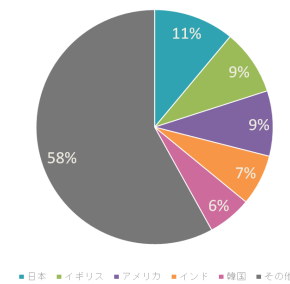


図 1: 仮想通貨の保有率 ([12] を参考に作成)

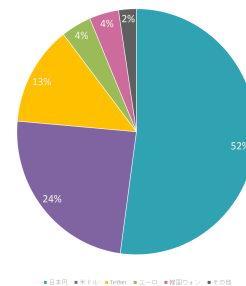


図 2: 国別ビットコイン取引量 ([13] を参考に作成)

コア、日本語のツイートの重み付き感情スコア、Google トレンドのデータを用いて、既存のビットコインの騰落予測手法の改良を目指す。

2 深層学習を用いたビットコイン価格の騰落予測モデル

2.1 データ概要

本研究で使用したデータは、CoinMarketCap²で提供されているビットコイン価格の日次データ、Bloomberg で提供されているビットコインに関する日本語のニュース記事、Twitter 社が提供している API を用いて取得した「ビットコイン」という語を含むツイート、それを投稿したアカウントに関する情報、Google 社が提供している Google トレンドの「ビットコイン」の検索数に関するデータである。いずれのデータも、データ期間は 2020 年 6 月 26 日から 2020 年 10 月 17 日までの 114 日間である。図 3 に 2020 年 6 月 26 日から 2020 年 10 月 17 日までのビットコインの日次価格の推移を示す。図 3 から、2020 年 7 月中旬まであまり大きな変動は見られなく停滞していたが、7 月末に上昇し始め 100 万を超えていることが見てとれる。8 月中旬には、2020 年最高額である 130 万円に達した。騰落に関しては、価格が上

²<https://coinmarketcap.com>

昇している日は 65 日、価格が下降している日は 49 日となっている。



図 3: ビットコイン価格の推移 (単位: 円)

ツイートのアカウントに関する情報は、ユーザー名、投稿日時、ツイート、いいね数、RT 数、フォロワー数、フォロワー数の 7 つの変数で構成されている。本研究では、投稿日時、ツイート、いいね数、RT 数、フォロワー数のみを用いる。

説明変数として扱うニュースとツイートのデータは Google Natural Language API を使って感情分析を行い、その出力である感情スコアを算出する。Google Natural Language API は機械学習で構築された強力な事前トレーニング済みモデルで、これを用いるとテキストの構造・感情を明らかにすることが可能である。感情スコアは、-1 から 1 までの値で算出される。-1 に近いほどネガティブで、1 に近いほどポジティブな感情を表す。算出した感情スコアは、次節で述べるように、ツイート情報で重み付けを行った上で、日毎に平均を求め、それらを説明変数として用いる。ニュースの感情スコアについても、次節で述べるように、その日のニュース数で重み付けを行った後に説明変数として用いる。

Google トレンドのデータは、「ビットコイン」の検索数が 0 から 100 に正規化された値をそのまま説明変数として用いる。

2.2 感情スコアの重み付け

先行研究では、1 日の全ニュース記事、全ツイートの感情スコアの平均をその日のニュース、ツイートの感情スコアとして説明変数に用いていた。本研究では、ツイートの感情スコアに関しては、インフルエンサーの影響を考慮するために、いいね数・RT 数・フォロワー数で重み付けを行う。いいね数・RT 数・フォロワー数はそれぞれ 0 以上 1 以下の値をとるように正規化する。重み付けの方法には以下の 2 パターンを考える。

$$\begin{aligned} & \text{ツイートの重み付き感情スコア 1} \\ &= \text{ツイートの感情スコア} \\ & \quad \times (\text{RT 数} + 1) \\ & \quad \times (\text{いいね数} + 1) \end{aligned}$$

$$\times (\text{フォロワー数} + 1) \quad (1)$$

$$\begin{aligned} & \text{ツイートの重み付き感情スコア 2} \\ &= \text{ツイートの感情スコア} \\ & \quad \times (\text{RT 数} + \text{いいね数}) \\ & \quad \times (\text{フォロワー数} + 1) \end{aligned} \quad (2)$$

1 日分の全ツイートに対して重み付き感情スコアを算出し、その平均値を、その日のツイートの重み付き感情スコアとする。

ニュース記事についても、その日のニュース数も投資家心理に影響すると考え、感情スコアをニュース数で重み付けしたものを説明変数に用いることを考えた。

$$\begin{aligned} & \text{ニュースの重み付き感情スコア} \\ &= \text{ニュースの感情スコア} \times \text{ニュース数} \end{aligned} \quad (3)$$

2.3 LSTM ネットワークを用いた騰落予測モデル

本稿では、過去 6 日分の価格と説明変数のデータを用いて、翌日のビットコイン価格の騰落を予測するモデルを構築する。目的変数は価格が 1 % 以上の上昇なら +1, 0 ~ 1 % の上昇なら +0, 0 ~ 1 % の下降なら -0, 1 % 以上の下降なら -1 とラベル化した価格の騰落とする。本稿で提案するモデルの説明変数は、過去 6 日分の価格、ニュース・ツイートの重み付き感情スコア、Google トレンドのデータであるが、ここでは、提案モデルの有用性を示すために、これらの一部を説明変数として用いたモデルと予測精度の比較を行う。こうしたモデルの中には、先行研究で用いられたモデルも含まれる。

2020 年 6 月 26 日から 9 月 25 日までの 92 日間のデータを訓練データ、9 月 26 日から 10 月 17 日までの 22 日間のデータをテストデータとし、活性化関数に softmax 関数、隠れ層数を 64 層、損失関数に交差エントロピー誤差を適用した LSTM ニューラルネットワークによってモデルを構築した。本稿で利用した深層学習のモデル構築には python の機械学習ライブラリである TensorFlow を用いた。

ここでは、22 日間のテストデータに対する騰落の正解率を騰落予測精度とする。また、価格変動率に関わらず上昇・下降のみの正解率を単純騰落予測精度と呼ぶことにする。

Kinderis et al.[10] では、ニュース記事の感情スコアとツイートの感情スコアの他に、説明変数として S&P500 の株価、イリジウムの価格、パラジウムの価格、アルミニウムの価格、コバルトの価格、Length Lumber Futures 社の株価を用いていたが、これらの変数は本研究のデータ期間では予測精度を向上させなかったため、本稿で

は、これらの変数を用いないモデルによる予測結果のみを示す。また、多層パーセプトロンによるモデルを構築も行ったが、LSTM ニューラルネットワークを用いた方が高精度だったため、本稿では LSTM ニューラルネットワークを用いた予測結果のみを示す。

3 騰落予測結果

3.1 予測結果

本節では、6 パターンの説明変数の組に対して学習した LSTM ニューラルネットワークで予測を行った結果を示す。以下では、各モデルの定義を順に述べ、図 4 から図 9 で、予測結果の混同行列を示す。混同行列の図の縦軸は価格変動のラベル、横軸は予測値を表し、色の濃さは度数を表す。

モデル 1 はビットコインの価格データのみを説明変数として用いたもので、22 日分のテストデータに対する騰落予測精度は 23 %、単純騰落予測精度は 50 % という結果であった。図 4 から、どのラベルも予測精度が高くないことが見て取れる。

モデル 2 は価格データに加えて Kinderis et al.[10] でも用いられていたニュースとツイートの重みなしの感情スコアを説明変数として組み込んだもので、騰落予測精度が 27 %、単純騰落予測精度は 55 % となった。価格データだけで構築したモデル 1 と比較すると、若干予測

精度が向上しているが、図 5 を見てもわかるように、モデル 1 と同様に高い精度で予測できているラベルはないことが見てとれる。

モデル 3 は、モデル 2 の価格データ、ニュース・ツイートの感情スコアに加えて Google トレンドのデータを説明変数として組み込んだもので、騰落予測精度が 32 %、単純騰落予測精度は 64 % となった。モデル 2 に Google トレンドのデータを説明変数に加えることで、騰落予測精度が 5 ポイント、単純騰落予測精度は 9 ポイント向上している。図 6 から、0~1 % の価格上昇と 1 % 以上の価格上昇ともに 50 % の確率で予測できていることがわかる。

モデル 4 は価格データ、ニュースの感情スコア、(1) 式のツイートの重み付き感情スコア 1、Google トレンドのデータを説明変数として組み込んだもので、騰落予測精度が 41 %、単純騰落予測精度は 73 % となった。ツイートの感情スコアにいいねの数・RT 数・フォロワー数で重み付けすることで、モデル 3 から騰落予測精度が 9 ポイント、単純騰落予測は 9 ポイント向上している。図 7 から、1 % 以上の価格上昇は 83 %、価格上昇は 100 % の精度で予測できていることがわかる。この結果は、インフルエンサーの影響を考慮に入れることで、予測精度が向上することを示唆している。

モデル 5 は、モデル 4 のニュースの感情スコアを、(3) の重み付き感情スコアで置き換えたものである。このモデルは、騰落予測精度が 41 %、単純騰落予測精度は

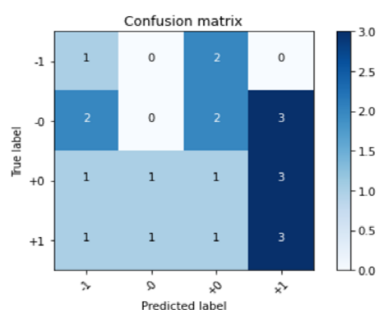


図 4: モデル 1 の予測結果

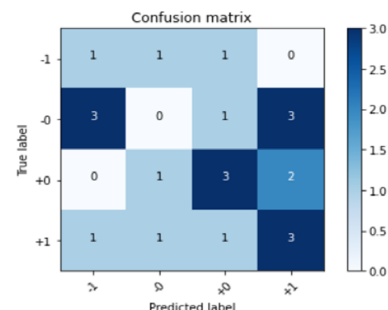


図 6: モデル 3 の予測結果

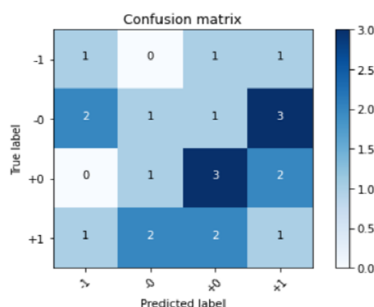


図 5: モデル 2 の予測結果

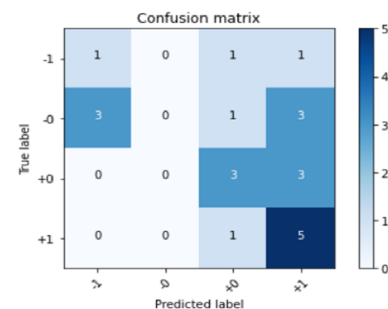


図 7: モデル 4 の予測結果

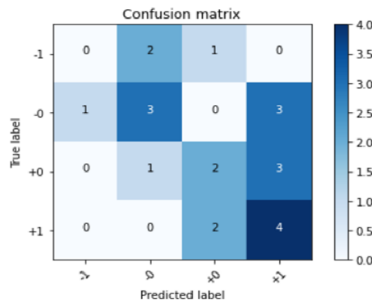


図 8: モデル 5 の予測結果

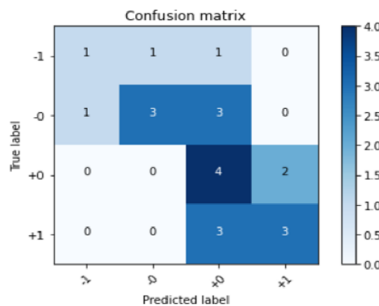


図 9: モデル 6 の予測結果

82 %となった。ニュースの感情スコアをニュース数で重み付けすることで、モデル 3 から騰落予測精度が 9 ポイント、単純騰落予測は 18 ポイント向上した。図 8 から、1 %以上の価格上昇は 67 %、価格上昇は 92 %、価格下降は 60 %の精度で予測できていることがわかる。

モデル 6 は、モデル 5 のツイートの重み付き感情スコアを (2) 式の重み付き感情スコアで置き換えたもので、騰落予測精度が 50 %、単純騰落予測精度は 82 %となった。ツイートの感情スコアにいいねの数と RT 数の和とフォロワー数の積で重み付けすることで、モデル 3 から騰落予測精度が 18 ポイント、単純騰落予測は 18 ポイント向上している。図 9 から、0~1 %の価格上昇は 67 %、価格上昇は 100 %、価格下降は 60 %の精度で予測できていることがわかる。モデル 6 が騰落予測精度、単純騰落予測精度ともに最も精度の高いモデルとなった。

3.2 考察

説明変数としてニュースの感情スコアとツイートの感情スコアを組み込んだ Kinderis et al.[10] のモデルは、価格のみのモデルよりも騰落予測精度が 4 ポイント、単純騰落予測精度は 5 ポイント向上するという結果になり、先行研究と同様に精度が向上した。このモデルに Google トレンドのデータを加えることによってさらに騰落予測精度が 5 ポイント、単純騰落予測精度は 9 ポイント向上するという結果になった。この結果

は、ビットコインの騰落予測における Google トレンドデータの有用性を示していると言える。ニュースの感情スコアとツイートの感情スコアで投資家の感情を捉えようと考えられる。一方、Google トレンドでの検索数は投資家のビットコインに対する興味・関心の量を表すと考えられるため、ビットコインの価格が上昇したときだけでなく、価格が下降したときにも増加する。つまり、Google トレンドのデータは、投資家の感情の強さを捉えようと考えられる。Google トレンドのデータを用いることで、より正確に人々の感情を捉えることが可能になり、ビットコインの騰落予測精度が向上したと考えられる。

また、Google トレンドデータに加え、重み付けしたニュース・ツイート感情スコアを組み込んだモデルは、より精度の高いものとなり、重み付けなしのニュース・ツイートの感情スコア、Google トレンドのデータを組み込んだモデルより、騰落予測精度が 9 ポイント、単純騰落予測精度は 18 ポイント向上するという結果になった。これは、仮説でもあったインフルエンサーの影響が大きいことと、数で重み付けすることでより正確に投資家の感情の強さを捉えられているためであると考えられる。このことから、ツイートの感情スコアをいいねの数・RT 数・フォロワー数で重み付けすることの有用性を示すことができたと言える。

ツイートの感情スコアも (1) 式、(2) 式の 2 パターン考えたが、(2) 式の重みの方が騰落予測精度が 9 ポイント向上するという結果になった。一般に、インフルエンサーとはフォロワー数が多い人のことを指す。しかし、人々の感情に影響を与えるという点においては、フォロワー数だけでなく、エンゲージメント率も重要となってくる。エンゲージメント率は以下の式で定義される。

エンゲージメント率

$$= (\text{クリック数} + \text{RT 数} + \text{フォロワー数} + \text{返信数} + \text{いいね数}) / \text{インプレッション} \times 100.$$

Twitter API では RT 数・いいねの数しか収集することができないため、本研究では RT 数といいねの数の和をそのツイートのエンゲージメントと定義する。この重みは、感情スコアにインフルエンサーを表すフォロワー数を掛け、さらにエンゲージメントを表す RT 数といいねの数の和を掛け合わせるによって、インフルエンサーによる拡散力に加えて、インフルエンサーのツイートに対する人々の反応をエンゲージメントとして表現したものと解釈できる。結果として、従来法に比べ精度が飛躍的に向上している点は興味深い。また、日本ではビットコインに関する詐欺アカウントが多く、そのようなアカウントはフォロワーを買っていて、本研究ではインフルエンサーとしての重み付けがされるが、それらはエンゲージメントが極めて低いため、(2) 式の

重みは悪質なツイートを除去するのにも役立ったとも考えられる。

4 おわりに

本稿では, Kinderis et al.[10] で扱われていた説明変数で構築したモデルのツイートの感情スコアに, いいね数・被リツイート数・フォロワー数で重み付けを行い, さらに Google トレンドのデータを説明変数に加えることで予測精度が向上するという結果を示した。この結果は, ツイートの感情スコアへの重み付けと Google トレンドのデータがビットコインの騰落予測において有用であることを示していると言える。

本稿では, ビットコインを対象としたが同じような予測手法が用いられている株などの金融商品にも本研究の手法は有用であると考えられる。

謝辞

本研究の一部は本研究は JSPS 科研費 JP17K00061 の助成を受けたものです。また, 本研究を行うにあたり, 同志社大学大学院文化情報学研究科の内藤宏明氏に有益なコメントをいただきました。ここに記して感謝の意を表します。

参考文献

- [1] 和泉潔, 後藤卓, 松井藤五郎 (2010). テキスト情報による金融市場変動の要因分析. 人工知能学会論文誌, **25**, 3, pp.383–387.
- [2] 和泉潔, 後藤卓, 松井藤五郎 (2011). 経済テキスト情報を用いた長期的な市場動向推定. 情報処理学会論文誌, **52**, 12, pp.3309–3315.
- [3] 五島圭一, 高橋大志 (2016). ニュースと株価に関する実証分析: ディープラーニングによるニュース記事の評判分析. 証券アナリストジャーナル, **54**, 3, pp.76–86.
- [4] Rakhi Batra and Sher Muhammad Daudpota (2018). Integrating StockTwits with sentiment analysis for better prediction of stock price movement. In *2018 International Conference on Computing, Mathematics and Engineering Technologies*, pp.1–5.
- [5] Joshi Kalyani, H.N.Bharathi and Rao Jyothi (2014). Stock trend prediction using news sentiment analysis. <https://arxiv.org/ftp/arxiv/papers/1607/1607.01958.pdf>, 2020 年 12 月 2 日閲覧.
- [6] Aparna Nayak, M. M. Manohara Pai and Radhika M. Pai (2016). Prediction Models for Indian Stock Market. *Procedia Computer Science*, **89**, pp.441–449.
- [7] Connor Lamon, Eric Nielsen and Eric Redondo (2017). Cryptocurrency Price Prediction Using News and Social Media Sentiment. *SMU Data Sci. Rev.*, **1**, pp.1–22.
- [8] Vytautas Karalevicius, Niels Degrande and Jochen De Weerd (2018). Using sentiment analysis to predict interday Bitcoin price movements. *Journal of Risk Finance*, **19**, pp.56–75.
- [9] Svitlana Galeshchuk, Oleksandra Vasylyshyn and Andriy Krysovaty (2017). Bitcoin Response to Twitter Sentiments. In *ICT in Education, Research and Industrial Applications*, Springer, pp.160–168.
- [10] Marius Kinderis, Marija Bezbradica and Martin Crane (2018). Bitcoin Currency Fluctuation. In *COMPLEXIS 2018 –3rd International Conference on Complexity, Future Information Systems and Risk*, pp.31–41.
- [11] Ahmed Baig, Benjamin M. Blau and Nasim Sabah (2019). Price clustering and sentiment in bitcoin. *Finance Research Letters*, **29**, pp.111–116.
- [12] Coincheck 「仮想通貨の日本人保有率は? 世界や国別の状況も紹介」, <https://coincheck.com/ja/article/130>, 2020 年 6 月 25 日アクセス.
- [13] Liquid 「ビットコイン (Bitcoin) の取引量は? 国別にみる取引量推移や日本が一位の理由」, <https://blog.liquid.com/ja/knowledge/bitcoin-world-trading-volume/>, 2021 年 2 月 2 日アクセス.

FinMegatron: Large Financial Domain Language Models

Xianchao Wu
NVIDIA

xianchaow@nvidia.com, wuxianchao@gmail.com

Abstract

General domain pretrained large-scale language models, such as BERT and GPT3, have achieved state-of-the-art results among numerous NLP classification and generation applications. This pretraining technology is also willing to be used in vertical domains, such as finance. The downstream applications include financial event extraction from news, summarization, and causal inferencing. In this paper, we propose large-scale pretrained BERT models for financial domain in English and Japanese languages. The original datasets come from professional financial news. We empirically study the factors of sub-word vocabulary set, model size and their impacts to the downstream financial NLP applications. The code and pretrained models are released from <https://github.com/NVIDIA/Megatron-LM>.

1 Introduction

Large-scale pretrained contextual language models (Qiu et al., 2020), such as ELMo (Peters et al., 2018), BERT (Devlin et al., 2019), GPT (Radford et al., 2018), ALBERT (Lan et al., 2020), Megatron (Shoeybi et al., 2019), and GPT3 (Brown et al., 2020), have achieved outstanding results in numerous classification and generation applications of NLP field. Leveraging the encoder blocks of *transformers* (Vaswani et al., 2017), BERT (Devlin et al., 2019) proposed the prediction of masked words or subwords and next sentences. These pretrained models only require large-scale textual documents as inputs without any additional manual annotations. Parameters include the number of heads used in multi-head self-attentions and the hidden size of point-wise feed-forward networks used in transformer’s encoder blocks. Comparing with this non-autoregressive framework of utilizing the encoder blocks, GPT models employ the decoder blocks of transformers. That is, a sentence is taken as input and the model is trained by predicting its word

step-by-step, one-by-one. In order to predict the next word, mask tensors are constructed by limiting the usage of left-hand-side visible words in a sentence. A list of NLP applications, such as machine translation (Stahlberg, 2020), machine reading comprehension (Zhang et al., 2020), document summarization (El-Kassas et al., 2021), sentiment analysis (Zhang et al., 2018), have achieved significantly better results in recent years, thank to these large-scale pretrained models (such as NVIDIA’s Megatron GPT2 which contains 8.3 billion parameters and OpenAI’s GPT3 which contains 160 billion parameters).

On the other hand, vertical domains, such as bioinformatics and finance, have large collections of textual documents, such as technical research papers, analysis reports and financial news. These domains require domain specific knowledge to tackle down-stream applications. In bioinformatics domain, there are researches such as NVIDIA’s BioMegatron (Shin et al., 2020) and Microsoft’s domain-specific language models for Biomedical NLP (Gu et al., 2021). We focus on building finance-domain pretrained language models in this paper. In this finance domain, there is limited work on quite limited languages. FinBERT (Liu et al., 2020) was proposed for financial text-mining in English language. 61GB textual data are collected from English wikipedia and BooksCorpus, FinancialWeb (13GB, 3.31B words), FinancialWeb (24GB, 6.38B words), YahooFinance (19GB, 4.71B words), and RedditFinanceQA (5GB, 1.62B words). In addition, FinBERT has two versions, *LARGE* and *BASE*, which share the same model settings of transformer and pre-training hyperparameters as BERT. NTT-data is one of the few companies that are known of developing financial BERT for Japanese language¹. NTT-data used 12.7GB textual data with a vocabulary size of 32K

¹<https://www.nttdata.com/jp/ja/news/release/2020/071000/>

	FinBERT	NTT-Data	FinMegatron
Language	En	Jp	En, Jp
Vocab	30K	32K	customize
Parameter	330M	330M	customize
Model open	No	No	Yes
Code open	No	No	Yes
Model parallel	NA	NA	Yes

Table 1: Comparison of BERT in financial domain.

that includes both words and sub-words (such as Japanese characters, kana). The default setting of BERT-LARGE is used in NTT-data’s financial BERT as well.

2 FinMegatron Configurations

We name our pretrained language models in financial domain as *FinMegatron*, in which *Fin* stands for financial domain and *Megatron*² (Shoeybi et al., 2019) is our general pretrained language models with model parallelism with open-source code written in PyTorch³, support mixed precision training (fp16 and fp32)⁴ under Apex⁵ and parallelism of data, pipeline and tensor.

Table 1 lists the differences of our FinMegatron and two baselines, FinBERT (Liu et al., 2020) and NTT-Data’s BERT. We support both English and Japanese languages in which Japanese tokenizer such as Mecab with IPAdict is integrated in our code. In addition, we support customizing vocabulary set so that vertical domain’s named entities can be better covered by training Byte-Pair Encoding (BPE)⁶ and WordPiece (Kudo, 2018) under the given large-scale textual dataset.

In order to support customized vocabulary size, we reuse the code *build_vocab.py* in Tohoku-U’s BERT⁷ for general domain. For example, in our Japanese BERT model, we set the vocabulary size to be 40k.

Besides the traditional usage of mini-batch for data-parallelism, we also leverage the tensor and pipeline parallelism. Tensor parallelism means that we can cut one tensor by its last dimension into several tensors and each tensor allocated in one

GPU. For example, if one tensor’s shape is (batch size=b, sequence length=s, hidden size=h) and we use two GPUs to respectively store half of it. Then each GPU will have a tensor of (b, s, h/2). We can define an activate function to implement this forward cutting and backward gathering.

For example, to implement the multi-layer perceptron (MLP) layer (i.e., affine linear projection of from h to 4*h and then back to hidden-size used in point-wise feed-forward layer) in transformer, we define two types of parallelism, column-parallel and row-parallel. For an input tensor X , column-parallel linear layer performs $Y = XA + b$ where matrix A is separated by columns alike $A = [A_1, A_2]$ (when p=2). Then we have $Y_1 = XA_1 + b$ and $Y_2 = XA_2 + b$. In forward process, X is not changed and broadcast to each GPU, in backward process, PyTorch’s *all-reduce* is performed to sum-up tensor values from different GPUs and sync their values to be the same and up-to-date. Typically, in the first affine linear projection from h to 4h, A_i will take a shape of (h, 4h/p) where p is the number of parallel process or number of GPUs for tensor parallelism. So we will have Y_i ’s shape of (b, s, h) * (h, 4h/p) = (b, s, 4h/p). The second affine linear projection is a row-parallel linear layer with original A to be (4h, h) and here cut by rows to be (4h/p, h). This means that now A is cut by rows and $A = [A_1, A_2]^T$ (when p=2).

To implement the tensor-parallelism for the multi-head self-attention layer, we can first separate by the number of heads and then each head’s tensor can be further paralleled by being separated following the final hidden size dimension. There are four linear layers in the multi-head self-attention layer, the first three linear layers are to project the input X into Q , K , and V tensors and the forth is a simple linear projection of the output of the multi-head self-attention tensor. We can combine the first three linear layers into one layer with weight matrix to be from the original 3*(h, h) to now (h, 3h). Then we can make use of the column-parallelism linear layer used in MLP layer. For the forth layer, it is actually a row-parallelism layer used in MLP layer.

There are 24 encoder blocks in BERT-LARGE and each block contains a multi-head self-attention layer followed by a point-wise feed-forward layer. We can further perform pipeline-parallelism to separate each For example, splitting a model with 24 transformer layers across 4 stages would mean each stage gets 6 transformer layers each. Through this

²<https://github.com/NVIDIA/Megatron-LM>

³<https://pytorch.org/>

⁴<https://docs.nvidia.com/deeplearning/performance/mixed-precision-training/index.html>

⁵<https://github.com/NVIDIA/apex>

⁶https://en.wikipedia.org/wiki/Byte_pair_encoding

⁷<https://github.com/cl-tohoku/bert-japanese>

	English	Japanese
file size	15GB	13GB
documents	5,516,610	5,221,226
sentences	97,122,204	40,467,789
tokens	2.38B	1.06B
vocab	7.18M	364K
vocab-bert	30K	40K
model size	335M	346M

Table 2: Statistical Information of the English and Japanese datasets for pre-training.

way, the forward and backward computing can be paralleled across each stage.

In our BERT models, we also follow the default settings of BERT-LARGE. The network contains 24 layers in which each layer contains a multi-head self-attention layer and a point-wise feed-forward layer with residual adding and layer normalization included as well. We set the max position embedding and maximum sequence length to be 512, number of attention heads to be 16, hidden size to be 1024. The whole dataset is separated into three subsets: 94.9% for training, 5% for validating and 1% for testing.

3 Experiments

3.1 Datasets

We extract finance related documents from large-scale multilingual datasets, such as Google’s C4⁸ and Facebook’s CC-100⁹ which are all originally collected from the web. In order to perform the filtering, we keep a large finance word list for each language. The word lists are collected from the web.

Table 2 lists the major statistical information of the English and Japanese datasets of financial domain. An independent document stands for a complete web page or a whole news or even a whole analysis report.

3.2 Training Details

Figure 1 and 2 illustrates the train/validation losses during training the BERT-LARGE models. We set 1 million iterations (with time costing to be 190ms/iteration) for the English dataset and 1 million iterations (with time costing to be 200ms/iteration) for the Japanese dataset.

For the English training, we employ a NVIDIA

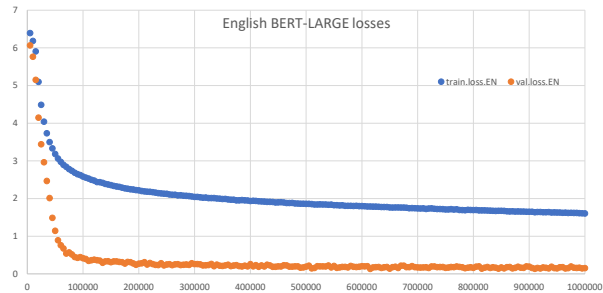


Figure 1: Training and validating losses for the English financial BERT-LARGE.

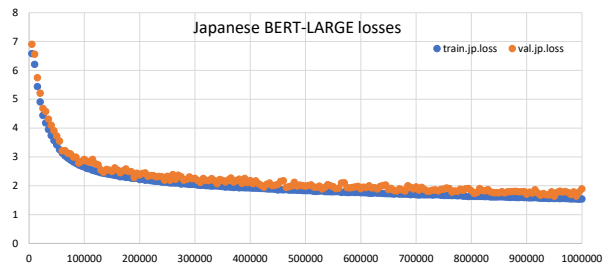


Figure 2: Training and validating losses for the Japanese financial BERT-LARGE.

DGX-1 workstation with 8 V100 cards¹⁰ (16GB memory each, micro-batch size to be 4 and global batch size to be 32), and for the Japanese training, we employ a NVIDIA DGX-1 workstation with 4 V100 cards (16GB memory each, micro-batch size to be 4 and global batch size to be 16). All these workstations come from a large-scale DGX-SUPERPOD¹¹ in NVIDIA. In addition, note that our open-source code is suitable of using large-scale textual datasets as well to train on NVIDIA’s large-scale clusters of DGX-SUPERPOD with as many as 3,072 A100 GPUs¹². The losses in English models drop from 6.5 around to less than 2 and finally at around 1.6 for training and 0.2 for validating. Finally, the 1% test set’s loss is 0.224. On the other hand, the losses in Japanese models drop from 7 around to around 2 in both the training and validating. Finally, the 1% test set’s loss is 1.648.

4 Conclusion

We have trained two BERT-LARGE BERT models in English and Japanese for financial domain using the Megatron open-source toolkit in NVIDIA. Our

⁸<https://www.tensorflow.org/datasets/catalog/c4>

⁹<http://data.statmt.org/cc-100/>

¹⁰<https://www.nvidia.com/en-us/data-center/v100/>

¹¹<https://www.nvidia.com/en-us/data-center/dgx-superpod/>

¹²<https://github.com/NVIDIA/Megatron-LM>

code and model support both customized vocabulary sizes and model paralleling among multi-GPU of single-node and even multi-nodes. Considering that this is still a work in progress, we will include the training of GPT models as well for English and Japanese financial domain. Finally, down-stream NLP applications such as market understanding (Wu, 2020) are also to be fine-tuned again with these pretrained BERT models.

References

- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. [Language models are few-shot learners](#). In *Advances in Neural Information Processing Systems*, volume 33, pages 1877–1901. Curran Associates, Inc.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Wafaa S. El-Kassas, Cherif R. Salama, Ahmed A. Rafea, and Hoda K. Mohamed. 2021. [Automatic text summarization: A comprehensive survey](#). *Expert Systems with Applications*, 165:113679.
- Yu Gu, Robert Tinn, Hao Cheng, Michael Lucas, Naoto Usuyama, Xiaodong Liu, Tristan Naumann, Jianfeng Gao, and Hoifung Poon. 2021. [Domain-specific language model pretraining for biomedical natural language processing](#).
- Taku Kudo. 2018. [Subword regularization: Improving neural network translation models with multiple subword candidates](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 66–75, Melbourne, Australia. Association for Computational Linguistics.
- Zhenzhong Lan, Mingda Chen, Sebastian Goodman, Kevin Gimpel, Piyush Sharma, and Radu Soricut. 2020. [Albert: A lite bert for self-supervised learning of language representations](#). In *International Conference on Learning Representations*.
- Zhuang Liu, Degen Huang, Kaiyu Huang, Zhuang Li, and Jun Zhao. 2020. [Finbert: A pre-trained financial language representation model for financial text mining](#). In *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI-20*, pages 4513–4519. International Joint Conferences on Artificial Intelligence Organization. Special Track on AI in FinTech.
- Matthew Peters, Mark Neumann, Mohit Iyyer, Matt Gardner, Christopher Clark, Kenton Lee, and Luke Zettlemoyer. 2018. [Deep contextualized word representations](#). In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pages 2227–2237, New Orleans, Louisiana. Association for Computational Linguistics.
- Xipeng Qiu, Tianxiang Sun, Yige Xu, Yunfan Shao, Ning Dai, and Xuanjing Huang. 2020. [Pre-trained models for natural language processing: A survey](#).
- Alec Radford, Karthik Narasimhan, Tim Salimans, and Ilya Sutskever. 2018. [Improving language understanding by generative pre-training](#).
- Hoo-Chang Shin, Yang Zhang, Evelina Bakhturina, Raul Puri, Mostofa Patwary, Mohammad Shoeybi, and Raghav Mani. 2020. [Biomegatron: Larger biomedical domain language model](#).
- Mohammad Shoeybi, Mostofa Patwary, Raul Puri, Patrick LeGresley, Jared Casper, and Bryan Catanzaro. 2019. [Megatron-lm: Training multi-billion parameter language models using model parallelism](#). *CoRR*, abs/1909.08053.
- Felix Stahlberg. 2020. [Neural machine translation: A review and survey](#).
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. [Attention is all you need](#). In I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems 30*, pages 5998–6008. Curran Associates, Inc.
- Xianchao Wu. 2020. [Event-driven learning of systematic behaviours in stock markets](#). In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 2434–2444, Online. Association for Computational Linguistics.
- Lei Zhang, Shuai Wang, and Bing Liu. 2018. [Deep learning for sentiment analysis : A survey](#).
- Zhuosheng Zhang, Hai Zhao, and Rui Wang. 2020. [Machine reading comprehension: The role of contextualized language models and beyond](#).

株価の残差リターンに注目した深層学習ポートフォリオ最適化

Deep Portfolio Optimization of Residual Factors in the Stock Market

今城 健太郎^{1*} 南 賢太郎¹ 伊藤 克哉¹ 中川 慧²

Kentaro Imajo, Kentaro Minami, Katsuya Ito, Kei Nakagawa

¹ 株式会社 Preferred Networks

¹ Preferred Networks, Inc.

² 野村アセットマネジメント株式会社

² Nomura Asset Management Co., Ltd.

Abstract: Recent developments in deep learning techniques have motivated intensive research in machine learning-aided stock trading strategies. However, since the financial market has a highly non-stationary nature hindering the application of typical data-hungry machine learning methods, leveraging financial inductive biases is important to ensure better sample efficiency and robustness. In this study, we propose a novel method of constructing a portfolio based on predicting the distribution of a financial quantity called residual factors, which is known to be generally useful for hedging the risk exposure to common market factors. The key technical ingredients are twofold. First, we introduce a computationally efficient extraction method for the residual information, which can be easily combined with various prediction algorithms. Second, we propose a novel neural network architecture that allows us to incorporate widely acknowledged financial inductive biases such as amplitude invariance and time-scale invariance. We demonstrate the efficacy of our method on U.S. and Japanese stock market data. Through ablation experiments, we also verify that each individual technique contributes to improving the performance of trading strategies. We anticipate our techniques may have wide applications in various financial problems.

1 はじめに

定量的な取引戦略の開発は金融分野における中心的課題である。近年の機械学習技術の発展に伴い、機械学習技術に基づく取引戦略の開発が注目されている [1, 2, 3]。一方、金融時系列の予測は一般には非常に難しい問題であると考えられている [4]。特に、金融市場のダイナミクスは非常に短い期間で敵対的に変動するということが指摘されており (例えば効率市場仮説 [5])、多くの機械学習手法の動作条件となる、定常かつ十分なサンプルサイズを伴ったデータを得ることは困難である。

一般に、データ効率のよい機械学習手法を開発するためには、データに特有の帰納バイアスを考慮したモデルを利用することが有効である。そこで、本研究では、金融分野における既存研究によって報告されている、金融時系列に関する以下の3つの性質に着目する。

1. 市場の共通因子に対するヘッジ: 金融分野におい

ては、株式投資のリターンの挙動は因子モデルによって説明されることが多い [6]。因子モデルによれば、特定の株式のリターンは、市場全体の指数などに連動する共通因子と、その銘柄特有の挙動をあらわす残差因子に分解される。取引戦略の構築においては、モメンタム効果などに基づく古典的な戦略は共通因子と高い相関をもち、経済危機など市場全体がもつリスクに晒されることが指摘されている [7]。一方、残差因子に基づく取引戦略を構築することで、共通因子の影響を削減しつつ正のリターンが得られうることが報告されている [7, 8]。

2. 分布仮定に基づくポートフォリオ構築: ポートフォリオの「良さ」を評価するために、与えられたリスク回避度のもとでリターンを最大化するという規準を考える。ポートフォリオ理論では、リターンの分布の性質に基づいて最適ポートフォリオを構築するための枠組みが与えられている [9]。特に、将来リターンの平均および分散の予測が与

*連絡先: 東京都千代田区大手町 1-6-1 大手町ビル
E-mail: imos@preferred.jp

えられれば、それに基づいたポートフォリオを客観的に構築することが可能である。

3. **金融時系列におけるスケール不変性:** 金融時系列は、ボラティリティ・クラスタリング [10] やフラクタル性 [11] といった性質を持つことが指摘されている。これらは、それぞれ時系列データの振幅および時間スケールに対するある種の不変性として解釈することができる。

本論文では、ニューラルネットワークによる残差因子の分布予測に基づいてポートフォリオを構築する手法を提案する。提案手法のアイデアは次の通りである。上記 1 の性質から、リターンの残差因子の情報を利用することで、市場の共通因子に対して相関の低いポートフォリオが構築できることが期待される。また、上記 2 の性質から、分布予測を経由したポートフォリオ構築の枠組みを考えることで、ある種の最適性をもったポートフォリオを客観的に構築することが可能である。さらに、分布予測においては、上記 3 の性質を考慮したネットワーク構造を提案することで、効率的な学習をめざす。また、本論文では、日本の株式市場データを用いた実証分析によって、様々なベースラインの手法に対する提案手法の有効性を示す¹。

2 問題設定

過去の株価の観測に基づいてポートフォリオを構築する問題を考える。\$S\$ 個の銘柄に対応する株価

$$\mathbf{p}^{(i)} = \{p_1^{(i)}, p_2^{(i)}, \dots, p_t^{(i)}, \dots\}$$

が観測できるとする。ただし、\$p_t^{(i)}\$ は銘柄 \$i\$ の時刻 \$t\$ における価格である。また、銘柄 \$i\$ の時刻 \$t\$ におけるリターンを

$$r_t^{(i)} = \frac{p_{t+1}^{(i)}}{p_t^{(i)}} - 1$$

と定義する。以下本論文では、株価そのものではなく、リターンの性質について考察する。

ポートフォリオとは、\$S\$ 個の銘柄に対する (時間依存する) 重みベクトル \$\mathbf{b}_t = (b_t^{(1)}, \dots, b_t^{(i)}, \dots, b_t^{(S)})\$ のこととする。ただし、\$b_t^{(i)}\$ は銘柄 \$i\$ に対する投資比率を表し、\$\sum_{i=1}^S |b_t^{(i)}| = 1\$ を満たすものとする。ポートフォリオ \$\mathbf{b}_t\$ が与えられたとき、時刻 \$t\$ におけるポートフォリオのリターンを

$$R_t = \sum_{i=1}^S b_t^{(i)} r_t^{(i)}$$

とする。本論文で考察する問題は、過去の時刻におけるリターンの観測に基づいてポートフォリオ \$\mathbf{b}_t\$ の値を決定し、\$R_t\$ を最大化することとして定式化される。また、本論文では正味の投資金額がゼロであるという \$\sum_{i=1}^S b_t^{(i)} = 0\$ を満たすゼロ金額投資ポートフォリオのパフォーマンスを用いて評価を行う。

3 提案手法

本節では提案システムの概要について説明する。図 1 に提案システムの概要を示す。提案システムでは、まず (i) スペクトラル残差によって残差因子の情報を抽出し (3.1 節)、次に (ii) ニューラルネットワークによる残差因子の分布予測を行い (3.2 節)、最後に (iii) 分布情報に基づいた最適ポートフォリオの構築を行う (3.4 節)。また、(ii) の分布予測においては、金融時系列特有の構造を考慮したネットワーク構造を提案する。これについては 3.3 節で解説する。

3.1 残差因子の抽出

本論文では、株式のリターンの残差因子に基づく取引戦略を開発する。冒頭で説明したように、残差因子は複数銘柄間で共有される要因をヘッジした残りの情報に基づいた因子であるため、市場がもつリスクに対してロバストなポートフォリオを構築できることが期待される。

以下では、観測データから残差因子を抽出する方法を簡単に説明する。時刻 \$t\$ における株 \$S\$ のリターンを

$$\mathbf{r}_t = (r_t^{(1)}, \dots, r_t^{(S)})^\top$$

とする。リターンの挙動を説明するモデルとして、因子モデル

$$\mathbf{r}_t = \mathbf{B} \mathbf{f}_t + \boldsymbol{\epsilon}_t \quad (1)$$

を考える。ただし、\$\mathbf{f}_t \in \mathbb{R}^K\$ は \$K\$ 個の共通因子を表すベクトル、\$\mathbf{B}\$ は \$S \times K\$ 行列、\$\boldsymbol{\epsilon}_t\$ は残差因子を表す。我々の目的は、リターン \$\mathbf{r}_t\$ の観測に基づいて残差因子 \$\boldsymbol{\epsilon}_t\$ を復元することである。

残差因子の抽出に関する既存手法について簡単に触れておく。(1) は動的因子モデルと呼ばれることがあり [13]、共通因子 \$\mathbf{f}_t\$ や残差因子 \$\boldsymbol{\epsilon}_t\$ の分布に関する様々な仮定のもとで因子分解の手法が複数提案されている。多くの場合、このような分解は因子分析 (FA) や主成分分析 (PCA) と類似のスペクトル分解法に対して時刻 \$t\$ に関する相関構造を考慮した拡張を行うことで得られる。しかし本論文では、時系列の相関構造の複雑なモデル化を避け、単純化したスペクトル分解に基づく抽出手法について考察する。

¹本稿は [12] の要約であり、詳細は <https://arxiv.org/abs/2012.07245> を参照。

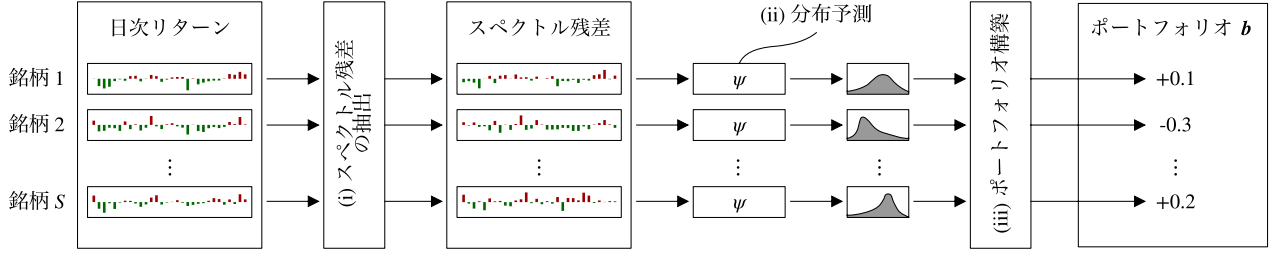


図 1: 提案システムの概要. 提案システムは (i) 残差因子の抽出 (ii) ニューラルネットワークに基づく残差因子の分布予測 (iii) ポートフォリオ理論に基づく最適ポートフォリオの構築の 3 つの機能からなる.

$H > 1$ を窓幅として, 行列

$$\mathbf{X}_t = [\mathbf{r}_{t-H}, \dots, \mathbf{r}_{t-1}] = \begin{bmatrix} r_{t-H}^{(1)} & \cdots & r_{t-1}^{(1)} \\ \vdots & \ddots & \vdots \\ r_{t-H}^{(S)} & \cdots & r_{t-1}^{(S)} \end{bmatrix}$$

を定義する. また, $\tilde{\mathbf{X}}_t$ を行ごとの平均を引いて中心化した行列とする. $C (< S)$ 共通因子の数に相当するパラメータとする. 以上のもとで, スペクトル残差 (spectral residual) を以下のように定義する. まず, $\tilde{\mathbf{X}}_t$ の特異値分解を

$$\tilde{\mathbf{X}}_t = \mathbf{V}_t \text{diag}(\lambda_1, \dots, \lambda_S) \mathbf{U}_t$$

とする. ただし, $\lambda_1 \geq \dots \geq \lambda_S$ は降順に並べた特異値である. このとき, スペクトル残差を小さい方から数えた $S-C$ 主成分とする. すなわち, 時刻 $t-H \leq s \leq t-1$ に対して, 残差因子ベクトルを

$$\tilde{\epsilon}_s = \mathbf{A}_t \mathbf{r}_s \quad (2)$$

と定義する. ただし,

$$\mathbf{A}_t = \mathbf{V}_t \text{diag}(\underbrace{0, \dots, 0}_C, \underbrace{1, \dots, 1}_{S-C}) \mathbf{V}_t^\top$$

は「残差空間」への射影行列である. この定義は, 直感的には, 時間について局所的には $\tilde{\mathbf{X}}_t$ の C 個の主成分で張られる空間に共通因子に関する情報が多く含まれ, 残りの成分であるスペクトル残差には残差因子に関する情報が多く含まれることを期待したものである.

3.2 残差因子の分布予測

次に, 過去の残差因子の系列 $\tilde{\epsilon}_{i,t-H}, \dots, \tilde{\epsilon}_{i,t-1}$ が抽出されたもとで, 次の時刻における $\tilde{\epsilon}_{i,t}$ の分布を予測する問題を考える. ここでは, 分位点回帰 [14] を利用した手法について説明する. 分位点回帰は, 条件付き分位点を推定するために広く利用されている手法である. 直感的には, 未来の時刻における残差因子 $\tilde{\epsilon}_{i,t}$ に

対して, 十分多い数の分位点を予測することができれば, 分布に関する様々な情報を復元することが可能である.

まず, 与えられた $\alpha \in (0, 1)$ に対して, 条件付き α -分位点を予測する方法について述べる. 簡単のため $y_t = \tilde{\epsilon}_{i,t}$ および $\mathbf{x}_t = (\tilde{\epsilon}_{i,t-H}, \dots, \tilde{\epsilon}_{i,t-1})$ と表す. y_t の条件付き分位点を予測する関数 $\psi_\alpha: \mathbb{R}^H \rightarrow \mathbb{R}$ は, 次の損失関数の最小化によって得られる:

$$\mathbb{E}_{y_t, \mathbf{x}_t} [\ell_\alpha(y_t, \psi_\alpha(\mathbf{x}_t))].$$

ただし $\ell_\alpha(y, y')$ は

$$\ell_\alpha(y, y') = \max\{(\alpha - 1)(y - y'), \alpha(y - y')\}$$

によって定義される pinball loss である. 次に, $\alpha_j = j/Q$ ($j = 1, \dots, Q-1$) を与えられたグリッドとして, α_j -分位点を同時に予測したい. そこで, 関数 $\psi: \mathbb{R}^H \rightarrow \mathbb{R}^{Q-1}$ を, 以下の目的関数を最小化するように訓練する

$$\mathcal{L}_Q(y_t, \psi(\mathbf{x}_t)) = \sum_{j=1}^{Q-1} \ell_{\alpha_j}(y_t, \psi_j(\mathbf{x}_t)).$$

上記の方法によって $Q-1$ 個の条件付き分位点 $\tilde{y}_t^{(j)} = \psi_j(\mathbf{x}_t)$ ($j = 1, \dots, Q-1$) が得られたもとでは, y_t の分布に関する情報を以下のように推定することができる. たとえば, 平均は

$$\hat{\mu}_t = \hat{\mu}(\tilde{\mathbf{y}}_t) = \frac{1}{Q-1} \sum_{j=1}^{Q-1} \tilde{y}_t^{(j)} \quad (3)$$

によって推定できる. また, 分散は, 分位点の不偏分散

$$\begin{aligned} \hat{\sigma}_t &= \hat{\sigma}(\tilde{\mathbf{y}}_t) \\ &= \text{Variance}(\tilde{\mathbf{y}}_t) \\ &= \frac{1}{Q-1} \sum (\tilde{y}_t^{(j)} - \hat{\mu}_t)^2 \end{aligned} \quad (4)$$

や, ロバストな推定量 (Median Absolute Deviation など) によって推定することができる.

3.3 金融時系列を考慮したネットワーク構造

分位点の予測器 $\psi: \mathbb{R}^H \rightarrow \mathbb{R}^{Q-1}$ として、ニューラルネットワークに基づくモデルを利用することを考える。本論文では、金融時系列について知られている2つの不変性を考慮したネットワーク構造を導入する。

3.3.1 ボラティリティ不変性

まず、振幅に関する不変性について考察する。金融時系列においてはボラティリティ・クラスタリングと呼ばれる性質があり、リターンの変動の後には同程度の大きさの変動が追従しやすいということが知られている [15]。このため、ある信号が金融時系列として観測されるならば、正の定数倍をして得られる信号も金融時系列として尤もらしいものになると考えられる。

この不変性を効率的に利用するために、予測器が positive homogeneous であるという制約を導入する。ただし、関数 $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ が positive homogeneous であるとは、任意の入力 $\mathbf{x} \in \mathbb{R}^n$ と正の定数 $a > 0$ に対して $f(a\mathbf{x}) = af(\mathbf{x})$ となることをいう。一般に、positive homogeneous な関数は以下のような正規化によって構成できる。 $\tilde{\psi}: S^{H-1} \rightarrow \mathbb{R}^{Q-1}$ を球面 $S^{H-1} = \{\mathbf{x} \in \mathbb{R}^H : \|\mathbf{x}\| = 1\}$ 上の任意の関数とすると、

$$\psi(\mathbf{x}) = \|\mathbf{x}\| \tilde{\psi}\left(\frac{\mathbf{x}}{\|\mathbf{x}\|}\right)$$

によって定義される関数は positive homogeneous になる。よって、この正規化によって、球面上の関数を表現する任意のモデルを、振幅に関する不変性をもつモデルに変換することができる。

3.3.2 フラクタル性

次に時間スケール不変性について考える。株価の時系列にはフラクタル構造があることが知られている [11]。ここで、フラクタル構造とは特徴的な時間スケールがないことであり、ある時系列を異なるサンプリングレートで観測した場合にそれぞれの形状を区別することができないことを意味する。

本論文では、予測問題においてフラクタル構造を活用するために、フラクタルネットワークと呼ばれるネットワーク構造を提案する。図2にフラクタルネットワークの構造を示す。提案手法の基本的なアイデアは、複数の異なるサンプリングレートのもとで観測した信号に対して同一の操作を適用することで、時間スケール不変性を考慮した予測を行うというものである。提案手法は、(a) リサンプリング機構および (b) 2つのニューラルネットワーク ψ_1, ψ_2 からなる。リサンプリング機構 $\text{Resample}(\mathbf{x}, \tau)$ では、与えられたサンプリング

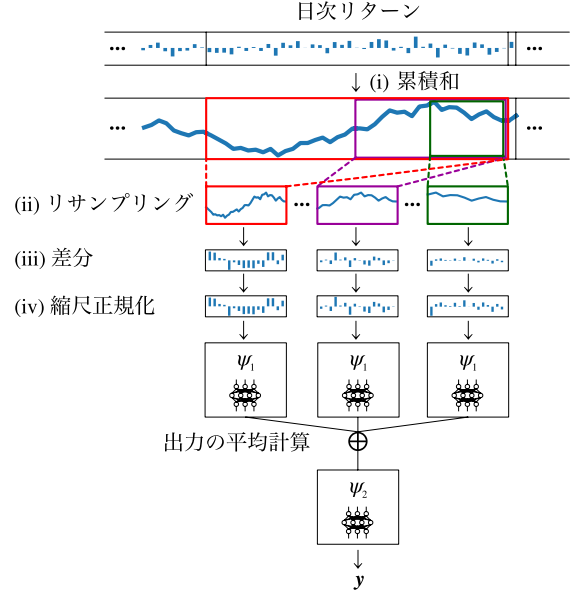


図2: フラクタルネットワークにおける処理。

レート τ に対応する系列をサンプリングと正規化によって生成する (図2の (i)–(iv))。まず、与えられたサンプリングレートの列 $\tau_1 < \dots < \tau_L (= 1)$ に対して $\text{Resample}(\mathbf{x}, \tau_i)$ を適用し、 L 個の系列を生成する。次に、 L 個の系列それぞれに対して共通の変換 ψ_1 を適用する。最後に、これらの系列の平均をとることによって情報を集約し、変換 ψ_2 を適用した結果を出力する。以上を数式としてまとめると、フラクタルネットワークの処理は以下のように表せる

$$\psi(\mathbf{x}) = \psi_2 \left(\frac{1}{L} \sum_{i=1}^L \psi_1(\text{Resample}(\mathbf{x}, \tau_i)) \right). \quad (5)$$

3.4 ポートフォリオ構築

ここでは、分布予測の結果に基づいて最適なポートフォリオを構築する方法について説明する。期待リターンと共分散行列の推定値をそれぞれ $\hat{\mu}_t$ と $\hat{\Sigma}_t$ とする。このとき、 $\lambda > 0$ をリスク回避度として、最適なポートフォリオは

$$\hat{\mathbf{b}}_t = \frac{1}{\lambda} \hat{\Sigma}_t^{-1} \hat{\mu}_t$$

と書ける [16]。

今回は、残差リターンの各成分はほぼ無相関であるという仮定に基づき、共分散行列は対角行列によって近似するものとする。すると、 j 番目の残差リターンのウェイトは

$$\hat{b}_j^{\text{res}} = \frac{\hat{\mu}_{t,j}}{\lambda \hat{\sigma}_{t,j}}$$

によって与えられる。ただし、 $\hat{\mu}_{t,j}$ は式 (3), および $\hat{\Sigma}_t$ は式 (4) によって与えられる平均および分散の推定値である。このように得られた残差リターンに対応するポートフォリオ $\mathbf{b}^{\text{res}}$ を元の株式に対するポートフォリオ $\mathbf{b}$ に変換するためには、関係式 (2) に基づいて

$$\mathbf{b}_t = \mathbf{A}_t^\top \mathbf{b}_t^{\text{res}}$$

とすればよい。さらに、ゼロ金額投資ポートフォリオに変換するには、ポートフォリオの平均

$$\bar{b}_t = \frac{1}{S} \sum_{i=1}^S b_{t,i}$$

を各ウェイトから引き、ウェイトの合計が 1 となるように正規化すればよい。

4 数値実験

アメリカ株式市場のデータとして、S&P 500 指数構成銘柄を使用した実証分析を行った。2000 年 1 月から 2007 年 12 月までをトレーニング期間として用い、2008 年 1 月から 2020 年 4 月までをテスト期間として使用した。ベースラインの手法として、単純バイアンドホールド (Market), AR(1) モデル, リッジ回帰 (Linear), ニューラルネットワーク (MLP), State Frequency Memory RNNs を元にした株価予測の最先端モデル (SFM) を用いた。

また、提案手法 (DPO) の各要素の寄与を調べるため以下の切除モデルとの比較も行った。

- DPO-NQ … 提案手法の分布予測モデルを用いずに平均を L_2 誤差最小化によって求めたモデル。
- DPO-NF … 提案手法のフラクタルネットワークを多層パーセプトロンに置き換えたモデル。
- DPO-NV … 提案手法のボラティリティ不変性のための正規化を行わないモデル。

評価指標としては、各手法を用いて構築されたポートフォリオのトータルリターン (CW), 年率リターン (AR), リスクを表す標準偏差 (AVOL), リスクー単位あたりのリターンを表す Sharpe レシオ ($\text{ASR} := \text{AR} / \text{AVOL}$), ポートフォリオの全期間での最大の下落幅を表す最大ドローダウン (MDD) を用いた。

表 1 が評価指標のサマリーである。提案手法が AR を除くすべての評価指標においてベースラインを上回っている。図 3 が CW の推移を各手法で比較したものである。CW もベースラインを安定的に上回り、かつ全期間を通じて右上がりの推移となっていることから、提案手法の有効性が確認できる。

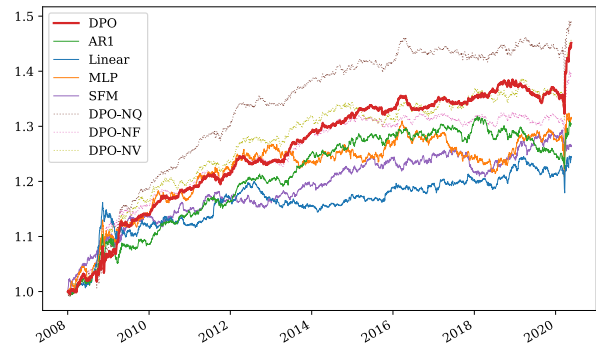


図 3: アメリカ株式市場における各手法の CW の推移

表 1: アメリカ株式市場における性能評価

Method	ASR	AR[%]	AVOL[%]	MDD[%]
Market	+0.607	+0.130	0.215	0.496
AR(1)	+0.858	+0.021	0.025	0.072
Linear	+0.724	+0.017	0.024	0.059
MLP	+0.728	+0.022	0.030	0.077
SFM	+0.709	+0.019	0.026	0.058
DPO-NQ	+1.237	+0.032	0.026	0.063
DPO-NF	+1.284	+0.027	0.021	0.042
DPO-NV	+1.154	+0.030	0.026	0.053
DPO (Proposed)	+1.393	+0.030	0.021	0.045

また、日本株式市場のデータに対しても同様の実験を行った。日本株式市場のデータとしては、TOPIX 500 指数構成銘柄を使用し実証分析を行った。2005 年 1 月から 2012 年 12 月までをトレーニング期間とし、2013 年 1 月から 2018 年 12 月までをテスト期間として使用した。表 2 の通り提案手法が AR を除くすべての評価指標においてベースラインを上回る同様の結果が得られた。図 4 が CW の推移を各手法で比較したものである。CW もベースラインを安定的に上回り、かつ全期間を通じて右上がりの推移となっていることから、提案手法の有効性が確認できる。

表 2: 日本株式市場における性能評価

Method	ASR	AR[%]	AVOL[%]	MDD[%]
Market	+0.819	+0.158	0.193	0.248
AR(1)	+1.835	+0.034	0.019	0.031
Linear	+1.380	+0.023	0.017	0.024
MLP	+1.802	+0.030	0.017	0.032
SFM	+0.250	+0.005	0.019	0.042
DPO-NQ	+1.369	+0.024	0.018	0.032
DPO-NF	+1.770	+0.030	0.017	0.025
DPO-NV	+1.979	+0.039	0.020	0.026
DPO	+2.171	+0.036	0.017	0.025

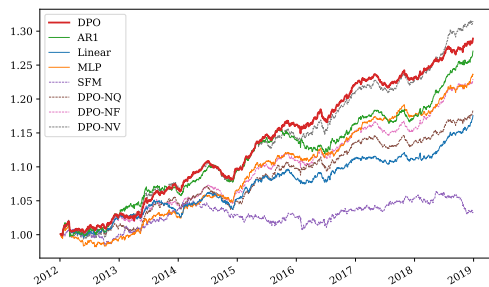


図 4: 日本株式市場における各手法の CW の推移

5 結論

本論文では、残差因子の分布予測に基づく新しいポートフォリオ構築手法を提案した。提案手法では、単純なスペクトル分解法によって残差因子の情報を抽出し(スペクトル残差)、金融時系列の不変性を考慮した新しいネットワーク構造を利用して分布予測を行う。日本株式市場を対象とした実証分析から、提案手法が一定の有効性をもつことが確認された。

参考文献

- [1] Jingyuan Wang, Yang Zhang, Ke Tang, Junjie Wu, and Zhang Xiong. Alphastock: A buying-winners-and-selling-losers investment strategy using interpretable deep reinforcement attention networks. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pp. 1900–1908, 2019.
- [2] Rohit Choudhry and Kumkum Garg. A hybrid machine learning system for stock market forecasting. *World Academy of Science, Engineering and Technology*, Vol. 39, No. 3, pp. 315–318, 2008.
- [3] Vatsal H Shah. Machine learning techniques for stock prediction. *Foundations of Machine Learning—Spring*, Vol. 1, No. 1, pp. 6–12, 2007.
- [4] Christopher Krauss, Xuan Anh Do, and Nicolas Huck. Deep neural networks, gradient-boosted trees, random forests: Statistical arbitrage on the s&p 500. *European Journal of Operational Research*, Vol. 259, No. 2, pp. 689–702, 2017.
- [5] Burton G Malkiel and Eugene F Fama. Efficient capital markets: A review of theory and empirical work. *The journal of Finance*, Vol. 25, No. 2, pp. 383–417, 1970.
- [6] Eugene F. Fama. Efficient capital markets: A review of theory and empirical work. *The Journal of Finance*, Vol. 25, No. 2, p. 383, May 1970.
- [7] David Blitz, Joop Huij, and Martin Martens. Residual momentum. *Journal of Empirical Finance*, Vol. 18, No. 3, pp. 506–521, 2011.
- [8] David Blitz, Joop Huij, Simon Lansdorp, and Marno Verbeek. Short-term residual reversal. *Journal of Financial Markets*, Vol. 16, No. 3, pp. 477–504, 2013.
- [9] Harry Markowitz. Portfolio selection. *The journal of finance*, Vol. 7, No. 1, pp. 77–91, 1952.
- [10] Thomas Lux and Michele Marchesi. Volatility clustering in financial markets: a microsimulation of interacting agents. *International journal of theoretical and applied finance*, Vol. 3, No. 04, pp. 675–702, 2000.
- [11] Edgar E Peters. *Fractal market analysis: applying chaos theory to investment and economics*, Vol. 24. John Wiley & Sons, 1994.
- [12] Kentaro Imajo, Kentaro Minami, Katsuya Ito, and Kei Nakagawa. Deep portfolio optimization via distributional prediction of residual factors. *Proceedings of the AAAI Conference on Artificial Intelligence (to be appear)*, 2021.
- [13] James H. Stock and Mark W. Watson. *Dynamic Factor Models*. Oxford University Press, Jul 2011.
- [14] Roger Koenker. *Quantile Regression*. Econometric Society Monographs. Cambridge University Press, 2005.
- [15] Benoit B Mandelbrot. The variation of certain speculative prices. In *Fractals and scaling in finance*, pp. 371–418. Springer, 1997.
- [16] Raymond Kan and Guofu Zhou. Optimal portfolio choice with parameter uncertainty. *Journal of Financial and Quantitative Analysis*, Vol. 42, No. 3, pp. 621–656, 2007.

状態空間モデルを用いたバランス型投資信託のリバランス推定

Rebalance Estimation of Multi-Asset Mutual Fund using State Space Model

今井 崇公^{1*}
Takahiro Imai¹

¹ 野村アセットマネジメント株式会社

¹ Nomura Asset Management Co.,Ltd.

Abstract: ベンチマークのないバランス型運用では、しばしば自身と類似した運用を行うプレイヤーとのパフォーマンス比較が重視されるため、他のプレイヤーのポジションを把握することが相対的な運用リスク管理の観点から有用である。しかし多くのヘッジファンド複製に関する先行研究で示されるように、運用の自由度の高いポートフォリオのポジションを高精度に推定することは容易ではない。そこで本研究ではヘッジファンドより投資制約の強いバランス型投資信託をケーススタディとして、高精度なウェイト推定を実現する手法を提案し、人工データ実験によりその有効性を示す。具体的には、一般的なショート・レバレッジ制約付きの投資信託のウェイト変動が時価変動要因とリバランス要因から構成されることを、状態空間モデルを用いて明示的に組み込むことにより、ウェイトひいてはリバランスまで高精度に推定できることを明らかにする。

1 はじめに

バランス型運用の客観的なパフォーマンス評価は、多くの場合にベンチマークが与えられていないことから容易ではない。絶対リターンやシャープレシオなどを用いて絶対的な評価をすることも可能ではあるが、データ期間の相場つきの影響を強く受けることから複数の期間のパフォーマンスを中立的に評価できているとは言い難い。この問題は、ウェイト制約などによりベータの影響を受けやすい運用ほど顕著になる。こうした状況から、バランス型運用ではしばしば自身と類似した運用を行う他のプレイヤーとのパフォーマンス比較が重視される。しかしベンチマーク対比でのパフォーマンス尺度の場合と異なり、比較対象である競合の運用商品のポジションを即時的に知ることができないことから、競合比を前提とした運用のリスク管理は行えず、事後的なパフォーマンス評価に終始してしまうという問題が存在する。

そこで本研究では、バランス型運用の相対的なリスク管理を企図して、他のプレイヤーのポジションを高精度に推定する手法を提案する。ここで相対的なリスク管理とは、一般的な株式運用におけるトラッキングエラーやアクティブシェアの管理に相当する概念である。

しかし、多くのヘッジファンド複製に関する先行研究で示されるように、ヘッジファンドのような運用の自由度が高く、透明性の低い運用商品のポジションを高精度に推定することは容易ではない。そこで本研究では、ヘッジファンドより投資制約が強く透明性の高いバランス型投資信託をケーススタディとして、高精度なウェイト推定の実現を目指す。本ケースでは、ショート・レバレッジ制約¹付きの投資信託を想定しており、デリバティブが用いられないことからウェイト変動は時価変動要因とリバランス要因に明確に分離される。この性質を状態空間モデルを用いて明示的に組み込むことにより、高い精度のウェイト推定ひいてはリバランス推定を実現する手法を提案する。

本稿の構成は次の通りである。2章では、本研究に関連するヘッジファンド複製に関する先行研究について簡単に述べる。3章では、線形回帰、ベイズ線形回帰、カルマンフィルタ、粒子フィルタを用いたウェイトの推定手法を導入する。そして4章において、バランス型投資信託を想定した人工データを用いた数値実験によってそれらの有効性を評価する。最後に、結論を述べる。

*連絡先：野村アセットマネジメント株式会社
〒135-0061 東京都江東区豊洲二丁目2番1号
E-mail: takahiro.imai.biz@gmail.com

本稿の内容は全て筆者に属し、所属する組織としての見解を示すものではない。また本稿にあり得べき誤りは、全て筆者の責に帰する。

¹ウェイトが非負で、ウェイト和が1となるような投資制約で、最適化の分野では単体制約と呼ばれる。投資信託では、このようなウェイト制約が課されることが多い。

表 1: 一般のヘッジファンドと本研究の対象の比較

	一般のヘッジファンド	本研究の対象
投資対象の可変性	有	無
投資対象の非線形性（オプションなど）	有	無
レバレッジ	有	無
ショート	有	無
リターンの開示頻度	四半期, 月次	日次

2 先行研究

ポートフォリオのウェイト推定問題の関連研究として、ヘッジファンド複製に関する研究が古くから数多く行われている。ヘッジファンド複製の主なモチベーションは、ヘッジファンドのパフォーマンスを低コストで複製することであり、大きく3つのアプローチに分かれる[1]。第一のアプローチは、ヘッジファンドの公表データ、金融指標データなどから、ファクターモデルなどを用いてヘッジファンドのリターンを複製する手法であり、帰納的に毎期のリターンを複製する。第二のアプローチは、対象とするヘッジファンドが行っていると想定される運用戦略をボトムアップに構築することでリターンを複製する手法であり、演繹的に毎期のリターンを複製する。第三のアプローチは、ヘッジファンドのリターンの統計的性質を規定し、その特性を満たすような戦略を構築する手法であり、リターン分布を複製する。これらの中でポートフォリオのウェイト推定に直接的に関連するのは、第一のアプローチであり、多くの場合にファンドのエクスポージャー²推定が行われている。

リターンの帰納的な逐次複製は、回帰をベースに行われることが多い。古典的な手法として、OLSを用いた複製はリスクファクターでの複製[2, 3, 4]、流動性資産での複製[5]など数多く行われている。なお[6]では、S&P500の個別銘柄に投資するファンドを対象に逐次的に保有銘柄候補を変えながらOLSを用いてウェイト推定をしているが、保有銘柄の組合せを当てることに焦点が当てられており、ウェイト推定を高精度にできているわけではない。この他にも、OLSを拡張した手法として、Ridge回帰を用いて多重共線性の問題に対応する手法[7]、LASSO回帰を用いて大規模な投資ユニバースから複製に用いるべき資産を選択する手法[8, 9]、それらを組み合わせたElastic Net回帰を用いた手法[10]など、OLSに正則化項を付加した手法も提案されている。[10]によれば、ステップワイズ回帰、主成分回帰、部分最小二乗回帰などの古典的な手法と

比較して、正則化項をつけた縮小推定量の方が実証的に高いリターンの複製精度を得られている。

より精緻なモデリングを実現しようとした試みとして、状態空間モデルを導入した研究もある。状態空間モデルを用いたアプローチでは、エクスポージャーを状態とみなして状態方程式で記述し、リターンが観測方程式に従ってエクスポージャーから生成されると考える。状態空間モデルの推定は、ベイズフィルタを通してオンラインで行われるため、エクスポージャー推定の遅延性を改善できる。具体的な手法として、状態方程式と観測方程式が共に線形で、誤差を正規分布とした場合にはカルマンフィルタを用いたアプローチ[10, 11, 12]、より一般の分布の場合には粒子フィルタを用いたアプローチ[12, 13]が提案されている。粒子フィルタを用いることで、リターンのファットテール性やオプション等の非線形なリターンを有する資産などを考慮することも可能となる。その他に、ヘッジファンド指数をショート・レバレッジ制約付きのポートフォリオで複製しようとした試みもある[13]。この研究ではエクスポージャー分布をディリクレ分布とした状態空間モデルでモデリングしており、粒子フィルタを用いた推定により、相対的に優れたリターンの複製精度が得られている。しかし、投資対象が縛られないヘッジファンド指数と複製ポートフォリオの投資対象が一致しないこと、ヘッジファンド指数のリターンの開示頻度が低いことから、ウェイト推定ができていないわけではない。

以上に挙げたヘッジファンド複製に関する研究のほとんどのモチベーションは、低コストでヘッジファンドと同一のパフォーマンスを得ることであり、正確なウェイト推定を目的としているわけではない。この点がヘッジファンド複製に関する先行研究と本研究の大きな違いである。一般のヘッジファンドと本研究の対象には表1のような違いがあり、本研究の問題の構造のシンプルさと透明性をうまく活かすことで、精度の良いウェイト推定が可能となる。

²ファクター等に対するファンドの感応度を指し、ファンドの保有資産のウェイトを指しているわけではないことに注意されたい。

3 ウェイトの推定手法

ウェイト推定問題の基本構造は回帰問題として捉えられ、一種の逆問題になっている [14]。多くの順問題の演算が和、積、積分から構成されることから、逆問題はそれらの逆演算である差、商、微分から構成されることが多い。逆問題における商演算は不安定性を生む要因となり得て、線形回帰においては逆行列に起因する多重共線性の問題が存在する。

以上を背景に、ウェイト推定問題は一般に不良設定問題となるが、不良設定問題を良設定問題に変換するためにしばしば正則化が施される。正則化はモデルのバイアスを上げ、バリエーションを下げる効果があるが、問題の構造に合わせた正則化を施すことで汎化性能を向上させることも可能である。ここで幾つかの正則化はベイズ的に解釈可能³であることが知られている [15]。

本研究では自由度の高いベイズモデルを用いることで問題の構造に合わせた正則化を図る。具体的には線形回帰から始めて、ベイズモデルにおける事前分布（グラフィカルモデルおよび各ノードの分布形）をうまく設定することで、効率的なウェイト推定の実現を目指す。

3.1 線形回帰

T, K をそれぞれデータ期間、ファンドのキャッシュを除く投資対象資産数とする。ここでは、組み入れ資産のリターン行列 $\mathbf{R} \in \mathbb{R}^{T \times K}$ 、ファンドのリターンベクトル $\mathbf{r} \in \mathbb{R}^T$ を入力値としてウェイトベクトル $\mathbf{w} \in \mathbb{R}^K$ を求める問題を考える。なおリターン誤差ベクトルを $\boldsymbol{\varepsilon} = \mathbf{r} - \mathbf{R}\mathbf{w}$ とする。

線形回帰を用いたウェイト推定問題は、ショート・レバレッジ制約下でリターンの二乗誤差 $\boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon}$ を最小化すればよく、以下の最小化問題に帰着する：

$$\begin{aligned} \min_{\mathbf{w}} \quad & \frac{1}{2} \mathbf{w}^\top \mathbf{R}^\top \mathbf{R} \mathbf{w} - \mathbf{r}^\top \mathbf{R} \mathbf{w} \\ \text{s.t.} \quad & \mathbf{w} \geq \mathbf{0}, \mathbf{w}^\top \mathbf{1} \leq 1. \end{aligned} \quad (1)$$

3.2 ベイズ線形回帰

ナイーブな線形回帰を用いたウェイト推定の問題点として、「正則化されていないため、即時性を高めようとしてデータの窓幅を短くすると外れ値の影響が大きくなること」、「ファンドの事前情報を全く利用していないため、推定の非効率性が残ること」が挙げられる。

そこで本節では線形回帰をベイズ化し、ファンドに関する事前情報を事前分布の形で与えることで問題の解決を図る。具体的には、ウェイトの事前分布として、

³例えば、 L^1 正則化はラプラス分布を、 L^2 正則化は $\mathcal{N}(\mathbf{0}, \eta^2 \mathbf{I})$ を事前分布に置いた場合の最大事後確率 (MAP) 推定と解釈できる。

平均および共分散行列を前回推定したウェイト平均（前回の推定ウェイト）およびウェイト共分散行列とした正規分布を用いる。また尤度として、誤差が正規分布に従うと仮定する⁴。事前分布をこのように置くことで、「今回推定するウェイトは前回推定したウェイトの近傍にあること」、「ファンドのウェイト変動には特徴（例えば、外国株式比率が上がる時は国内債券比率が下がるなど）があること」を踏まえることができる。

ただし今回の問題設定には、ショート・レバレッジ制約が課されているため、事後分布の平均値が必ずしも実行可能領域内に存在するとは限らない。そこで事後分布の平均値が、実行可能領域外に出てしまった場合、実行可能領域内で最大事後確率をとる点⁵ が平均値になるように事後分布を平行移動させる手順を加える。以上のベイズ更新の手順について、2 資産 + キャッシュの場合を図 1 に示す。

以上より、正定値行列 $\mathbf{S}_0 \in \mathbb{R}^{K \times K}$ 、正実数 $\beta \in \mathbb{R}_{>0}$ を用いて、ウェイトの事前分布 $p(\mathbf{w}) = \mathcal{N}(\mathbf{w}_0, \mathbf{S}_0)$ 、リターン誤差 $\varepsilon_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, \beta^{-1})$ を仮定すると、ベイズ線形回帰を用いたウェイト推定問題は、以下の最小化問題⁶に帰着する：

$$\begin{aligned} \min_{\mathbf{w}} \quad & \frac{1}{2} \mathbf{w}^\top \mathbf{S}^{-1} \mathbf{w} - (\mathbf{w}_0^\top \mathbf{S}_0^{-1} + \beta \mathbf{r}^\top \mathbf{R}) \mathbf{w} \\ \text{where} \quad & \mathbf{S}^{-1} = \mathbf{S}_0^{-1} + \beta \mathbf{R}^\top \mathbf{R} \\ \text{s.t.} \quad & \mathbf{w} \geq \mathbf{0}, \mathbf{w}^\top \mathbf{1} \leq 1. \end{aligned} \quad (2)$$

3.3 カルマンフィルタ

本研究対象のようにファンドが現物資産にのみ投資している場合には、ウェイトは時価変動要因とリバランス要因の 2 つの要因で変動する。したがって実際のウェイトは、前回更新ウェイトではなく、前回更新ウェイトから「時価変動した」ウェイトの近傍にある。この情報をウェイトを状態とみなした状態空間モデルを用いて組み込むことを本節と次節では考える。

ウェイトの時価変動とファンドリターンは共に、ウェイトと組み入れ資産のリターンの線形式で記述できる。仮にリバランスとファンドリターン誤差が正規分布に従うと仮定した場合、状態空間モデルは線形・正規分布型となり、カルマンフィルタを用いて状態を推定することができる。

⁴尤度が正規分布の場合、正規分布は共役事前分布となることから、事後分布も正規分布となる。

⁵実行可能領域内において、元の事後分布の平均値からのマハラノビス距離最小の点である。

⁶目的関数は $(\mathbf{r} - \mathbf{R}\mathbf{w})^\top (\mathbf{r} - \mathbf{R}\mathbf{w}) + \frac{1}{\beta} (\mathbf{w} - \mathbf{w}_0)^\top \mathbf{S}_0^{-1} (\mathbf{w} - \mathbf{w}_0)$ のように二乗誤差 + マハラノビス距離で記述された正則化項の形に変形できる。 $\mathbf{w}_0 = \mathbf{0}, \mathbf{S}_0 = \mathbf{I}$ の場合には通常の L^2 正則化となることから、問題の構造を踏まえて L^2 正則化の拡張をしていると解釈できる。

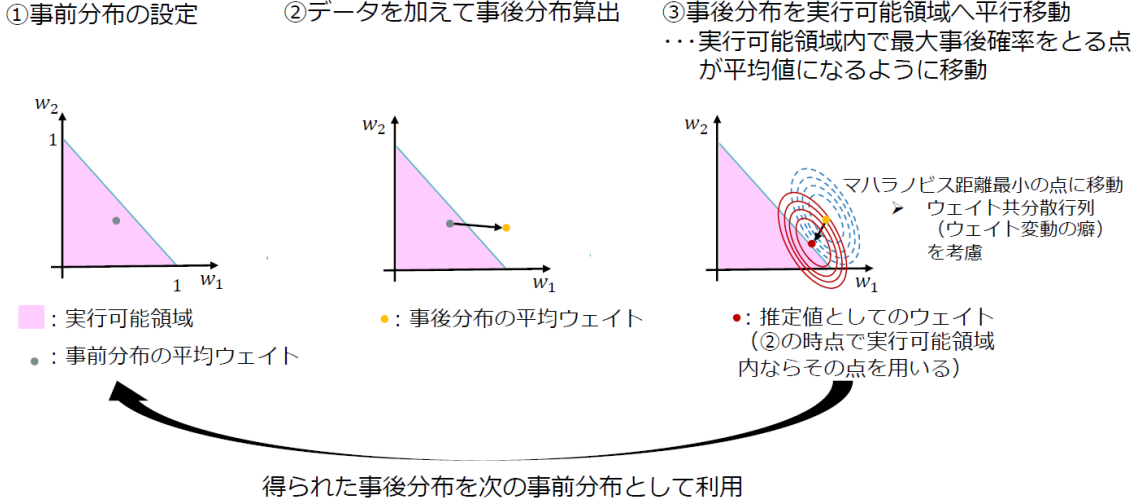


図 1: 2 資産 + キャッシュの場合のベイズ更新の手順

ここでは、推定した $t-1$ 時点の期初ウェイトベクトル $\mathbf{w}_{t-1} \in \mathbb{R}^K$ 、キャッシュウェイト $w_{t-1, \text{cash}} \in \mathbb{R}$ 、 $t-1$ 時点および t 時点の組み入れ資産リターンベクトル $\mathbf{r}_{a,t-1}, \mathbf{r}_{a,t} \in \mathbb{R}^K$ 、 t 時点のファンドリターン $r_t \in \mathbb{R}$ を入力値として、 t 時点の期初ウェイトベクトル $\mathbf{w}_t \in \mathbb{R}^K$ を求める問題を考える。

この場合、状態方程式（ウェイトの更新式）は

$$\begin{aligned} \mathbf{w}_t &= \mathbf{G}_t \mathbf{w}_{t-1} + \mathbf{e}_t \quad \mathbf{e}_t \sim \mathcal{N}(\mathbf{0}, \mathbf{E}_t) \\ \mathbf{G}_t &= \frac{\text{diag}(\mathbf{1} + \mathbf{r}_{a,t-1})}{\|\text{diag}(\mathbf{1} + \mathbf{r}_{a,t-1}) \mathbf{w}_{t-1}\|_1 + w_{t-1, \text{cash}}} \end{aligned} \quad (3)$$

で表される⁷。ここで $\mathbf{G}_t$ は時価変動を表し、 $\mathbf{e}_t$ はリバランスを表す。通常、誤差項として取り扱われる $\mathbf{e}_t$ に経済的意味付けをしている点が本モデルの要所である。これにより、ウェイトのみならずリバランスを含めた推定が可能となる。

また観測方程式（ファンドリターンの更新式）は

$$\begin{aligned} r_t &= \mathbf{F}_t \mathbf{w}_t + v_t \quad v_t \sim \mathcal{N}(0, V_t) \\ \mathbf{F}_t &= \mathbf{r}_{a,t}^\top \end{aligned} \quad (4)$$

で表される。ここで v_t はファンドリターン誤差を表す。なお本モデルにおいて、 $\mathbf{E}_t$ と V_t は既知とする。

ショート・レバレッジ制約については、フィルタリング分布についてベイズ線形回帰の場合（図 1）と同様の対応をする。

以上より、カルマンフィルタを用いたウェイト推定問題では、キャッシュを除く資産について初期ウェイト分布 $\mathcal{N}(\mathbf{w}_0, \mathbf{S}_0)$ を置いてフィルタリング分布の制約付き MAP 推定を行えばよく、以下の最小化問題に帰着する：

⁷ $\text{diag}(\mathbf{x})$ はベクトル $\mathbf{x}$ の各要素を対角成分とする対角行列を返す。

$$\begin{aligned} \min_{\mathbf{w}_t} \quad & \frac{1}{2} \mathbf{w}_t^\top (\mathbf{A}_t^{-1} + \mathbf{F}_t^\top \mathbf{V}_t^{-1} \mathbf{F}_t) \mathbf{w}_t \\ & - (\mathbf{F}_t^\top \mathbf{V}_t^{-1} r_t + \mathbf{A}_t^{-1} \mathbf{a}_t)^\top \mathbf{w}_t \\ \text{where} \quad & \mathbf{a}_t = \mathbf{G}_t \mathbf{w}_{t-1} \\ & \mathbf{A}_t = \mathbf{E}_t + \mathbf{G}_t \mathbf{S}_{t-1} \mathbf{G}_t^\top \\ & \mathbf{S}_{t-1}^{-1} = \mathbf{A}_{t-1}^{-1} + \mathbf{F}_{t-1}^\top \mathbf{V}_{t-1}^{-1} \mathbf{F}_{t-1} \\ \text{s.t.} \quad & \mathbf{w}_t \geq \mathbf{0}, \mathbf{w}_t^\top \mathbf{1} \leq 1. \end{aligned} \quad (5)$$

ここでリバランスは、 $\mathbf{w}_t - \mathbf{a}_t$ で求めることができる。

3.4 粒子フィルタ

前節までは、ショート・レバレッジ制約を最適化問題の制約条件に入れることで考慮していた。本節では、単体を定義域とするディリクレ分布でモデリングすることで、より自然にウェイト制約を考慮する。なおディリクレ分布を導入したことで解析的な推定が困難になったため、本節では数値計算的な推定手法である粒子フィルタを用いる。

この場合、状態方程式（キャッシュを含むウェイトの更新式）は以下で表される：

$$\mathbf{w}_t \sim \text{Dirichlet}(\boldsymbol{\alpha}_t). \quad (6)$$

ここで $\boldsymbol{\alpha}_t \in \mathbb{R}^{K+1}$ はディリクレ分布のパラメータであり、分布の期待値が前回の推定ウェイトから時価変動した値となるように決定する。また観測方程式は、式 (4) と同様とする。

粒子フィルタを用いたウェイト推定アルゴリズムを、Algorithm 1 に示す。ここで $\alpha_{\text{scale}} \in \mathbb{R}_{>0}$ はディリクレ分布の分散に影響を与えるパラメータで、値が大きいくほど分散が小さくなる。また $N \in \mathbb{Z}_{>0}$ は粒子数である。

Algorithm 1 粒子フィルタを用いたウェイト推定アルゴリズム

Input: α_{scale} , N

Output: W

```

1: for  $t = 1$  to  $T$  do
2:   if  $t = 1$  then
3:      $\alpha_t \leftarrow 1$ 
4:   else
5:      $w_{t|t-1}^{(tmp)} \leftarrow G_t w_{t-1|t-1}$ 
6:      $\alpha_t \leftarrow \alpha_{scale} w_{t|t-1}^{(tmp)}$ 
7:   end if
8:   for  $i = 1$  to  $N$  do
9:      $w_{t|t-1}^{(i)} \sim \text{Dirichlet}(\alpha_t)$ 
10:  end for
11:   $W_{t|t-1} = \{w_{t|t-1}^{(1)}, w_{t|t-1}^{(2)}, \dots, w_{t|t-1}^{(N)}\}$  を尤度比で復元抽出し,  $W_{t|t} = \{w_{t|t}^{(1)}, w_{t|t}^{(2)}, \dots, w_{t|t}^{(N)}\}$  を生成
12:   $w_{t|t} \leftarrow \mathbb{E}(W_{t|t})$ 
13: end for
14:  $W \leftarrow \{w_{1|1}, w_{2|2}, \dots, w_{T|T}\}$ 

```

り、値が大きいほど高い精度の推定が可能となる。ここでリバランスは、 $w_{t|t} - w_{t|t-1}^{(tmp)}$ で求めることができる。

4 数値実験

4.1 人工データ

ここでは、仮想的な株式、債券、キャッシュを保有するショート・レバレッジ制約付きのファンドのウェイト推定問題を考える。入力値は株式、債券、ファンドのリターンであり、株式、債券の1期間のリターンはそれぞれ $\mathcal{N}(0, 0.06^2)$, $\mathcal{N}(0, 0.03^2)$ に独立に従うとする。データ期間は300期間で、このファンドは50期毎に以下の戦略に切り替えるようリバランスすると仮定する。

- 通常戦略 (0, 50, 150, 250 期):
株式 40%, 債券 40%, キャッシュ 20%
- リスクオフ戦略 (100 期):
株式 10%, 債券 20%, キャッシュ 70%
- リスクオン戦略 (200 期):
株式 70%, 債券 20%, キャッシュ 10%

推定すべき真のウェイト推移を図2に、真の株式・債券のリバランス行動を図3⁸に示す。

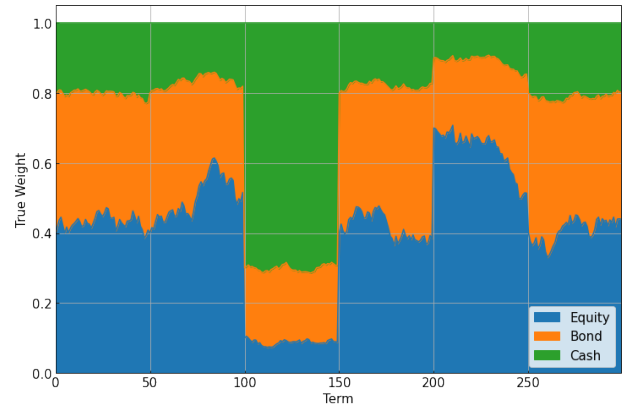


図 2: 真のウェイト推移

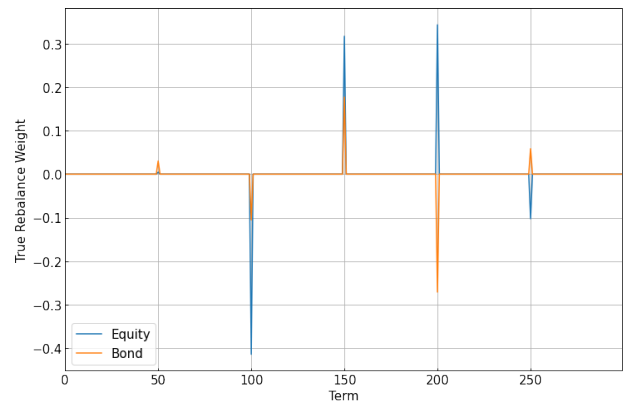


図 3: 真の株式・債券のリバランス行動

⁸煩雑になるため掲載していないが、ショート・レバレッジ制約下でのリバランス比率はネットで0であるため、キャッシュリバランス比率は $-\text{株式リバランス比率} - \text{債券リバランス比率}$ となる。

4.2 ハイパーパラメータ

線形回帰では、窓幅 20 のローリングウィンドウで推定を行った。ベイズ線形回帰では、初期ウェイト分布を $w_0 = (0.5, 0.5)^T$, $S_0 = 10^{-2}I$ とし, $\beta = 10^2$ で 1 期間ずつデータを与えて更新した。カルマンフィルタでは、初期ウェイト分布を $w_0 = (0.5, 0.5)^T$, $S_0 = 10^{-2}I$, 誤差共分散を $E_t = 10^{-2}I$, $V_t = 10^{-6}$ とした。粒子フィルタでは、粒子数を 10 万, ディリクレ分布のパラメータを $\alpha_{scale} = 10^2$, 誤差分散を $V_t = 10^{-6}$ とした。

4.3 分析結果と考察

4 つの手法を用いて推定したウェイト推移を図 4 に示す。線形回帰を用いたウェイト推定では、100 期までの比較的緩やかなウェイト変動の期間では、水準を概ね当てられている。これは、窓幅内のウェイトが均一に近い局面では最良線形不偏推定量である線形回帰が機能することを示唆している⁹。一方で、100 期以降の戦略のスイッチングによる激しいウェイト変動には応答の遅行が目立つ。線形回帰は正則化していないバッチ学習であり、窓幅内の急激なウェイト変動を前提としていないモデルであることが原因と考えられる。ベイズ線形回帰を用いたウェイト推定では、オンライン学習していることから、通期にわたり遅効性が少ない。一方で L^2 正則化に相当する最適化を行っていることから、推定ウェイトがスムージングされてしまっており、不連続なウェイト変動を推定できていない。カルマンフィルタや粒子フィルタといった状態空間モデルを用いたウェイト推定では、高い精度で遅行なくウェイト推定ができています。ウェイトの変動構造を基礎方程式として陽にモデルに入れたことが、推定精度を向上させたと考えられる。なお線形回帰以外の手法では、真のウェイトと異なる初期値を与えているが、短期間で正しい水準に収束しており、手法の頑健性も窺える。

50 期以降のウェイトとリバランスの推定精度を表 2 に、状態空間モデルを用いて推定した株式・債券のリバランス行動を図 5 に示す。表 2 において、コモナリティ平均とは、真のウェイトと推定ウェイトの共通部分の平均ウェイトを表し、値が大きいほど精度が高いことを示す。誤リバランス比率平均とは、真のリバランスと推定したリバランスのウェイト差の絶対和の平均を表し、値が小さいほど精度が高いことを示す。通期のコモナリティ平均は、状態空間モデルを用いた手法の精度が高く、いずれも 95% 以上の精度であった。リバランス後 5 期間に絞ったコモナリティ平均は、カルマンフィルタが最良の結果で、95% 以上の精度を保っている。図 5 によれば、カルマンフィルタは 50 期毎

のリバランスの発生を明確に特定できており、表 2 の誤リバランス比率平均が低いことから、大きなリバランスを迅速かつ正確に捉えられていることが窺える。相対リスクである他のプレイヤーとのトラッキングエラーはリバランスに伴い変動するため、この結果の意義は大きい。一方で図 5 によると、粒子フィルタを用いた手法はスパースなリバランス推定にはなっていない。これはカルマンフィルタでは MAP 推定を行っているのに対し、粒子フィルタでは期待事後 (EAP) 推定を行っており、安定的だが大きな変化を伴わない解を出しやすいことが影響していると考えられる。

最後に、線形回帰、ベイズ線形回帰、カルマンフィルタを用いた手法はいずれも線形制約付き凸二次計画問題に帰着するため、大域的最適解が一意的に存在すること、計算量が少ないというメリットがある。一方で、粒子フィルタを用いたアプローチは汎用性は高いが、サンプリングベースであるため、資産数が増加した際に十分な粒子数がなければ実用的に利用可能な解が得られないことに注意が必要である。

5 まとめ

本研究ではバランス型運用の相対的な運用リスク管理の導入を企図しながら、線形回帰、ベイズ線形回帰、カルマンフィルタ、粒子フィルタをベースとした 4 つの手法を用いて、一般的なバランス型投資信託を想定した人工データのウェイト推定を行った。結果として、状態空間モデルを用いてウェイトの時価変動構造を明示的に入れる提案手法により、高い精度のウェイト推定、ひいてはリバランス推定ができることが明らかとなった。特にカルマンフィルタを用いた推定ではリバランスを非常に高い精度で推定できており、解析的なアプローチであることから計算量の観点からも実的に価値が高いことが示唆された。

今後の課題として二点ほど触れておく。第一に粒子フィルタを用いた手法の改良である。単体制約を自然に反映させるために導入したディリクレ分布はパラメータ数が資産数しかなく、記述力が低いため十分な事前情報を入れられない。例えば共分散行列に相当する情報を組み入れられるようにモデルを拡張することで、より精緻なリバランス推定ができる可能性がある。粒子フィルタは分布形が限定されない汎用的な手法であることから、本研究で取り扱った以上のクラスの問題に拡張する際にも利便性が高く、手法の改良が望まれる。第二に実証分析への応用である。具体的には公募バランス型投資信託のウェイト・リバランス推定を行い、バランス型投資信託全体の投資行動の特性 (個々のリバランス特性、ファンド間の類似性、金融危機時の応答など) を考察することが挙げられる。

⁹ リバランス頻度が高くウェイトが一定に維持される場合は、線形回帰が最も有効であるのは自明である。



図 4: 4つの手法を用いて推定したウェイト推移

表 2: 50期以降のウェイトとリバランスの推定精度

	線形回帰	ベイズ線形回帰	カルマンフィルタ	粒子フィルタ
コモナリティ平均 (%)	90.3	92.1	99.5	97.5
リバランス後5期間のコモナリティ平均 (%)	70.6	76.8	95.2	85.9
誤リバランス比率平均 (%)	-	-	1.3	2.5

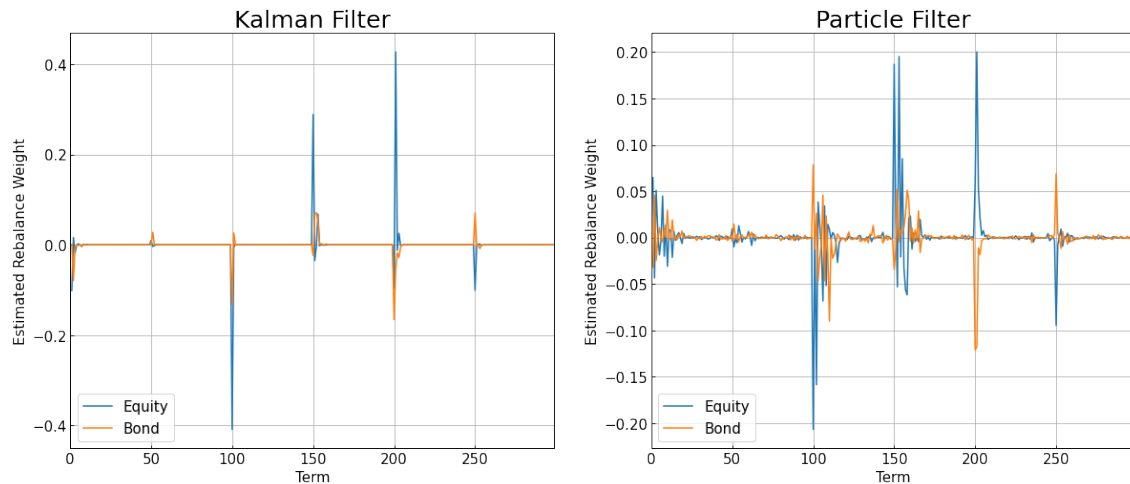


図 5: 状態空間モデルを用いて推定した株式・債券のリバランス行動

参考文献

- [1] Akihiko Takahashi, Kyo Yamamoto, et al. Hedge fund replication. *The Recent Trend of Hedge Fund Strategies*, pages 57–95, 2008.
- [2] William F Sharpe. Asset allocation: Management style and performance measurement. *Journal of portfolio Management*, 18(2):7–19, 1992.
- [3] William Fung and David A Hsieh. Hedge fund benchmarks: A risk-based approach. *Financial Analysts Journal*, 60(5):65–80, 2004.
- [4] Stein Frydenberg, Kjartan Hrafnkelsson, Vegard Strand Bromseth, and Sjur Westgaard. Hedge fund strategies and time-varying alphas and betas. *The Journal of Wealth Management*, 19(4):44–60, 2017.
- [5] Jasmina Hasanhodzic and Andrew W Lo. Can hedge-fund returns be replicated?: The linear case. *The Linear Case (August 16, 2006)*, 2006.
- [6] David Byrd, Sourabh Bajaj, and Tucker Hybnette Balch. Fund asset inference using machine learning methods: What’s in that portfolio? *The Journal of Financial Data Science*, 1(3):98–107, 2019.
- [7] Joseph Simonian and Chenwei Wu. Factors in time: Fine-tuning hedge fund replication. *The Journal of Portfolio Management*, 45(3):159–164, 2019.
- [8] Jun Duanmu, Yongjia Li, and Alexey Malakhov. In search of missing risk factors: Hedge fund return replication with etfs. *Available at SSRN 2411910*, 2018.
- [9] Lars Stentoft and Sha Wang. Consistent and efficient dynamic portfolio replication with many factors. *The Journal of Portfolio Management*, 46(2):79–91, 2019.
- [10] Jiaqi Chen and Michael L Tindall. Hedge fund replication using shrinkage methodologies. *The Journal of Alternative Investments*, 17(2):26–49, 2014.
- [11] Noël Amenc, Lionel Martellini, Jean-Christophe Meyfredi, and Volker Ziemann. Passive hedge fund replication—beyond the linear case. *European Financial Management*, 16(2):191–210, 2010.
- [12] Thierry Roncalli and Guillaume Weisang. Tracking problems, hedge fund replication and alternative beta. *Hedge Fund Replication and Alternative Beta (January 9, 2009)*, 2009.
- [13] Laszlo F Korsos. The dirichlet portfolio model: Uncovering the hidden composition of hedge fund investments. *arXiv preprint arXiv:1306.0938*, 2013.
- [14] Finbarr O’Sullivan. A statistical perspective on ill-posed inverse problems. *Statistical science*, pages 502–518, 1986.
- [15] Nicholas G Polson and Vadim Sokolov. Bayesian regularization: From tikhonov to horseshoe. *Wiley Interdisciplinary Reviews: Computational Statistics*, 11(4):e1463, 2019.

シュレーディンガー・リスクパリティポートフォリオ

Schrödinger Risk Parity Portfolio

内山 祐介^{1*} 中川 慧²
Yusuke Uchiyama¹ Kei Nakagawa²

¹ 株式会社 MAZIN

¹ MAZIN Inc.

² 野村アセットマネジメント株式会社

² Nomura Asset Management Co, Ltd.

Abstract: 複数資産からなるポートフォリオのアセット・アロケーションの問題において、期待リターンとリスクのトレードオフを考慮した平均分散法が使用されてきた。しかし、期待リターンの推定は困難であり、アウトオブサンプルのパフォーマンスが優れないことなどから、リスクのみに焦点を当てたリスクベースのポートフォリオ構築手法が複数提案されており、実務を中心に注目されている。加えて、ポートフォリオを構成する資産変動同士は背後に共通するファクターを持つと考えられ、これらの情報を抽出ために次元削減の手法である主成分分析が応用されている。本研究では、量子力学にあらわれるシュレーディンガー方程式を応用したシュレーディンガー主成分分析を用いたリスク分散ポートフォリオとして、シュレーディンガー・リスクパリティポートフォリオを提案する。これによりサンプル点が不等間隔や少数のケースであっても主成分と相互相関が精度良く推定でき、効率的なリスク分散が可能になると考えられる。提案手法を既存のリスク分散ポートフォリオと比較し、有効性と課題を検証する。

1 はじめに

複数資産からなるポートフォリオのアセット・アロケーションの問題において、期待リターンとリスクのトレードオフを考慮した平均分散法が使用されてきた。Markowitz[1]により提案されて以来長らく、平均分散法は投資意思決定において極めて重要なフレームワークである。その理由として、平均分散法は最適化問題が性質の良い二次計画問題として表すことができることから、実務上扱いやすいだけでなく、期待効用最大化原理と整合的であるため、経済的な正当性が得られることが挙げられる。平均分散法は期待リターン、分散および資産間の相関をそれぞれ推計し、それらをインプットとして与え、様々な投資制約のもとでポートフォリオの平均と分散のトレードオフを考慮し資産配分を決定する手法である。しかしながら、平均分散法にはいくつかの実務上の課題がある。最もよく指摘される点は、そもそもインプットとなる期待リターンの正確な値を予測することが非常に困難であり[2]、アウトオブサンプルのパフォーマンスが非常に悪いという点である[3]。リーマンショックや直近のコロナショック等を例に、市場の

不確実性は年々高まっており、このような市場環境下で期待収益率の推定はより困難になっていると考えられる。さらに期待リターンの小さな変化に対して、最適化の結果として得られる最適なポートフォリオが大きく異なってしまうという点も課題として挙げられる[4]。

そこで、期待リターンを用いずにリスクのみに焦点を当てたリスクベースのポートフォリオ構築手法が複数提案されており、実務を中心に注目されている。このようなリスクベース・ポートフォリオ構築手法として、最小分散ポートフォリオ[5]、リスクパリティ・ポートフォリオ[6, 7]、最大分散度ポートフォリオ[8]などが代表的なものとして提案されており、さらにそれぞれを発展させたリスクベースポートフォリオ[9, 10, 11]が提案されている。実際に、株式ポートフォリオや資産配分を対象とした様々な実証分析やバックテストにおいて、これらのリスクベースのポートフォリオ構築手法は平均分散ポートフォリオや時価総額加重型のポートフォリオと比較して良好なパフォーマンスを示している[12]。

一方で、ポートフォリオを構成する資産変動同士は背後に共通するファクターを持つと考えられており、これらの情報を抽出ために次元削減の手法である主成分分析が応用されている。これらは資産ベースでのリスクを均等に配分する通常のリスクパリティポートフォリオに対して、リスクの源泉であるファクターについての

*連絡先：株式会社 MAZIN
〒111-0035 東京都台東区西浅草 3 丁目 29-14
E-mail: uchiyama@mazin.tech

リスクを均等に配分するため、リスク分散ポートフォリオとも呼ばれている。また、資産ベースではなくファクターベースでのリスクベースのポートフォリオ構築手法として、主成分分析によって抽出されたファクターに対してリスクを均等に配分する主成分リスクパリティポートフォリオ [13] が提案されている。さらに、動的なリスクファクターを抽出することを目的に、ヒルベルト主成分分析により抽出したファクターに対して均等にリスクを配分する複素主成分リスクパリティポートフォリオが提案され [10], これをさらに発展させた四元数リスクパリティポートフォリオも提案されている [14]。

本研究では、量子力学にあらわれるシュレーディンガー方程式を応用した、シュレーディンガー主成分分析 [15] を用いたリスク分散ポートフォリオとして、シュレーディンガー・リスクパリティポートフォリオを提案する。シュレーディンガー主成分分析によりサンプル点が不等間隔や少数のケースであっても主成分と相互相関が精度良く推定でき、効率的なリスク分散が可能になると考えられる。提案手法を既存のリスク分散ポートフォリオと比較し、有効性と課題を検証する。

2 提案手法

2.1 ガウス場とシュレーディンガー方程式

空間座標 $\mathbf{x} \in \mathbb{R}^d$ に対してスカラー場 $\phi: \mathbb{R}^n \rightarrow \mathbb{R}$ が与えられているとする。任意の N 個の座標ベクトル $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N$ に対応する場の量 $\phi(\mathbf{x}_1), \phi(\mathbf{x}_2), \dots, \phi(\mathbf{x}_N)$ が N 次元ガウス分布に従う確率変数であるとき、 $\phi(\cdot)$ をガウス過程という。特に、 $\phi(\cdot)$ は空間座標 $\mathbf{x}$ に対する場の量であることから、ガウス場とも呼ばれる。ガウス場から平均値 $\mathbb{E}[\phi(\mathbf{x})]$ を引くことにより、一般性を損なうことなく、 $\mathbb{E}[\phi(\mathbf{x})] = 0$ とすることが可能である。このとき、

$$k(\mathbf{x}, \mathbf{x}') = \mathbb{E}[\phi(\mathbf{x})\phi(\mathbf{x}')] \quad (1)$$

で与えられる正定値対称関数を共分散関数またはカーネル関数という。したがって、任意の N 個の座標ベクトル $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N$ に対してはカーネル関数によって N 次元ガウス分布の共分散行列

$$K = \begin{bmatrix} k(\mathbf{x}_1, \mathbf{x}_1) & k(\mathbf{x}_1, \mathbf{x}_2) & \dots & k(\mathbf{x}_1, \mathbf{x}_N) \\ k(\mathbf{x}_2, \mathbf{x}_1) & k(\mathbf{x}_2, \mathbf{x}_2) & \dots & k(\mathbf{x}_2, \mathbf{x}_N) \\ \vdots & \vdots & \ddots & \vdots \\ k(\mathbf{x}_N, \mathbf{x}_1) & k(\mathbf{x}_N, \mathbf{x}_2) & \dots & k(\mathbf{x}_N, \mathbf{x}_N) \end{bmatrix} \quad (2)$$

が与えられる。ガウス場の共分散関数に対しては、

$$k(\mathbf{x}, \mathbf{x}') = \sqrt{A(\mathbf{x})A(\mathbf{x}')} \times \exp \left[-\frac{1}{2}(\mathbf{x} - \mathbf{x}')^T \Sigma^{-1}(\mathbf{x} - \mathbf{x}') \right] \quad (3)$$

をみたす実関数 $A(\cdot)$ と正定値対称行列 Σ が存在する。

一方、カーネル関数を積分核に持つ積分作用素に対して、固有方程式

$$\int_{\Omega_{\mathbf{x}}} k(\mathbf{x}, \mathbf{x}') \varphi_i(\mathbf{x}') d\mathbf{x}' = \lambda_i \varphi_i(\mathbf{x}') \quad (4)$$

および、直交条件

$$\int_{\Omega_{\mathbf{x}}} \varphi_i(\mathbf{x}) \varphi_j(\mathbf{x}) d\mathbf{x} = \delta_{i,j} \quad (5)$$

をみたす固有値 λ_i と固有関数 φ_i が存在する。ここで、 $\delta_{i,j}$ はクロネッカーのデルタである。カーネル関数の正値対称性から、固有値は正の実数である。固有値と固有関数の添字 i および j は $\lambda_1 \geq \lambda_2 \geq \dots \geq 0$ をみたすようにとる。式 (4) は共分散行列の固有値問題の無限次元ベクトルに対応することから、適当な有限個の固有関数で打ち切られた固有関数展開は主成分分析に対応する。

量子力学におけるシュレーディンガー方程式の固有値問題は以下のように与えられる。

$$-\frac{1}{2}\Delta_{\Sigma_m}\psi(\mathbf{x}) + V(\mathbf{x})\psi(\mathbf{x}) = E\psi(\mathbf{x}) \quad (6)$$

ここで、 $\psi(\cdot)$ は波動関数であり、 Δ_{Σ_m} は正定値対称行列 Σ_m によって重み付けられた 2 階の偏微分作用素

$$\Delta_{\Sigma_m} = \sum_{i,j} \frac{\partial}{\partial x_i} (\Sigma_m)_{i,j} \frac{\partial}{\partial x_j} \quad (7)$$

である。実関数 $V(\cdot)$ はポテンシャル関数であり、 E はエネルギー固有値である。式 (6) の左辺に現れる線形作用素はハミルトン演算子やハミルトニアンともよばれ、

$$H = -\frac{1}{2}\Delta_{\Sigma_m} + V(\mathbf{x}) \quad (8)$$

と表記される。ハミルトニアンの正値対称性から、エネルギー固有値は非負の実数値をとる。

2.2 シュレーディンガー主成分分析

式 (4) における積分方程式と式 (6) で表されるシュレーディンガー方程式は、以下の対応関係

$$\phi(\mathbf{x}) \Leftrightarrow \psi(\mathbf{x}) \quad (9)$$

$$-A(\mathbf{x}) \Leftrightarrow V(\mathbf{x}) \quad (10)$$

$$A(\mathbf{x})\Sigma \Leftrightarrow \Sigma_m \quad (11)$$

によって2次までの精度で一致する。したがって、もとのガウス場をシュレーディンガー方程式の固有関数で展開した際に適当なところで打ち切ることにより主成分分析を行うことが可能となる。これをシュレーディンガー主成分分析という。通常の主成分分析は有限のデータセットを用いて行われうるため、少数サンプルしか得られない場合の推定精度が悪化することが知られている。一方で、シュレーディンガー主成分分析はデータのサンプル数に依存しないため、2次までの近似精度の範囲内においては安定した推定結果を得ることができる。

2.3 シュレーディンガー・リスクパリティポートフォリオ

シュレーディンガー主成分分析は場の量、すなわち空間座標の各点において変動する量を対象としている。一方で、アセットアロケーションの問題では空間座標には対応付けられない個々の資産価格が変動する系が対象となる。そのため、 i 番目の資産を便宜的に空間座標値 x_i と対応付けることによって、複数資産の価格変動からなる多変量時系列に対してシュレーディンガー主成分分析を適用する。その結果を利用して構成するリスク分散ポートフォリオをシュレーディンガー・リスクパリティポートフォリオとして定式化する。

シュレーディンガー・リスクパリティポートフォリオの構成手順を以下に示す。

1. N 個の資産の T 期分のリターンの系列 $\{r_{t,i}\} (i = 1, 2, \dots, N; t = 1, 2, \dots, T)$ に対して、

$$\mu_i = \frac{1}{T} \sum_{t=1}^T r_{t,i}, \quad (12)$$

$$\sigma_i^2 = \frac{1}{T-1} \sum_{t=1}^T (r_{t,i} - \mu_{r,i})^2 \quad (13)$$

を求める。

2. ポテンシャル関数の関数形を仮定し、 $V(x_{j(i)}) = \sigma_i^2$ となるような置換 $i \rightarrow j(i)$ を行い、 $x_j = x_0 + j\Delta x$ における x_0 および Δx と、ポテンシャル関数のパラメータを推定する。
3. 空間1次元のシュレーディンガー方程式の固有値問題

$$-\frac{1}{2} \frac{d^2}{dx^2} \psi(x) + V(x) \psi(x) = E \psi(x) \quad (14)$$

を解くことで、固有値と固有関数の組 $\{E_l, \psi_l\} (l = 1, 2, \dots, L)$ を求める。ここで、 L はシュレーディンガー主成分分析の打ち切り次数に対応する。

4. 固有関数から行列

$$\Psi_{l,i} = \psi_l(x_i) \quad (i = 1, 2, \dots, N; l = 1, 2, \dots, L) \quad (15)$$

を生成し、ポートフォリオのウェイト w_i ($i = 1, 2, \dots, N$) に対して

$$\tilde{w}_l = \sum_{i=1}^N \Psi_{l,i} w_i \quad (16)$$

を求める。

5. 各 l に対して、固有値 E_l と $\tilde{w}_l$ から

$$v_l = \frac{E_l \tilde{w}_l^2}{\sum_{l=1}^L E_l \tilde{w}_l^2} \quad (17)$$

を求める。

6. エントロピー

$$\mathcal{H} = - \sum_{l=1}^L v_l \exp v_l \quad (18)$$

を、制約条件

$$\sum_{i=1}^N w_i = 1 \quad (19)$$

のもとで最大化することで、最適なウェイトを求める。

空間1次元の場合には、適当なポテンシャル関数のもとでのシュレーディンガー方程式は解析解を持つことが知られている。たとえば、ポテンシャル関数が

$$V(x) = -\frac{1}{2} k x^2 \quad (20)$$

で与えられる調和振動子の場合には、 $y = x/\sqrt{k}$, $\lambda = 2E/\sqrt{k}$ とすると2階の常微分方程式

$$\frac{d^2 \psi}{dy^2} + (\lambda - y^2) \psi = 0 \quad (21)$$

が得られる。これを解くと、

$$\lambda_l = 2l + 1, \quad (22)$$

$$\psi_l(y) = H_l(y) e^{-y^2/2} \quad (23)$$

が得られる。ここで、関数 $H_l(\cdot)$ はエルミート多項式

$$H_l(y) = (-1)^l e^{y^2} \frac{d^l}{dy^l} e^{-y^2} \quad (24)$$

である。これにより、求めたい固有値と固有関数が得られる。このほかにも解析解が得られるポテンシャル関数は存在し、問題に応じて適宜使い分ければよい。

3 実証分析

ここでは提案手法を含む複数のポートフォリオのパフォーマンスを比較し、提案手法の有効性を確認する。

3.1 データセット

各ポートフォリオのパフォーマンスを評価するために、表 1 に示す 6 つの債券先物、6 つの株価指数先物、5 つの通貨を対象とするデータを用いる。これらのデータは全て Bloomberg 端末より取得した。データのサンプル期間は 2000 年 5 月から 2017 年 4 月であり、記述統計を表 1 に示す。

3.2 分析手法

等ウェイトポートフォリオ (EW), 主成分リスクパリティポートフォリオ (PCA), 複素主成分リスクパリティポートフォリオ (HPCA), シュレディンガー・リスクパリティポートフォリオ (SPCA) の 4 つを比較する。それぞれ詳細は次の通りである。また、各ポートフォリオは 20 日 (月次) でパラメータを再推定し、リバランスを行う。

EW: 全ての資産を均等に保有する等ウェイトポートフォリオ [3]。

PCA: 過去 60 日分のリターンデータに対して主成分分析を行い、主成分リスクが均等になるように投資する主成分リスクパリティ・ポートフォリオ [13]

HPCA: 過去 60 日分のリターンデータから複素主成分分析を行い、複素主成分リスクが均等になるように投資する複素主成分リスクパリティ・ポートフォリオ [13]

SPCA: 過去 60 日分のリターンデータから調和振動子ポテンシャルを用いたシュレディンガー主成分分析を行い、主成分リスクが均等になるように投資するシュレディンガー・リスクパリティポートフォリオ (提案手法)

3.3 評価指標

評価指標として、広く使用されている次の 3 つの指標を使用する [16]。年率収益率 (AR) は、ポートフォリオの年率換算した収益率を表し、年率リスクは収益率の年率換算した標準偏差として定義され、リターン/リ

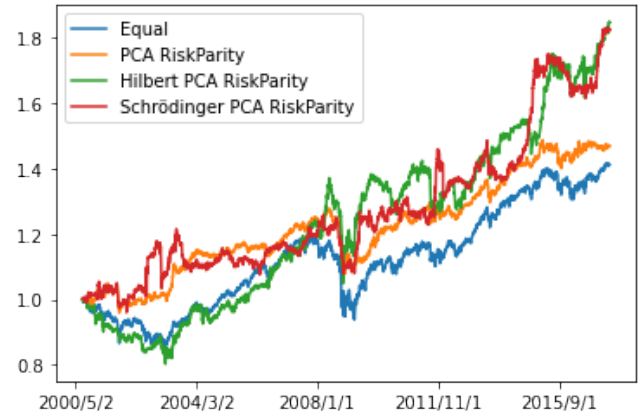


図 1: 各ポートフォリオの累積収益率の推移

スク比 (R/R) はリスクで基準化したリターンで定義される。

$$AR = \frac{250}{T} \sum_{t=1}^T R_t \quad (25)$$

$$RISK = \sqrt{\frac{250}{T-1} \times \sum_{t=1}^T (R_t - \mu)^2} \quad (26)$$

$$R/R = AR/RISK \quad (27)$$

ここで、 $\mu = (1/T) \sum_{t=1}^T R_t$ である。さらに、リスク指標として広く使用されている最大ドローダウン (MaxDD)[17] を評価する。MaxDD はある期間において、収益のピークから減少する幅のうち最大のものを指す。

$$MaxDD = \min_{k \in [1, T]} \left(0, \frac{W_k^{Port}}{\max_{j \in [1, k]} W_j^{Port}} - 1 \right) \quad (28)$$

$$W_k^{Port} = \sum_{i=1}^k (1 + R_i). \quad (29)$$

3.4 分析結果

各ポートフォリオの評価指標をまとめたものが表 2 である。表から、AR は HPCA が最も良く、RISK は PCA が最も低い一方で、RR は SPCA が最も良い。さらに下方リスク指標である MaxDD も SPCA が最も低い。したがって、提案手法はファクターのリスクを適切に分散することで効率的にリスク・リターンを確保でき、さらに下方リスクも制御できていることがわかる。図 1 が各ポートフォリオの累積収益率の推移を示したものである。この図からも提案手法はドローダウンが少なく安定した収益獲得ができていることがわかる。

表 1: 対象資産の詳細と記述統計

Ticker	資産タイプ	説明	平均	標準偏差	歪度	尖度
TX1	債券先物	10 Year T-Note Futures	0.0066	0.4650	-0.2840	4.56
XM1	債券先物	Australian 10 Year Bond	0.0009	0.0621	-0.1400	1.81
CN1	債券先物	Canadian Government 10 Year Note	0.0094	0.4440	-0.2220	2.87
RX1	債券先物	Eurex Euro Bund	0.0129	0.4640	-0.8210	6.95
G1	債券先物	Gilt UK	0.0033	0.5320	-7.3500	214
JB1	債券先物	Japan 10 Year Bond Futures	0.0041	0.2830	-0.5710	5.78
SP1	株式指数先物	S&P500	0.2100	14.7	-0.2500	5.37
XP1	株式指数先物	S&P/ASX 200(Austraria)	0.6240	46.9	-0.2890	5.60
PT1	株式指数先物	S&P/TSX 60 Index(Canada)	0.0776	7.42	-0.6280	7.42
GX1	株式指数先物	DAX(German)	1.1000	90.8	-0.2270	4.01
Z1	株式指数先物	FTSE100(UK)	0.1660	61.3	-0.2230	3.61
NK1	株式指数先物	Nikkei225(Japan)	0.1650	194	-0.3020	5.44
AUD	通貨	AUD/USD	0.0000	0.0064	-0.3610	6.11
CAD	通貨	CAD/USD	0.0000	0.0069	0.0726	2.92
EUR	通貨	EUR/USD	0.0000	0.0077	0.0346	1.85
GBP	通貨	GBP/USD	-0.0001	0.0094	-0.6650	7.55
JPY	通貨	JPY/USD	0.0007	0.6750	-0.1470	3.19

表 2: 各ポートフォリオのパフォーマンス比較

	EW	PCA	HPCA	SPCA
AR	2.11	2.29	3.77	3.67
RISK	5.29	4.28	7.06	6.70
RR	0.40	0.53	0.53	0.55
MaxDD	-26.25	-19.91	-32.28	-17.99

4 まとめ

複数資産からなるポートフォリオのアセット・アロケーションの問題に対して、量子力学にあらわれるシュレーディンガー方程式を応用したシュレーディンガー・リスクパリティポートフォリオを提案し、既存のリスクベースポートフォリオに対する性能評価を行った。提案手法は少数サンプル点しか得られないケースにおいてもリスクの推定精度が安定する効果があることから、既存手法に対して高いリスク分散効果が得られることが期待される。実際にリスク・リターンおよび最大ドローダウンを比較したところ、既存手法に対して提案手法が上回っていたことから、期待どおりの効果が得られることを確認した。

本研究での実証分析においては、提案手法で使用するポテンシャル関数を調和振動子ポテンシャルとし、資産同士の空間配置を等間隔と仮定したが、シュレーディンガー主成分分析の実施においては、ポテンシャル関数の選択と空間配置を不当間隔にする自由度が存在する。また、[10]と同様にポートフォリオを構成する各資産の

リターンの時系列をヒルベルト変換によって複素化することによって、提案手法を HPCA に拡張することも可能である。これらの方向で提案手法を発展させた際の効果を検証することが今後の課題である。

参考文献

- [1] Harry Markowitz. Portfolio selection. *The Journal of Finance*, 1952.
- [2] Robert C Merton. On estimating the expected return on the market: An exploratory investigation. *Journal of financial economics*, Vol. 8, No. 4, pp. 323–361, 1980.
- [3] Victor DeMiguel, Lorenzo Garlappi, and Raman Uppal. Optimal versus naive diversification: How inefficient is the 1/n portfolio strategy? *The review of Financial studies*, Vol. 22, No. 5, pp. 1915–1953, 2009.
- [4] Richard O Michaud. The markowitz optimization enigma: Is ‘optimized’ optimal? *Financial analysts journal*, Vol. 45, No. 1, pp. 31–42, 1989.
- [5] Roger Clarke, Harindra De Silva, and Steven Thorley. Minimum-variance portfolio composition. *The Journal of Portfolio Management*, Vol. 37, No. 2, pp. 31–45, 2011.

- [6] Edward Qian. Risk parity portfolios: Efficient portfolios through true diversification. *Panagora Asset Management*, 2005.
- [7] Sébastien Maillard, Thierry Roncalli, and Jérôme Teïletche. The properties of equally weighted risk contribution portfolios. *The Journal of Portfolio Management*, Vol. 36, No. 4, pp. 60–70, 2010.
- [8] Yves Choueifaty and Yves Coignard. Toward maximum diversification. *The Journal of Portfolio Management*, Vol. 35, No. 1, pp. 40–51, 2008.
- [9] 中川慧. リスクベース・ポートフォリオの高次モーメントへの拡張. リスク管理・保険とヘッジ (ジャーナル・ジャーナル) , pp. 49–71, 2017.
- [10] Yusuke Uchiyama, Takanori Kadoya, and Kei Nakagawa. Complex valued risk diversification. *Entropy*, Vol. 21, No. 2, p. 119, 2019.
- [11] Kei Nakagawa and Yusuke Uchiyama. Go-gjrsk model with application to higher order risk-based portfolio. *Mathematics*, Vol. 8, No. 11, p. 1990, 2020.
- [12] Thorsten Poddig and Albina Unger. On the robustness of risk-based asset allocations. *Financial Markets and Portfolio Management*, Vol. 26, No. 3, pp. 369–401, 2012.
- [13] Attilio Meucci. Managing diversification. *Risk*, Vol. 22, No. 5, p. 74, 2009.
- [14] Seisuke Sugitomo and Keiichi Maeta. Quaternion valued risk diversification. *Entropy*, Vol. 22, No. 4, p. 390, 2020.
- [15] Ziming Liu, Sitian Qian, Yixuan Wang, Yuxuan Yan, and Tianyi Yang. Schrödinger pca: You only need variances for eigenmodes. *arXiv preprint arXiv:2006.04379*, 2020.
- [16] Michael W Brandt. Portfolio choice problems. In *Handbook of financial econometrics: Tools and techniques*, pp. 269–336. Elsevier, 2010.
- [17] Malik Magdon-Ismail and Amir F Atiya. Maximum drawdown. *Risk Magazine*, Vol. 17, No. 10, pp. 99–102, 2004.

人工市場を用いたリベートが取引シェアに与える影響調査

Investigation of impact by rebates on market share using artificial markets

星野 真広^{1*} 水田 孝信² 八木 勲³
Mahiro Hoshino¹, Takanobu Mizuta², Isao Yagi³

¹ 神奈川工科大学大学院工学研究科情報工学専攻

¹ Information and Computer Sciences, Graduate School of Kanagawa Institute of Technology

² スパークス・アセット・マネジメント株式会社

² SPARX Asset Management Co., Ltd.

³ 神奈川工科大学情報学部工学科

³ Department of Information and Computer Sciences, Faculty of Information Technology,
Kanagawa Institute of Technology

Abstract: 取引所が提供する手数料体系の1つとして指値注文者（メイカー）にリベート（報酬）を支払う、メイカー・テイカー制が挙げられる。メイカー・テイカー制を採用した市場では、効率的な市場形成が見込まれるため、他の取引所に対して取引シェアの向上が期待できるとされているが、取引シェアについてはまだ十分な議論がなされていない。そこで本研究では、メイカー・テイカー制を採用した市場と採用していない2つの人工市場を構築し、市場間での取引シェアの変化の様子を調査した。その結果、取引所がリベートを十分に提供した場合、メイカー・テイカー制を採用した市場の取引シェアが向上することが確認できた。しかし、リベートを十分に提供できていない場合、メイカー・テイカー制を採用していない市場に取引シェアが奪われることも確認できた。

1 はじめに

取引所が提供する手数料体系の1つとして指値注文者（メイカー）にリベート（報酬）を支払い、成行注文者（テイカー）にその取引を行った手数料を請求するメイカー・テイカー制が挙げられる。メイカー・テイカー制は、効率的な市場形成や取引シェアの獲得が期待できると言われており、それらに関する調査が進められている [Foucault 13, Battalio 16, 岡田 17, Cox 19, Yagi 20, Brolley 20]。実際に、Yagi et al.[Yagi 20] の研究により、メイカー・テイカー制を採用した市場では効率的な市場形成が見込まれることが既に確認されている。しかし、もう1つの目的である、取引シェアの向上についてはまだ十分な議論がなされていない。そのため、本研究ではメイカー・テイカー制を採用した市場と採用していない2つの市場間での取引シェアの検討を行うことを目的にする。

本研究のように、ある制度が採用された市場とそうでない市場の比較調査を、現実市場データを用いて行おうとすると、その制度以外にも市場に影響を与える外的

要因が含まれることが多い。そのため実証研究では比較検討が難しくなる。実証研究では議論が困難な課題を分析する方法の1つとして、人工市場を用いる方法がある。人工市場とは、計算機上に仮想的に構築された金融市場マルチエージェントシステムである [Chiarella 09, Chen 12, Yeh 13]。人工市場では、エージェントにそれぞれ独自の売買手法を与え、それらを投資家として金融資産の取引をさせると、市場がどのような振る舞いをするかを確認することができる。

そこで、本研究ではメイカー・テイカー制を採用した人工市場と、採用していない人工市場の2つを構築し、エージェントに市場選択を行わせ各市場の取引シェアを調査した。またエージェントの市場選択が市場へ与える影響を確認するため、各市場のボラティリティと市場非効率性の確認も行った。

2 人工市場モデル

本研究では Yagi et al.[Yagi 20] の人工市場モデルを基に、一般投資家エージェント、アルゴリズムエージェント、ポジションマーケットメイカーを導入した。Yagi

* 神奈川工科大学大学院工学研究科情報工学専攻
〒 243-0292 神奈川県厚木市下荻野 1030
E-mail: s2085013@cce.kanagawa-it.ac.jp

et al. の市場は単一市場であったが、本研究では市場間取引シェアを確認するため取引できる市場を2つに拡張する。1つはメイカー・テイカー制を採用した市場（以下、採用市場とする）、もう1つはメイカー・テイカーを採用していない市場（以下、非採用市場とする）である。なお、2つの市場では同一資産が取引されているものとする。メイカー・テイカー制の手数料体系は2.1節で詳しく述べる。また、取引シェアを測定するため、一般投資家エージェントに市場選択を行う要素を追加した。

一般投資家エージェントは n 体、アルゴリズムエージェントは m 体（ただし、 $1 \leq m \leq n$ ）とする。一般投資家エージェント j は $j = 1$ から順に注文を出していく。そして $j = n$ まで注文を出し終えたあと $j = 1$ に戻る。一般投資家エージェントが n/m 体（小数点以下切り捨て）注文を出すごとに、アルゴリズムエージェントが1体注文を出す。アルゴリズムエージェント k は $k = 1$ から順に注文を出し、 $k = m$ まで到達すると $k = 1$ に戻る。一般投資家エージェントは注文を行う直前に注文を出す市場を選択する。アルゴリズムエージェントは k が偶数なら採用市場に、 k が奇数なら非採用市場に注文を出す。ポジションマーケットメイカーは各市場に1体ずつ存在し、一般投資家エージェントおよびアルゴリズムエージェントが注文を行う前に売りと買いの注文両方をそれぞれ出す。

時刻 t は一般投資家エージェントおよびアルゴリズムエージェント1体が注文を出すたびに1だけ増える。注文をただけで取引が成立しない場合も時刻 t は1ステップ進む。ポジションマーケットメイカーの注文で時刻 t は進まない。このモデルでの価格決定メカニズムは、買い手と売り手が価格を提示し、両者の提示価格が合致するとその価格での取引が成立するザラ場方式（連続ダブルオークション方式）とした。

2.1 手数料体系

取引所の市場運営は営利事業であり、その利益は各投資家が取引を行った際の手数料でまかなっている。メイカー・テイカー制を採用した取引所の利益は、式 (1) で表すことができる。

$$R_{EX} = C_T - R_M \quad (1)$$

R_{EX} は取引所の必要利益、 R_M は取引所がメイカーへ支払うリベート（負の手数料）、 C_T はテイカーが取引所に支払う手数料を表す。Yagi et al.[Yagi 20] と同じく、 $R_{EX} = 0.1\%$ とし、 R_{EX} 、 R_M 、 C_T は後述するファンダメンタル価格に対する比で示す。

2.2 一般投資家エージェント

一般投資家エージェントは、一般的な投資戦略に基づいて取引を行う投資家を想定したエージェントである。一般投資家エージェントは、ファンダメンタル価格を参照し投資判断を行うファンダメンタル戦略、過去の価格推移を利用して投資行動を行うテクニカル戦略、試行錯誤的な投資判断を表すノイズ戦略からなる。また、市場状況の変化にあわせて学習することで、ファンダメンタル戦略とテクニカル戦略の比重を適宜切り替えていく。以下に一般投資家エージェントの注文プロセスを記す。一般投資家エージェントは以下の手順に従い、買いと売りの判断を行う。一般投資家エージェント j が時刻 t の時に予想する価格の変化率（予想リターン） $r_{e_j}^t$ は式 (2) から求められる。

$$r_{e_j}^t = \frac{w_{1_j}^t r_{1_j}^t + w_{2_j}^t r_{2_j}^t + u_j \epsilon_j^t}{w_{1_j}^t + w_{2_j}^t + u_j} \quad (2)$$

ここで、 $w_{i_j}^t$ は時刻 t における一般投資家エージェント j の i 項目の重みであり、シミュレーション開始時にそれぞれ0から w_{imax} までの一様乱数で決める。右辺の分子の1項目の $w_{1_j}^t$ はファンダメンタル戦略の成分の重み、2項目の $w_{2_j}^t$ はテクニカル戦略の成分の重みである。 u_j はノイズ戦略の成分の重みであり、シミュレーション開始時にそれぞれ0から u_{max} までの一様乱数で決められ、シミュレーション中は変化しない。これらの重みは互いに独立して変化する。これら3つの重みからくる影響は式 (2) の右辺の分母にて正規化することで平準化している。

$r_{i_j}^t$ は時刻 t における一般投資家エージェント j の i 項目の予想リターンである。1項目の $r_{1_j}^t$ はファンダメンタル成分のリターンであり、 $\ln(P_f/P^{t-1})$ とする。これは、ファンダメンタル価格と1期前の取引価格を比較し、取引価格の方が低ければ正、高ければ負の予想リターンを意味する。 P_f は時間で変化しない一定のファンダメンタル価格である。ファンダメンタル価格とは、株式を発行する企業自体がもっている実態価値に基づいた価格のことである。 P^t は時刻 t における取引価格（取引されなかった時刻では直近取引された価格であり、 $t = 0$ では $P^t = P_f$ とする）である。2項目の $r_{2_j}^t$ はテクニカル成分の予想リターンであり、 $\ln(P^{t-1}/P^{t-1-\tau_j})$ とする。これは、過去のリターンが正なら正、負なら負の予想リターンを意味している。 τ_j は1から τ_{max} までの一様乱数でエージェントごとに決める。 ϵ_j^t は時刻 t におけるエージェント j のノイズ成分であり、平均0、標準偏差 σ_ϵ の正規分布乱数である。

式 (2) で導いた予想リターンを基に予想価格 $P_{e_j}^t$ を式 (3) で求める¹。

$$P_{e_j}^t = P^{t-1} \exp(r_{e_j}^t) \quad (3)$$

¹本研究では対数リターンを使用している。そのため予想リターン

注文価格 P_{oj}^t は平均 P_{ej}^t , 標準偏差 P_{σ}^t の正規分布乱数で決める。ただし, $P_{\sigma}^t = P_{ej}^t \cdot est$ とする。 est ($0 < est \leq 1$) を便宜上, 「ばらつき係数」と呼ぶ。そして, P_{oj}^t が P_{ej}^t より小さければ, リスク資産 1 単位の買い注文を出し, P_{oj}^t が P_{ej}^t より大ければ, リスク資産 1 単位の売り注文を出す。

一般投資家エージェントは注文価格と注文種別（買いまたは売り）を決定したのちに注文を出す市場を選択する。市場を選択するため, それぞれの市場の注文板から最も安い売り注文（最良売り気配注文）または最も高い買い注文（最良買い気配注文）を確認する。エージェントが決定した注文が採用市場または非採用市場のいずれかで約定する場合, 注文が約定する方の市場に（成行）注文を出す。どちらの市場でも約定しない場合, 採用市場 : 非採用市場 = 50 : 50 で（指値）注文を出す。どちらの市場でも約定する場合, 市場の手数料を考慮した上でより安く（高く）取引できる市場に買い（売り）の（成行）注文を出す。

2.3 アルゴリズムエージェント

アルゴリズムエージェントはアルゴリズムトレードを行う機関投資家を想定したエージェントである。ここでいうアルゴリズムトレードとは予め決められた大口の注文を小口に分けて少しづつ自動的に執行する取引を指す。本モデルではこの戦略を規則的に注文数 1 の成行買い注文でモデル化している。

アルゴリズムエージェントは発注の際に, 最良売り気配注文を確認する。最良売り気配注文が存在すればその注文の価格からティックサイズ ΔP を加えた価格で買い注文を出す。注文板に最良売り気配注文が存在しない場合, 注文は行わない。

2.4 ポジションマーケットメイカー

ポジションマーケットメイカーはマーケットメイク戦略をとる機関投資家を想定したエージェントである。自身のポジション（保有しているリスク資産数, 買いなら正, 空売りなら負）を考慮に入れ, 最良買い気配値と最良売り気配値から注文基準価格を求め, この価格に, 提示スプレッド θ_{pm} (1 取引あたりの期待利益率) を加えた価格で売り注文を, 減じた価格で買い注文を同時に出す [Nakajima 04, 草田 15, Yagi 20]。時刻 t における取引市場の最良売り気配値 $P^{t,sell}$, 最良買い気配値 $P^{t,buy}$, ポジションマーケットメイカーの提示スプレッドを θ_{pm} , 時刻 t と $t+1$ の間にポジションマー

ケットメイカーが抱えるポジションを s_{pm}^t , ポジション考慮度を w_{pm} とすると, 買い注文価格 $P_{o,pm}^{t,buy}$ と売り注文価格 $P_{o,pm}^{t,sell}$ は式 (4) で基準注文価格 $P_{fv,pm}^t$ を求めた後に, 式 (5), 式 (6) で決定される。

$$P_{fv,pm}^t = \frac{P^{t,buy} + P^{t,sell}}{2} \left(1 - w_{pm} (s_{pm}^t)^3\right) \quad (4)$$

$$P_{o,pm}^{t,buy} = P_{fv,pm}^t - \frac{1}{2} P_f \theta_{pm} \quad (5)$$

$$P_{o,pm}^{t,sell} = P_{fv,pm}^t + \frac{1}{2} P_f \theta_{pm} \quad (6)$$

s_{pm}^t は, 買い（売り）注文の取引が成立するごとに 1 増加（減少）する。 s_{pm}^t が正（負）の状態のとき買い（売り）ポジションであることを意味する。 w_{pm} はポジションが買いまたは売りに偏らないようにするために用いる係数である。 w_{pm} を高く設定するほど, ポジションマーケットメイカーのポジションが買い（売り）のとき, 式 (4) の $P_{fv,pm}^t$ がより低く（高く）なるよう設定される。その結果, ポジションマーケットメイカーの買い（売り）注文は最良買い（売り）注文になりにくくなるため取引も成立しなくなるが, 売り（買い）注文は最良売り（買い）注文になりやすくなるため取引が成立しやすくなる。

$P_{o,pm}^{t,buy}$ と $P_{o,pm}^{t,sell}$ にはポジションマーケットメイカーの成行注文を防ぐため価格に制約を加える。その制約を式 (7) に示す。

$$\begin{aligned} P_{o,pm}^{t,buy} &\geq P^{t,sell} \\ P_{o,pm}^{t,sell} &\leq P^{t,buy} \end{aligned} \quad (7)$$

これらの制約を満たすときのポジションマーケットメイカーの発注価格は式 (8), 式 (9) のようになる。これにより買い注文と売り注文の価格が逆転することも防ぐことができる。

$P_{o,pm}^{t,buy} \geq P^{t,sell}$ のとき,

$$\begin{aligned} P_{o,pm}^{t,buy} &= P^{t,sell} - \Delta P, \\ P_{o,pm}^{t,sell} &= (P^{t,sell} - \Delta P) + P_f \cdot \theta_{pm} \end{aligned} \quad (8)$$

$P_{o,pm}^{t,sell} \leq P^{t,buy}$ のとき,

$$\begin{aligned} P_{o,pm}^{t,buy} &= (P^{t,buy} + \Delta P) - P_f \cdot \theta_{pm}, \\ P_{o,pm}^{t,sell} &= P^{t,buy} + \Delta P \end{aligned} \quad (9)$$

本研究ではポジションマーケットメイカーに, 1 取引 (1 単位の売り注文と買い注文) で生じる期待利益を設ける。ポジションマーケットメイカーは期待利益とメイカーのリバート R_M を考慮に入れ提示スプレッド θ_{pm} を調節していく。 θ_{pm} の求め方を式 (10) に示す。

$$\theta_{pm} = Re_M - 2R_M \quad (10)$$

は現在の価格の対数と予想価格の対数の差である。すなわち, $r_{ej}^t = \log P_{ej}^t - \log P^t = \log P_{ej}^t / P^t$ であり, これより式 (3) が導き出される。

表 1: リバートが変化したときの提示スプレッドとテイカーの手数料（ただし, $Re_M = 0.300\%$, $Re_X = 0.100\%$ ）

R_M	θ_{pm}	C_T
-0.050%	0.400%	0.050%
-0.025%	0.350%	0.075%
0.000%	0.300%	0.100%
0.025%	0.250%	0.125%
0.050%	0.200%	0.150%
0.075%	0.150%	0.175%
0.100%	0.100%	0.200%
0.125%	0.050%	0.225%
0.140%	0.020%	0.240%
0.145%	0.010%	0.245%

Re_M はポジションマーケットメイカーが1取引あたりに必要な期待利益である。本研究ではYagi et al.[Yagi 20]と同じく, $Re_M = 0.3\%$ の固定値とし, 対ファンダメンタル価格比で表す。 R_M に係数2がついていることに注意する。これは, メイカーが1取引行ったとき（売り注文と買い注文の両方が約定したとき）に, 取引所から2注文に対するリバートが支払われることを意味している。以上の条件下で, ポジションマーケットメイカーのリバートが変化したときの提示スプレッドとテイカーの手数料の変化の関係を表1に示す。なお, 表1より $R_M = -0.050\%$ のとき, $C_T = 0.050\%$ となり, メイカーとテイカーともに 0.050% の取引手数料を支払うこととなる。そのため, 本研究では $R_M = -0.050\%$ をメイカー・テイカー制非採用時とする。

3 実験

本研究では採用市場と非採用市場の2つの人工市場を構築し, 各市場間での取引シェアの議論を行う。それに加え, 市場選択が各市場へ与える影響を確認するため, 各市場のボラティリティと市場非効率性の確認も行う。採用市場のリバートは表1のようにリバート R_M を -0.050% から 0.125% までの 0.025% 刻みと 0.140% , 0.145% と変化させる。非採用市場のリバートは2.4節で述べたように, $R_M = -0.050\%$ で一定とする。

各パラメータの値は $n = 980$, $m = 20$, $w_{1max} = 1$, $w_{2max} = 10$, $u_{max} = 1$, $\tau_{max} = 10,000$, $\sigma_\epsilon = 0.06$, $est = 0.003$, $\Delta P = 1.0$, $P_f = 10,000.0$, $t_e = 1,000,000$, $\delta_l = 0.01$, $w_{pm} = 0.00000005$ とする。 t_e はシミュレーション終了時の時刻である²。シミュレー

² $t_e = 1,000,000$ とした理由は, この値で実験の傾向を十分に把握することができ, 期間を延ばしても傾向に差異は生じなかったからである。

ションは各条件の下でそれぞれ50回ずつ試行し, その結果をもとに議論を行った。

3.1 取引シェア

各市場間の取引シェアを測定するために, 市場の出来高（取引が成立した回数）を使用する。シミュレーション終了時の各市場の出来高を計測し, 出来高の比率を取引シェアとして算出する。採用市場の出来高を V_A , 非採用市場の出来高を V_B とした場合の, 採用市場の取引シェア S_A と非採用市場の取引シェア S_B の求め方を式(11), 式(12)に示す。

$$S_A = \frac{V_A}{V_A + V_B} \quad (11)$$

$$S_B = \frac{V_B}{V_A + V_B} \quad (12)$$

3.2 市場非効率性

市場の効率性を測定する指標として市場非効率性 M_{ie} を用いる[Mizuta 16]。式(13)に市場非効率性の求め方を示す。

$$M_{ie} = \frac{1}{t_e} \sum_{t=0}^{t_e} \frac{|P_i^t - P_f|}{P_f} \quad (13)$$

P_i^t は時刻 t における i 市場での取引価格（取引されなかった時刻では直近取引された価格であり, $t = 0$ では $P_i^t = P_f$ とする）である。 M_{ie} は0以上の値をとり, 0ならば完全に効率的, 値が大きくなるにつれて非効率性であることを表している。市場非効率性はファンダメンタル価格に対する市場価格の平均乖離度で表している。

4 結果と考察

実験の結果, 採用市場と非採用市場の取引シェア比較すると, 採用市場のリバートが低い間は, 非採用市場の取引シェアの方が高くなった。しかし, 採用市場のリバートが一定に達すると今度は採用市場の取引シェアの方が高くなった。ボラティリティは採用市場のリバートが増加するにつれて, 採用市場では常に低下し, 非採用市場では常に上昇し続けた。市場非効率性は採用市場と非採用市場ともに, 採用市場のリバートが増加につれて減少し, 2市場間での差異はほぼ見られなかった。

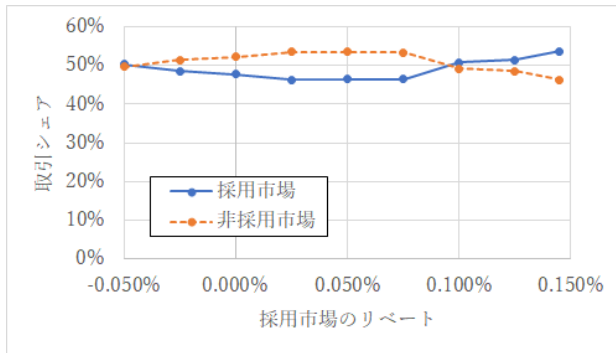


図 1: 各市場の取引シェア

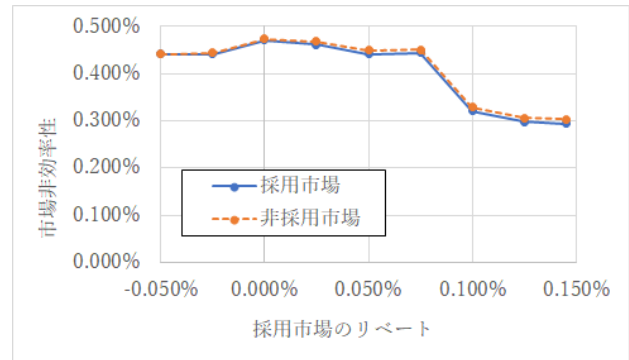


図 3: 各市場の市場非効率性

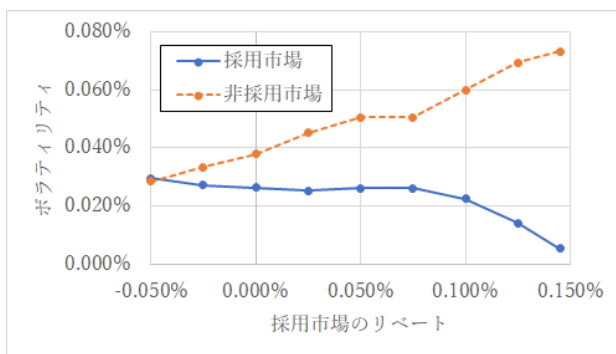


図 2: 各市場のボラティリティ

4.1 取引シェア

各市場間の取引シェアの結果を図 1 に示す。

採用市場のリベートが-0.050%から0.075%の間は、非採用市場の取引シェアの方が高くなった。考えられることの1つとして、この区間では採用市場がメーカーに対して以下に示すような十分なリベートを提供できていないことが挙げられる。

市場がメーカーに提供するリベートを増加させると、それに伴いメーカーは提示スプレッドを狭めることができる。そして、ある地点で市場のビット・アスク・スプレッドより提示スプレッドが小さくなる。すると、メーカーの注文が最良気配値となり、次第にテイカーの取引コストが少なくなることが確認されている [Yagi 20]。取引コストが低くなれば、結果的に採用市場での取引の方が安く買える（高く売れる）ことが多くなり、その結果、採用市場での取引が活発になり取引シェアが高くなる。つまり、十分なリベートとはメーカーの注文が最良気配値となりテイカーの取引コストが減少する金額である。リベートが-0.050%から0.075%の間は採用市場の取引シェアが低くなっているのは、メーカーに十分なリベートを提供できていないためだと言える。

同様の理由で、採用市場のリベートが0.100%から0.145%の間は、メーカーに十分なリベートが提供でき

たため、採用市場の取引シェアが高くなったと考えられる。

4.2 ボラティリティ

各市場のボラティリティの結果を図 2 に示す。本研究のボラティリティはリターン（騰落率）の標準偏差を示している。

採用市場のボラティリティは、リベートがある値より大きくなると低下している。この理由は、4.1 節で述べた取引コストと同様に市場のビット・アスク・スプレッドによることが知られている [Yagi 20]。市場の価格変動幅は最良買い（売り）気配値で決まるため、メーカーの提示スプレッドが小さくなりメーカーの注文が最良気配値になるにつれボラティリティも小さくなる。

採用市場のリベートが増加するにつれて、非採用市場のボラティリティは上昇した。非採用市場のボラティリティは採用市場のリベートが-0.050%から0.075%の間とそれ以降の採用市場のリベート0.100%から0.145%の間の、2段階で上昇したと考察する。

まずは採用市場のリベートが-0.050%から0.075%の間について述べる。この区間の非採用市場の取引シェア（図 1 参照）は採用市場よりも多くなっており活発に取引されている。取引量が多くなればその分価格が頻繁に動くことになり、価格が上下に大きく変化する可能性も増加する。この期間は取引量が多くなったため、ボラティリティが上昇した。

次に採用市場のリベート0.100%から0.145%の間について述べる。非採用市場ではメーカーが提示スプレッドを狭めることはなく、一定で取引をする。採用市場では市場のビット・アスク・スプレッドよりもポジションマーケットメーカーが形成するスプレッドが小さくなるため、メーカーの提示スプレッド付近に注文が集中し、板が安定して形成される。それに比べ、非採用市場では提示スプレッドは市場のビット・アスク・スプレッドより広く、最良気配値付近の注文が採用市場

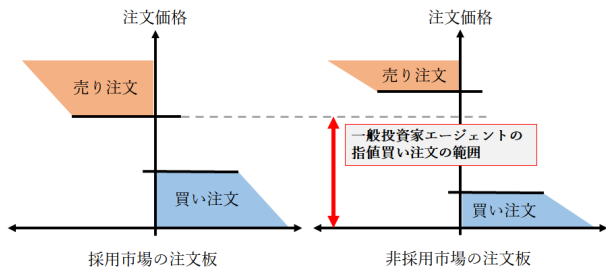


図 4: 一般投資家エージェントの指値注文の範囲（買い注文の場合）

ほど多くない。これは、非採用市場では注文板上の注文が採用市場よりまばらに配置されることを意味する。非採用市場では注文板上の注文がばらけているため、1回の取引で大きく価格が動く可能性が上昇する。これにより、取引量は少ないものの、1回の取引で大きく価格が動くため、ボラティリティが上昇した。

4.3 市場非効率性

各市場の市場非効率性を図3に示す。市場非効率性は採用市場、非採用市場ともにリベートが上昇するにつれ減少していった。

採用市場のリベート-0.050%から0.075%の間、両市場の市場非効率性は上下に多少ぶれているがおおむね一定である。一方で、採用市場のリベート0.100%から0.145%の間、両市場の市場非効率性は減少していることが確認できる。

採用市場の市場非効率性が減少する結果は既に確認されている[Yagi 20]。市場価格がファンダメンタル価格から離れるような値動きをしたとしても、ポジションマーケットメイカーの形成する狭い提示スプレッド付近で注文が約定され、市場価格の変動幅は小さいものになるためである。

一方で、非採用市場ではポジションマーケットメイカーが提示スプレッドを狭めることはないが、採用市場と同様に市場非効率性は減少していった。これは、一般投資家エージェントが市場選択を行うと、採用市場の値動きが非採用市場の値動きに影響を与えるためである。

図4を用いて説明する。図4の通り、採用市場がリベートを十分に提供した場合、採用市場のビット・アスク・スプレッドはポジションマーケットメイカーの提示スプレッドの影響で、非採用市場のそれより小さくなる。この条件の下で、一般投資家エージェントの注文が指値注文になる範囲を考える。一般投資家エージェントは、注文が非採用市場で約定せずとも採用市場で約定するのであれば、採用市場に成行で発注する。

そのため、一般投資家エージェントが買い注文を行う際、その注文が指値注文になるのは図4中の矢印の範囲となる。これは、採用市場の最良売り注文が上限となるため、非採用市場は採用市場より価格が上昇することがなくなることを示している。

指値売り注文でも同様に採用市場の最良買い注文が下限となり、非採用市場は採用市場より価格が低くなることはない。これが要因となって、非採用市場の値動きが採用市場の値動きと近くなる。そのため、採用市場と非採用市場の市場非効率性は同じように減少する結果が得られた。

5 まとめ

本研究ではメイカー・テイカー制を採用した人工市場（採用市場）と、採用していない人工市場（非採用市場）の2つを構築し、その2つの市場間での取引シェアの検討を行った。また、市場選択が各市場へ与える影響を確認するため、各市場のボラティリティと市場非効率性の確認も行った。

その結果、採用市場の取引シェアが向上することが確認できた。しかし、リベートを十分に提供できていない場合、非採用市場に取引シェアが奪われることも明らかになった。ボラティリティは採用市場では低下し、非採用市場では上昇することが確認できた。市場非効率性は採用市場のリベートが上昇すると、どちらの市場でも低下し、またそれらに大きな差がないことも確認できた。

以上より、メイカー・テイカー制は、その制度を採用している市場のボラティリティの低下、市場効率化に加え、他市場の市場効率化も図ることが確認できた。しかし、リベートを適切に提供できていない場合、メイカー・テイカー制を採用した市場が他の市場に取引シェアを奪われてしまう扱いづらい一面があることが今回の実験で明らかになった。取引所がメイカー・テイカー制を採用する場合、闇雲にリベートを提供するのではなく、その市場ではメイカーがどの程度のリベートを望んでいるか、そのリベートを提供できる余裕はあるか、それに見合った見返りを得られるかなどを十分に精査する必要があるだろう。

今回は取引シェアの推移と、市場選択がどのような影響を与えているのかを、主に市場目線で議論した。今後の課題として、メイカー・テイカー制を含む市場選択が、市場参加者の取引コストにどのような影響を与えるのかを調査することが挙げられる。

留意事項

本論文はスパークス・アセット・マネジメント株式会社の公式見解を表すものではありません。すべては個人的見解です。

参考文献

- [Battalio 16] Battalio, R., Corwin, S. A., and Jennings, R.: Can brokers have it all? On the relation between make-take fees and limit order execution quality, *The Journal of Finance*, Vol. 71, No. 5, pp. 2193–2238 (2016)
- [Brolley 20] Brolley, M. and Malinova, K.: Maker-Taker Fees and Liquidity: The Role of Commission Structures, *SSRN Working Paper Series* (2020)
- [Chen 12] Chen, S.-H., Chang, C.-L., and Du, Y.-R.: Agent-based economic models and econometrics, *The Knowledge Engineering Review*, Vol. 27, No. 2, pp. 187–219 (2012)
- [Chiarella 09] Chiarella, C., Iori, G., and Perelló, J.: The impact of heterogeneous trading rules on the limit order book and order flows, *Journal of Economic Dynamics and Control*, Vol. 33, No. 3, pp. 525–537 (2009)
- [Cox 19] Cox, J., Van Ness, B., and Van Ness, R.: Increasing the Tick: Examining the Impact of the Tick Size Change on Maker-Taker and Taker-Maker Market Models, *Financial Review*, Vol. 54, No. 3, pp. 417–449 (2019)
- [Foucault 13] Foucault, T., Kadan, O., and Kandel, E.: Liquidity cycles and make/take fees in electronic markets, *The Journal of Finance*, Vol. 68, No. 1, pp. 299–341 (2013)
- [Mizuta 16] Mizuta, T., Noritake, Y., Hayakawa, S., and Izumi, K.: Affecting market efficiency by increasing speed of order matching systems on financial exchanges-investigation using agent based model, in *2016 IEEE Symposium Series on Computational Intelligence (SSCI)*, pp. 1–8IEEE (2016)
- [Nakajima 04] Nakajima, Y. and Shiozawa, Y.: Usefulness and feasibility of market maker in a thin market, in *the International Conference Experiments in Economic Sciences: New Approaches to Solving Real-world Problems*, pp. 1000–1003 (2004), <https://www.cc.kyoto-su.ac.jp/project/orc/execo/EES2004/proceedings.html>
- [Yagi 20] Yagi, I., Hoshino, M., and Mizuta, T.: Analysis of the impact of maker-taker fees on the stock market using agent-based simulation, *arXiv preprint in 2020 ACM International Conference on AI in Finance*, *arXiv:2010.08992* (2020)
- [Yeh 13] Yeh, C.-H. and Yang, C.-Y.: Do price limits hurt the market?, *Journal of Economic Interaction and Coordination*, Vol. 8, No. 1, pp. 125–153 (2013)
- [岡田 17] 岡田 功太, 齋藤 芳充: 米国株式市場のメイカー・テイカー・モデルを巡る議論-流動性向上策としてのリベートの功罪-, *野村資本市場クォーターリー*, Vol. 21-2, pp. 35–53 (2017)
- [草田 15] 草田裕紀, 水田孝信, 早川聡, 和泉潔: 保有資産を考慮したマーケットメイク戦略が取引所間競争に与える影響: 人工市場アプローチによる分析, *人工知能学会論文誌*, Vol. 30, No. 5, pp. 675–682 (2015)

金融データにおける VC 相関の応用

Application of VC Correlation in Financial Data

落合 友四郎^{1*} ホセ・ナチエル^{2†}
Tomoshiro Ochiai¹ Jose C. Nacher²

¹ 大妻女子大学社会情報学部

¹ Faculty of Social Information Studies, Otsuma Women's University

² 東邦大学理学部情報科学科

² Department of Information Science, Faculty of Science, Toho University

Abstract: 通常の相関係数は要素間の相関はわかるが、その因果関係を含めた方向性まではわからない。ところが VC 相関を用いると、その要素間の相関とともにその方向性まで推定することができる。これまで、この VC 相関を金融やその他分野に応用してきた。日米株式相関・為替レートの間方向性、日中と夜間の収益率の相関、また金融以外への応用として遺伝子制御関係の推定にもこの VC 相関を応用した。最新の財務データへの応用も含めて、VC 相関のメトリックとして優れた点と展望を検証する。

1 はじめに

これまで様々な分野の研究において、時系列データの相関係数を計算することにより、その背後のネットワークを推定する努力が進められてきた。これら時系列データから、背後のネットワークを同定する研究は、いわゆる通常のピアソン相関係数 (Pearson product-moment correlation coefficient) を用いられることが多い。ところが、通常の相関係数では、相関の有無はわかって、その方向性まではわからない。つまり時系列データ A, B に対して、相関係数 $\text{Cor}(A, B)$ と、 $\text{Cor}(B, A)$ では同じ値をとるために、A と B のどちらが制御側で、どちらが被制御側かわからないのである。そこで、我々は Volatility-Constrained-Correlation (VC-correlation) と呼ぶ新しいタイプの相関を計測する手法を開発した。この VC-correlation を用いることにより、2つの要素 A, B の間の相関のみならず因果の方向性まで検出できるようになった。日米株式相関・為替レートの間方向性、日中と夜間の収益率の相関、また金融以外への応用として遺伝子制御関係の推定にこの VC 相関を応用した。

2 提案方法

$P_i(t)$ ($i = 1, 2$) を金融資産 i の時刻 t ($t_i \leq t \leq t_f$) における値とする。すると対数収益率は以下で定義さ

れる。

$$R_i(t) = \ln P_i(t+1) - \ln P_i(t).$$

すると、平均、標準偏差は通常通り以下で定義される。

$$E(R_i(t)) = \frac{1}{(t_f - t_i)} \sum_{t_i \leq t < t_f} R_i(t),$$

$$\sigma(R_i(t)) = \sqrt{\frac{1}{(t_f - t_i)} \sum_{t_i \leq t < t_f} (R_i(t) - E(R_i(t)))^2}.$$

ここで、 t_i と t_f は、データセットの最初および最後の時刻である。さらに、時系列データペア $\{(R_1(t), R_2(t))\}$ ($t_i \leq t < t_f$) に対して、相関係数 (Pearson product-moment correlation coefficient) は以下で定義される。

$$C(R_1(t), R_2(t)) = \frac{1}{(t_f - t_i)} \times \sum_{t_i \leq t < t_f} \frac{(R_1(t) - E(R_1(t)))}{\sigma(R_1(t))} \frac{(R_2(t) - E(R_2(t)))}{\sigma(R_2(t))}.$$

金融市場では、2つの金融資産が相関する現象は頻繁に観察される。さらに、資産 1 から資産 2 への影響の強さが、その逆方向 (資産 2 から資産 1) の影響の強さと異なることも良く観察される。そこで、この現象をとらえるために、以下のように volatility-constrained correlation (VC 相関) と呼ぶメトリックを導入する。

一般に Ω を、時系列 $\{t | t_i \leq t < t_f\}$ の部分集合とする。この部分時系列 Ω に対して、フィルター (制限)

*連絡先: E-mail: ochiai@otsuma.ac.jp

†連絡先: E-mail: nacher@is.sci.toho-u.ac.jp

された平均、標準偏差、相関係数を以下で定義する。

$$E(R(t), \Omega) = \frac{1}{\#\Omega} \sum_{t \in \Omega} R(t),$$

$$\sigma(R(t), \Omega) = \sqrt{\frac{1}{\#\Omega} \sum_{t \in \Omega} (R(t) - E(R(t), \Omega))^2},$$

$$C(R_1(t), R_2(t), \Omega) = \frac{1}{\#\Omega} \sum_{t \in \Omega} \frac{(R_1(t) - E(R_1(t), \Omega))}{\sigma(R_1(t), \Omega)} \times \frac{(R_2(t) - E(R_2(t), \Omega))}{\sigma(R_2(t), \Omega)},$$

ここで、 $\#\Omega$ は、 Ω の要素の数である。

時系列部分集合 Ω の特別な場合として、以下のよう
におく。

$$\Omega_{[t_1, t_2; \alpha, \beta]} = \{t \in [t_i, t_f] \mid t_1 \leq t < t_2 \text{ and } \alpha \cdot \sigma(R_1(t)) \leq |R_1(t)| < \beta \cdot \sigma(R_1(t))\},$$

ここで、 $t_i \leq t_1 < t_2 \leq t_f$ である。

そこで、次のように volatility-constrained correlation を定義する。

$$F[\alpha, \beta](s) = C(R_1(t), R_2(t), \Omega_{[s, s+\Delta s; \alpha, \beta]}),$$

言い換えると、volatility-constrained correlation $F[\alpha, \beta](s)$

は、一方の資産の $|R_1(t)|$ が特定のレンジに制限されているときの $(\alpha \cdot \sigma(R_1(t)) \leq |R_1(t)| < \beta \cdot \sigma(R_1(t)))$ 、2つの金融資産 $(R_1(t), R_2(t))$ の間の特定時期の相関係数を表す。制限を課す方の資産（影響の起点となる方） $R_1(t)$ を base asset と呼ぶことにする。

$F[\alpha, \beta](s)$ は、 $R_1(t)$ のボラティリティーによって制限された相関係数であり、もし、二つの資産 $R_1(t)$ と $R_2(t)$ を交換すると、 $F[\alpha, \beta](s)$ は異なる値をとる。いかえると、この定義は、2つの資産 $R_1(t)$ と $R_2(t)$ の交換について非対称である。この非対称性は、2つの資産価格間の影響の方向性を検知するのに重要な役割を果たす。ここで、この2種類の VC 相関の値を比べることにより影響の方向性を推定することができる。例えば、2つの資産 $R_1(t)$ と $R_2(t)$ を交換したとき、 $F[\alpha, \beta](s)$ が減少したら、資産 $R_1(t)$ から資産 $R_2(t)$ の方向に影響を与えていることがわかる。逆に、2つの資産 $R_1(t)$ と $R_2(t)$ を交換したとき、 $F[\alpha, \beta](s)$ が増大したら、資産 $R_2(t)$ から資産 $R_1(t)$ の方向に影響を与えていることがわかる。（図1）

3 VC 相関の応用 1

日経平均株価と他の金融資産との影響伝播の方向性（VC 相関）を分析した [1]。具体的には、日経平均株価、

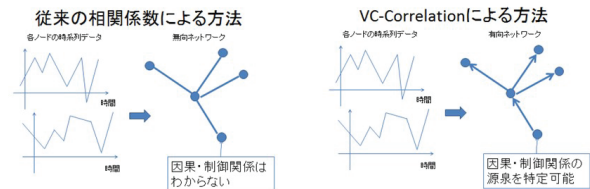


図 1: VC 相関。

米ドル為替レート、DJIA(ダウ平均株価) の間の VC 相関の非対称性をしらべた。この分析により、日経平均株価が米ドル円と DJIA の影響を受けていることが統計的に示すことができた。これは、日米間の国際的な経済力のバランスを反映していると解釈できる。

4 VC 相関の応用 2

VC 相関は、金融のみならず他分野に応用でき、実績を上げている [2]。疾患分子経路やシグナル伝達ネットワークを明らかにするためには、無方向性ネットワークだけでなく細胞構成要素間の機能的相互作用や物理的相互作用の方向性を推論する必要がある。遺伝子間の機能的相互作用を同定するための一般的な方法は、実験的な遺伝子発現測定値間の相関関係にある程度依存している。しかし、標準的なピアソンまたはスピアマン相関に基づくアプローチでは、細胞成分間の無方向性の相関関係を決定することしかできない。論文 [2] では、遺伝子発現プロファイルに対して、遺伝子間の相互作用の方向性を捉えるための新しいメトリックであるボラティリティー制約相関法を適用する。予測結果を評価するために、DREAM5 ネットワークの4つのデータセットを用いた。VC 相関による解析結果は、実験的に検証された遺伝的調節リンクの方向性のデータと比較した。その結果、本手法は高い統計的有意性を持って遺伝的相互作用の方向性を予測することに成功していることがわかった。

5 VC 関連の応用 3

金融工学や経済物理学の研究では、日中の取引に焦点を当てた研究が多いが、非取引時間帯に焦点を当てた研究は少ない。論文 [3] では、前日のオーバーナイト・リターンのと日中のリターンの相関関係（相関 ND）と、日中のリターンと翌日のオーバーナイト・リターンの相関関係（相関 DF）を調査した結果、いくつかの知見が得られた。第一に、日本の株式市場では、オーバーナイト・リターンと日中リターンの間に弱い負の相関（相関 ND）が観測された。第二に、VC 相関法を適用することで、このシグナルが有意に増幅され、標準的な相関に比べて日中リターンの予測可能性が高まることがわかった。さらに、VC 相関から得られた増幅シグナルを各銘柄ごとに分析したところ、標準相関と VC 相関の間には直線的なスケールの関係が見られた。また、VC 相関から得られる影響の方向性は時間の方向性と一致することが確かめられた。

6 VC 関連の応用 4

財務諸表などの財務データには多くの変数がありますが、その中から因果関係、すなわち、各変数の方向性を見出すために相関（VC 相関）法を用いて、2 つの変数の間の方向関係を予測した [4]。正確には、VC 相関法を、1990 年から 2018 年までの 28 年間の東京上場企業 2321 社の売上高、当期純利益、営業利益、自己資本、時価総額の時系列データに応用した。

7 終わりに

VC 相関を用いて、日米株式相関・為替レートの間の方方向性、日中と夜間の収益率の相関、また金融以外への応用として遺伝子制御関係の推定にもこの VC 相関を応用した。これらの結果により、VC 相関は時系列データの要素間の方向性を推定するのに優れたメトリックであることがわかり、金融やそれ以外への分野へのさらなる応用が期待できる。

謝辞

T.O. was partially supported by JSPS Grants-in-Aid for Scientific Research (Grant Number 15K01200).

参考文献

- [1] T. Ochiai, J.C. Nacher, “Volatility-constrained correlation identifies the directionality of the influence between Japan’s Nikkei 225 and other financial markets”, *Physica A* 393, 364–375 (2014)
- [2] T. Ochiai, J.C. Nacher, “Predicting link directionality in gene regulation from gene expression profiles using volatility-constrained correlation”, *Biosystems*, Volume 145, 9–18 (2016)
- [3] T. Ochiai, J.C. Nacher, “VC correlation analysis on the overnight and daytime return in Japanese stock market” *Physica A*, Volume 515, 537-545(2019)
- [4] T. Ochiai, J.C. Nacher, “Unveiling the directional network behind the financial statements data using volatility constraint correlation”, preprint, arXiv:2008.07836 [q-fin.GN]

非同期時系列の Lead-Lag 効果推定のための新しい推定量 Mining Lead-Lag Relationship from Non-synchronous Data

伊藤 克哉^{1*} 中川 慧²
Katsuya Ito¹ Kei Nakagawa²

¹ Preferred Networks 株式会社

¹ Preferred Networks, Inc.

² 野村アセットマネジメント株式会社

² Nomura Asset Management Co, Ltd.

Abstract: リードラグ効果は、金融市場のいたるところで観察され、特に高頻度データを用いた投資戦略を策定する上で重要な要因である。一方で、Lead-Lag 効果を推定するために、次の3つの課題が存在する。(1) 高頻度データにおいては、2つの時系列が常に同時に観測できるとは限らない(非同期)。(2) 高頻度データのサイズは大きく、応用のためには短期間で推定を完了する必要がある。(3) Lead-Lag 効果は時变的であり、かつ短期間しか持続せず、しばしば外部要因の影響を受ける。本研究では、これらの課題をすべて解決する、新しい Lead-Lag 効果の推定量 (NAPLES: Negative And Positive lead-lag EStimator) を提案する。人工データセットと実際の金融市場のデータセットを用いた実験の結果、NAPLES は、重要なマクロ経済のアナウンス (外部要因) によって引き起こされたものを含め、Lead-Lag 効果と強い相関関係が確認できた。

1 はじめに

裁定機会 (arbitrage opportunity) の解析は数理ファイナンスにおいて重要な課題の一つである。裁定機会の代表的なものに一つの銘柄の値動きが他の銘柄の値動きよりも先行または遅延する Lead-Lag 効果がある [19, 21, 12, 9]。先行している銘柄の値動きに合わせて遅延している銘柄を売買することによりリスクを取ることなくリターンを発生することができるため、Lead-Lag 効果については古くから計量経済学や金融工学の分野で研究がされてきた [21, 12]。

投資判断・市場構造分析において、より高速・正確に Lead-Lag 効果を推定することの重要性が近年増化している。その理由の1つは、近年、高速な情報通信と計算機を用いた電子取引が可能となり、市場はより効率的になり Lead-Lag 効果の探索は難しくなったからである [8, 6, 24]。2つ目の理由として、直近ではミリ秒・ナノ秒単位で電子取引を行う高頻度取引が登場したため、より短い時間の Lead-Lag 効果を利用して収益を得ることが可能となった [3, 25, 22]。

はじめに Lead-Lag 効果の推定の方法論を数学的に定式化する。この推定は、2つ時系列を観測し、一方のタイムスタンプをずらしながら相関係数を最大化す

る方法が一般的である。資産価格 X_t を $0 = s_0 < s_1 < s_2 < \dots < s_n$ という時刻で観測し、資産価格 Y_t を $0 = t_0 < t_1 < t_2 < \dots < t_m$ という時刻で観測する。 X_t が Y_t を **Lead-Lag** パラメタ θ で先行しているとは、任意の $\theta' \neq \theta$ について、 X_t と $Y_{t+\theta}$ の相関が X_t と $Y_{t+\theta'}$ の相関よりも強いことをいう。この Lead-Lag パラメタを推定するには、 $\hat{\text{Corr}}$ という相関係数の推定量を定義し、以下の最大化問題を解く。

$$\hat{\theta} = \operatorname{argmax}_{\theta \in \Theta} \hat{\text{Corr}}(\{X_{s_i}\}_{i=1}^n, \{Y_{t_j+\theta}\}_{j=1}^m) \quad (1)$$

高頻度取引における Lead-Lag 効果の推定における重要かつ難解な課題は大きく3つあり、それらは、(1) タイムスタンプの非同期性、(2) 計算量の膨大性、(3) Lead-Lag 効果の時变性である。

(1) タイムスタンプの非同期性とは、tick データに対して Lead-Lag 効果を計算するときの課題である。tick データは市場で起こった全てのイベントを保存するので、タイムスタンプが揃っていない。この課題に対しては Hayashi & Yoshida(2005) が非同期に観測される時系列データに対して共分散推定量を定義し [16]、Hoffman, Rosenbaum, & Yoshida(2013) がその推定量を最小化することによって得られる Lead-Lag 効果の推定量を定義したことにより、解決方法が提案された [17]。しかし、その手法を実世界に適用するために様々な理論的解析 [10, 14, 15] および実証的解析 [1, 5, 7, 18] が現時点でも行われている。

*連絡先: Preferred Networks 株式会社
〒100-0004 東京都千代田区大手町 1-6-1 大手町ビル 2F
E-mail: katsuyalito@gmail.com

(2) 計算量の膨大性とは、高頻度取引が一般的になるにつれて、計算量が膨大となる課題である。データ数 T 、銘柄数 N 、ラグ候補数 L としたときに現在の計算量は、 $O(T^2 N^2 L)$ のオーダーでの計算量が多く、これ以上の計算オーダーの計算アルゴリズムは開発されていない。

(3) Lead-Lag 効果の時変性とは、Lead-Lag パラメタが時間に応じて変化することで推定が困難になるという課題である。基本的に Lead-Lag 効果は、情報の非対称性 [24] やトレーディング・ボリュームの差 [11] によって価格変更が遅れていること等に起因して発生し、常に発生しているわけではない [2, 4, 23]。そのために時間変化するラグを扱うことのできる方法が必要とされている [20]。本論文では、これらの 3 つの問題を解決するための Lead-Lag 効果の新しい推定量を提案する。提案方法は、非同期に観測される時系列のセットを使用して実行できるだけでなく、観測数に対して線形な計算複雑性で実行が可能である。本論文の主な貢献は下記の通りである。

- 3 つの課題を解決する Lead-Lag 効果の新しい推定法を提案した。
- 推定法の数学的な性質を与えるとともにその解釈を行った。
- 人工データ・実データを用いた実験により、提案手法の有効性を示すとともに、マクロ指標の発表前にリードラグの推定を行い、その考察を行った。

2 問題設定と先行研究

ここでは、リードラグ効果の推定の定式化と先行研究において提案された手法を概観する。

2.1 Notation

2 つの実数 $a, b \in \mathbb{R}$ の区間を次のように書く。 $[a, b] = \{x | a \leq x \leq b\}$, $]a, b[= \{x | a < x < b\}$, $[a, b[= \{x | a \leq x < b\}$, $]a, b] = \{x | a < x \leq b\}$ また、 $a \vee b$ と $a \wedge b$ はそれぞれ $\max\{a, b\}$ と $\min\{a, b\}$ を表す。 1_A は条件 A についての指示関数である。

2.2 問題設定

X_t, Y_t を 2 つの確率過程 (e.g., 資産価格) とする。 X_t は時間間隔 $0 = s_1 < s_2 < \dots < s_n = T$ で観察され、 Y_t は時間間隔 $0 = t_1 < t_2 < \dots < t_m = T$ でそれぞれ観察されるとする。はじめに Lead-Lag 効果と非同期観測の定義を行う。

Definition 1 X_t が Y_t に対して、Lead-Lag パラメタ θ の Lead-Lag 効果を持つとは、任意の α に対して、 $X_t, Y_{t+\alpha}$ の相関係数の絶対値が $X_t, Y_{t+\theta}$ よりも大きいことをいう

Definition 2 X_t と Y_t が、すべての i に対して、 $s_i = t_i$ が成立する場合、同期観測であるといい、同期観測でない場合、非同期観測という。

本研究における問題設定は、非同期観測されたデータから Lead-Lag 効果を推測することである。つまり、

Problem 1 X_t が Lead-Lag パラメタ θ を持ち、 Y_t を先行すると仮定する。このとき、非同期観測データ $\{X_{s_i}\}_{i=0}^n$ および $\{Y_{t_j}\}_{j=0}^m$ から、Lead-Lag パラメタ θ を推定する。

2.3 先行研究

高頻度取引が盛んになる前は、同期観測のデータに対する Lead-Lag の推定量が Lo & MacKinlay(1990)[21] や de Jong & Nijman [12] によって考案されていた。

高頻度取引が盛んになり非同期観測場合における Lead-Lag の推定のために Hayashi & Yoshida(2005) は次のような共分散推定量を提案した、

Definition 3 資産価格 X_t を $s_1 < s_2 < \dots < s_n$ という時刻で観測し、資産価格 Y_t を $t_1 < t_2 < \dots < t_m$ という時刻で非同期的に観測するとき、Hayashi-Yoshida の共分散推定量は $HY(X_s, Y_t, \{s_i\}_{i=1}^n, \{t_j\}_{j=1}^m)$

$$= \sum_{i=1}^{n-1} \sum_{j=1}^{m-1} (X_{s_{i+1}} - X_{s_i})(Y_{t_{j+1}} - Y_{t_j}) 1_{[t_{i-1}, t_i] \cap [s_{j-1}, s_j] \neq \emptyset}$$

のように定義されるものである。この共分散推定量を用いて、

$$U^n(\tilde{\theta}) := 1_{\tilde{\theta} \geq 0} HY(X_{s-\tilde{\theta}}, Y_t, \{s_i - \tilde{\theta}\}_{i=1}^n, \{t_j\}_{j=1}^m)$$

$$+ 1_{\tilde{\theta} < 0} HY(X_{s+\tilde{\theta}}, Y_t, \{s_i + \tilde{\theta}\}_{i=1}^n, \{t_j\}_{j=1}^m)$$

のように時系列をシフトした共分散推定量を作ることによって Lead-Lag 効果の推定する。

この推定量の理論的拡張として、Chiba(2017) はこの推定量を非整数ブラウン運動の場合に拡張した。また、Hayashi & Koike(2017) [14] はフーリエ変換を用いてこの推定量の改良を行い、Hayashi & Koike(2018) [15] はウェーブレット変換を用いてこの推定量の改良を行った。

一方でこれらの研究と独立した形で、Dobrev & Schaumburg は次のような Lead-Lag 推定のための数値を考案している。

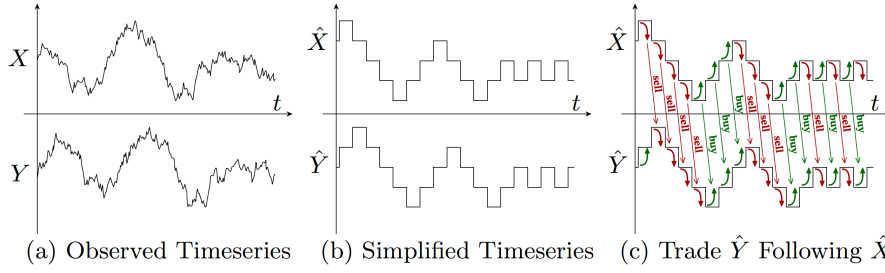


図 1: アルゴリズムの概念図

Definition 4 (式 (2) in [13]) 2 資産 X, Y が観測されている状況を考える。確率過程 $\hat{Z}_t$ は、ある資産 $Z = X, Y$ が時刻 t において取引されているときに 1, それ以外の時に 0 を返すような確率過程である。このとき、Dobrev-Schaumburg の指数 DS_t とは次のように計算されるものである。

$$DS_t := \frac{\sum_{i=|t|}^{N-|t|} \hat{X}_{i+t} \hat{Y}_i}{\sum_{i=|t|}^{N-|t|} \hat{X}_{i+t} \wedge \sum_{i=|t|}^{N-|t|} \hat{Y}_i}.$$

Dobrev-Schaumburg の指数は Hoffman-Rosenbaum-Yoshida と同様に t に関して探索を行うことによって Lead-Lag 効果の推定を行うことができる。その統計的性質は明らかになっていないが、実証的研究によってその有効性が明らかにされている。

3 提案手法:NAPLES

提案手法は、(1) 時系列の符号による単純化と (2) 単純化した時系列をもちいた投資戦略の計算をする。(1) によって、時系列を単純化し、Lead-Lag に機敏に反応する推定量を作ることができる。また (2) によって、タイムスタンプの非同期性や Lead-Lag パラメタの時間変化を考慮した、高速な計算をすることができる。

ここで、 $\{Z_{u_k}\}_{k=1}^n \in \{\{X_{s_i}\}_i, \{Y_{t_j}\}_j\}$ を観測したとき、その対数リターンとリターンの符号をそれぞれ $r_{u_k}^{(Z)} = \log(Z_{u_k}/Z_{u_{k-1}})$, $b_{u_k}^{(Z)} = \text{sign}\{r_{u_k}^{(Z)}\}$ と書く。このとき、 $\hat{Z}_t$ を $b_{u_k}^{(Z)}$ の累積和 $\hat{Z}_t = \sum_{k=1}^n 1_{u_k < t} b_{u_k}^{(Z)}$ として定義する。

Definition 5 (NAPLES; Negative And Positive lead-lag EStimator) X_t を時刻 $0 = s_1 < s_2 < \dots < s_n = T$ で観察し、 Y_t を $0 = t_1 < t_2 < \dots < t_m = T$ で観察する。このとき、指数 NAPLES, $R(t; X_t, Y_t)$ は次のように定義される。

$$R(t; X_t, Y_t) := \sum_{i=1}^{n-1} \left(b_{s_i}^{(X)} \hat{Y}_{s_{i+1}} - b_{s_i}^{(X)} \hat{Y}_{s_i} \right) 1_{s_{i+1} < t} - \sum_{j=1}^{m-1} \left(b_{t_j}^{(Y)} \hat{X}_{t_{j+1}} - b_{t_j}^{(Y)} \hat{X}_{t_j} \right) 1_{t_{j+1} < t}$$

この定義は一見すると複雑に見えるが、以下のような単純な発想に基づくものである。つまり、 X が Y に対して先行しているとき、 Y を X の価格変動に応じて取引したときの利益または損失が $R(t; X_t, Y_t)$ である。具体的には、 $R(t; X_t, Y_t)$ の第一項は、時刻 s_i で X の価格が上昇(下落)したときに、時刻 s_i で $\hat{Y}$ を買い(売り)、時刻 s_{i+1} で $\hat{Y}$ を売る(買う)戦略のリターンを表している。一方で、第二項は、時刻 s_i で Y の価格が上昇(下落)したときに、時刻 s_i で $\hat{X}$ を買い(売り)、時刻 s_{i+1} で $\hat{X}$ を売る(買う)戦略のリターンを表している。図 1 が我々のアルゴリズムの概念図である。これを用いると、Lead-Lag パラメタ θ は次のように計算することができる。

$$\hat{\theta} = \operatorname{argmax}_{\theta \in \Theta} R(T, X_t, Y_{t+\theta}) \quad (2)$$

NAPLES は、上記の 3 つの問題を解決する。すなわち、(1) 非同期観測の時系列に対して計算可能であり、(2) 推定量は単に N 項の合計であるため、 $O(N)$ のオーダーで計算可能であり、(3) 取引戦略のリターンであるため、各時点ごとに Lead-Lag を推定でき、Lead-Lag の時変的な性質を捉えることができる。

さらに、 $R(t; X_t, Y_t)$ の性質を表す下記の定理を示す。

Theorem 1 二つの資産価格 X, Y が相関 ρ 、Lead-Lag パラメタ θ を持つ幾何ブラウン運動に従うとする。さらに観測時刻 s_i と t_j が常に等間隔 Δ であると仮定する。このとき、 $R(T; X_t, Y_t)$ は次を満たす。

$$E[R(T; X_t, Y_t)] = \begin{cases} L \arcsin\left(\frac{\rho|\theta|}{\Delta}\right) & 0 < |\theta| < \Delta \text{ のとき} \\ L \arcsin\left(\frac{\rho(2\Delta-|\theta|)}{\Delta}\right) & \Delta < |\theta| < 2\Delta \text{ のとき} \\ 0 & \text{それ以外の場合} \end{cases}$$

ただし l は $\theta > 0$ のとき $m-1$ で、 $\theta < 0$ のとき $n-1$ であるような定数であり、 $L = \frac{l}{\pi} \text{sign}(\theta)$ である。

この定理の証明を一般化した定理とともに Appendix に示す。

3.1 提案手法の定性的解釈

前節では提案手法の数学的性質を明らかにした。この節では提案手法の具体例を計算することによってそ

の原理を定性的に明らかにする。Theorem 1から等間隔に観測される場合において我々の指数の期待値は以下になることがわかっている。

$$E[R(T)] = 1_{0 < |\theta| < \Delta} \frac{l}{\pi} \text{sign}(\theta) \arcsin\left(\frac{\rho|\theta|}{\Delta}\right) + 1_{\Delta < |\theta| < 2\Delta} \frac{l}{\pi} \text{sign}(\theta) \arcsin\left(\frac{\rho(2\Delta - |\theta|)}{\Delta}\right).$$

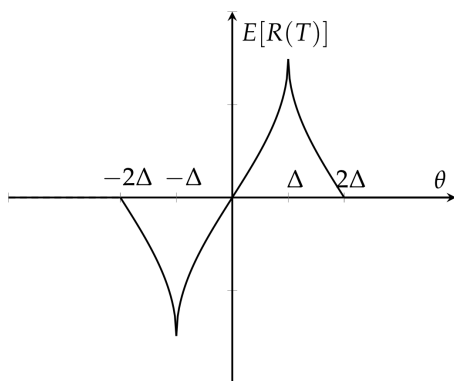


図 2: θ を変化させたときの $E[R(T)]$ の変化

図 2は、Lead-Lag 効果を表す θ のみを動かしたときの我々の推定量の値である。その定性的性質を考察する。まず、この手法において絶対値が 2Δ よりも大きなラグを検出することは、 $E[R(T)]$ が常に 0 であるためできない。これはこの手法の技術的限界を示すわけではなく、 X_t, Y_t という系列があったときにも $|\theta| > 2\Delta$ であったとしても $X_t, Y_{t+\alpha}$ について $R(T)$ を計算することで $|\theta - \alpha| < 2\Delta$ とすることで Lead-Lag 効果の検出が可能になるということを示している。これによって $|\theta - \alpha| > 2\Delta$ を満たすような α に関しては全く反応がないので検出の精度を上げることができる。

また、 $|\theta| < \Delta$ に関しては $E[R(T)]$ が単調増加している。これはラグが大きければ大きいほど早く先行する銘柄の動きを捉えることができ、次の Δ 秒後のブロックで売買を行うことによって利益を得ることができることを示している。一方で $|\theta| > \Delta$ に関しては $E[R(T)]$ の絶対値が単調減少している。これはラグが観測間隔よりも大きいときは、検出できたとしてもその売買を行うブロックが遅れた 2 つ遅れたブロックになってしまうことから、ブラウン運動の独立増分性からその先行者から得られる情報を活用できなくなることを示している。

一般的に Δ は秒からミリ秒やナノ秒単位に縮小する傾向があり、この関数は 2 つの正と負の δ 関数を組み合わせたような形になることが考えられる。そのような場合において、 α による探索を行うことでより良い Lead-Lag 効果の推定量を得ることが考えられる。一方で Δ が比較的大きな場合においても絶対値が Δ の中では

指数が単調増加することから良い推定量が得られる事がわかる。

4 実験と考察

本章では提案手法の有効性を示すために、人工データ及び実データによる実験を行う。本章では、先行研究と同様に、Hoffman-Rosenbaum-Yoshida の手法 (HRY) と Dobrev-Schaumburg の手法 (DS) との比較実験を行う [15, 13]。

4.1 人工データの問題設定

実験のために次のような確率微分方程式に従って生成したデータを用いる。

$$\begin{cases} X_t = x_0 \exp\left(\sigma_1 B_t^{(1)}\right) \\ Y_t = y_0 \exp\left(\rho\sigma_2 B_{t-\vartheta}^{(1)} + \sigma_2 (1 - \rho^2)^{1/2} W_{t-\vartheta}\right) \end{cases}.$$

ここで各実験の数値はそれぞれ、観測する時刻 $0 \leq t \leq 10000$ 、各資産の初期値 $x_0 = y_0 = 100$ 、各資産のボラティリティ $\sigma_1 = \sigma_2 = 1.0$ 、Lead-Lag パラメータ $\theta = 10$ 、共分散 $\rho \in \{1.0, 0.9, 0.8, 0.7, 0.6, 0.5\}$ とした。

4.2 Lead-Lag の検出力の実験

本実験では、Lead-Lag を検出できる能力について考察する。非同期観測の幾何ブラウン運動での人工データを用いた実験結果を示す。一般的に時系列の相関係数が低くなっているときの Lead-Lag 効果は測定することが難しいので本実験では相関係数を動かし、それ以外のパラメータは固定した場合での正答率を見る。ここでの共分散 $\rho \in \{1.0, 0.9, 0.8, 0.7, 0.6, 0.5\}$ とした。推定する Lead-Lag はある候補の中から選ばれるがその項は $\{-100, -99, \dots, 99, 100\}$ とした。Lead-Lag を検出できたか否かについては、Lead-Lag の真値である $\theta = 10$ から上下 2 の間である $\{8, 9, 10, 11, 12\}$ のうちのどれかが推定された場合それは正解であるとし、それ以外は不正解とした。実験は 1000 回乱数を発生させてその中での正答率を見るというセットを 1 セットとし、そのセットを 100 セット行う事によって、平均と標準偏差を計算した。

表 1 はその正答率を示す。表から HRY は共分散が 1 以外の場合においては極端に悪い性質を示していることがわかる。また DS はハイパーパラメータが最も良いものをチューニングして使っており、これは同期観測においては我々の手法を $\{0, 1\}$ に直しただけである。故に同様の結果が得られていることがわかる。

表 1: 非同期観測の人工データにおける Lead-Lag 効果推定

Methods	$\rho = 1.0$	$\rho = 0.9$	$\rho = 0.8$	$\rho = 0.7$	$\rho = 0.6$	$\rho = 0.5$
HRY[17]	0.974 ± 0.004	0.944 ± 0.006	0.939 ± 0.006	0.931 ± 0.006	0.919 ± 0.008	0.898 ± 0.009
DS[13]	1.000 ± 0.000	0.984 ± 0.004	0.964 ± 0.005	0.939 ± 0.007	0.918 ± 0.007	0.887 ± 0.008
提案手法	1.000 ± 0.000	1.000 ± 0.000	1.000 ± 0.000	0.999 ± 0.001	0.999 ± 0.001	0.994 ± 0.002

4.3 Lead-Lag の推定における収束速度の比較

本実験では人工データの観測時間を増やすことによって各アルゴリズムのが非同期観測の条件下においてどのような速度で収束するのかを見る。各モデルとパラメータはと次のようにして行う。共分散 $\rho = 0.9$ 観測間隔は平均 10 標準偏差 2 の正規分布に従うとした。推定する Lead-Lag はある候補の中から選ばれるがその項は $\{-100, -99, \dots, 99, 100\}$ とした。本実験においては観測する時刻 $t = 0, 1, \dots, T$ の最大値である T を動かすことによって、各アルゴリズムが収束するまでにどれだけの時間を要するかを調べる。 T に関しては、 $T \in \{10^{2.5}, 10^{2.6}, 10^{2.7}, \dots, 10^{4.9}, 10^{5.0}\}$ を用いる。 T が大きくなるにつれて相対的に観測間隔は小さくなるため、これは HRY の収束の仮定を満たすことにもなる。実験に関しては 1,000 回実験を行った際の推定値の平均絶対誤差をもって評価する。つまり、乱数を $n = 1,000$ 回発生させ、 i 番目の実験において各アルゴリズムによって推定された Lead-Lag を $\hat{\theta}_i$ とし、真値を θ_i とするとき、そのアルゴリズムの平均絶対誤差は $\frac{1}{n} \sum_{i=1}^n |\hat{\theta}_i - \theta_i|$ となる。

図 3 はその収束の様子を示すグラフである。図からわかる通り、HRY は収束速度について早いだが、正確な値に収束することはない。DS はそれに比べると良い値に収束するが、それでも収束が遅いことがわかる。一方で提案手法は DS と同様の収束速度を保ちながら、有意に正確な値を推定することができる。

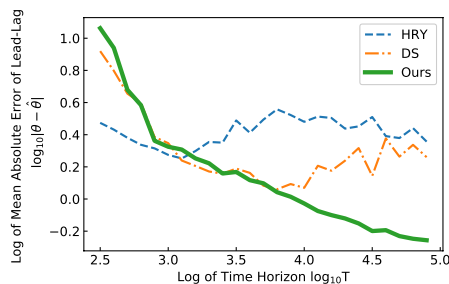


図 3: 各アルゴリズムの収束の比較

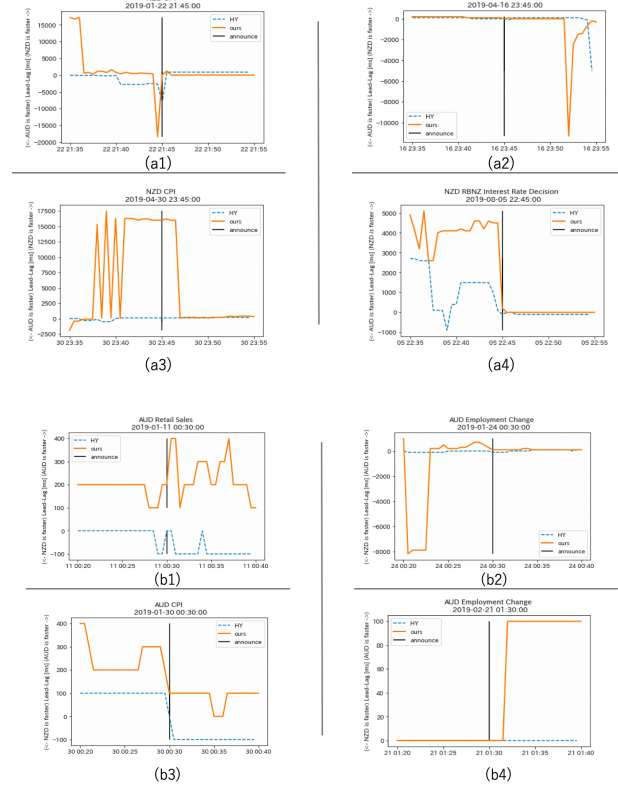


図 4: 重要なマクロ指標発表前の Lead-lag パラメータ。

4.4 実データにおける Lead-Lag の推定

本節では、提案手法が実データにおいても有効であることを示す。本実験では、AUD/USD と NZD/USD という 2 つの地理的に近く似た通貨とされる通貨ペアを用いた。ただし、AUD はオーストラリアドル、USD はアメリカドル、NZD はニュージーランドドルの略称である。我々は、2019 年 1 月から 2019 年 9 月 30 日までのデータを Dukascopy からダウンロードした¹。

また、オーストラリアとニュージーランドの重要な指標発表については、“Economic calendar Investing.com Forex (2011-2019)”²を用いた。図 4.3 が重要な指標発表前後の Lead-Lag パラメータを示している。

¹<https://www.dukascopy.com/>

²<https://www.kaggle.com/devorvant/economic-calendar>

4.5 考察

本論文では、Lead-Lag 効果推定のための指数 NAPLES を提案した。実験 4.2 から、単なる共分散推定量よりも、符合を用いた DS と提案手法がより効果的であることがわかる。実験 4.3 から、定性的に定理 1 を考察が示すように提案手法は Lead-Lag に鋭敏に反応するため、収束が良く精度も良いことがわかる。実験 4.4 からは、以下のような考察が得られる。

- Lead-lag 効果は重要な指標発表前後で変化する。
- NAPLES は HRY よりも早く Lead-Lag 効果を推定できる (a1, a2, b4)。
- NAPLES と HRY が示す Lead-Lag 効果は符合としては同じである。
- NAPLES は HRY が発見できない Lead-Lag 効果を発見できる。 (a3 and b2 and b4)。

5 まとめ

本論文では、高頻度取引における Lead-Lag 効果の推定における 3 つの課題である、(1) タイムスタンプの非同期性、(2) 計算量の膨大性、(3) Lead-Lag 効果の時変性を解決するための Lead-Lag 効果の新しい推定量である NAPLES を提案し、推定法の数学的な性質を与えるとともにその解釈を行った。また人工データ・実データを用いた実験により、提案手法の有効性を示すとともに、マクロ指標の発表前にリードラグの推定を行い、その考察を行った。今後の課題としては、提案手法の統計的性質を明らかにすること、および具体的な裁定取引戦略への応用により収益性を検証することが挙げられる。

参考文献

- [1] Hamad Alsayed and Frank McGroarty. Ultra-high-frequency algorithmic arbitrage across international index futures. *Journal of Forecasting*, Vol. 33, No. 6, pp. 391–408, jun 2014.
- [2] Torben G. Andersen and Tim Bollerslev. Intraday periodicity and volatility persistence in financial markets. *Journal of Empirical Finance*, Vol. 4, No. 2-3, pp. 115–158, June 1997.
- [3] Bruno Biais, Thierry Foucault, and Sophie Moinas. Equilibrium fast trading. *Journal of Financial Economics*, Vol. 116, No. 2, pp. 292–313, May 2015.
- [4] Markus Bibinger, Nikolaus Hautsch, Peter Malec, and Markus Reiss. Estimating the spot covariation of asset prices—statistical theory and empirical evidence. *Journal of Business & Economic Statistics*, Vol. 37, No. 3, pp. 419–435, December 2017.
- [5] Nicolas P.B. Bollen, Michael J. O’Neill, and Robert E. Whaley. Tail wags dog: Intraday price discovery in VIX markets. *Journal of Futures Markets*, Vol. 37, No. 5, pp. 431–451, aug 2016.
- [6] Bradford Case, Yawei Yang, and Yildirim Yildirim. Dynamic correlations among asset classes: REIT and stock returns. *The Journal of Real Estate Finance and Economics*, Vol. 44, No. 3, pp. 298–318, March 2010.
- [7] Andrea Ceron, Luigi Curini, and Stefano M. Iacus. First- and second-level agenda setting in the twittersphere: An application to the italian political debate. *Journal of Information Technology & Politics*, Vol. 13, No. 2, pp. 159–174, mar 2016.
- [8] Alain P. Chaboud, Benjamin Chiquoine, Erik Hjalmarsson, and Clara Vega. Rise of the machines: Algorithmic trading in the foreign exchange market. *The Journal of Finance*, Vol. 69, No. 5, pp. 2045–2084, Sep 2014.
- [9] Kalok Chan. A further analysis of the lead-lag relationship between the cash market and stock index futures market. *Review of Financial Studies*, Vol. 5, No. 1, pp. 123–152, January 1992.
- [10] Kohei Chiba. Estimation of the lead-lag parameter between two stochastic processes driven by fractional brownian motions. *Statistical Inference for Stochastic Processes*, Vol. 22, No. 3, pp. 323–357, January 2019.
- [11] Tarun Chordia and Bhaskaran Swaminathan. Trading volume and cross-autocorrelations in stock returns. *The Journal of Finance*, Vol. 55, No. 2, pp. 913–935, 2000.
- [12] Frank de Jong and Theo Nijman. High frequency analysis of lead-lag relationships between financial markets. *Journal of Empirical Finance*, Vol. 4, No. 2-3, pp. 259–277, June 1997.
- [13] Dobrislav Dobрева and Ernst Schaumburgb. High-frequency cross-market trading: Model free

- measurement and applications, 2017. Working paper.
- [14] Takaki Hayashi and Yuta Koike. Multi-scale analysis of lead-lag relationships in high-frequency financial markets, 2017.
- [15] Takaki Hayashi and Yuta Koike. Wavelet-based methods for high-frequency lead-lag analysis. *SIAM Journal on Financial Mathematics*, Vol. 9, No. 4, pp. 1208–1248, January 2018.
- [16] Takaki Hayashi and Nakahiro Yoshida. On covariance estimation of non-synchronously observed diffusion processes. *Bernoulli*, Vol. 11, No. 2, pp. 359–379, April 2005.
- [17] M. Hoffmann, M. Rosenbaum, and N. Yoshida. Estimation of the lead-lag parameter from non-synchronous data. *Bernoulli*, Vol. 19, No. 2, pp. 426–461, May 2013.
- [18] Nicolas Huth and Frédéric Abergel. High frequency lead/lag relationships—empirical facts. *Journal of Empirical Finance*, Vol. 26, pp. 41–58, 2014.
- [19] Ira Kawaller, Paul Koch, and Timothy W Koch. The temporal price relationship between s& p 500 futures and the s& p 500 index. *Journal of Finance*, Vol. 42, pp. 1309–29, 02 1987.
- [20] Yuta Koike. Inference for time-varying lead-lag relationships from ultra high frequency data. *SSRN Electronic Journal*, 2017.
- [21] Andrew W. Lo and A. Craig MacKinlay. An econometric analysis of nonsynchronous trading. *Journal of Econometrics*, Vol. 45, No. 1-2, pp. 181–211, July 1990.
- [22] Victor Hugo Martinez and Ioanid Rosu. High frequency traders, news and volatility. *SSRN Electronic Journal*, 2012.
- [23] Sait Ozturk, Michel van der Wel, and Dick van Dijk. Intraday price discovery in fragmented markets. Tinbergen Institute Discussion Paper 14-027/III, Amsterdam and Rotterdam, 2014.
- [24] Bhaskar Sinha and Sumati Sharma. Lead - lag relationship in indian stock market: Empirical evidence. *SSRN Electronic Journal*, 2008.
- [25] Xiaoli Wang. An empirical analysis of lead-lag relationship among various financial markets. *Accounting and Finance Research*, Vol. 4, No. 2, mar 2015.

6 Appendix

次の定理を同期観測の場合に制限すれば直ちに主定理が得られる。

Theorem 2 2つの資産価格 X, Y が幾何ブラウン運動 $dX_t = \sigma_1 dW_t^{(1)}$, $X_0 = x_0$, $dY_t = \sigma_2 dW_t^{(2)}$, $Y_0 = y_0$ に従うとする。ただし、 $W_t^{(1)}, W_{t-\theta}^{(2)}$ は $W_t^{(1)}$ と独立な標準ブラウン運動 $W_t^{(3)}$ を用いて $W_{t-\theta}^{(2)} = \rho W_t^{(1)} + \sqrt{1-\rho^2} W_t^{(3)}$ とあらわされる相関係数が ρ の標準ブラウン運動である。先行遅行関係は θ であらわす。資産価格 X_t を $0 = s_0 < s_1 < s_2 < \dots < s_n$ という時刻で観測し、資産価格 Y_t を $0 = t_0 < t_1 < t_2 < \dots < t_m$ という時刻で非同期的に観測する。これに対して $\bar{t}_i := \min\{t_j \mid s_i \leq t_j\}$, $\underline{t}_i := \max\{t_j \mid s_i > t_j\}$, $\bar{s}_j := \min\{s_i \mid t_j \leq s_i\}$, $\underline{s}_j := \max\{s_i \mid s_i < t_j\}$ として、 s_i, t_j に一番近いような t_j, s_i のうち大きいものと小さいものを1つずつ選んで対応付して、 $T = \max\{s_n, t_m\}$ とおく。タイムスタンプについて次のような条件を仮定する。(1) 任意の i について $s_i < t_j < s_{i+1}$ を満たすような t_j は1個または0個しか存在しない。(2) 任意の j について $t_j < s_i < t_{j+1}$ を満たすような s_i は1個または0個しか存在しない。このとき $R(T)$ の期待値は次のようになる。

$$\begin{aligned} E[R(T)] &= \frac{1}{\pi} \sum_{i=1}^{n-1} 1_{s_{i-1} < \bar{t}_i + \theta < s_i} \arcsin\left(\frac{\rho(s_i - \bar{t}_i - \theta)}{\sqrt{(\underline{t}_{i+1} - \bar{t}_i)(s_i - s_{i-1})}}\right) \\ &+ \frac{1}{\pi} \sum_{j=1}^{m-1} 1_{\bar{s}_j < t_j + \theta < \underline{s}_{j+1}} \arcsin\left(\frac{\rho(t_j + \theta - \bar{s}_j)}{\sqrt{(\underline{s}_{j+1} - \bar{s}_j)(t_j - t_{j-1})}}\right) \\ &- \frac{1}{\pi} \sum_{i=1}^{n-1} 1_{s_{i-1} < \underline{t}_{i+1} + \theta < s_i} \arcsin\left(\frac{\rho(\underline{t}_{i+1} + \theta - s_{i-1})}{\sqrt{(\underline{t}_{i+1} - \bar{t}_i)(s_i - s_{i-1})}}\right) \\ &- \frac{1}{\pi} \sum_{j=1}^{m-1} 1_{\bar{s}_j < t_{j-1} + \theta < \underline{s}_{j+1}} \arcsin\left(\frac{\rho(\underline{s}_{j+1} - t_{j-1} - \theta)}{\sqrt{(\underline{s}_{j+1} - \bar{s}_j)(t_j - t_{j-1})}}\right). \end{aligned}$$

Lemma 3 N, M をともに平均が0で分散が1の標準正規分布に従い、相関係数が ρ であるような確率変数であるとする。このとき次が成り立つ

$$P(N > 0, M > 0) = \frac{1}{4} + \frac{\arcsin(\rho)}{2\pi} = \frac{\arccos(-\rho)}{2\pi}.$$

Proof 1 (補題の証明) 2つ目の等式は明らかである。故に1つ目の等式を示す。これは *Box-Muller* 変換を用いればわかる。実際 V, W を独立な標準正規分布に従う確率変数とする。 U を $[-\pi, \pi)$ 上一様分布にしたが

う確率変数とする。このとき $R = \sqrt{V^2 + W^2}$ とおき $\phi = \arcsin \rho$ とおく。ここで

$$(N, M) = (\sqrt{2}(V \cos \phi + W \sin \phi), W),$$

$$(V, W) = (R \cos U, R \sin U)$$

が同時分布が同じという意味で成り立つことがわかる。よって以下のように補題が証明される。

$$\begin{aligned} P[N > 0, M > 0] &= P[\cos U \cos \phi + \sin U \sin \phi > 0, \sin U > 0] \\ &= P[U \in (\phi - \pi/2, \phi + \pi/2), U \in (0, \pi)] \\ &= \frac{1}{4} + \frac{\phi}{2\pi} = \frac{1}{4} + \frac{\arcsin(\rho)}{2\pi} = \frac{\arccos(-\rho)}{2\pi}. \end{aligned}$$

Proof 2 (定理の証明) 最後の時刻である T まで観測しているので $R(T)$ は以下のようになる。

$$R(T) = \sum_{i=1}^{n-1} (b_{s_i}^{(X)} \hat{Y}_{s_{i+1}} - b_{s_i}^{(X)} \hat{Y}_{s_i}) - \sum_{j=1}^{m-1} (b_{t_j}^{(X)} \hat{X}_{t_{j+1}} - b_{t_j}^{(Y)} \hat{X}_{t_j}).$$

まず第 1 項 $\sum_{i=1}^{n-1} (b_{s_i}^{(X)} \hat{Y}_{s_{i+1}} - b_{s_i}^{(X)} \hat{Y}_{s_i})$ について考える。 $b_{s_{i+1}}^{(X)}$ と $b_{s_i}^{(X)}$ が同符号であるときは $(b_{s_{i+1}}^{(X)} \hat{Y}_{s_{i+2}} - b_{s_{i+1}}^{(X)} \hat{Y}_{s_{i+1}}) + (b_{s_i}^{(X)} \hat{Y}_{s_{i+1}} - b_{s_i}^{(X)} \hat{Y}_{s_i}) = (b_{s_{i+2}}^{(X)} \hat{Y}_{s_{i+1}} - b_{s_i}^{(X)} \hat{Y}_{s_i})$ となり、または $b_{s_i}^{(X)}$ が 0 であるときは $(b_{s_{i+1}}^{(X)} \hat{Y}_{s_{i+1}} - b_{s_i}^{(X)} \hat{Y}_{s_i}) = 0$ となるので同符号または 0 となるような s_i を省略することによって、 $b_{s_{i+1}}^{(X)} = -b_{s_i}^{(X)}$ かつ $b_{s_i}^{(X)} \neq 0$ がすべての i に成り立つように s_i を取り直して計算しても $R(t)$ の数値は同じであるので、以下はそうに取り直した s_i を用いる。

また第 1 項は先述の通り

$$\hat{b}_{s_{i+1}}^{(Y)} := (\hat{Y}_{s_{i+1}} - \hat{Y}_{s_i}), \hat{b}_{t_{j+1}}^{(X)} := (\hat{X}_{t_{j+1}} - \hat{X}_{t_j})$$

と定義することによって

$$\sum_{i=1}^{n-1} b_{s_i}^{(X)} (\hat{Y}_{s_{i+1}} - \hat{Y}_{s_i}) = \sum_{i=1}^{n-1} b_{s_i}^{(X)} \hat{b}_{s_{i+1}}^{(Y)}$$

と書き直すことができる。ここで仮定より

$$\hat{b}_{s_{i+1}}^{(Y)} = \text{sign}(Y_{t_{i+1}} - Y_{t_i}), \hat{b}_{t_{j+1}}^{(X)} = \text{sign}(Y_{s_{j+1}} - Y_{s_j}),$$

が成り立つので確率変数 $b_{s_i}^{(X)} \hat{b}_{s_{i+1}}^{(Y)}$ はこの 2 つの確率変数が同符号のとき 1、異符号のとき -1 となるような確率変数である。ゆえに、その期待値を計算するときは、同符号・異符号であるときの確率をそれぞれ計算して、

$$\begin{aligned} E[\sum_{i=1}^{n-1} b_{s_i}^{(X)} \hat{b}_{s_{i+1}}^{(Y)}] &= \sum_{i=1}^{n-1} E[b_{s_i}^{(X)} \hat{b}_{s_{i+1}}^{(Y)}] \\ &= \sum_{i=1}^{n-1} P((W_{s_i}^{(1)} - W_{s_{i-1}}^{(1)})(W_{t_{i+1}}^{(2)} - W_{t_i}^{(2)}) > 0) \\ &\quad - \sum_{i=1}^{n-1} P((W_{s_i}^{(1)} - W_{s_{i-1}}^{(1)})(W_{t_{i+1}}^{(2)} - W_{t_i}^{(2)}) < 0) \\ &= \sum_{i=1}^{n-1} 2P((W_{s_i}^{(1)} - W_{s_{i-1}}^{(1)})(W_{t_{i+1}}^{(2)} - W_{t_i}^{(2)}) > 0) + 1 \end{aligned}$$

とすることができる。最後の確率を計算するために

$$p_{s_i, s_{i-1}} := (W_{s_i}^{(1)} - W_{s_{i-1}}^{(1)}), q_{t_{i+1}, t_i} := (W_{t_{i+1}}^{(2)} - W_{t_i}^{(2)})$$

という確率変数を定めて、

$$P(p_{s_i, s_{i-1}} q_{t_{i+1}, t_i} > 0)$$

を計算する。 $W^{(1)}, W^{(2)}$ はブラウン運動であるので $p_{s_i, s_{i-1}} q_{t_{i+1}, t_i}$ はともに正規分布に従う。よってこの確率を計算するためには、補

題 3 から $p_{s_i, s_{i-1}} q_{t_{i+1}, t_i}$ の相関係数を計算すればよい。ここで、 $W^{(1)}, W^{(2)}$ は先行運行関係があるブラウン運動がであったので

$$\begin{aligned} q_{t_{i+1}, t_i} &= (W_{t_{i+1}}^{(2)} - W_{t_i}^{(2)}) \\ &= \rho(W_{t_{i+1}+\theta}^{(1)} - W_{t_i+\theta}^{(1)}) \\ &\quad + \sqrt{1-\rho^2}(W_{t_{i+1}+\theta}^{(3)} - W_{t_i+\theta}^{(3)}) \\ &= \rho p_{t_{i+1}+\theta, t_i+\theta} + \sqrt{1-\rho^2} r \end{aligned}$$

と書ける。ただし、 $r = (W_{t_{i+1}+\theta}^{(3)} - W_{t_i+\theta}^{(3)})$ であり p とは独立な確率変数である。ゆえに求めるべき共分散は共分散の双線形性と r の独立性によって

$$\begin{aligned} \text{Cov}(p_{s_i, s_{i-1}}, q_{t_{i+1}, t_i}) &= \text{Cov}(p_{s_i, s_{i-1}}, \rho p_{t_{i+1}+\theta, t_i+\theta} + \sqrt{1-\rho^2} r) \\ &= \rho \text{Cov}(p_{s_i, s_{i-1}}, p_{t_{i+1}+\theta, t_i+\theta}) \\ &\quad + \text{Cov}(p_{s_i, s_{i-1}}, \sqrt{1-\rho^2} r) \\ &= \rho \text{Cov}(p_{s_i, s_{i-1}}, p_{t_{i+1}+\theta, t_i+\theta}) \\ &= \rho \text{Cov}((W_{s_i}^{(1)} - W_{s_{i-1}}^{(1)}), (W_{t_{i+1}+\theta}^{(1)} - W_{t_i+\theta}^{(1)})) \end{aligned}$$

となるので、ブラウン運動の独立増分性から

$$\begin{aligned} \text{Cov}(p_{s_i, s_{i-1}}, q_{t_{i+1}, t_i}) &= \begin{cases} \rho(s_i - t_i - \theta) & s_{i-1} < t_i + \theta < s_i \text{ のとき} \\ \rho(t_{i+1} + \theta - s_{i-1}) & s_{i-1} < t_{i+1} + \theta < s_i \text{ のとき} \\ 0 & \text{それ以外のとき} \end{cases} \end{aligned}$$

であることがわかる。ゆえにこれらの確率変数の相関係数は以下のようになる。 $\text{Corr}(p_{s_i, s_{i-1}}, q_{t_{i+1}, t_i}) =$

$$\begin{cases} \frac{\rho(s_i - t_i - \theta)}{\sqrt{(t_{i+1} - t_i)(s_i - s_{i-1})}} & s_{i-1} < t_i + \theta < s_i \text{ のとき} \\ \frac{\rho(t_{i+1} + \theta - s_{i-1})}{\sqrt{(t_{i+1} - t_i)(s_i - s_{i-1})}} & s_{i-1} < t_{i+1} + \theta < s_i \text{ のとき} \\ 0 & \text{それ以外のとき} \end{cases}$$

ここで $P(p_{s_i, s_{i-1}} q_{t_{i+1}, t_i} > 0)$ を計算するために補題 3 を用いる。求めるべきは、 $P(p_{s_i, s_{i-1}} q_{t_{i+1}, t_i} > 0) = P(p_{s_i, s_{i-1}} > 0, q_{t_{i+1}, t_i} > 0) + P(p_{s_i, s_{i-1}} < 0, q_{t_{i+1}, t_i} < 0) = 2P(p_{s_i, s_{i-1}} > 0, q_{t_{i+1}, t_i} > 0)$ である。ここで補題から $P(p_{s_i, s_{i-1}} q_{t_{i+1}, t_i} > 0) =$

$$\begin{cases} \frac{1}{\pi} \arcsin\left(\frac{\rho(s_i - t_i - \theta)}{\sqrt{(t_{i+1} - t_i)(s_i - s_{i-1})}}\right) + \frac{1}{2} & s_{i-1} < t_i + \theta < s_i \text{ のとき} \\ \frac{1}{\pi} \arcsin\left(\frac{\rho(t_{i+1} + \theta - s_{i-1})}{\sqrt{(t_{i+1} - t_i)(s_i - s_{i-1})}}\right) + \frac{1}{2} & s_{i-1} < t_{i+1} + \theta < s_i \text{ のとき} \\ \frac{1}{2} & \text{それ以外のとき} \end{cases}$$

が得られる。 $R(t)$ の第 2 項の

$$\sum_{j=1}^{m-1} (b_{t_j}^{(Y)} \hat{X}_{t_{j+1}} - b_{t_j}^{(Y)} \hat{X}_{t_j}) = \sum_{j=1}^{m-1} b_{t_j}^{(Y)} \hat{b}_{t_{j+1}}^{(X)}$$

についてもこれまでの議論を Y, X を置き換えて同様にブラウン運動の共分散を計算することによって

$$P(p_{s_{j+1}, s_j} q_{t_j, t_{j-1}} > 0) =$$

$$\begin{cases} \frac{1}{\pi} \arcsin\left(\frac{\rho(t_j + \theta - s_j)}{\sqrt{(s_{j+1} - s_j)(t_j - t_{j-1})}}\right) + \frac{1}{2} & s_j < t_j + \theta < s_{j+1} \text{ のとき} \\ \frac{1}{\pi} \arcsin\left(\frac{\rho(s_{j+1} - t_{j-1} - \theta)}{\sqrt{(s_{j+1} - s_j)(t_j - t_{j-1})}}\right) + \frac{1}{2} & s_j < t_{j-1} + \theta < s_{j+1} \text{ のとき} \\ \frac{1}{2} & \text{それ以外のとき} \end{cases}$$

が成り立つことがわかる。これらを足し合わせて定理が証明される。

予約価格を用いた指値・成行注文選択による約定寿命の 冪分布の再現

Recovering power-law distributions of lifetimes of orders both cancelled and executed using an agent-based model containing reservation prices

吉村勇志¹ 陳昱¹

Yushi Yoishimura¹, Yu Chen^{1,2}

¹ 東京大学大学院新領域創成科学研究科人間環境学専攻

¹Department of Human and Engineered Environmental Studies,
Graduate School of Frontier Science, The University of Tokyo

Abstract: A lifetime of limit order is defined as the elapsed time from appear (submission) to disappear (cancellation or execution). Lifetime is a key of decision-making of traders because traders submit an order based on the trade-off between execution cost (how much price will be executed at) and delay risk (how long does it takes to be executed) and execution lifetime means the waiting time to execution and cancellation lifetime means the limit time of patience. Therefore, recovering power-law distribution of lifetimes of orders by an agent-based model (ABM) is a benchmark of time related decision-making of agents and contributes to constructing more advanced models. In this study, we created an ABM reproducing both of cancellation and execution lifetime distributions by extending our previous ABM doing only the distribution of cancelled orders.

1. 背景

個々の市場参加者の意思決定、市場全体の挙動の何れを考察するにせよ、時間という概念は重要である。

市場参加者が注文を出す際、それを指値にするか成行にするか、指値ならば価格は幾らを指定するかといった意思決定は、取引価格と待ち時間のトレードオフの関係にある。キャンセル注文の提出は時間経過による市況の変化、或いは待ち時間が許容範囲を超えたことによって行われる筈である。このように、市場に対する代表的な注文方法である指値・成行・キャンセル注文の全てがその意思決定において時間という要素を持っている。

市場の時間発展について考えてみると、誰が何時注文を出したか、市場参加者間の注文の順序が大きく影響することは明らかであるように思われる。というのも、市場参加者は板の状態を見て意思決定を行うと考えられるが、板の状態は市場参加者の注文によって変化する。即ち、ある市場参加者が板を変え、それに応じて他の市場参加者の意思決定、注文が変化し、また板の状態が変わるというフィードバ

ック構造が存在する。従って、ある瞬間の市場に対して誰が真っ先に反応し、注文を出したかによってその後の流れが変わり得る筈である。長い時間スケールが興味の対象で時間的に平均化された統計性質を問題とする場合にはこの違いは重要ではないであろうが、人工市場がより個別の市場の詳細な性質、時間発展をも研究対象として含むことになるならば、注文の順序等の市場の時間構造の再現は必要になってくるであろう(現在のところ、市場の個別の時間発展の研究としては 2010 年アメリカのフラッシュクラッシュを模擬した人工市場モデル[1]が存在する為、一般的な市場の統計性質のみならず個別、具体的な市場の展開の研究に対する潜在的ニーズも存在する筈である)。

人工市場においても現実の市場の時間構造を模倣出来ることが望ましいということを論じてきたが、時間構造とは何を指しているのか、何を再現すれば十分なのかは曖昧であり、何らかの形で捨象して再現すべき性質を定量的に表現する必要がある。本研究ではそれとして注文の寿命の冪分布を選択する。

注文の寿命とは、指値注文が板上に出されてからキャンセル若しくは約定によって消えるまでの経過

時間のことである。これが市場の時間構造において重要な意味を持つことは、次のように説明出来る。

約定によって消えた注文の寿命(約定寿命)は、定義上当然ながら指値注文を出してから約定するまでの経過時間である。即ち、これの分布を知ることは自分が指値注文を出した場合に約定までにどれくらい待つ必要があるのか検討を付けられるということであり、即座に約定する成行注文との選択において意味を持つ(これは指値注文の価格を無視した極めて簡易な議論であることは言うまでもない)。そしてキャンセルによって消えた注文の寿命(キャンセル寿命)は忍耐力が切れてキャンセルするまでの経過時間であり、同時にある市場参加者が注文を出してからもう一度注文を出すまでの経過時間、即ち、1人の市場参加者が市場に参画する頻度の一側面の表現でもある。

このように市場の時間構造の一部である注文の寿命であるが、これの分布を取ると冪分布になるということが知られている。Challet [2]の観測においてはキャンセル寿命の冪指数が $\alpha_C = 2.1$ 、約定寿命が $\alpha_M = 1.5$ であった。この性質は寿命を実時間で測定してもティック時間(注文が1つ出される度に1進む仮想的な時間)で測定しても変わらないようであるが、本研究では簡単の為に以降後者で考える。

2. 先行研究

人工市場研究における注文の寿命の扱いについて述べる。キャンセル寿命に関して議論する為にはキャンセルがモデルに組み込まれていなければならない。キャンセルが導入された人工市場には、例えば水田[3,4]や Bartolozzi[5]のように、指値/成行注文が出される度に板上の全注文に対しキャンセルされるかどうかの判定を行い、1時間ステップとするものが多い。そしてキャンセルの発生も単純に経過時間が一定値に到達しているかどうか[3,4]、或いはボラティリティにのみ依存する[5]等、単純なものが多くキャンセル寿命の分布を議論すること自体出来ない。例外的に low intelligence model (エージェントは仮定せず確率的に板に注文が到達するが、その確率が板の状態に応じて変化するモデル)の1つである Mike[6](と恐らくはその拡張形[7,8])は Challet の観測とは異なる値でキャンセル寿命の冪分布を再現するものの、約定寿命の冪分布は再現していない。

対して、本研究の前身となった吉村[9,10]は指値/成行注文とキャンセル注文の発注タイミングを同一とした上でキャンセル寿命の冪分布を Challet の指数通りに再現することに成功した。指値/成行/キャンセル注文の発注タイミング統一はそれらの注文間の時間的競合を発生させ、取引に即時性をどれ程求め

ているか、その強さ(緊急性と以下呼称)の比較を可能にし、将来的なモデル拡張にも有益であると考えている。

3. モデル開発

3.1 PR モデル

3.2 でモデルの詳細を述べる前に本項でその原形である吉村[9,10](以下 PR(previous)モデルと呼ぶ)を解説する。

エージェントの注文順序が固定でも完全ランダムでもなく各人の必要性に応じて行われ、尚且つ競争があるような状況を表現する為に緊急性を用いたルーレット選択によってその時間ステップに注文を出すエージェントを決定するものとする。即ち、各エージェントは自分の緊急性 $u_i(t)$ を計算し、エージェント i が選ばれる確率は $u_i(t)/\sum_j u_j(t)$ となる。

PR モデルでは緊急性 $u_i(t)$ を解析的に導出することで注文の寿命の冪分布の再現を試みた。まず、1人のエージェントが同時に複数の注文発注を考えるのを許容すると緊急性の計算が煩雑になるので、以下のようにエージェントの行動を簡略化した。エージェントに N 状態(板上に自分の指値注文が無い)、L 状態(板上に自分が指値注文が残っている)の区分を与え、N エージェントは指値/成行、L エージェントはキャンセル注文しか出さないものとし、緊急性も $u^{LM}(=1)$ と $u^C(\tau_i(t))$ と異なる数式を用いることとする。N 状態の緊急性を定数とした理由は2つあり、指値/成行注文とキャンセル注文の比率のみに着目したことと、N 状態では指値/成行は確率 q_L 、 q_M で出し分けるのみ、注文内容は不定であるので計算する理由が無いことである。L 状態の緊急性はエージェントが注文を出してから経過時間に依存する関数で決まるとしている。これは忍耐がキャンセルに影響するという考えと、多変数に依存すると解析的に解けないという事情に基づく。

この設定の下、板上の注文の個数や経過時間の分布の時間的安定性を仮定して $u^C(\tau_i(t))$ の解析解を解くと、次のようになる。

$$u^C(\tau) = \frac{1}{P(\tau)} (N - N_{LOB}^*) \frac{q_L - q_M}{q_L + q_M} \frac{\tau^{-\alpha_C}}{\sum_{\tau=1}^{\tau_{max}} \tau^{-\alpha_C}}$$

$$P(\tau) = \frac{1}{2} - p_C^* \frac{\sum_{\tau=1}^{\tau-1} \tau^{-\alpha_C}}{\sum_{\tau=1}^{\tau_{max}} \tau^{-\alpha_C}} - p_M^* \frac{\sum_{\tau=1}^{\tau-1} \tau^{-\alpha_M}}{\sum_{\tau=1}^{\tau_{max}} \tau^{-\alpha_M}}$$

$$N_{LOB}^* = p_C^* \frac{\sum_{\tau=1}^{\tau_{max}} \tau^{-\alpha_C+1}}{\sum_{\tau=1}^{\tau_{max}} \tau^{-\alpha_C}} + p_M^* \frac{\sum_{\tau=1}^{\tau_{max}} \tau^{-\alpha_M+1}}{\sum_{\tau=1}^{\tau_{max}} \tau^{-\alpha_M}}$$

式の詳細についての説明は省略する。

得られた解析解を用いてシミュレーションを行うと、キャンセル寿命の冪分布は綺麗に再現することが出来たが約定寿命の冪分布に関しては、寿命が大きい領域では再現することが出来たが小さい領域では冪分布より少ない頻度の注文しか観測されなかった。即ち、指値注文がだされてからすぐに成行注文によって約定するような展開が現実の市場と比べて少ないということである。

3.2 DREC-2 モデル

そこで、市場に出されてから短時間で注文が約定されるという動きをもたらすトレーダーの行動性質を考察し、それをモデルに導入することを考える。

短時間による約定を引き起こすメカニズムとして、指値注文が最良気配値を更新すると、それによって成行注文で得られる利益が増加、成行注文が出され易くなり、その指値注文が結果として短時間で約定されるというような流れが考えられる。これが人工市場において再現されるには、エージェントが最良気配値を意思決定に用い、取引を行う際にその利益を計算出来ること、PR モデルと異なり指値/成行注文の使い分けが単純な確率的選択ではなくそれらに依存することが必要である。

それを達成する為、各エージェントに予約価格 $P_{res}^i(t)$ を与え、時間変化させる。予約価格とは 1 取引単位だけ金融商品を買った(売った)場合に自身の総資産の価値が変化しない価格のことであり[11]、これより安く(高く)買う(売る)ことが出来ればその分だけ総資産が増える、即ち利益を得られるということになる。その意味で予約価格は本人が思っている金融商品の本質的価値、ファンダメンタルに近いとは言えるが、自分のポジション等にも依存し、また安全マージンを取って買いと売りの予約価格が異なることもあり得る。

エージェントに予約価格を設定したことによって成行注文を出した時の利益が最良気配値との差を取るによって計算可能になり、N エージェントが指値と成行注文を使い分ける為の基盤が整った。リアリティに極限まで拘るならば、価格の流れや板の状態から指値注文の価格と予期される待ち時間の関係を推測し、それらと比較した上で成行注文発注の妥当性を考えることになるが、これは煩雑に過ぎる。本研究では指値注文の価格は PR モデルと同様冪分布に従ったランダムのもとし、指値/成行注文の確率を成行注文の利益によって変化させるのみとする。即ち、次のように注文を出すとする。その時刻に注文を出すエージェントとして N 状態のそれが選ばれた時、彼の予約価格がスプレッド内であれば、成行注文では決して利益を得られないのだから必ず指値

注文を出す。予約価格が最良売り気配値より高い場合は、エージェントが金融商品の価値を高く見ているので買い注文を出すものとし、成行注文の選択確率は予約価格と最良気配値の差 x に対する増加関数であるとする。増加関数と仮定する理由は、 $x = 0$ では成行注文の利益が 0 である為成行注文の確率 0、 $x \rightarrow \infty$ では指値注文と成行注文の利益差が消失する、即ち約定速度のみ問題となるので成行注文の確率 1 となることからである。予約価格が最良買い気配値より低い場合も売買方向は逆で同様に設定する。これを視覚的に表すと図 1 のようになる。

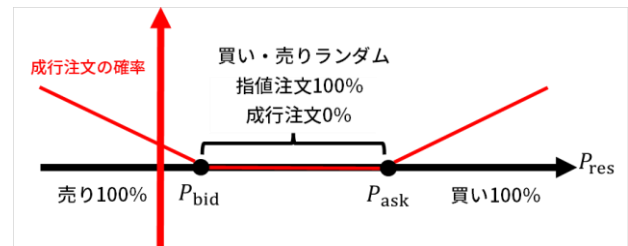


図 1 N エージェントの行動の視覚的表現

またキャンセルも経過時間のみに依存するのではなく、利益と約定速度のトレードオフを組み込むこととする。即ちキャンセルの緊急性を

$$v^C(\tau) = w_{DR} w_{EC} u^C(\tau)$$

という式に改める。 w_{DR} 及び w_{EC} はそれぞれ約定までの時間、利益の効果を表している。

w_{DR}

$$= \begin{cases} \exp \left[\frac{P_{ask}^i - P_{mid}(t)}{P_{ask}^i - P_{mid}(t_0)} - 1 \right] & (\text{売り注文の場合}) \\ \exp \left[\frac{P_{mid}(t) - P_{bid}^i}{P_{mid}(t_0) - P_{bid}^i} - 1 \right] & (\text{買い注文の場合}) \end{cases}$$

$$w_{EC} = \begin{cases} \exp \left[1 - \frac{P_{ask}^i - P_{res}^i(t)}{P_{ask}^i - P_{res}^i(t_0)} \right] & (\text{売り注文の場合}) \\ \exp \left[1 - \frac{P_{res}^i(t) - P_{bid}^i}{P_{res}^i(t_0) - P_{bid}^i} \right] & (\text{買い注文の場合}) \end{cases}$$

この数式の意味は次のようになる。1 つ目の式は指値注文を出した当初と比べ約定までにかかる時間の見込みが増加しているならばキャンセルの緊急性を上げる、減少しているならば下げる効果を表している。エージェント i が出している売り(買い)注文の価格を $P_{ask}^i(P_{bid}^i)$ とし、それと板の中価格との距離の変化によって約定までの見込み時間の変化を表現している。2 つ目の式は自分の注文が約定した際の利益が指値注文を出した当初と比べ増加しているならばキャンセルの緊急性を下げ、減少しているならば上げる効果を表している。利益に関しては自分の注文

の価格と予約価格の差によって表現している。

以上により利益と約定速度を考慮した注文行動を行うエージェントベースモデルの枠組みは完成した訳だが、シミュレーションを実行するにはまだ2つ決定しなければならないことがある。予約価格の時間発展の数式と、最良気配値と予約価格の差に対する成行注文の選択確率の増加関数の具体的な形状である。

前者に関しては、本研究では

$$P_{\text{res}}^i(t+1) = P_{\text{res}}^i(t) + b(P(t) - P_{\text{res}}^i(t)) + \sigma \varepsilon_i(t)$$

を用いる。予約価格の時間発展を最終約定価格へ近づく・離れる項とランダム項の和でモデル化している。尚、式中の b は2つのパターンで計算する。1つは定数で、もう1つは時刻、エージェント毎に異なる値を取る乱数であり

$$b = 1 - \Lambda(m, s^2)$$

によって表現される。ここで $\Lambda(m, s^2)$ は対数正規分布である。 $\varepsilon_i(t)$ は $N(0,1)$ に従う乱数である。

後者に関しては、遺伝的プログラミングを用いて生成するものとする。 $x=0$ で0、 $x \rightarrow \infty$ で1になる増加関数であるから、この要件を満たす為に $f(x)/\{f(x)+10\}$ の $f(x)$ を $x=0$ で0、 $x \rightarrow \infty$ で ∞ になるように生成する。関数の値域を変更したのは計算の都合である。このようにすることによって、正数

x に対して $+$ 、 $\times$ 、 $\sqrt{\quad}$ 、 $\log(x+1)$ の4つの演算をどのように組み合わせて $f(x)$ を生成しても値域が条件を満たすようになる。指数関数 e^x も数学的には条件を満たすが、数値計算上では発散し易い為今回は使用しない。

ここで遺伝的プログラミングを用いる意味は、予約価格の更新式の多様性を吸収する為である。現実世界で市場参加者の予約価格を観測することは不可能、従ってその更新式のモデル化の仕方にも無数に考えられる。更新式の違いに応じて人工市場が適切に動作する成行注文の選択確率の増加関数の形状も変化すると想定し、2つの数式を同時に変更するのではなく、1つのみ変更しもう1つはそれに応じて遺伝的プログラミングで生成することで条件を満たす人工市場モデルの構築を容易にしている。

4. シミュレーション結果

パラメータはエージェント数 $N=300$ 、注文の経過時間の最大値 $\tau_{\text{max}}=8.01 \cdot 10^6$ 、再現すべき注文の寿命の冪分布を $\alpha_C=2.1$ 、 $\alpha_M=1.5$ とし、1つの人工市場モデルの計算は100万時間ステップ実行する。遺伝的プログラミングは50個体の計算を5世代行う。これは個体数、世代数共にかなり少ないが、値

域の分かっている単調増加関数というかなり条件の限定された関数の生成に用いる為、1つ1つの個体が人工市場であり計算量を要求することからこのように設定した。ランダム生成される木の深さ上限は4で、初期解候補はramped half-and-halfで生成し、新世代は交叉70%、突然変異20%、再生産10%で生成する。

その結果、予約価格更新式中の b が定数の時には注文の寿命の冪分布の再現に失敗したが、乱数の場合にはキャンセル、約定寿命共に指数含めその再現に成功した。それを図2、3に示す。

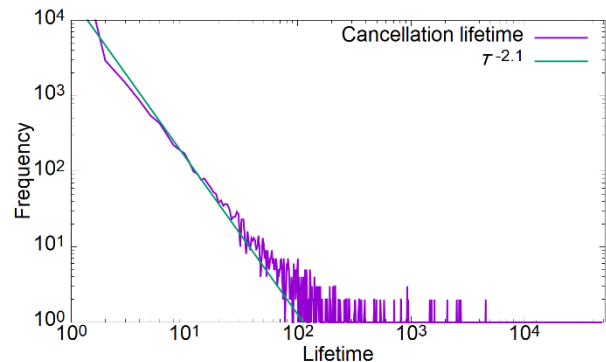


図2 シミュレーションのキャンセル寿命の分布

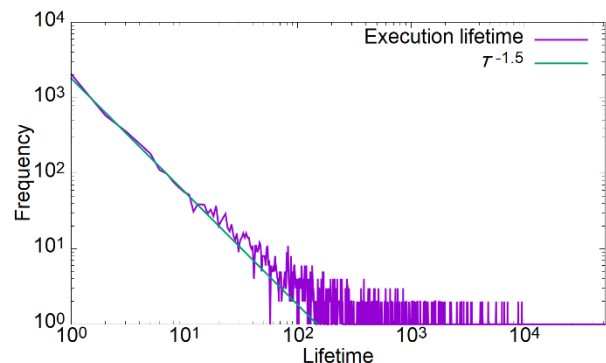


図3 シミュレーションの約定寿命の分布

尚、この時遺伝的プログラミングによって生成された数式は $f(x) = 90 \times 0.3 \times \sqrt{\log(x+1)} + 40$ であった。これを用いた場合、 N エージェントは成行注文が可能である場合には高い確率で出すということになる。図2、3を得た際のパラメータは $m=0.99$ 、 $s=0.25$ 、 $\sigma=1$ である。

5. 考察

N エージェントの成行注文の選択確率の増加関数を生成する解の任意性が高い遺伝的プログラミングで生成したにも関わらず予約価格の更新式の内容如

何によっては寿命の冪分布を再現出来ないということは、N エージェントが指値注文と成行注文の両方が選択可能な時にどう判断するかよりも、そもそも成行注文で収益を上げられる状態がいつ、どのように発生するかの方が重要であるということを意味している。しかしながら本研究では予約価格が具体的にどのような性質を持っていれば寿命の冪分布の再現に十分であるか、その条件の特定には至っていない為、それは今後の課題である。

予約価格更新式が理由で寿命分布が再現出来なかった場合に起きていたことを以下述べる。シミュレーションのパラメータや遺伝的プログラミングの適合度関数におけるキャンセル寿命と約定寿命の重み付けによってキャンセル寿命、約定寿命のどちらが再現出来なかったのかは異なっていたが、キャンセル寿命が1ばかり、即ち指値注文とそれに対するキャンセル注文がずっと交互に繰り返されていたり、市場において約定が殆ど起きず、分布がそもそも描けないということが起きていた。これは指値、成行、キャンセル注文の3つが、エージェントの意思決定により選択され時間的に競合する形で市場に出されるモデル特有の現象だと思われる。エージェントの行動の結果として市場の状態が決まり、それが今度ではエージェントの行動を左右するモデルで、エージェントの行動の自由度が高い場合には正常に動くモデルを得ることも容易でないということを示している。遺伝的プログラミング等の機械学習によってエージェントの行動規則を生成させる場合においても、正常に動くモデルが少なければ探索・搾取が上手く機能しないこともあり得る。エージェントの行動にどの程度制限を与えどの程度自由を与えるかは今後のモデル構築においても重要である。実際、本研究でも計算した時間ステップ数に比して観測された約定回数が少なく、基本的な市場のダイナミクスを捉え切れていない節がある(これは出した直後にキャンセルされた指値注文の多さが一因である)。

6. 結論

本研究では予約価格を導入し、エージェントが自分の注文が約定した際の利益と約定までにかかるであろう時間のトレードオフ関係を意識して指値、成行、キャンセル注文を出すモデルを構築することによって、キャンセル注文によるものだけでなく約定によって市場から消えた注文の寿命の冪分布も指数含めて再現することに成功した。その際に予約価格の時間発展が重要であることは分かったが、具体的に予約価格が満たすべき条件の特定には至らなかった。

謝辞

本研究は JSPS 科研費 JP19J12283 の助成を受けた。

参考文献

- [1] Vuorenmaa, T. A., & Wang, L. (2014). An agent-based model of the flash crash of may 6, 2010, with policy implications. *Available at SSRN 2336772*.
- [2] Challet, D., & Stinchcombe, R. (2003). Limit order market analysis and modelling: on a universal cause for over-diffusive prices. *Physica A: Statistical Mechanics and its Applications*, 324(1), 141-145.
- [3] Mizuta, T., Hayakawa, S., Izumi, K., & Yoshimura, S. (2013, September). Simulation study on effects of tick size difference in stock markets competition. In *International Workshop on Agent-based Approaches in Economic and Social Complex Systems*(Vol. 2013).
- [4] Mizuta, T., Matsumoto, W., Kosugi, S., Izumi, K., Kusumoto, T., & Yoshimura, S. (2014, March). Do dark pools stabilize markets and reduce market impacts? Investigations using multi-agent simulations. In *Computational Intelligence for Financial Engineering & Economics (CIFEr), 2104 IEEE Conference on*(pp. 71-76). IEEE.
- [5] Bartolozzi, M. (2010). A multi agent model for the limit order book dynamics. *The European Physical Journal B*, 78(2), 265-273.
- [6] Mike, S., & Farmer, J. D. (2008). An empirical behavioral model of liquidity and volatility. *Journal of Economic Dynamics and Control*, 32(1), 200-234.
- [7] Gu, G. F., & Zhou, W. X. (2009). Emergence of long memory in stock volatility from a modified Mike-Farmer model. *EPL (Europhysics Letters)*, 86(4), 48002.
- [8] Zhou, J., Gu, G. F., Jiang, Z. Q., Xiong, X., Chen, W., Zhang, W., & Zhou, W. X. (2017). Computational experiments successfully predict the emergence of autocorrelations in ultra-high-frequency stock returns. *Computational Economics*, 50(4), 579-594.
- [9] 吉村勇志, 陳昱(2007), 能動性キャンセルを含む人工市場によるキャンセル寿命の冪分布の再現, 第19回人工知能学会金融情報学研究会
- [10] Yoshimura, Y., Okuda, H., & Chen, Y. (2020). A mathematical formulation of order cancellation for the agent-based modelling of financial markets. *Physica A: Statistical Mechanics and its Applications*, 538, 122507.
- [11] Avellaneda, M., & Stoikov, S. (2008). High-frequency trading in a limit order book. *Quantitative Finance*, 8(3), 217-224.

高頻度注文情報を用いたデータの層化と 多段階事前学習による株価動向予測

Stock Price Prediction Using Limit Order Book Data with Data Stratification and Multi-Phase Pre-training

松原冬樹^{1*} 和泉潔¹ 坂地泰紀¹
Fuyuki Matsubara¹ Kiyoshi Izumi¹ Hiroki Sakaji¹

¹ 東京大学大学院工学系研究科

¹ Graduate School of Engineering, The University of Tokyo

Abstract: Predicting the movement of stock price is an important issue for market participants. Recently, there have been many attempts applying machine learning techniques in financial time series prediction. However, overfitting presents a huge challenge when machine learning approaches are used in financial time series prediction. In this paper, we propose a stock price prediction method utilizing limit order book data from stocks other than target stocks by stratifying the data and holding a multi-phase pre-training considering market liquidity. Experimental results shows that the proposed approach enhances prediction performance.

1 はじめに

金融商品の価格は市場参加者の行動によって変化し続けている。市場参加者は金融商品を購入価格より高い価格で売却、もしくは売却価格より低い価格で購入することによって収益を上げることができる。金融商品の将来における価格変動を過去の情報を用いることによって事前に予測することが可能であるならば、予測を基に取引をしリスクに見合う以上の収益を追求することができると考えられる。経済学、物理学などの学問分野の学者や金融市場に身を置く実務家によって「金融市場における将来の価格予測が可能であるか」という問いに関して数多くの理論研究、実証分析が行われてきた。

[Fama 70] は、「いかなる時も利用可能なすべての情報が証券価格に完全に反映される」とする効率的市場仮説を提唱した。実際に取引が行われる取引価格に対して、すべての市場参加者が情報を共有し情報の非対称性がない場合の均衡価格を本源的価値とすると、本源的価値はすべての市場参加者が価格付けの方法に関して合意している場合の証券価格となる。

一方で、効率的市場仮説への反証をあげる研究、事例も存在する [Cont 01] は経験的に知られている市場の定型化された事実 (stylized facts) として証券のリターンが正規分布ではなくべき乗分布に従うことやボラティ

リティ(資産のリターン標準偏差)に正の自己相関が存在するボラティリティクラスタリングといった現象等について言及している。これらの市場の経験則は効率的市場仮説とは相容れない性質である。また、2010年5月6日には米国市場においてダウ工業株30種平均が数分間で9% (約1000ドル) も下落したフラッシュクラッシュは、市場が効率的であるとする仮説とは非整合的な現象であると考えられる。

これらのことから、市場は必ずしも効率的であるとは言えないと考えられる。金融商品の価格変動に一定の傾向を見出すことができるとすると、そのパターンを利用して金融市場の動向を予測できる可能性はあると言える。近年では計算機性能の向上や取引の電子化に伴い入手可能な金融データが増えたことで、人工知能分野の技術を用いた金融市場の動向予測に関する研究も増えてきている。中でも、機械学習手法は大規模なデータの中から隠れたパターンを認識することに長けているため、金融市場予測において同手法を用いることの妥当性は高いと言える。

1.1 金融時系列情報への機械学習の応用に関する研究

高頻度注文情報には市場参加者の様々な行動が反映されており、価格のみに関する時系列情報と比較して市場参加者の行動に関する情報が豊富に織り込まれていると考えられる。近年、高頻度注文情報を機械学習

*連絡先：東京大学大学院工学系研究科
〒113-8656 東京都文京区本郷 7-3-1
Email: m2019fmatsubara@socsim.org

手法によってを分析し、証券価格動向を予測する数々の研究がなされてきた。

[Dixon 18] は仲値から上下五本の価格、注文数量に加え、成行注文の比率を RNN(Recurrent Neural Network) に入力し、仲値の動向予測を行った。RNN による予測はロジスティック回帰やカルマンフィルタによる予測精度を上回り、次の仲値変動と核の注文系列、流動性の偏りとの時間依存性を考慮に入れた非線形な関係を捉えることに適していると論じた。また、ニューラルネットワークを金融時系列予測へ適用する際には過学習への陥りやすいことを述べている。[Zhang 19] は CNN(Convolutional Neural Network) と LSTM(Long Short-Term Memory) を組み合わせ、注文情報の空間的な構造と時間依存性の 2 側面を捉える DeepLOB を提案した。DeepLOB は価格と注文数量を入力としたときの仲値動向予測において、既存手法の精度を上回ると述べた。また、データ量が少ない場合は過学習を起こし、十分な汎化性能をモデルが出来ない可能性について指摘している。金融時系列予測においてモデルが訓練データに過剰適合をする場合、その予測を基に取引を行う主体は損失を出す危険性が高まると考えられる。したがって、機械学習手法を金融時系列予測に用いる際、過学習を抑制するは主要な課題の一つであると言える。

1.2 銘柄を超えた学習による株価動向予測に関する研究

[Sirignano 19b] は米国株式の高頻度注文情報から得られる価格、注文数量を入力とし、LSTM を用いて次に仲値が変動する際の方向の予測を行った。その際、予測対象の銘柄のみでモデルを学習させたケースよりも予測対象以外の銘柄のデータも併せて学習させた場合の方が予測精度が高くなる結果を得ている。

1.3 投資家行動と価格形成に関する研究

マーケットマイクロストラクチャー研究は投資家行動と価格形成の関係に焦点を当て、在庫モデルや戦略的トレーダーモデルを扱う理論研究や、注文情報を様々な角度から分析、解釈する実証分析が存在する。[Maureen 15] はこうした市場環境下において、資産の価格形成に市場の流動性がより一層大きな役割を担っていると述べ、マーケットマイクロストラクチャーに着目し現実市場の理解を深めることの重要性について論じた。[David 11] は、2010 年 5 月 6 日のフラッシュクラッシュに関する研究を行った。特定の取引日の出来高から計測される売買の偏りから算出され、情報の非対称性を表す VPIN(Volume-Synchronized Probability of

Informed Trading) の累積密度関数が急激な価格変動に先行する可能性を指摘した。また、[David 12] では LFT(Low Frequency Trading) 戦略を用いる取引主体が TWAP(Time Weighted Average Price) といった執行アルゴリズム等を用いて物理時間を基に取引を行うことに記録されるデータの特徴から HFT(High Frequency Trader) に取引意図を推察され、利用されている可能性について言及した。また、[Cao 09] は豪州株式市場における注文情報の分析を行い、需給を表す注文の偏り (Order Imbalance) と証券の短期リターンの間に有意な関係性が存在すると述べている。これらの研究からマーケットマイクロストラクチャーを考慮に入れることが金融市場の動向予測に役立つ可能性が示唆され、データドリブンなアプローチによる金融時系列予測を行う際にもマーケットマイクロストラクチャーを反映させることで予測精度、解釈性の向上に繋がると考えられる。

金融商品の価格動向を予測することができれば、投資家にとってリターンの追求やリスクの低減が可能となる。効率的市場仮説は金融商品の動向の予測は不可能であるとするが、数々の実証研究が示唆するように、現実の市場は完全に効率的とは必ずしも言えない事例が多く存在する。金融市場予測を行うにあたって、高頻度注文情報は様々な市場参加者の注文行動が記録されているため、金融商品の価格予測に適していると考えられる。また、高頻度注文情報のデータ量は膨大であるため金融時系列の動向予測をする際に機械学習手法を用いることの妥当性は高いが、過学習の抑制は最も重要な課題の一つである。機械学習手法を用いた金融時系列予測における過学習の主な原因として、金融時系列の実データは時間軸方向の数に限りがあるため訓練データが不足してしまうことが挙げられる。

本研究では、機械学習手法を用いた株価動向予測を行うにあたり、予測対象銘柄のデータに加えて他銘柄のデータも併せて学習に用いることでデータ数の不足にアプローチを行う。予測対象銘柄以外の銘柄の注文情報を用いる際、市場参加者や価格形成の挙動といった性質の類似度が高い銘柄と低い銘柄が混在しており、予測対象銘柄と類似度が低い銘柄の情報で学習を行うと予測精度が低下すると考えられる。そのため、予測精度を向上させるためには予測対象銘柄と類似度が低い銘柄の割合を下げるのが妥当であると考えられる。したがって、銘柄の類似度を基に高頻度注文情報のサブグルーピングを行い、多段階の学習フェーズを設けることとする。この手続きにより、他銘柄のデータと併せて学習を行うことによって訓練データの数を補いつつ、学習データにおける予測対象銘柄と類似度が低い銘柄のデータの割合を段階的に下げていくことができると考える。また、この手法を用いる際データが層状に分割されるため、以降高頻度注文情報のサブグルー

ピングを行い、多段階の学習フェーズを設ける一連の手続きを「データの層化」と呼ぶ。データの層化手法の詳細は第2章で述べる。また、1.3節で見たように短期市場では特に投資家行動、価格形成にマーケットマイクロストラクチャーが影響していると考えられるため、本研究ではマーケットマイクロストラクチャーを考慮に入れた高頻度注文情報のサブグルーピングを提案する。高頻度注文情報を用いた機械学習手法による金融市場予測タスクにおいて、マーケットマイクロストラクチャーを考慮に入れた多段階事前学習を行う研究の例はなく、本研究における新規性となる。本目的の達成は投資家にとっての運用成績向上につながるだけでなく、金融市場という市場参加者の行動が複雑に集積する系における深層学習の有効性を示すことに貢献できると考える。

2 データの層化と多段階事前学習

本章では、高頻度注文情報の層化を行い多段階の学習フェーズを設ける提案手法の枠組みについて述べる。2.1節では本研究で用いるデータの概要について説明し、2.2節では枠組みの詳細を述べる。

2.1 使用データ

東京証券取引所は市場の売買取引に関する注文、約定、売買高、といった情報をリアルタイムで提供し、蓄積されたものはヒストリカルデータとして取得可能にしている。中でも、FLEX Standard は普通債・TOKYO PRO-BOND Market 銘柄を除く全上場銘柄に関する四本値、売買高、売買代金、寄前/複数気配値段・数量、成行注文数量といった時系列情報を提供している。

本研究では、この FLEX Standard のうち、東証株価指数 (TOPIX) ニューインデックスシリーズのうち TOPIX Core30 と TOPIX Large70 を合わせた計 100 銘柄 (TOPIX100) に関する高頻度注文情報 (2015 年 1 月 4 日から 2015 年 9 月 18 日まで) を用いる。TOPIX ニューインデックスシリーズは年に 1,2 回の頻度で銘柄の規模の変化に伴って定期的に区分の選定が行われる。本研究で用いる銘柄は 2014 年 10 月 31 日に実施された定期選定結果に従った。なお、2014 年 10 月 31 日以降は 2015 年 10 月 30 日までニューインデックスシリーズの入れ替えは行われていない。

2.2 手法

本節では、深層学習モデルに多段階の事前学習をさせる際の高頻度注文情報の層化の枠組みについて述べる。1.3.3 節で言及したように、市場参加者の取引は注文を

介するため高頻度注文情報には予測という観点で有用な情報が含まれている可能性がある。また、流動性投資家、利益追求投資家の需要と供給によって価格形成がなされることは株式市場において銘柄を問わず共通している。このことから、予測対象以外の銘柄の高頻度注文情報を用いることによって株価動向予測の精度向上が期待できる。1.3.3 節で述べたように、[Sirignano 19b] は、米国株式を対象に、銘柄を超えた学習による精度向上を示した。一方で、銘柄をセクターやティックサイズの大小によってデータ分割を行い学習させる場合、使用可能な全情報を用いて仲値変動の予測をする場合と比較して精度向上にはつながらなかった。これはデータの分割を行うことによって学習に用いられるデータの量が減少したことが原因の一つであると考えられる。したがって、用いることのできる全銘柄のデータで学習させた後、予測対象銘柄のデータを用いてファインチューニングをすることによって全銘柄の情報をしつつ、予測対象の銘柄の性質に即した学習を行うことができると考えられる。

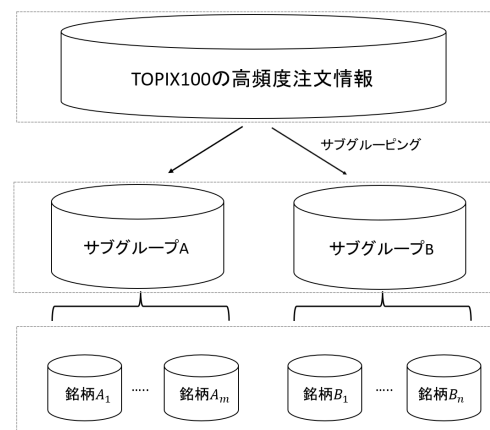


図 1: サブグルーピングによるデータの層化

図1では、サブグループを2種類 (サブグループA、サブグループB) とし、サブグループAの銘柄数が m 銘柄、サブグループBの銘柄数が n 銘柄 (ただし $m+n=100$) とした場合の高頻度注文情報の層化の概要を示した。

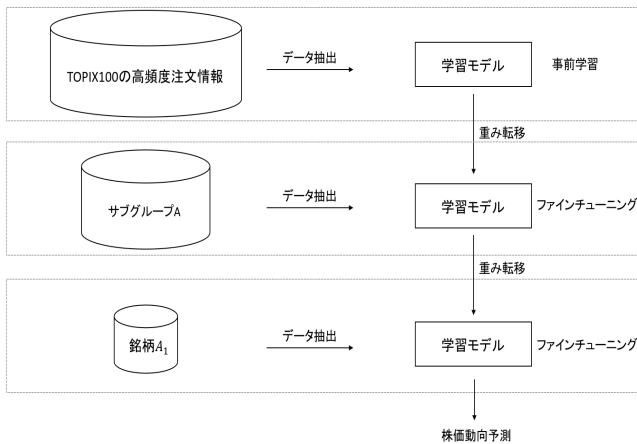


図 2: データ層化による多段階事前学習の概要図

図 2 では高頻度注文情報を用いて TOPIX100 の構成銘柄のサブグループを形成し、多段階の事前学習の枠組みの概要を示した。銘柄 A_1 を予測対象銘柄とするとき、TOPIX100 をグルーピングした際に銘柄 A_1 がサブグループ 1 に属する場合の学習手順の例である。

予測対象銘柄の比較した際に、注文情報や市場参加者の挙動及び性質が類似している銘柄のデータと合わせて学習させることが類似度の低い銘柄のデータと合わせて学習させるよりも予測精度向上に有効であると仮定すると、過去の注文情報及び市場参加者の性質が類似している銘柄同士でグルーピングを行いデータの層化を行った後、多段階の学習フェーズを設けることで予測精度の向上に繋がると考える。今回用いるデータセットの TOPIX100 は TOPIX100 ニューインデックス区分の TOPIX Core30 と TOPIX70 から構成されており、サブグルーピングを行うにあたってニューインデックス区分に従った 2 サブグループに分割させるサブグルーピングは自然な発想であると言える。次章では FLEX Standard から計測した指標を基に TOPIX100 のサブグルーピングを行うが、サブグループの数はニューインデックス区分を基に分割されるグループ数に合わせて 2 つのサブグループに分割する。分割手法の詳細は次章で述べる。また、データセットの銘柄数が N 銘柄から成る場合データは最大で N 層にまで層化することが可能であるが本研究では 3 層までとする。

3 マーケットマイクロストラクチャーを考慮に入れたサブグルーピング

第 1 章で述べた通り、近年金融市場環境は急速に変化している。特に短期で売買を行う高頻度取引、そして電子取引を活用した大口取引の執行アルゴリズムの存在感は年々増してきている。こうした取引戦略や執

行戦略は [Jarnecic 14]、[Kissell 06] で言及されているように収益機会の向上や取引コストの削減のため、市場の流動性に影響を受けると考えられる。したがって、市場の短期価格変動の分析をする際に流動性を考慮に入れることで現実に応じた予測アプローチを取ることができると考える。2.1 節で述べた通り TOPIX ニューインデックスシリーズは時価総額、流動性を基準に銘柄の選定を行っている。米国市場を襲ったフラッシュクラッシュは流動性の枯渇によって引き起こされたと考察されるが、発生直後に出来高自体は急増していること ([David 11]) を考慮に入れると売買代金のみで一元的に流動性を算出することは流動性の側面を捉える際に制限が大きいと言える。したがって、本研究では流動性の性質より詳細に捉えるため、流動性の概念を表す複数の側面に分解して考えサブグループを形成することとする。流動性指標の計測方法については 3.1 節で述べる。

3.1 流動性指標の計測

本節では [太田 11] を参考に流動性の概念を導入し、本研究で用いる流動性の計測方法について述べる。[太田 11] は流動性に 3 つの側面があるとし、それぞれ

1. 取引したいタイミングで取引できるか
2. より低い取引費用で取引できるか
3. 必要な数量を取引できるか

とした。それぞれの側面は順に即時性、タイトネス、デプスと呼ばれ、本研究では FLEX Standard を用いて各流動性を計測した。即時性を表す指標としては約定間隔 (Transaction Interval) を、タイトネスを表す指標として気配スプレッド率 (Quoted Spread Rate) [太田 11]、デプスを表す指標として仲値を中心とする上下 8 本の売り注文、買い注文の総量 (Order Volume) を選定した。1 取引日当たりの各指標を以下のように表す。

$$Transaction\ Interval = \frac{1}{n} \sum_{i=1}^n (Transaction\ Time_{i+1} - Transaction\ Time_i) \quad (1)$$

$$Quoted\ Spread\ Rate = \frac{1}{N} \sum_{i=1}^N \left(\frac{BidAsk\ Spread_i}{MidPrice_i} \right) \quad (2)$$

$$Order Volume = \frac{1}{N} \sum_{i=1}^N \left\{ \sum_{j=1}^8 (Ask Price_j^i \times Ask Volume_j^i + Bid Price_j^i \times Bid Volume_j^i) \right\} \quad (3)$$

3.2 クラスタリング

3.1 節で述べた方法により計測された流動性指標によって TOPIX100 の計 100 銘柄のクラスタリングを行いサブグループを形成する。ここで、クラスター数は TOPIX ニューインデックスシリーズにより TOPIX100 が Core30 と Large70 の 2 グループから構成されていることを考慮し、2 クラスターとする。本研究では k-mean 法を用いる。なお、各々の流動性指標に関してはヒストグラムに偏りが見られたため、自然対数を取るスケーリングを行ったのち標準化を行った。その後、100 銘柄をユークリッド距離を距離関数とする k-means 法によって 2 クラスターに分類する。

3.3 結果

本節では、3.1 節および 3.2 節で述べた流動性指標の計測方法にしたがって TOPIX100 の銘柄をクラスタリングした結果を示す。

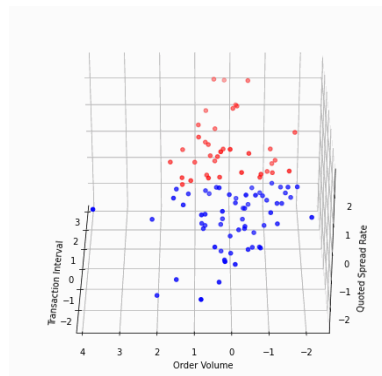


図 3: 流動性指標を用いた TOPIX100 のクラスタリング結果-その 1-

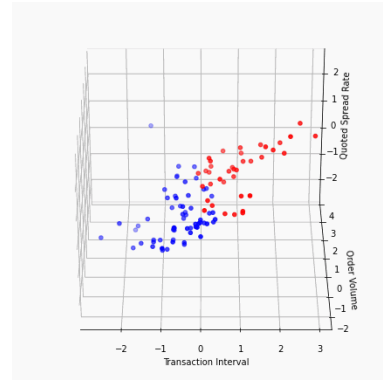


図 4: 流動性指標を用いた TOPIX100 のクラスタリング結果-その 2-

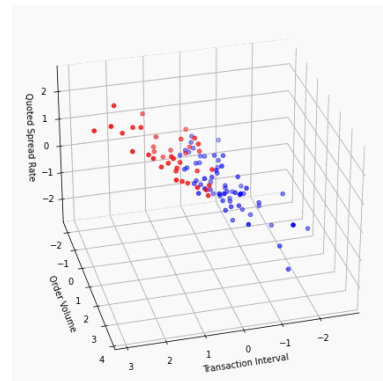


図 5: 流動性指標を用いた TOPIX100 のクラスタリング結果-その 3-

図 3、図 4、図 5 における青色に示された銘柄群をグループ 0、赤色に示された銘柄群をグループ 1 とするとグループ 0 は即時性、タイトネスの観点で流動性が平均的に高く、グループ 1 は即時性、タイトネスの観点で流動性が平均的に低い結果となった。

4 深層学習による株価動向予測

本章では、深層学習による株価動向予測を行い、予測精度を比較することにより、前章までに述べた多段階事前学習の枠組みの有効性を検証する。

4.1 実験

本研究では下記の 5 通りで TOPIX100 に対して次の仲値変動の方向の予測を行う。

1. 予測対象銘柄のデータのみを用いて学習
→ 予測

2. TOPIX100 の全銘柄のデータを用いて学習
→予測
3. TOPIX100 の全銘柄のデータを用いて事前学習
→予測対象銘柄のデータを用いてファインチューニング
→予測
4. TOPIX100 の全銘柄のデータを用いて事前学習
→予測対象銘柄がCore30に属する場合はCore30、
Large70に属する場合はLarge70のデータを用いて
ファインチューニング
→予測対象銘柄のデータを用いてファインチューニング
→予測
5. TOPIX100 の全銘柄のデータを用いて事前学習
→流動性を基準に分割したサブグループのデータ
を用いてファインチューニング
→予測対象銘柄のデータを用いてファインチューニング
→予測

本研究では以後上記の5手法を Target only、TOPIX100 only、TOPIX100+Target、TOPIX100+TOPIX New Index+Target、TOPIX100+Liquidity+Target と表記することとする。

本実験では、株価動向予測を行うにあたり FLEX Standard のうち、TOPIX100 に関する高頻度注文情報 (2015 年 1 月 4 日から 2015 年 9 月 18 日まで) を用いる。2015 年 1 月 4 日から 2015 年 3 月 31 日までを訓練用データ、2015 年 4 月 1 日から 2015 年 4 月 30 日までを検証用データ、2015 年 5 月 1 日から 2015 年 9 月 18 日までを評価用データとした。本研究における実験では、13 時 00 分 00 秒から 13 時 01 分 00 秒までの高頻度注文情報をサンプリングし、13 時 01 分 00 秒より後に初めて仲値が変動したときにそれが上昇であるか下落であるかの 2 クラス分類を行った。仲値変動が上昇である場合には教師ラベル 0 を、下落である場合には教師ラベル 1 を与えた。評価指標には F1 値を用いた。入力には、[Cao 09]、[Gould 16]、[Sirignano 19a]、[Stoikov 16] を参考に Order Book Imbalance を用いた。Order Book Imbalance は (4) 式のように表される。

$$OrderBookImbalance = \frac{V_k^b - V_k^a}{V_k^a + V_k^b} \quad (4)$$

これは売買注文の偏りを表し、 V_k^a を仲値から k 番目の売り注文数量 ($k = 1$ のとき V_k^a は最良売り気配)、 V_k^b を仲値から k 番目の買い注文数量 ($k = 1$ のとき V_k^b は最良買い気配) とすると以下のように表される。本研究で用いた高頻度注文情報は仲値を中心とした上下 8 本までの売り注文および買い注文が記録されているため、

入力の次元数は 8 次元となる。また、学習モデルには 1 層の LSTM を使い、LSTM の最終隠れ層の出力はユニット数 2 の全結合層へ入力され、ソフトマックス関数により得られた確率による分類を行った。

4.2 結果

各手法ごとの F1 値の TOPIX100 全体に関する平均値は以下ようになった。

表 1: 手法別 F1 値の TOPIX100 に関する平均値

Method	F1 Score
Target only	0.5115
TOPIX100 only	0.5713
TOPIX100 + Target	0.5819
TOPIX100 + TOPIX New Index + Target	0.5884
TOPIX100 + Liquidity + Target	0.5975

4.3 考察

4.2 節から、データを層化させ多段階の学習フェーズを設ける手法によって予測精度が平均的に向上していると言える。一方で、予測対象銘柄のデータのみで学習した場合など、データの層数が少ない方が予測精度が高い銘柄も存在した。他銘柄のデータと併せて学習を行うことにより予測精度の低下が起きる原因として、以下のことが考えられる。

まず、予測対象銘柄の市場参加者や価格形成の挙動といった性質が他銘柄のそれと大きく異なる場合が考えられる。この場合、他銘柄の注文情報と併せて学習する際にデータ数の増大による予測精度向上比べ、予測対象銘柄とは類似度の低いデータによる学習の割合が高いことによる予測精度の低下の度合いが大きくなる可能性があり、最終的には予測対象銘柄のデータのみで学習した場合のほうが予測精度が高くなることが考えられる。次に、訓練用データと評価用データでの市場参加者や価格形成の挙動といった性質の時間的変化を本実験では考慮していないことがこの結果につながったと考えられる。訓練用データと評価用データにおいて市場参加者や価格形成の挙動といった性質が変化した場合、訓練用データにおいて類似度が高い銘柄群が評価用データにおいては必ずしも類似度が高いとは言えず、評価用データにおける予測精度の低下を招いたと考えられる。

4.4 投資シミュレーション

本節では、投資シミュレーションの設定、結果について述べる。4.2 節で得られた結果の中で最も F1 値が

高い銘柄数が多かった TOPIX100+Liquidity+Target モデルによる投資シミュレーションのパフォーマンスをロングオンリー戦略、ショートオンリー戦略の2手法をベースラインとして比較する実験を行う。現実の市場において仲値での取引はできないが、同じ条件下でベースラインのパフォーマンスを上回るのであればモデルに予測力が存在することを示唆していると考ええる。ここでは、13時00分00秒から13時01分00秒までの高頻度注文情報をサンプリングし、13時01分00秒より後に初めて仲値が変動したときにそれが上昇とモデルが予測した場合には買いを、下落であると予測した場合には売りをサンプリング期間の直後に入れ、仲値が変動したときに反対売買を行う戦略に基づくパフォーマンスを比較する。なお、ロングオンリー戦略はモデルの予測結果に関わらずサンプリング期間の直後に買いを、ショートオンリー戦略は売りを入れ、仲値が変動した直後に反対売買を行う。また、各手法の取引は2015年5月1日から2015年9月18日まで一営業日に一度行い、初期状態の資金を1として複利で運用した。

TOPIX100に関する投資シミュレーションの結果、100銘柄中64銘柄において、提案手法によるパフォーマンスが最も高く、ロングオンリー戦略、ショートオンリー戦略はそれぞれ13銘柄、23銘柄においてパフォーマンスが最も高い結果となった。したがって、提案手法の取引パフォーマンスはベースラインを多くの場合上回っていると言える。このことから、提案手法は概ね実務への応用につながる予測力を有していると考えられる。

5 今後の課題

今後の課題としては、今回計測した流動性指標では説明できない予測精度の向上、低下も見られたことから、流動性指標計測の精緻化が考えられる。本研究では即時性、タイトネス、デプスをそれぞれ約定時間、気配スプレッド率、注文総量で代表させ日次で平均をとる処理を行ったが、各値の分散も考慮に入れることでより解釈性向上につながると考えられる。次に、市場参加者や価格形成の挙動といった市場構造の時間的変化への対応が挙げられる。具体的には、学習データを期間ごとに分割しデータの層化を行い古いデータを段階的に減少させるよう多段階の学習フェーズを設ける学習手法や生成ネットワークを用いて生成した注文情報によるデータ拡張といった手法が挙げられる。最後に、本研究ではマーケットマイクロストラクチャーを考慮に入れるべく流動性指標を計測し、それを基に銘柄のクラスタリングを行ったが、自己符号化器などを用いて高頻度注文情報を直接的に入力したクラスタリ

ングを行うことができれば人手を介する処理を省くことにつながる。人手を介することによる恣意性を排除できれば、バイアスを抑制でき予測精度の向上につながると考えられる。

参考文献

- [Cao 09] Cao, C., Hansch, O., and Wang, X.: The information content of an open limit-order book, *Journal of Futures Markets: Futures, Options, and Other Derivative Products*, Vol. 29, No. 1, pp. 16–41 (2009)
- [Cont 01] Cont, R.: Empirical properties of asset returns: stylized facts and statistical issues (2001)
- [David 11] Easley, D., De Prado, M. M. L., Maureen O’Hara: The microstructure of the “flash crash”: flow toxicity, liquidity crashes, and the probability of informed trading, *The Journal of Portfolio Management*, Vol. 37, No. 2, pp. 118–128 (2011)
- [David 12] Easley, D., Prado, de M. M. L., Maureen O’Hara: The volume clock: Insights into the high-frequency paradigm, *The Journal of Portfolio Management*, Vol. 39, No. 1, pp. 19–29 (2012)
- [Dixon 18] Dixon, M.: Sequence classification of the limit order book using recurrent neural networks, *Journal of computational science*, Vol. 24, pp. 277–286 (2018)
- [Fama 70] Fama, E. F.: Efficient capital markets: A review of theory and empirical work, *The journal of Finance*, Vol. 25, No. 2, pp. 383–417 (1970)
- [Gould 16] Gould, M. D. and Bonart, J.: Queue imbalance as a one-tick-ahead price predictor in a limit order book, *Market Microstructure and Liquidity*, Vol. 2, No. 02, p. 1650006 (2016)
- [Jarnecic 14] Jarnecic, E. and Snape, M.: The provision of liquidity by high-frequency participants, *Financial Review*, Vol. 49, No. 2, pp. 371–394 (2014)
- [Kissell 06] Kissell, R.: The expanded implementation shortfall: Understanding transaction cost components, *The Journal of Trading*, Vol. 1, No. 3, pp. 6–16 (2006)
- [Maureen 15] Maureen O’Hara: High frequency market microstructure, *Journal of Financial Economics*, Vol. 116, No. 2, pp. 257–270 (2015)

- [Sirignano 19a] Sirignano, J. A.: Deep learning for limit order books, *Quantitative Finance*, Vol. 19, No. 4, pp. 549–570 (2019)
- [Sirignano 19b] Sirignano, J. and Cont, R.: Universal features of price formation in financial markets: perspectives from deep learning, *Quantitative Finance*, Vol. 19, No. 9, pp. 1449–1459 (2019)
- [Stoikov 16] Stoikov, S. and Waeber, R.: Reducing transaction costs with low-latency trading algorithms, *Quantitative Finance*, Vol. 16, No. 9, pp. 1445–1451 (2016)
- [Zhang 19] Zhang, Z., Zohren, S., and Roberts, S.: Deeplob: Deep convolutional neural networks for limit order books, *IEEE Transactions on Signal Processing*, Vol. 67, No. 11, pp. 3001–3012 (2019)
- [太田 11] 太田亘, 宇野淳, 竹原均 F 株式市場の流動性と投資家行動, 中央経済社 (2011)

Restricted Genetic Network Programmingによる リスク回避型外国為替取引戦略の構築 Constructing a Risk-Averse Forex Trading Strategy Using Restricted Genetic Network Programming

内田 純平^{*1}
Jumpei Uchida

穴田 一^{*1}
Hajime Anada

^{*1} 東京都市大学
Tokyo City University#1

Recently, many researchers have studied foreign exchange trading using technical analysis. However, it is difficult to achieve profitability using this technique. Therefore, using Genetic Network Programming, we construct a model that considers the technical index signal strength for devising a profitable trading strategy. Finally, we confirmed the effectiveness of our model using historical data of the exchange market.

1. はじめに

近年、テクニカル分析を用いた株式売買や外国為替証拠金取引(Foreign exchange, FX)に関する研究が精力的に行われている。為替市場での分析方法は、各国や世界全体の財政面や景気の指標などを見るファンダメンタル分析と、過去の時系列データを数理的に扱うテクニカル分析に大きく分けることができる。

テクニカル分析を用いた投資戦略に関する研究では、遺伝的アルゴリズム(Genetic Algorithm ; GA)によってテクニカル指標のパラメーターを最適化する研究[平林 08]や、間普らによって考案された遺伝的ネットワークプログラミング (Genetic Network Programming ; GNP)[間普 11]を用いた株式売買に関する研究[間普 07]などがあり、これらの研究は相場のトレンドや転換点を判断するテクニカル指標を組み合わせることにより売買戦略を構築している。しかし、テクニカル指標の売買シグナルには、取引のタイミングではないにも関わらず誤って売買シグナルを出すといったダマシが存在し、テクニカル指標の売買シグナルのみを頼りにして利益を常に上げることは難しい。そこで我々は、テクニカル指標による売買シグナルのダマシで取引をしないための信頼度の 1 つとして売買シグナルの強弱を定義し、GNPを用いて為替取引戦略の進化モデル GNP with Signal Strengthを構築し、その有効性を確認した[内田 20]。しかし、GNPによる売買戦略構築ではテクニカル指標の組み合わせ候補数が多すぎるという問題があった。そこで、本研究では GNP の解表現と進化の方法に制限を付けた制限付き遺伝的ネットワークプログラミング(Restricted GNP ; R-GNP)を構築し、新たに損切り・利益確定機能、トレンド判定機能、空売り機能を追加し、その有効性を確認した。

2. 提案手法

それぞれの個体が売買戦略のネットワークと 2 進数で表されるオシレーター系指標の閾値のリスト、損切り利確価格リストを持ち、ネットワークで表された戦略に従って取引を行う。その取引結果から個体を評価した値である適応度を求め、ネットワークとオシレーター系指標の閾値のリストを個体の遺伝子として遺伝的操作を用いることでより適応度が高くなるように個体を進化させていく。

2.1 テクニカル指標

テクニカル指標は金融取引の売買タイミングを判断するために使われる指標であり、トレンド系、オシレーター系の 2 つがある。トレンド系は為替の推移からトレンドを判断する指標、オシレーター系は為替の推移からトレンドの転換点を判断する指標である。

2.2 R-GNP の構造

R-GNP は、用意した全ての判定ノードと処理ノードを機能ごとに 1 つずつ配置し、無作為に自分以外のノードに接続する。判定ノードは条件判定を行い、その判定結果に基づき次に実行するノードを決定し、処理ノードは決められた処理を行う。R-GNP のノード遷移は開始ノードから始まり、条件に従い遷移を行いノードを決定する役割のみを持つ。また、R-GNP はノードの遅れ時間、終了条件の意味を持つネットワークの総遅れ時間が定義されている。また、各オシレーター系指標の閾値、損切り価格、利益確定価格、単純移動平均によるトレンド判定日数(day_{SMA})、指数平滑平均によるトレンド判定日数(day_{EMA})を 2 進数でそれぞれ表現し、Binary GA で用いる遺伝子情報を保存している。その遺伝子情報を以下の図 1 に示す。

連絡先: 内田純平, 東京都市大学, 〒158-8557
東京都世田谷区玉堤 1-28-1



図 1 Binary GA の遺伝子構造

2.3 R-GNP のノード遷移および学習

本研究では、各判定ノードの遅れ時間を 1, 各処理ノードの遅れ時間を R , 総遅れ時間を R に設定した。ここで総遅れ時間 R は売買の意思決定を行なう際、1 日あたり最大何個のテクニカル指標を使用するかを意味する。よって、本研究における 1 日の取引は、 $R - 1$ 回以下の判定の後 1 回の処理を行って終了するか、 R 回の判定で終了する場合が考えられる。また、ノード遷移は開始ノードから始まり、ノード間の接続と判定ノードでの判定結果に従って行われる。

I. 売買シグナルの強弱

テクニカル指標による売買シグナルの強さはテクニカル指標の種類によって 2 つに分けて定義した。1 つ目は、設定された値をテクニカル指標によって計算された値が越える度合い、2 つ目は、短期日数で計算されたテクニカル指標と長期日数で計算されたテクニカル指標が交差する角度の大きさによる定義である。1 つ目の場合テクニカル指標がシグナルを出した値と設定された値との差分。2 つ目の場合テクニカル指標がシグナルを出した時の交差の角度を個体毎に記憶し、それらを利用することで売買シグナルの強弱の判断基準を計算し、ノード遷移の際に計算された差分や角度が基準を越えている時を強いシグナル、越えていないときを弱いシグナルとした。

売買シグナルの強弱の判断基準については以下の通りである。

① 差分と角度の記憶

テクニカル指標の種類によって差分か角度どちらかをテクニカル指標毎にメモリに保存する。表 1 に角度または差分を記憶するテクニカル指標を分類して示す。

表 1 テクニカル指標の分類

角度	差分
短期中期 SMA	RSI
短期長期 SMA	Williams %R
Perfect Order SMA	単純移動平均乖離率
短期中期 EMA	指数移動平均乖離率
短期長期 EMA	CMO
Perfect Order EMA	ROC
MACD	Psychological line

Fast stochastic	
Slow stochastic	
DMI (+DMI & -DMI)	
DMI (ADX & ADXR)	

判定ノードの各テクニカル指標で、シグナルが出た時の差分または角度を、買いサインと売りサインで分けて個体毎に損益が確定するまでメモリに保存する。Stochastic は、差分と角度の両方を扱うテクニカル指標であるが、両方を考慮することは難しいため、交差の角度のみを売買シグナル強弱に利用している。

② 基準差分と基準角度の更新

全てのテクニカル指標(全 18 個)に、買いシグナルのメモリと売りシグナルのメモリの 2 つのメモリが存在する。従って、全テクニカル指標で $36(18 \times 2)$ 個のメモリが存在する。決済が確定した時の各テクニカル指標のメモリ内の平均を計算し、 $J_{ij}(i = 1, 2, \dots, 18, j = 1, 2)$ とする。ここで i はテクニカル指標の種類を表し、 $j = 1$ のときに買いシグナル、 $j = 2$ のときに売りシグナルの判断基準を表している。そして、決済の結果、利益が出ている時、 J_{ij} が 0 ではないテクニカル指標において売買シグナルの強さの判断基準 $B_{ij,t}$ の更新を次式で定義する

$$B_{ij,t} = B_{ij,t-1} + \frac{2(J_{ij} - B_{ij,t-1})}{t + 1} \quad (1)$$

ここで $t(1, 2, \dots)$ は更新回数を表す。

②において決済の結果、損失が確定した場合、個体のメモリをリセットする。

II. 開始ノード

自分の所持するポジションの有無と種類によって遷移先を変更することで、多点スタート戦略を可能にした。

III. 判定ノード

各判定ノードが、1 つの判定条件を所持する。表 2 にノードの判定条件を示す。

表 2 ノードの判定条件

機能番号	判定条件
1~18	買いサインかつシグナル強度が強い
	買いサインかつシグナル強度が弱い
	売買サインが無い
	売りサインかつシグナル強度が弱い
	売りサインかつシグナル強度が強い
19	$3\sigma < Close$
	$2\sigma \leq Close \leq 3\sigma$
	$-2\sigma \leq Close \leq 2\sigma$

20	$-3\sigma \leq Close \leq -2\sigma$
	$-3\sigma > Close$
	$s2 < Close \leq s3$
	$s1 < Close \leq s2$
	$r1 \leq Close \leq s1$
	$r2 \leq Close < r1$
21	$r3 \leq Close < r2$
	$Long_{sma} < Short_{sma} < SMA_t$
	$\frac{Short_{sma} + Long_{sma}}{2} < SMA_t$
	$SMA_t < Short_{sma} < Long_{sma}$
	$SMA_t < \frac{Short_{sma} + Long_{sma}}{2}$
	else
22	$Long_{ema} < Short_{ema} < EMA_t$
	$\frac{Short_{ema} + Long_{ema}}{2} < EMA_t$
	$EMA_t < Short_{ema} < Long_{ema}$
	$EMA_t < \frac{Short_{ema} + Long_{ema}}{2}$
	else
	else
23	$profit < 2L$
	$2L \leq profit < L$
	else
	$2G \geq profit > G$
	$profit > 2G$

ここで、 $Close$ は終値を表している。また、ノード番号 19 は Bollinger Band, 20 は Pivot という指標を表し、売買シグナルの強弱による遷移を行わない為、遷移方法が異なる。また、ノード番号 21 は単純移動平均(SMA)によるトレンド判定ノード、22 は指数平滑平均(EMA)によるトレンド判定ノードを表し、 $Long$ と $Short$ を次式で定義する。

$$Long_{sma} = SMA_{t-long_day_{sma}} \quad (2)$$

$$Short_{sma} = SMA_{t-short_day_{sma}} \quad (3)$$

$$Long_day_{sma} = 2 \times Short_day_{sma} \quad (4)$$

$$Short_day_{sma} = 5 \times ([day_{sma}]_{2 \rightarrow 10} + 1) \quad (5)$$

ここで、 t は開始から現在までの日数、 $[\cdot]_{2 \rightarrow 10}$ は 2 進数から 10 進数への変換を表す。また、 day_{sma} は各個体が持つ Binary GA の遺伝子情報を用いる。EMA の式についても同様である。

ノード番号 23 は損切り・利益確定ノードを表し、 $profit$ は現在のポジションを解消したときの手数料を考慮した利益を表す。ここで、買いポジションを持っているときの L と G は次式で定義する。

$$L = -(6.25 \times ([loss_{buy}]_{2 \rightarrow 10} + 1))^V \quad (6)$$

$$G = (6.25 \times ([get_{buy}]_{2 \rightarrow 10} + 1))^V \quad (7)$$

$$V = \sqrt{\frac{1}{n} \sum_{i=1}^n (p_i - \bar{p})^2} \quad (n = 5) \quad (8)$$

ここで、 $loss_{buy}$ は買いポジションを持っているときの損切り価格、 get_{buy} は買いポジションをも

っているときの利益確定価格であり、遺伝子上に個体ごとで保存されている。また、 V は(8)式で算出されるボラティリティを用いる。(8)式は過去 n 日間の終値の標準偏差を表している。

IV. 処理ノード

処理ノードは買いポジション獲得、売りポジション獲得、ポジション解消、いずれかの処理機能を持ち、処理ノードに遷移した時に買いポジション獲得の機能を持っていたら買い、売りポジション獲得の機能を持っていたら空売りを行う。しかし、複数のポジションを持つことはできないため、既に関買ポジションを持っている状態のポジション獲得や、ポジションを持っていないにもかかわらずポジションの解消はできない。しかし、ポジション解消ノードの遷移先がポジション獲得ノードであった場合、総遅れ時間を超えていても機能に応じたポジションを獲得する。

2.4 遺伝的操作

2.4.1 初期個体生成

ネットワークはスタートノード 1 個、判定ノード m 個、処理ノード n 個の合計 $m + n + 1$ 個のノードで N 個体生成する。ノードの機能は表 2 を参考にする。また、ノードの接続は自分以外の他のノードに無作為に接続する。そして、Binary GA の初期生成時に各オシレーター系テクニカル指標の閾値に振り分けられる領域 (bits) を以下表 3 に示す。

表 3 各テクニカル指標の割り当て領域

テクニカル指標	下限	上限
RSI	3bits	3bits
Williams %R	3bits	3bits
単純移動平均乖離率	5bits	5bits
指数移動平均乖離率	5bits	5bits
CMO	3bits	3bits
Psychological line	3bits	3bits
Fast stochastic	3bits	3bits
Slow stochastic	3bits	3bits

また、上記の閾値をデコードするとき用いる式を次に示す。

$$\begin{cases} \text{上限} : 5 \times ([TI]_{2 \rightarrow 10} + 1) \\ \text{下限} : 100 - 5 \times ([TI]_{2 \rightarrow 10} + 1) \end{cases} \quad (9)$$

$$\begin{cases} \text{上限} : 0.3 \times ([TI]_{2 \rightarrow 10} + 1) \\ \text{下限} : -0.3 \times ([TI]_{2 \rightarrow 10} + 1) \end{cases} \quad (10)$$

ここで、 TI は 2 進数で表現されたテクニカル指標の閾値を表す。 TI が 3bits であるときは(9)式、5bits であるときは(10)式を用いる。

2.4.2 評価

個体の適応度fitnessを次式で定義する.

$$\text{fitness} = (1 - Q) + \frac{\text{profit}}{\text{allprofit}} \quad (11)$$

$$\begin{cases} px^{(k+1)} + 1 - p - x = 0 \\ Q = x^r \end{cases} \quad (12)$$

ここで, Q はバルサラの破産確率, profit は売買を行う期間の損益の合計(銭), allprofit は期間内の上り幅の合計から取引手数料を上昇回数分引いたもの, p は勝率, k は損益率, r は資本比率, x は式(12)の上の式(バルサラの破産確率の特性方程式)の解を表し, この解を用いて Nauzer J. Balsara によって考案された破産確率 Q を求める.

2.4.3 エリート1個体保存

適応度が最も高い個体を 1 つそのまま次世代に保存し, 残りの個体は交叉と突然変異によって新しく生成されたものと入れ替える.

2.4.4 進化的操作

(1) 現世代から 2 個体をサイズ T のトーナメント選択で選択

(2) 交叉操作

A) 交換するノードを選択

確率 P_c で $m + n + 1$ 個のノード番号をそれぞれが交叉番号となるか判定する.

B) 交換する遺伝子の選択

全遺伝子番号についてそれぞれ一様交叉で交叉遺伝子となるか判定する

(3) 突然変異

A) 変異する接続を選択

交叉番号と判定されたノードのそれぞれの接続において確率 P_m で接続先を変更するか判定し, 無作為に変更する.

B) 変異する遺伝子の選択

交叉遺伝子番号についてそれぞれ確率 P_c で突然変異をするか判定し, 無作為に変更する.

生成された個体が $N - 1$ 個になるまで繰り返す.

3. 結果

本研究では, 提案手法である R-GNP の優位性を確認するために, 比較手法として適応度を profit とし, 表 4 で示す判定条件による遷移を行い, 空売りができない GNP(以下比較手法を GNP とする)を用いた. そして, 日足ドル円レートを用いて表 5 に示す期間で学習とテストを行った. また, 進化と学習のパラメータを表 6, 各テクニカル指標の設定日数について表 7 に示す. 取引を行う売買手数料は標準的な FX 会社に合

わせて 0.4 銭に設定した. なお, パラメーターは結果が最も良いものを使用した.

表 4 ノードの判定条件

機能番号	判定条件
1~18	買いサイン
	売買サインが無い
	売りサイン
19	$3\sigma < \text{Close}$
	$2\sigma \leq \text{Close} \leq 3\sigma$
	$-2\sigma \leq \text{Close} \leq 2\sigma$
	$-3\sigma \leq \text{Close} \leq -2\sigma$
	$-3\sigma > \text{Close}$
20	$s2 < \text{Close} \leq s3$
	$s1 < \text{Close} \leq s2$
	$r1 \leq \text{Close} \leq s1$
	$r2 \leq \text{Close} < r1$
	$r3 \leq \text{Close} < r2$

表 5 取引期間

学習期間	2001 年 1 月 1 日 ~ 2002 年 12 月 31 日
テスト期間	2003 年 1 月 1 日 ~ 2018 年 12 月 31 日

表 6 進化と学習のパラメーター

世代数	2000
個体数 N	101
交叉確率 P_c	25.0(%)
突然変異確率 P_m	1.0(%)
トーナメントサイズ T	2
総遅れ時間 R	5
試行回数	50

表 7 テクニカル指標の設定日数

テクニカル指標	設定日数		
SMA	5	10	25
EMA	5	10	25
MACD	12	26	9
Stochastic	9	3	3
DMI	14	14	14
Williams %R	10		
Psychological line	8		
単純移動平均乖離率	5		
指数移動平均乖離率	5		
CMO	14		
ROC	10		
RSI	14		
Bollinger Band	20		

3.1 学習期間

各世代で最も適応度が高い個体の 1 年間の学習期間の 1 ドル当たりの平均利益の世代推移を図 2 に示す. この図は, 50 試行を平均したものであり, 縦軸は平均利益(銭), 横軸は世代数を表し, 青色の点線は比較手法である GNP, オレンジ色の実線は提案手法である R-GNP を表す.

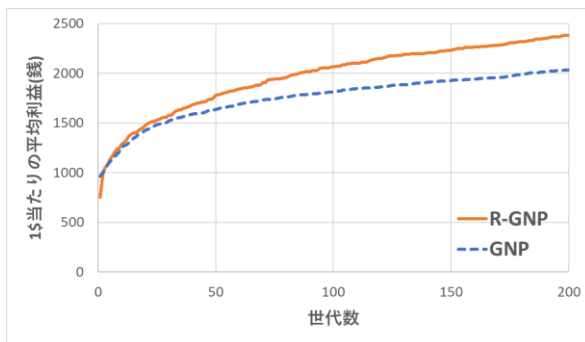


図2 最良個体の平均利益の世代推移(学習期間)

学習期間において提案手法である R-GNP の方が利益を上げることができている。これは、比較手法である GNP は空売りができないため提案手法よりも利益を出せなかったと考えられる。

3.2 テスト期間

各世代で最も適応度が高い個体の 16 年間各年の 1 ドル当たりの年間平均利益を図3に示す。この図は、50 試行を平均したもので、縦軸は 1 ドル当たりの平均利益(銭)、横軸は時間(年)を表し、青色の棒グラフは比較手法である GNP, オレンジ色の棒グラフは提案手法である R-GNP を表す。

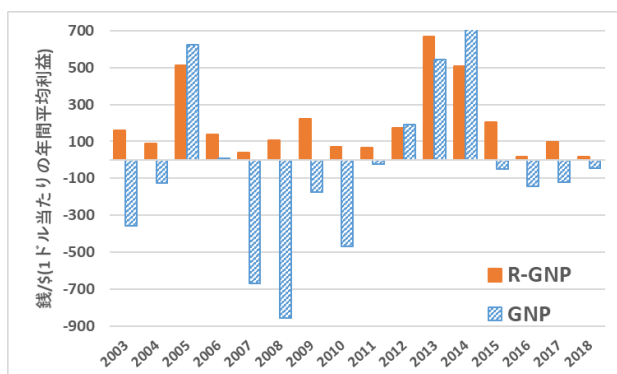


図3 最良個体の各年平均利益(テスト期間)

テスト期間において、提案手法である R-GNP は全ての期間で利益を出せていることがわかる。

4. 今後の課題

本研究では、複数のポジションを所持する事が出来ず、買い増しなどの戦略を取る事が出来ない。資産の 10%を購入などといったより複雑な売買ルールを導入する事を考えている。

参考文献

- [平林 08] 平林明憲, 伊庭齊志: 遺伝的アルゴリズムによる外国為替取引手法の最適化, 第 22 回人工知能学会全国大会, 2008
- [間普 11] 間普真吾, 平澤宏太郎: 遺伝的ネットワークプログラミングのアーキテクチャについて, シ

ステム制御情報学会誌, Vol. 55, No. 11, pp. 480-485, 2011

[間普 07] S. Mabu, K. Hirasawa, and T. Furuzuki: Trading Rules on Stock Markets Using Genetic Network Programming with Reinforcement Learning and Importance Index, IEEJ Trans. EIS, Vol. 127, No. 7, pp. 1061-1067, 2007

[内田 20] 内田純平, 穴田一: 売買シグナルの強弱を考慮した Genetic Network Programming による外国為替取引戦略の構築, 第 34 回人工知能学会全国大会, 2020

テンソルを用いたマルチモーダル学習による株式価格の変動予測

A tensor based multimodal learning for predicting stock price fluctuations

笹尾知広^{1*} 中田和秀¹
Tomohiro Sasao¹ Kazuhide Nakata¹

¹ 東京工業大学工学院経営工学系

¹ Graduate Major in Industrial Engineering and Economics, Tokyo Institute of Technology

Abstract: 近年では株価データと他のデータを組み合わせたマルチモーダル学習が頻繁に行われている。しかし、複数の情報源からそれぞれ特徴量ベクトルを作成した後、それらを単純に連結させたベクトルを用いて予測するのでは、情報源間の共起関係が十分に考慮されないという問題があった。このような問題に対して、高次元テンソルを用いて対処しようとする研究がされている。本研究はそれらの研究を発展させることで、複数の情報源を使ったときに、株式価格の変動をより精度よく予測するためのテンソルを用いたマルチモーダル学習フレームワークを提案する。そして、3種類の情報源としてトムソン・ロイター社が発信した記事、トムソン・ロイター市場心理指数 (TRMI) のセンチメントデータ、ニューヨーク証券取引所の株価データを用いて評価実験を行った。その結果、一日後の株価の予測・残差リターンの正負を予測する実験において、提案手法により予測性能の有意な改善が見られた。

1 はじめに

今日ではビックデータを用いた研究が注目されており、様々な研究や開発に適用されている。株式市場予測の分野でもインターネット上で取得することができる SNS やニュース記事などのオルタナティブデータを用いる研究がされてきた。一般的に複数の情報源を用いた機械学習のことをマルチモーダル学習といい、近年では株価データと他のデータを組み合わせたマルチモーダル学習による株価予測 (回帰問題) や株価の変動方向予測 (二値分類問題) が頻繁に行われている。しかし従来の手法は複数の情報源から特徴量ベクトルを作成した後、それらを一つのベクトルとして連結させて処理することがほとんどであり、そのような連結させたベクトルを用いた予測では、情報源間の共起関係が十分に考慮されないという問題があった。このような問題に対して、特徴量の組み合わせに高次元テンソルを用いて対処しようとする研究がさまざまな分野でされており、株式市場予測の分野では Li ら [6] が「テンソルを用いたマルチモーダル学習フレームワーク」という手法を提案した (図 1)。

本研究では複数の情報源を使ったときに株式価格の変動をより精度よく予測するための、「改良したテンソルを用いたマルチモーダル学習フレームワーク」を提案する。提案手法は既存手法と比べて、補助行列の作成と予測手法に工夫を加えたフレームワークになっている。具体的に予測手法に関して、学習時間が短い予測・二値分類問題を考慮した予測を行うために、画像・動画分析の分野で提案された2種類のテンソル予測モデル (リッジ回帰ベースのテンソル回帰モデル [3], ロジスティック回帰ベースのテンソル分類モデル [11]) を導入し、アルゴリズムの終了条件に改良を行った。フレームワーク中の処理に関して、補助行列を作成するために解く必要がある最小化問題において、新たに特殊な正則化項を加えた最小化問題を定式化した。

実験では株式データ、ニュースデータ、センチメントデータを用いた株価予測と株価の変動方向予測の実験を行った。

その結果それぞれの情報源から作成した特徴量を1つのベクトルとして連結して予測を行うのではなく、提案したテンソルを用いたマルチモーダル学習のフレームワークを用いることによって、より精度良く予測できることを示した。

2 関連研究

機械学習を用いた株式価格の変動予測において、さまざまなデータや機械学習手法が使われてきた。Liu ら [8] は株式データとニュースデータの2つを用いて、それぞれから特徴量を作成し LSTM モデルで、Apple 社の株価の変動方向予測を行った。Bollen ら [1] は株式データと Twitter の投稿で表現された国民感情を DJIA 指数の変動方向予測に利用した。Suwa ら [10] はヤフー株式掲示板のデータのトピック分析を行い、トピック分析から得られた特徴量と株式データを用いて、VI 指数の変動方向予測を行った。また Li ら [7] は株式データ、ニュースデータ、掲示板データの3つを組み合わせて個別銘柄の株価予測を行った。このようにさまざまなデータを用いた将来の株式価格やインデック指標、またその変動方向予測が行われている。

一般的に複数の情報源を用いた機械学習のことをマルチモーダル学習といい、近年では株式データと他のデータを組み合わせたマルチモーダル学習が頻繁に行われている。しかし従来の手法は複数の情報源から特徴量ベクトルを作成した後、それらを一つのベクトルとして連結させて処理することがほとんどであった。そのような連結させたベクトルを用いた予測では、情報源間の共起関係が十分に考慮されないという問題がある。

この複数の情報源間の組み合わせ方の問題に着目し、特徴量の組み合わせに高次元テンソルを用いた研究が様々な分野でされている。特に株式市場予測の分野では、Li ら [6] が株価データ、ニュースデータ、掲示板データから作成した特徴量を3次元テンソルを用いて組み合わせ予測を行う「テンソルを用いたマルチモーダル学習フレームワーク」を提案し、個別銘柄の株価予測で高次元テンソルによる組み合わせの有

*連絡先: 東京工業大学工学院経営工学系
〒152-8552 東京都目黒区大岡山 2-12-1
E-mail:sasao.t.aa@m.titech.ac.jp

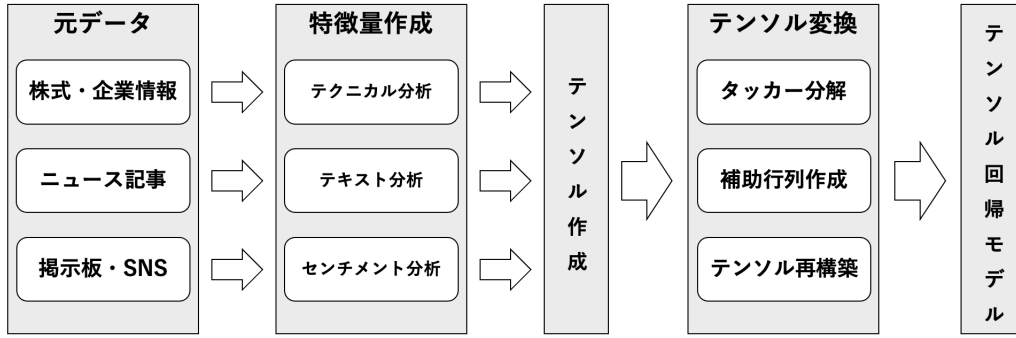


図 1: 「テンソルを用いたマルチモーダル学習フレームワーク」の概要図

効性を示した。さらに Huang ら [4] はこのモデルをさらに改良し、他の企業の情報も加味した手法を提案した。

しかし Li ら [6] と Huang ら [4] のモデルには、大きく分けて2つの問題点がある。1つ目は予測手法に関して、従来のモデルではサポートベクトル回帰をベースにしたモデルを用いており学習時間が遅いこと、また回帰問題しか想定しておらず二値分類問題が考えられていないことである。2つ目はフレームワークを実行する中で解く必要のある最小化問題において、求める最適解が一意に定まらず不安定なモデルであること、また制約条件が強すぎるためうまく学習できない場合があることである。

3 改良したテンソルを用いたマルチモーダル学習フレームワーク

本章では提案する「改良したテンソルを用いたマルチモーダル学習フレームワーク」について説明する(図2)。この手法は Li ら [6] の手法を発展させたフレームワークであり、大きく分けて「テンソル作成」、「テンソル変換」、「テンソル予測」の3つの部分から構成されている。

3.1 テンソル作成

テンソルを用いたマルチモーダル学習フレームワークでは、企業 s 、日付 t 毎に情報源からテンソル $\mathcal{X}_{s,t}$ を作成する。テンソルの階数は扱う情報源の数に対応している。本研究においては、企業データ、ニュースデータ、センチメントデータの3つの情報源を用いるため3階テンソルを用いており、以降3階のテンソルとして考える。テンソルは以下のような手順で作成する。

まず初めに企業データ、ニュースデータ、センチメントデータから企業 s 、日付 t における特徴量ベクトル $\mathbf{f}_1^{s,t} \in R^{I_1}$, $\mathbf{f}_2^{s,t} \in R^{I_2}$, $\mathbf{f}_3^{s,t} \in R^{I_3}$ を作成する。テンソル $\mathcal{X}_{s,t} \in R^{I_1 \times I_2 \times I_3}$ の各要素を $x_{i,j,k}^{s,t}$ とすると、各情報源の特徴量を次のようにテンソルに格納する。

- $x_{i,1,1}^{s,t}$, $1 \leq i \leq I_1$: 株価データの特徴量 $\mathbf{f}_1^{s,t}$
- $x_{2,j,2}^{s,t}$, $1 \leq j \leq I_2$: ニュースデータの特徴量 $\mathbf{f}_2^{s,t}$
- $x_{3,3,k}^{s,t}$, $1 \leq k \leq I_3$: センチメントデータの特徴量 $\mathbf{f}_3^{s,t}$
- それ以外は 0

このようなテンソル $\mathcal{X}_{s,t}$ を、各企業・日付毎に作成する。

3.2 テンソル変換

作成した企業 s 、日付 t のテンソル $\mathcal{X}_{s,t}$ から①タッカー分解、②補助行列作成、③テンソル再構築の3つの手順を介して、情報源間の共起関係の増強・ノイズが除去された新しいテンソル $\tilde{\mathcal{X}}_{s,t} \in R^{I_1 \times I_2 \times I_3}$ を作成する。具体的には、類似した結果になる2つのテンソルについて、それらのテンソルを分解して得られる因子行列同士が類似するような処理を行う。

まずはじめに、タッカー分解 [5] をもちいてテンソル $\mathcal{X}_{s,t}$ を次のように分解する。

$$\mathcal{X}_{s,t} = \mathcal{C}_{s,t} \times_1 U_1^{s,t} \times_2 U_2^{s,t} \times_3 U_3^{s,t} \quad (1)$$

ここで $U_k^{s,t} \in R^{I_k \times D_k}$ ($D_k \leq I_k$, $k = 1, 2, 3$) は企業データ、ニュースデータ、センチメントデータそれぞれを表現する因子行列になっており、 $\mathcal{C}_{s,t} \in R^{D_1 \times D_2 \times D_3}$ はそれぞれの情報源間の関係の強さを表すコアテンソルである。

続いて情報源間の共起関係の増強・ノイズを減少させるような、各因子に対応する補助行列 $V_{s,k} \in R^{I_k \times I_k}$ ($k = 1, 2, 3$) を導入する。Li ら [6] はこの補助行列 $V_{s,k}$ を求めるために次のような最小問題を定式化した。

$$\min_{V_{s,k}} J(V_{s,k}) = \sum_{i=1}^N \sum_{j=i}^N \|V_{s,k}^T U_k^{s,i} - V_{s,k}^T U_k^{s,j}\|^2 w_{s,i,j} \quad (2)$$

$$s.t. \sum_{i=1}^N \sum_{j=i}^N \|V_{s,k}^T U_k^{s,i}\|^2 w_{s,i,j} = 1 \quad (3)$$

$$w_{s,i,j} = \begin{cases} 1 & \text{if } i \leq j \text{ and } |y_{s,i} - y_{s,j}|/y_{s,j} \leq 5\% \\ 0 & \text{otherwise} \end{cases} \quad (4)$$

ここでの目的は、モード k の因子行列 $U_k^{s,i}|_{i=1}^N$ の本来の性質を保持しつつ、終値が類似している日付 i, j 同士の2つの行列 $V_{s,k}^T U_k^{s,i}$, $V_{s,k}^T U_k^{s,j}$ の値の差を小さくすることである。式(2)で時刻間の類似性を考慮しており、解の個数を制限させ $V_{s,k}$ のオーバーフィッティングを避けるために式(3)のような制約式を設けている。

本研究ではこの最小化問題に2つの問題があることを数式的に示すことができた。詳しい証明は割愛するが具体的に1つ目は、最小化問題の解が集合の形で与えられ、一意の解に定まらず不安定なものであること。2つ目は、制約式(3)がデータ数 N や $w_{s,i,j}$ の密度によって変化しないため、データ数 N が多い場合や $w_{s,i,j}$ が密な行列の場合、 $V_{s,k}$ が非常に小さい数値になることである。

そこで本研究ではこの補助行列 $V_{s,k}$ を求めるために、次のような制約なしの凸関数最小化問題を新たに定式化した。

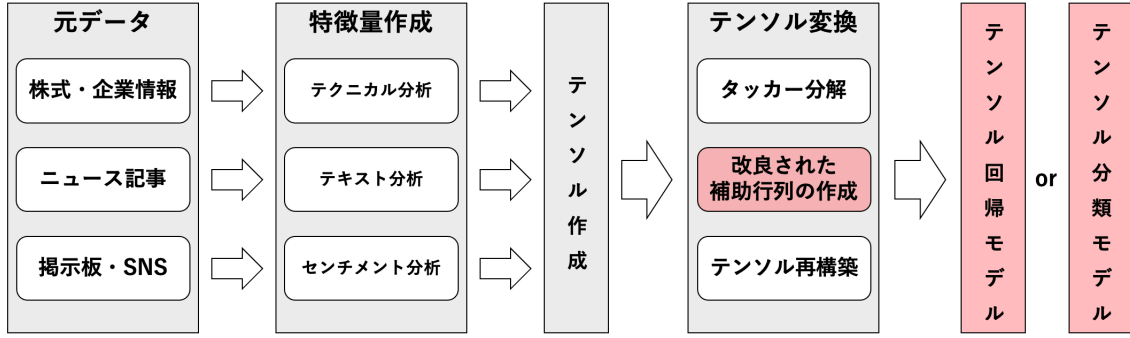


図 2: 「改良したテンソルを用いたマルチモーダル学習フレームワーク」の概要図

$$\begin{aligned} \min_{V_{s,k}} J(V_{s,k}) = & \frac{\mu}{\sum W_s} \sum_{i=1}^N \sum_{j=1}^N \|V_{s,k}^T U_k^{s,i} - V_{s,k}^T U_k^{s,j}\|^2 w_{s,i,j} \\ & + \frac{\gamma}{\sum E_s} \sum_{t=1}^N \sum_{c=1}^C \|V_{s,k}^T U_k^{s,t} - V_{s,k}^T U_k^{c,t}\|^2 z_{s,c} e_{t,s,c} \\ & + \frac{1}{\sum W_s} \sum_{i=1}^N \sum_{j=1}^N \|V_{s,k}^T U_k^{s,i} - I_{s,k} U_k^{s,i}\|^2 w_{s,i,j} \quad (5) \end{aligned}$$

$$w_{s,i,j} = \begin{cases} 1 & \text{if } i \leq j \text{ and } |y_{s,i} - y_{s,j}|/y_{s,j} \leq 5\% \\ 1 & \text{if } i > j \text{ and } |y_{s,i} - y_{s,j}|/y_{s,i} \leq 5\% \\ 0 & \text{otherwise} \end{cases} \quad (6)$$

または二値分類問題の場合、予測する二値数 $y_{s,i}^* = \{+1, -1\}$ に対して

$$w_{s,i,j} = \begin{cases} 1 & \text{if } y_{s,i}^* = +1 \text{ and } y_{s,j}^* = +1 \\ 0 & \text{otherwise} \end{cases} \quad (7)$$

$$z_{s,c} : \text{企業 } s, c \text{ の類似度} \quad (8)$$

$$e_{s,t,c} = \begin{cases} 1 & \text{時刻 } t \text{ の銘柄 } s, c \text{ のデータが存在} \\ 0 & \text{otherwise} \end{cases} \quad (9)$$

式 (5) の第 1,2 項目で時刻間・企業間の類似性を考慮しており、第 3 項目で $V_{s,k}$ が単位行列に近くなるという正則化項を加えている。さらに式 (5) は制約なしの凸関数最小化問題になっているため、 $(D_{s,k} - G_{s,k})$, $G_{s,k}$, $Q_{s,k}$ のいずれかが正定値行列という仮定をおけば、以下のように解析的に一意の解を求めることができる。

$$\begin{aligned} V_{s,k} = & \frac{1}{\sum W_s} \left\{ \frac{2\mu}{\sum W_s} (D_{s,k} - G_{s,k}) \right. \\ & \left. + \frac{1}{\sum W_s} D_{s,k} + \frac{\gamma}{\sum E_s} Q_{s,k} \right\}^{-1} D_{s,k} \quad (10) \end{aligned}$$

$$D_{s,k} = \sum_{i=1}^N \sum_{j=1}^N w_{s,i,j} U_k^{s,i} U_k^{s,jT}$$

$$G_{s,k} = \sum_{i=1}^N \sum_{j=1}^N w_{s,i,j} U_k^{s,i} U_k^{s,jT}$$

$$Q_{s,k} = \sum_{t=1}^N \sum_{c=1}^C (U_k^{s,t} - U_k^{c,t}) (U_k^{s,t} - U_k^{c,t})^T e_{t,s,c} z_{s,c}$$

最後に、タッカー分解で得られたコアテンソルと各情報源の因子行列、補助行列を用いて、次のようにテンソルを再構築する。

$$\tilde{\mathcal{X}}_{s,t} = \mathcal{C}_{s,t} \times_1 (V_{s,1}^T U_1^{s,t}) \times_2 (V_{s,2}^T U_2^{s,t}) \times_3 (V_{s,3}^T U_3^{s,t}) \quad (11)$$

このように各企業 s 、各日付 t について新しいテンソル $\tilde{\mathcal{X}}_{s,t}$ を作成する。新しく作られたテンソル $\tilde{\mathcal{X}}_{s,t} \in R^{I_1 \times I_2 \times I_3}$ は、補助行列 $V_{s,k}$ によって、元のテンソル $\mathcal{X}_{s,t}$ よりも情報源間の共起関係の増強・ノイズ除去されたテンソルとなる。

3.3 テンソル予測

提案する「改良したテンソルを用いたマルチモーダル学習フレームワーク」ではテンソルを共変量とする回帰問題と二値分類問題を扱う。各企業 s 、各観測値 $t = 1, \dots, N$ ごとに終値である実数 $y_{s,t}$ と共変量 $\tilde{\mathcal{X}}_{s,t}$ が与えられているとし、次のような回帰問題を考える。

$$\hat{y}_{s,t} = \langle \tilde{\mathcal{X}}_{s,t}, \mathcal{W}_s \rangle + b_s \quad (12)$$

$$= \tilde{\mathcal{X}}_{s,t} \times_1 \mathbf{w}_s^{(1)} \times_2 \mathbf{w}_s^{(2)} \times_3 \mathbf{w}_s^{(3)} + b_s \quad (13)$$

また二値分類問題であれば、日次収益率の正負である二値数 $y_{s,t}^* \in \{+1, -1\}$ と共変量 $\tilde{\mathcal{X}}_{s,t}$ が与えられているとし、次のような二値数の条件つき確率を考える。

$$\Pr(y_{s,t}^* = +1 | \tilde{\mathcal{X}}_{s,t}) = \frac{1}{1 + \exp(-\langle \mathcal{W}_s, \tilde{\mathcal{X}}_{s,t} \rangle - b_s)} \quad (14)$$

$$= \frac{1}{1 + \exp(-\tilde{\mathcal{X}}_{s,t} \times_1 \mathbf{w}_s^{(1)} \times_2 \mathbf{w}_s^{(2)} \times_3 \mathbf{w}_s^{(3)} - b_s)} \quad (15)$$

ここでの目的は、回帰係数テンソル $\mathcal{W}_s \in R^{I_1 \times I_2 \times I_3}$ 、切片 $b_s \in R$ を推定し、新しい入力 $\tilde{\mathcal{X}}$ に対して y , $\Pr(y_{s,t}^* = +1 | \tilde{\mathcal{X}}_{s,t})$ を予測することである。

従来の「テンソルを用いたマルチモーダル学習フレームワーク」ではこのサポートベクトル回帰をベースとした予測モデルしか考えられていなかった。本研究では学習時間が短いリッジ回帰をベースにしたテンソル回帰モデル [3] と、二値分類問題を考慮するためのロジスティック回帰をベースにしたテンソル分類モデル [11] をこのフレームワークに導入する。これらのテンソル予測モデルは、画像・動画分析の分野で提案された手法である。回帰式は式 (12), (14) であり、リッジ回帰ベースの手法では損失関数式 (16)、ロジスティック回帰ベースの手法では損失関数式 (17) を用いて、最適な $\mathcal{W}_s, b_s$ を求める。

$$L = \frac{1}{2} \sum_{t=1}^N \left(y_{s,t} - \tilde{x}_{s,t} \times_1 \mathbf{w}_s^{(1)} \times_2 \mathbf{w}_s^{(2)} \times_3 \mathbf{w}_s^{(3)} - b_s \right)^2 + \frac{\lambda}{2} \left(\|\mathbf{w}_s^{(1)}\|^2 + \|\mathbf{w}_s^{(2)}\|^2 + \|\mathbf{w}_s^{(3)}\|^2 + \|b_s\|^2 \right) \quad (16)$$

$$L = \sum_{t=1}^N \log \left[1 + \exp \left\{ -y_{s,t}^* \left(\tilde{x}_{s,t} \times_1 \mathbf{w}_s^{(1)} \times_2 \mathbf{w}_s^{(2)} \times_3 \mathbf{w}_s^{(3)} + b_s \right) \right\} \right] + \frac{\lambda}{2} \left(\|\mathbf{w}_s^{(1)}\|^2 + \|\mathbf{w}_s^{(2)}\|^2 + \|\mathbf{w}_s^{(3)}\|^2 + \|b_s\|^2 \right) \quad (17)$$

これらの損失関数は変数が多いため、変数 $\mathbf{w}_s^{(1)}, \mathbf{w}_s^{(2)}, \mathbf{w}_s^{(3)}$ のうち2つを固定し通常のリッジ回帰・ロジスティック回帰の問題に帰着させ、交互に最適な $\mathbf{w}_s, b_s$ を求めていく。また今回2つの予測モデルをフレームワークに導入するにあたって、終了条件の改良を行った。具体的には従来の終了条件は「 $\mathbf{W}_s$ の収束または、一定回数以上の反復」であったが、「損失関数の値の収束または、一定回数以上の反復」という終了条件にすることで、過学習が起きづらい学習モデルにすることができた。

4 数値実験

4.1 使用データ

今回使用するデータは企業データ、ニュースデータ、センチメントデータの3種類である。

「企業データ」：米国のニューヨーク証券取引所に上場している10業界82社のデータ。取引日・企業ごとに「終値」、「出来高」が記録されており、この二つとテクニカル指標指標の「ストキャスティクス%K」、「RSI」の4種類の特徴量を用いる。またインデックス指標としてS&P500の「終値」を用いる。

「ニュースデータ」：トムソン・ロイター社が提供しているニュースデータを用いる。「記事ID」、「タイトル」、「記事本文」、「公開日時」、「更新日時」、「企業ID」などが1セットとなったデータである。これらのデータからZhangら[12]の手法を参考に、タイトルから動詞を抽出し、それを学習済みのword2vec[9]と主成分分析を用いて20次元の分散表現を得る。

「センチメントデータ」：トムソン・ロイターマーケットサーイク指数(TRMI)を用いる。トムソン・ロイター社が世界中のニュースや公開情報、ソーシャルメディアの掲示板、ツイッターから情報集め、各企業や各国通貨、資源などの資産クラスごとに様々な種類のTRMI指数を算出している。今回はニュースデータと区別するために、TRMIの中でもニュース以外のインターネット掲示板やツイッターをソースとして計算されたTRMI指数を用いる。1日に一度(ニューヨーク時間の15:30)、各企業の過去24時間に対してのセンチメントの値が取得される。今回は「buzz」とTRMI指数である「sentiment」、「emotionVsFact」の3つの指標を用いる。「buzz」はその銘柄の「どれくらい話題になっているかの情報量を図る」指標であり[2]、「sentiment」は「ポジティブセンチメントとネガティブセンチメントの差」、「emotionVsFact」は「すべての感情的なセンチメントとすべての事実や時事問題の差」を図る指標である[13]。

4.2 実験概要

本研究はニューヨーク証券取引所に上場している10業界82社を対象に実験を行う。ニューヨーク証券取引所は現地時間の9:30-16:00が取引時間であり、また実験に用いるTRMIデータがNY時間の15:30に過去24時間のセンチメントの値を算出しているため、「各企業ごとに前日の15:30から当日の15:30までのデータを用いて、当日の終値または残差リターンの正負を予測する」というタスクを用いて手法の比較を行う。二値分類の予測対象に残差リターンの正負を用いる理由は、各銘柄は市場全体の流れで価格が変動することがあり、市場全体の影響を取り除いた残差リターンを考えると、今回実験に用いる3つのデータからその銘柄の予測しやすくなると考えたからである。

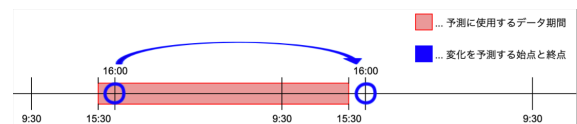


図3: 予測に用いるデータ期間と対象

4.3 結果

表1: サポートベクトル回帰ベースの予測

決定係数(R ²)				
μ	γ	類似度		
		(1)	(2)	(3)
1	1	-0.48123	0.92536	0.91818
1	0.1	0.89571	0.92155	0.92084
1	0.01	0.91906	0.92127	0.92116
1	0	0.92121		
0.1	1	0.97357	0.97344	0.97344
0.1	0.1	0.97347	0.97342	0.97342
0.1	0.01	0.97342	0.97342	0.97342
0.1	0	0.97342		
0.01	1	0.97357	0.97336	0.97336
0.01	0.1	0.97340	0.97336	0.97336
0.01	0.01	0.97337	0.97336	0.97337
0.01	0	0.97336		
0	1	0.97357	0.97337	0.97336
0	0.1	0.97339	0.97335	0.97335
0	0.01	0.97336	0.97335	0.97334
テンソル変換なし		0.97336		
(ベクトルで予測)		0.97323		

表1, 表2, 表3において、テンソル変換を行って予測精度が向上したものを赤く塗っている。それぞれの表の $\mu, \gamma, (1), (2), (3)$ はテンソル変換のハイパーパラメータを表しており、特に(1)(2)(3)は類似度 $z_{s,c}$ の種類を表している。(1)は銘柄同士の残差リターン正負の一致度、(2)は残差リターンの相関係

表 2: リッジ回帰ベースの予測

決定係数(R ²)				
μ	γ	類似度		
		(1)	(2)	(3)
1	1	0.97241	0.97248	0.97247
1	0.1	0.97242	0.97228	0.97237
1	0.01	0.97243	0.97235	0.97246
1	0	0.97252		
0.1	1	0.97240	0.97239	0.97224
0.1	0.1	0.97238	0.97233	0.97248
0.1	0.01	0.97241	0.97243	0.97230
0.1	0	0.97232		
0.01	1	0.97246	0.97237	0.97239
0.01	0.1	0.97245	0.97228	0.97242
0.01	0.01	0.97244	0.97235	0.97235
0.01	0	0.97241		
0	1	0.97251	0.97248	0.97237
0	0.1	0.97233	0.97232	0.97230
0	0.01	0.97238	0.97238	0.97228
テンソル変換なし		0.97244		
(ベクトルで予測)		0.97186		

表 3: ロジスティック回帰ベースの予測

正解率(AC)				
μ	γ	類似度		
		(1)	(2)	(3)
1	1	0.51293	0.52166	0.52189
1	0.1	0.52165	0.52191	0.52202
1	0.01	0.52185	0.52201	0.52202
1	0	0.52202		
0.1	1	0.52511	0.52303	0.52294
0.1	0.1	0.52376	0.52263	0.52280
0.1	0.01	0.52269	0.52279	0.52280
0.1	0	0.52278		
0.01	1	0.52531	0.52326	0.52353
0.01	0.1	0.52264	0.52420	0.52415
0.01	0.01	0.52396	0.52413	0.52409
0.01	0	0.52409		
0	1	0.52530	0.52345	0.52374
0	0.1	0.52225	0.52418	0.52424
0	0.01	0.52426	0.52415	0.52412
テンソル変換なし		0.52409		
(ベクトルで予測)		0.52210		

表 4: 既存のサポートベクトル回帰ベースのテンソル回帰モデルと提案したリッジ回帰ベースのテンソル回帰モデルの比較

トレーニング データ数	企業数	決定係数(R ²)		学習時間[s]	
		SVR ベース	リッジ回帰 ベース	SVR ベース	リッジ回帰 ベース
300~400	5	0.97724	0.97555	2.366	0.164
400~500	11	0.97329	0.97213	3.657	0.237
500~600	15	0.97285	0.96963	5.023	0.281
600~700	11	0.97305	0.97272	6.032	0.307
700~800	8	0.98187	0.98196	7.751	0.305
800~900	10	0.97881	0.97858	9.395	0.307
900~1000	7	0.97232	0.97191	11.812	0.346
1000~1100	2	0.98415	0.98371	16.653	0.447
1100~1200	9	0.95960	0.96003	19.666	0.371
1200~1300	4	0.96821	0.96805	17.905	0.334
平均	82	0.97336	0.97244	8.711	0.299

数, (3) は銘柄同士が同じ記事に共に出了割合を銘柄間の類似度として。

これらの結果からまず、ただ単に特徴量を1つのベクトルに連結することや3次テンソルに格納したデータで予測を行うのではなく、テンソル変換を用いてテンソル再構築したデータで予測を行った時の方が精度よく学習できていることがわかる。これは提案した補助行列が時刻間・企業間の類似性を考慮したテンソル変換を行い、情報源間の共起関係の増強・ノイズ除去された新しいテンソルを作れたからだと考えられる。さらにハイパーパラメータごとの精度の違いを見ると、 μ が小さく、 γ が大きいほど精度よく予測できる傾向があることがわかる。このことから同一銘柄の時刻間の類似性を考慮するよりも、同一時刻の他企業の類似性を考慮したほうが精度が良いテンソル変換が行えるということがわかる。企業間の類似度については、終値予測タスク・残差リターンの正負予測タスクの両方において(1)の「残差リターンの正負の一致度」を用いたものが全体的に精度が良いことがわかる。

表4は82企業をトレーニングデータ数で分けて結果をまとめたものである。この結果から全体の株価予測精度は、既存手法のサポートベクトル回帰のテンソル回帰モデルの方が優れていることがわかる。しかし学習時間を見ると、トレーニングデータ数が増えていくごとにサポートベクトル回帰ベースの予測モデルの学習時間が増えていることがわかる。これはサポートベクトル回帰の学習の遅さによるものであり、実際にこれらのモデルを用いる際は予測精度と学習時間のトレードオフでどちらのモデルを用いるかを検討すべきである。

5 おわりに

本研究の目的は、株式価格の変動予測のためのマルチモーダル学習において、複数の情報源の情報を十分に考慮しより精度の高い予測を行うということであった。そこで本研究では「改良したテンソルを用いたマルチモーダル学習フレームワーク」を提案し、株価データ、ニュースデータ、センチメントを用いた株価予測(回帰問題)と残差リターンの正負予測

(二値分類問題)の実験を行った。新たに導入したリッジ回帰ベースのテンソル回帰モデルは、既存のサポートベクトル回帰ベースのテンソル回帰モデルよりも予測精度は劣るものの学習時間を大幅に削減した予測モデルとなった。また株価予測・残差リターンの正負予測の両方において、それぞれの情報源から作成した特徴量を1つのベクトルとして連結して予測を行うのではなく、提案したテンソルを用いたマルチモーダル学習のフレームワークを用いることによって、より精度良く予測できることが示された。

今後の課題として次の2つが考えられる。1つ目の課題はさらなる予測精度の向上である。今回の手法・実験ではわずかな精度の向上しか確認できなかった。これは従来より株式市場の予測は難しいタスクであるとともに、今回の問題設定が1日間隔の予測を行っていたからだと考える。一日の間に複数の出来事があり、相場のモメンタムが激しく変わる場合などは予測が難しくなる。そこで実験の時間間隔をより短くすることで精度が向上の可能性があると考え。2つ目の課題はテンソルの各軸の次数を削減できるようなテンソル変換手法の考案である。従来の手法・今回提案した手法では高次元テンソルを扱うため、学習・予測の際に利用するリソースが必要になる。従来の「テンソルを用いたマルチモーダル学習フレームワーク」ではテンソル変換の際に各軸の次元削減が可能であったが、今回の提案手法では $V_{s,k}$ が正方行列という条件を追加したため、次元削減の要素を失ってしまった。よって次元削減と予測精度向上の両立できるテンソル変換モデルを考案することが望ましい。

謝辞

本研究はみずほ第一フィナンシャルテクノロジー株式会社にデータや実験環境の提供、研究の相談など多大なるご協力をいただきました。この場を借りて厚く御礼申し上げます。

参考文献

- [1] Johan Bollen, Huina Mao and Xiaojun Zeng, “Twitter Mood Predicts the Stock Market,” *Journal of Computational Science*, 2(1), pp.1-8, 2011.
- [2] Baoqing Gan, Vitali Alexeev, Ron Bird and Danny Yeung, “Sensitivity to Sentiment: News vs Social Media,” *International Review of Financial Analysis*, 67(101390), 2020.
- [3] Weiwei Guo, Irene Kotsia and Ioannis Patras, “Tensor Learning for Regression,” *IEEE Transactions on Image Processing*, 21(2), pp.816-827, 2012.
- [4] Jieyun Huang, Yunjia Zhang, Jialai Zhang and Xi Zhang, “A Tensor-Based Sub-Mode Coordinate Algorithm for Stock Prediction,” *Proceedings of 2018 IEEE Third International Conference on Data Science in Cyberspace (DSC)*, pp.716-721, 2018.
- [5] Tamara G. Kolda and Brett W. Bader, “Tensor Decompositions and Applications,” *SIAM Review*, 51(3), pp.455-500, 2009.
- [6] Qing Li, Yuanzhu Chen, Li L.Jiang, Ping Li and Hsinchun Chen, “A Tensor-Based Information Framework for Predicting the Stock Market,” *ACM Transactions on Information Systems*, 34(2), 2016.
- [7] Qing Li, Tiejun Wang, Ping Li, Ling Liu, Qixu Gong and Yuanzhu Chen, “The Effect of News and Public Mood on Stock Movements,” *Information Sciences*, 278, pp.826-840, 2014.
- [8] Yang Liu, Qingguo Zeng, Huanrui Yang and Adrian Carriro, “Stock Price Movement Prediction from Financial News with Deep Learning and Knowledge Graph Embedding,” *Proceedings of Knowledge Management and Acquisition for Intelligent Systems*, pp.102-113, 2018.
- [9] Radim Řehůřek and Petr Sojka, “Software Framework for Topic Modelling with Large Corpora,” *Proceedings of the LREC 2010 Workshop on New Challenges for NLP Frameworks*, pp.45-50, 2010.
- [10] Hiroki Suwa, Yuki Ogawa, Eiichi Umehara, Kento Kakigi, Keiichi Yasumoto, Tatsuto Yamashita and Kota Tsubouchi, “Develop Method to Predict the Increase in the Nikkei VI Index,” *Proceedings of 2017 IEEE International Conference on Big Data (Big Data)*, pp.3133-3138, 2017.
- [11] Xu Tan, Yin Zhang, Siliang Tang, Jian Shao, Fei Wu and Yueting Zhuang, “Logistic Tensor Regression for Classification,” *Proceedings of the Third Sino-Foreign-Interchange Conference on Intelligent Science and Intelligent Data Engineering*, pp.573-581, 2012.
- [12] Xi Zhang, Yunjia Zhang, Senzhang Wang, Yuntao Yao, Binxing Fang and Philip Yu, “Improving Stock Market Prediction via Heterogeneous Information Fusion,” *Knowledge-Based Systems*, pp.236-247, 2017.
- [13] リチャード・L・ピーターソン, [監修] 長尾慎太郎, [訳] 井田京子, 『市場心理とトレード -ビッグデータによるセンチメント分析』, パンローリング株式会社, 東京都, 2017.

効率的な Deep Hedging のためのニューラルネットワーク構造

Neural Network Architecture for Efficient Deep Hedging

今木 翔太^{1*} 今城 健太郎² 伊藤 克哉² 南 賢太郎² 中川 慧³
Shota Imaki, Kentaro Imajo, Katsuya Ito, Kentaro Minami, Kei Nakagawa

¹ 東京大学理学系研究科物理学専攻

¹ Department of Physics, The University of Tokyo

² 株式会社 Preferred Networks

² Preferred Networks, Inc.

³ 野村アセットマネジメント株式会社

³ Nomura Asset Management Co., Ltd.

Abstract: Deep Hedging [1] は、深層学習を用いて、不完全市場でのオプションの最適ヘッジ戦略を計算する汎用的な枠組みである。しかし、最適ヘッジ戦略は過去のヘッジ行動に依存するため、訓練が難しい。この課題を解決するため、我々は、No-Transaction Band 戦略のアイデアを活用する。No-Transaction Band 戦略は、指数型効用関数のもとでヨーロッパ・オプションをヘッジする最適戦略として知られる。我々は、この戦略がより一般の効用関数・エキゾチックを含む幅広いオプションに対しても最適ヘッジであることを証明する。この結果に基づき、我々は No-Transaction Band ネットワークを提案する。このネットワークは、過去のヘッジ行動をニューラルネットのインプットに使用せず、No-Transaction Band を出力する。そのため、高速な訓練が可能となり、かつ理論的結果により最適ヘッジ戦略をより正確に求めることが期待できる。数値実験の結果、ヨーロッパ・オプションおよびルックバック・オプションに対し、我々のネットワークが通常の順伝播型ニューラルネットより高速に、より優れたヘッジ戦略を実現することを示した。

1 はじめに

オプションのヘッジは、数兆ドル規模といわれるデリバティブ産業において不可欠な取引である。完全市場（つまり、取引コストなどの摩擦のない市場）においては、ブラック・ショールズの複製ポートフォリオ [2] を構築することで最適ヘッジができる。しかし、完全市場ではオプション価格の「Catch-22」として知られる逆説的な状況が生まれる [3]。これは、ブラック・ショールズモデルの前提が満たされるならオプションは冗長であり、前提が満たさないなら価格付けが正しくできないという、ブラック・ショールズモデルのジレンマを指す。実際の市場は常に摩擦、とりわけ取引コストの存在によってこのジレンマが解消される一方で、これによって最適ヘッジ戦略の構成がはるかに困難になる。

取引コストの存在下でのオプションの研究は [4] にさかのぼる。[4] は、ブラック・ショールズの複製ポートフォリオを離散的にリバランスするポートフォリオに拡張し、その後 [5] は、平均分散の意味で最適ヘッジ

を求めた。「最適」ヘッジの概念は、無差別効用価格付けの枠組みにより明確になった [6, 7]。この枠組みに基づき、多くの研究が最適ヘッジを追求してきた [8, 9]。

そのなかでも、[1] は、深層学習を用いて不完全市場でのオプションの最適ヘッジ戦略を計算する汎用的な枠組みである Deep Hedging を提案した。これは深層学習の高い表現力 [10] および最適化アルゴリズム [11] によって、最適ヘッジを達成することを狙いとする。実際、取引コストの存在下でのヨーロッパ・オプションのヘッジ戦略に関する数値実験は、この枠組みの効率性とスケーラビリティを示している。

このように Deep Hedging は有効であるものの、最適ヘッジ戦略は過去のヘッジ行動に依存するため、訓練が困難であるという問題がある。つまり、次時刻の最適なヘッジ行動が現在以前のヘッジ行動に依存するため、ニューラルネットの訓練が難しい。実際、我々が数値実験で示すように、ナイーブに設計したニューラルネットでは、大量の訓練を経ても収束しないことが確認できる（図 4, 6）。

そこで、本研究では、Deep Hedging よりも効率的に訓練でき、かつより広い設定で最適ヘッジ戦略を達成

*連絡先：東京大学理学系研究科物理学専攻 原子核理論研究室
東京都文京区本郷 7-3-1
E-mail: shota.imaki.0801@gmail.com

する方法論を開発すべく、No-Transaction Band 戦略に着目する。この戦略は、ヘッジ比率の許容幅（No-Transaction Band）内では取引を控えることで、不要な取引を抑え、取引コストを節約する。また、この戦略は指数型効用関数のもとでヨーロッパ・オプションをヘッジする最適戦略としても知られる [3]。

本研究は、これを利用した効率的な Deep Hedging のための No-Transaction Band ネットワークを提案する。このネットワーク構造は、Deep Hedging と比べて二つの特徴がある。第一に、No-Transaction Band ネットワークは、人間のトレーダーが用いるであろう様々な情報を用いることができる一方で、現在までのヘッジ行動をインプットに使用しない。これにより、前述の訓練の困難を回避し、効率的にヘッジ戦略を求めることができる。第二に、No-Transaction Band ネットワークは No-Transaction Band 戦略を出力する。我々はこの戦略が、指数型効用関数のもとでのヨーロッパ・オプションだけでなく、より一般の効用関数や、エキゾチックを含む幅広いオプションに対しても最適ヘッジであることを証明する。この結果は、No-Transaction Band を出力するニューラルネットを用いることで、最適ヘッジを達成できることを示唆する。No-Transaction Band ネットワークは、図 2 のように、任意のニューラルネットにある層を一層加えるだけで実装できる。

さらに、数値実験により、No-Transaction Band ネットワークが、エキゾチックを含むオプションに対し、通常のネットワーク構造に比べてより早くより優れたヘッジ戦略を実現することを示す。

本論文の構成は次の通りである。まず、第 2 章でオプションのヘッジ最適化問題を定式化する。次に、第 3 章では、Deep Hedging がどのように深層学習でヘッジ戦略を最適化するか説明し、訓練の難しさを指摘する。第 4 章で No-Transaction Band ネットワークを提案し、これが Deep Hedging における訓練の難しさを軽減する理由を説明する。第 5 章では、エキゾチックを含むオプションに対して数値実験を行い、No-Transaction Band ネットワークが高速に優れたヘッジ戦略を実現することを示す。最後に結論を述べる。

2 オプションのヘッジ最適化問題

本章では、市場の取引コストおよびオプションを説明し、オプションのヘッジ最適化問題を定式化する。

離散時間 $t \in \{t_0 = 0, t_1, \dots, t_n = T\}$ に取引可能な資産 $S = \{S_{t_i} | S_{t_i} > 0\}_{0 \leq t_i \leq T}$ からなる市場を考える。本研究では、資産価格ダイナミクスとして幾何ブラウン運動を仮定する。

$$\Delta S_{t_i} = \sigma S_{t_i} \Delta W_{t_i} \quad (1)$$

ここで $\Delta S_{t_i} \equiv S_{t_{i+1}} - S_{t_i}$ であり、 ΔW_{t_i} は i.i.d. の標準正規分布である。エージェントの各時刻における資産の保有単位を $\delta = \{\delta_{t_i} | \delta_{t_i} \in \mathbf{R}\}_{0 \leq t_i \leq T}$ で表す。ただし、 δ_{t_i} は時刻 t_i より過去に入手できる情報から決定する。また、無リスク利子はゼロとする。

取引には、取引額に対し割合 c の取引コストがかかるとする。つまり、時刻 t_i までに支払う累積コスト $C_{t_i}(S, \delta)$ は $C_0 = 0$ および $\Delta C_{t_i} = c S_{t_i} |\Delta \delta_{t_i}|$ をみたす。

資産 S のオプション Z は、満期 $t = T$ に、資産価格の経路に依存するペイオフ $Z(S)$ を払い出す。のちの数値実験では、ヨーロッパ・オプション $Z(S) = \max(S_T - K, 0)$ およびルックバック・オプション $Z(S) = \max(\max(S) - K, 0)$ を扱う。

オプションのディーラーは、顧客にオプションを販売し、ペイオフを支払う義務を負う。ディーラーは原資産を売買してこのリスクをヘッジする。その結果、満期におけるディーラーのポートフォリオは次の確率変数でかける。

$$P(-Z, S, \delta) = -Z(S) + (\delta \cdot S)_T - C_T(S, \delta) \quad (2)$$

$$(\delta \cdot S)_T \equiv \sum_{i=0}^{n-1} \delta_{t_i} \Delta S_{t_i}$$

各項は、顧客へのペイオフの支払い、原資産の価格変動による損益、取引コストを表す。ディーラーは、確率的なポートフォリオ (2) を効用関数 u について最適化する。すなわち、ヘッジ戦略 δ により、期待効用のマイナスで与えられる次の損失関数の最小化を目指す。

$$\ell(\delta) = -\mathbf{E}[u(P(-Z, S, \delta))] \quad (3)$$

のちの数値実験では、リスク回避度 1 の指数型効用関数 $u(x) = -\exp(-x)$ を扱う。

3 Deep Hedging

本章では、Deep Hedging [1] がいかに深層学習を用いて最適ヘッジ戦略を追求するか説明し、最適ヘッジ戦略の訓練が難しい原因を指摘する。

Deep Hedging の着想は、ヘッジ戦略 δ をニューラルネットで表現することである。[1] の提案するネットワークは、現在のヘッジ行動をインプットに用いる。

$$\delta_{t_{i+1}} = \text{NN}(I_{t_i}, \delta_{t_i}) \quad (4)$$

模式図を図 1 に示す。ここで NN は順伝播型ニューラルネット、 I_{t_i} は時刻 t_i 以前に入手できるヘッジに有用な情報の集合を表す。たとえば、価格が幾何ブラウン運動に従う資産のヨーロッパ・オプションの場合、 $I_{t_i} = \{S_{t_i}, t_i\}$ などが該当する。ニューラルネットの普

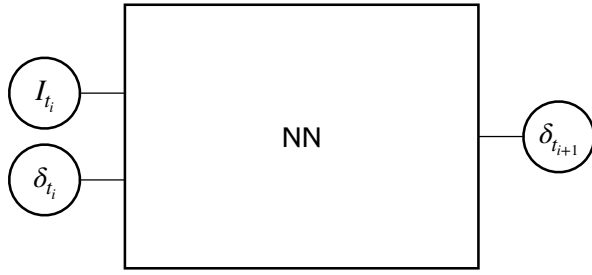


図 1: 順伝播型ニューラルネットワーク

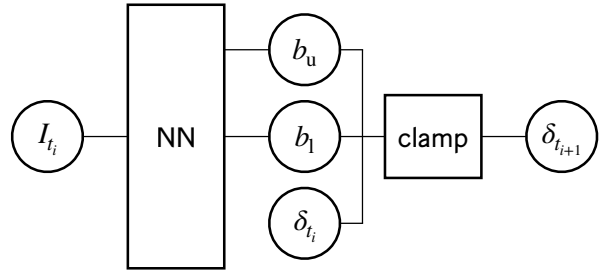


図 2: No-Transaction Band ネットワーク

逼近定理 [10] は、十分な数のパラメタをもつニューラルネットワーク表現 (4) が最適ヘッジ戦略をよく近似できることを示唆する。それゆえ、資産価格のシミュレーションを多数回行い、損失関数 (3) を最小化すべくニューラルネットワークのパラメタを調整すれば、最適に近いヘッジ戦略が得られると期待できる。

ところが、実際には、Deep Hedging におけるニューラルネットワークの訓練は難しい。なぜなら、ニューラルネットワーク (4) の入力値が現在のヘッジ行動 δ_{t_i} を含むためである。実際、第 5 章の数値実験で、ナイーブに実装したニューラルネットワークは、大量の訓練を経ても収束しないことが確認できる。

4 提案手法

本章では、No-Transaction Band ネットワークを提案し、この構造が訓練を効率化する理由を説明する。さらに、広いクラスの効用関数・エキゾチックを含むオプションに対して、No-Transaction Band 戦略が最適ヘッジとなることを証明する。

4.1 ネットワーク構造

No-Transaction Band ネットワークの構造は単純である。ニューラルネットワークは区間 $[b_l, b_u]$ を出力し、現在のヘッジ比率をこの区間に収めることで次時刻のヘッジ比率を得る。

$$\begin{aligned} (b_l, b_u) &= \text{NN}(I_{t_i}) \\ \delta_{t_{i+1}} &= \text{clamp}(\delta_{t_i}, b_l, b_u) \end{aligned} \quad (5)$$

模式図を図 2 に示す。ここで、**clamp** は次の関数である。

$$\text{clamp}(\delta_{t_i}, b_l, b_u) \equiv \begin{cases} b_l & \text{if } \delta_{t_i} < b_l \\ \delta_{t_i} & \text{if } b_l \leq \delta_{t_i} \leq b_u \\ b_u & \text{if } \delta_{t_i} > b_u \end{cases} \quad (6)$$

この戦略では、ヘッジ比率は常に区間 $[b_l, b_u]$ 内に保たれる一方、この区間内では取引を行わない。

No-Transaction Band ネットワークには二つの利点がある。第一に、ニューラルネットワーク (5) の入力値が現在のヘッジ比率を含まないため、前述の訓練の難しさを回避できる。第二に、この構造はヘッジ最適化のための適切な推論バイアスを組み込んでいる。なぜなら、No-Transaction Band を用いる戦略は、十分滑らかな効用関数のもとでペイオフが十分滑らかなオプションをヘッジする最適戦略だからである。最適性の証明は次に与えるが、次のような直感的な理解を述べる。仮に原資産価格が少々変動しても、次の時刻にもとの価格に戻りヘッジが不要になるかもしれないので、ヘッジ比率の許容幅内では取引を控え、取引コストを節約すべきである。従って、No-Transaction Band 戦略が取引コストの存在下で最適戦略となる。

4.2 No-Transaction Band 戦略の最適性

No-Transaction Band 戦略の効率性を正当化するため、[3] の設定（指数型効用関数の元でのヨーロッパン・オプション）における No-Transaction Band 戦略の最適性を、より広いクラスの効用関数とオプションに対して拡張する。

そのために、次のような損失関数を定義する。

$$\begin{aligned} \ell_{t_i}(-Z; \{S_{t_j}\}_{t_j \leq t_i}, \delta_{t_i}, C_{t_i}) \\ = \inf_{\{\Delta \delta_{t_j}\}_{t_j \geq t_i}} \mathbf{E}[-u(P(-Z, S, \delta)) | \{S_{t_j}\}_{t_j \leq t_i}, \delta_{t_i}, C_{t_i}] \end{aligned} \quad (7)$$

ここで、 u は十分に滑らかな効用関数とする。関数 (7) を最小化する戦略 δ は、損失関数 (7) を最小化することがわかる。この関数 (7) について、次の主張が成り立つ。証明は補遺 A に譲る。

Proposition 4.1 (No-Transaction Band 戦略の最適性). u を十分に滑らかな効用関数、 Z を十分に滑らかなペイオフをもつオプションとする¹。すると、損失関数 (7) は S_t に関して 2 階微分可能である。このとき、 Z の最適ヘッジは No-Transaction Band 戦略によって達成される。つまり、損失関数 (7) を最小化する δ^* は $\delta_{t_{i+1}}^* = \text{clamp}(\delta_{t_i}^*, b_l, b_u)$ によって与えられる。ここで、 b_l と b_u は $b_l \leq b_u$ を満たす $t, \{S_{t'}\}_{t' \leq t}$ の関数である。

5 数値実験

本章では、No-Transaction Band ネットワークが、通常の順伝播型ニューラルネットに比べ、より早くより優れたヘッジ戦略を与えることを、数値実験で実証する。

ヘッジ対象として、最も単純なオプションであるヨーロッパ・オプション (第 5.1 節) および、エキゾチックオプションの例であるルックバック・オプション (第 5.2 節) の二つを扱う。原資産価格は初期値 $S_0 = 1$ 、ボラティリティ $\sigma = 0.2$ の幾何ブラウン運動 (1) に従う。オプションの満期は $T = 30/365$ 、時間ステップは $t = 0, 1/365, \dots, 30/365$ である。取引コスト率 $c = 0$ および $c = e^k$ ($k = -10.00, -9.75, \dots, -5.00$) のそれぞれに対してヘッジ戦略を求め、達成される期待効用を計算した。

ヘッジ方法として、次の三つを比較する。

- (i) **No-Transaction Band ネットワーク** (図 2) : ニューラルネットは 32 ノードの隠れ層 4 層からなり、活性化関数は ReLU 関数である。ニューラルネットの二つの出力それぞれに LeakyReLU 関数を適用し、ブラック・ショールズのデルタに足す・引くことで No-Transaction Band の上限・下限を得る。インプット I_{t_i} はオプションにより異なるので、のちに説明する。
- (ii) **順伝播型ニューラルネット** (図 1) : ニューラルネットは (i) と同様に 4×32 ノードからなり、活性化関数は ReLU 関数である。ニューラルネットの出力に tanh 関数を適用し、ブラック・ショールズのデルタに足して次時刻のヘッジ比率を得る。
- (iii) **取引コストが小さい極限での最適解** : 取引コストが小さい極限で最適ヘッジとなる No-Transaction Band の値が [9] で解析的に求められている。この値を用いた No-Transaction Band 戦略である。

ニューラルネットの訓練は次のように行う。各モンテカルロシミュレーションで、原資産価格の経路を 50,000 通り生成し、各経路についてヘッジレポートフォリオ (2)

を計算する。そして、損失関数 (3) を最小化すべく Adam アルゴリズム [11] でパラメタを最適化する。各オプション・取引コスト率・ヘッジ方法のそれぞれについて、1,000 回のシミュレーションを行う。実験にあたって作成した PyTorch に基づくライブラリは、オープンソースで公開することを予定している²。

5.1 ヨーロピアン・オプション

まず、ヨーロッパ・オプションを扱う。行使価格は $K = 1$ とし、特徴量として対数 Moneyness、満期までの時間、ボラティリティ $I_{t_i} = \{\log(S_{t_i}/K), T - t_i, \sigma\}$ を用いる。

図 3 に示すように、No-Transaction Band ネットワークは、幅広いコスト率にわたり最も高い期待効用を達成する。また、コストがある閾値を超えると期待効用が下げ止まる。これは、No-Transaction Band ネットワークが取引コストが高すぎる場合にヘッジしないことを学習するためであり、この利点は他の手法ではみられない。また、取引コストが小さい領域における No-Transaction Band ネットワークの期待効用は (iii) の解に漸近しており、この領域で最適に近いヘッジ戦略を実現していることが推察できる。

さらに、図 4 の学習曲線が示すように、No-Transaction Band ネットワークの訓練は高速である。通常の順伝播型ニューラルネットが 1,000 回のシミュレーションを経ても収束しない一方、No-Transaction Band ネットワークは 100 回未満のシミュレーションで収束する。

5.2 ルックバック・オプション

次に、ルックバック・オプションを扱う。行使価格は $K = 1.03$ とし、特徴量として対数 Moneyness、満期までの時間、ボラティリティおよび対数 Moneyness の累積最大値 $M_{t_i} \equiv \max\{\{\log(S_{t_j}/K)\}_{t_j \leq t_i}\}$ を用いる。

図 5 に示すように、No-Transaction Band ネットワークは最も高い期待効用を達成する。また、前節 5.1 と同様、No-Transaction Band ネットワークはコストが高すぎる場合はヘッジしないことを学習し、期待効用の減少を抑える。

さらに、図 6 が示すように、No-Transaction-Band ネットワークの訓練は高速である。ルックバック・オプションは、経路依存性がありより多くの特徴量を要する点でヨーロッパ・オプションよりも複雑であるにもかかわらず、同程度の時間で収束することは特筆に値する。このことは、我々の手法が、より多くの特徴量を要する複雑なエキゾチックオプションにも適用できることを示唆する。

¹具体的には、有限のデルタとガンマを持つことを仮定する。

²GitHub: [pfnet-research/pfhedge](https://github.com/pfnet-research/pfhedge)

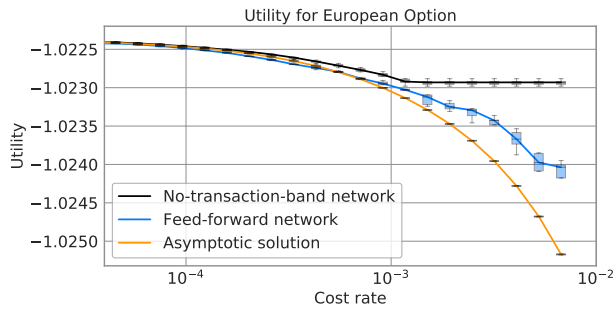


図 3: ユーロピアン・オプションに対する期待効用。20 回のシミュレーションで得た期待効用の平均値を示す。誤差は異なる乱数シードの五回の実験から算出した。

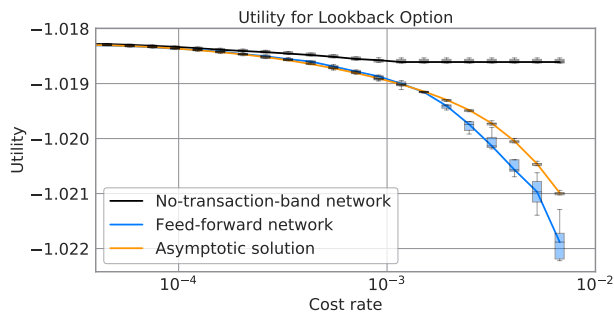


図 5: ルックバック・オプションに対する期待効用。20 回のシミュレーションで得た期待効用の平均値を示す。誤差は異なる乱数シードの五回の実験から算出した。

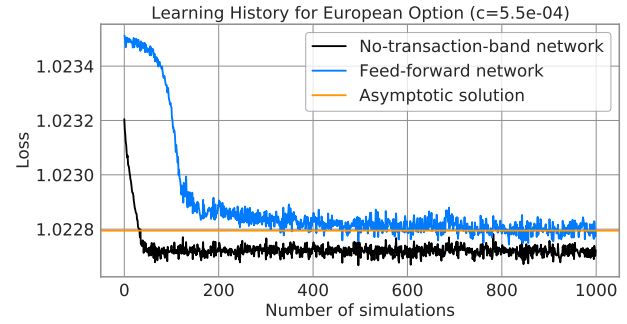


図 4: ユーロピアン・オプションに対する学習曲線。横軸はシミュレーション回数、縦軸は 20 回のシミュレーションで得た損失関数 (3) の平均値を表す。

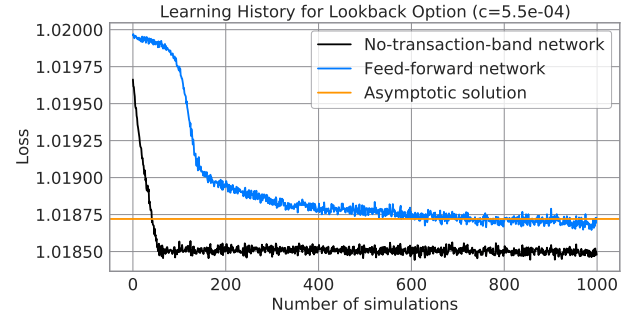


図 6: ルックバック・オプションに対する学習曲線。横軸はシミュレーション回数、縦軸は 20 回のシミュレーションで得た損失関数 (3) の平均値を表す。

6 おわりに

本研究では、シンプルなニューラルネットワークによる効率的なヘッジ戦略のための枠組みである、No-Transaction Band ネットワークを提案した。このネットワークは、インプットに現在までのヘッジ行動を必要とせず、No-Transaction Band を出力するという特徴がある。また、この戦略がより一般の効用関数・エキゾチックを含む幅広いオプションに対しても最適ヘッジであることを証明した。数値実験の結果、ユーロピアン・オプションおよびルックバック・オプションに対し、我々のネットワークが通常の順伝播型ネットワークよりも高速により優れたヘッジ戦略を実現することを示した。

本研究では資産価格ダイナミクスは幾何ブラウン運動としたが、実用のため、ボラティリティ・スマイルや期間構造、確率ボラティリティなどへの拡張が必要である。また、価格が確率ボラティリティに従う資産のオプションのヘッジのため、複数のヘッジ手段を組み込む必要がある。これらは今後の重要な研究課題としたい。

参考文献

- [1] Hans Buehler, Lukas Gonon, Josef Teichmann, and Ben Wood. Deep hedging. *Quantitative Finance*, Vol. 19, No. 8, pp. 1271–1291, 2019.
- [2] Fischer Black and Myron Scholes. The pricing of options and corporate liabilities. *Journal of political economy*, Vol. 81, No. 3, pp. 637–654, 1973.
- [3] Mark HA Davis, Vassilios G Panas, and Thaleia Zariphopoulou. European option pricing with transaction costs. *SIAM Journal on Control and Optimization*, Vol. 31, No. 2, pp. 470–493, 1993.
- [4] Hayne E Leland. Option pricing and replication with transactions costs. *The journal of finance*, Vol. 40, No. 5, pp. 1283–1301, 1985.
- [5] Erik R Grannan and Glen H Swindle. Minimizing transaction costs of option hedging strategies. *Mathematical finance*, Vol. 6, No. 4, pp. 341–364, 1996.

- [6] Stewart Hodges and Anthony Neuberger. Optimal replication of contingent claims under transaction costs. *Review Futures Market*, Vol. 8, pp. 222–239, 1989.
- [7] Vicky Henderson and David Hobson. Utility indifference pricing-an overview. *Volume on Indifference Pricing*, 2004.
- [8] Michael Monoyios. Option pricing with transaction costs using a markov chain approximation. *Journal of Economic Dynamics and Control*, Vol. 28, No. 5, pp. 889–913, 2004.
- [9] A Elizabeth Whalley and Paul Wilmott. An asymptotic analysis of an optimal hedging model for option pricing with transaction costs. *Mathematical Finance*, Vol. 7, No. 3, pp. 307–324, 1997.
- [10] Kurt Hornik. Approximation capabilities of multilayer feedforward networks. *Neural networks*, Vol. 4, No. 2, pp. 251–257, 1991.
- [11] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.

A 最適性の証明

Proposition 4.1 を証明する。時点 t_i での原資産の購入単位・売却単位をそれぞれ $y_{t_i}^{\text{buy}} \Delta t_i \geq 0$, $y_{t_i}^{\text{sell}} \Delta t_i \geq 0$ で表す。すると、最小化問題 (7) の Δt_i に関する Hamilton-Jacobi-Bellman 方程式は次のように書ける。

$$\inf_{y_t^{\text{buy}}, y_t^{\text{sell}}} \left(AS_t y_t^{\text{buy}} + BS_t y_t^{\text{sell}} \right) + \frac{\partial \ell_t}{\partial t} + \frac{1}{2} \sigma^2 S_t^2 \frac{\partial^2 \ell_t}{\partial S_t^2} = 0, \quad (8)$$

$$A = \frac{\partial \ell_t}{\partial \delta_t} + c \frac{\partial \ell_t}{\partial C_t}, \quad B = -\frac{\partial \ell_t}{\partial \delta_t} + c \frac{\partial \ell_t}{\partial C_t}. \quad (9)$$

ここで $t = t_i$ であり、 Δt の高次の項を無視した。最適なヘッジ比率 δ^* は式 (9) の最小化問題から得られる。いま、 $y_t^{\text{buy}}, y_t^{\text{sell}} \leq k < \infty$ と制限する。[3] と同様の議論から、 $A + B > 0$ であるため、係数 A, B の符号によって次の三つの場合に分けられる。

1. $A < 0, B > 0$: このとき、最小は $y_t^{\text{buy}} = k, y_t^{\text{sell}} = 0$ で得られる。すなわち、可能な限り原資産を購入するのが最適である。
2. $A > 0, B < 0$: このとき、最小は $y_t^{\text{buy}} = 0, y_t^{\text{sell}} = k$ で得られる。すなわち、可能な限り原資産を売却するのが最適である。

3. $A \geq 0, B \geq 0$: このとき、最小は $y_t^{\text{buy}} = 0, y_t^{\text{sell}} = 0$ で得られる。すなわち、何も取引しない (No Transaction) のが最適である。

ここで、 A と B がそれぞれ δ_t に関して単調増加、単調減少であると仮定する。これは損失関数が δ_t に関して凸であるため、小さな c に対しては正当化される。すると方程式 $\pm \partial \ell_t / \partial \delta + c \partial \ell_t / \partial C = 0$ はそれぞれ一つの解を持ち、それを $\delta_t = b_l, b_u, b_l \leq b_u$ とかく。極限 $k \rightarrow \infty$ を考えることにより、 $\delta_{t+\Delta t} = \delta_t + y_t^{\text{buy}} \Delta t - y_t^{\text{sell}} \Delta t = \text{clamp}(\delta_t, b_l, b_u)$ 、つまり $\delta_{t_{i+1}} = \text{clamp}(\delta_{t_i}, b_l, b_u)$ が成立し、No-Transaction Band が最適であることがわかる。