



Sony CSL

人工知能 (AI) によるファンド行動学習に
ついての委託研究業務
最終報告書 (要約版)

株式会社ソニーコンピュータサイエンス研究所

2020 年 9 月

Abstract

本誌は 2018 年 10 月に株式会社ソニーコンピュータサイエンス研究所が年金積立金管理運用独立法人（以下 GPIF）より受託した「人工知能（AI）によるファンド行動学習についての委託研究業務」に関する報告書の要約版である*1。

インプリケーション：

- 運用機関の選定で人工知能（以下 AI）がファンド間の違いの決め手を明確にすることができれば、個人の経験や能力によらない平準化された選定眼を持つことが可能となる。
- 運用機関の選定で AI が候補ファンドの運用行動の変化を検知。その変化の確認プロセスを重ねることで、AI の信頼性を高めることができる。

Case Study

運用機関の選定：「既存委託先 A と選定候補先 B の違いの決め手は何か？（国内株式）」

AI が対象ファンドを銘柄保有ウェイトに基づきマップ上に位置づけ、運用機関の選定に最適化された説明可能 AI（以下 XAI）によりその相違点を解釈した。1 つのサンプルとして既存委託先 A と選定候補先の違いの決め手を見たところ、限定的なサンプル期間においては選定候補先 B はある業種を強く選好することが判明。AI においても一般的な保有銘柄分析が示唆する内容を示していた。

運用機関の選定：「選定候補先 C の投資行動の変化（国内株式）」

月次レベルの粒度データにより選定候補先 C の運用行動を AI により検証した結果、ある時期に大きく変化したことを検知し、当該時期に担当 PM 変更があったことを確認した。この AI を活用することにより、PM 変更前後の投資行動の変化を AI においても確認することができた。

既存委託先の管理：「なぜ既存委託先 D はコロナショック相場で運用行動を変えたか？（外国株式）」

2020 年 2-4 月のいわゆるコロナショック相場で、既存委託先 D が運用行動を変化させていたことを AI が検知した。相場急変を踏まえた緊急リスクコントロール的対応等を想定したが、確認の結果、過去より継続していた運用のリスク管理モードの切替という変化であったり、重大事案ではなかった。AI によって重大事案のみを検知することは現段階では困難ではあるが、既存委託先の行動変化の検知を繰り返し検知事例を継続的に蓄積すれば、重大事案の捕捉も可能となるかもしれ

*1 本研究メンバー：北野宏明（代表取締役社長、所長）、田尻貴夫（プロジェクトリーダー）、佐々木貴宏（リサーチャー）、橋戸公則（非常勤リサーチャー）、森田匠（同）、村上朝成（同）、石川貴大（同）およびリサーチアシスタント

ない。

今後の研究

GPIF は第 4 期中期計画で業務運営の高度化・効率化のため AI、RPA 等の積極的活用を掲げている。本調査研究当初からの課題設定としてきた運用機関選定での定性項目評価に対する「客観性に欠ける」等の解決を目指す場合、運用機関選定での AI 活用は補助的なツールになりうる。この場合、「運用機関選定での事例蓄積」・「多くの選定候補先を評価できるよう試験活用範囲の拡大」・「選定における AI の位置づけと目標の整理」をより深掘りして実用に耐えうるモデルにしていくことが重要である。

GPIF は運用機関選定や既存委託先管理で AI を活用するという独創的視点で調査研究を行い、本取組みは EQ Derivatives 誌および Asianinvestor 誌より受賞し、多くの関係者やメディアより注目されてきた。基礎研究から試験活用まで行う重要性を理解し、このような機会を与えてくれた GPIF には改めて多謝申し上げる。

本論

1 調査研究の進展

これまで次の3つのフェーズで調査研究を進展させてきた：

- 第一フェーズ「人工知能（AI）が運用に与える影響についての調査研究業務 [1, 2]」（2017年-2018年）

AI技術の1つである深層学習 [3] を活用して「運用スタイル分析器（Style Detector Array、以下 SDA）」の原理試作を行った。これは、日本株式 100 銘柄のスモールスケールで、何かを検知する仕組みを探索したものである。何かについては GPIF および金融関係者に馴染みがあるバリューやモメンタムといった運用スタイルを検知対象とした。

- 第二フェーズ前半「人工知能（AI）によるファンド行動学習についての委託研究業務」（2018年10月-2019年10月）

第一フェーズの原理試作結果を受け、SDA のコア技術をそのままに、従来定量評価が難しいとされた投資方針や運用プロセスといった運用会社の「クセ」という定性情報を検知可能にする“Resembling”の調査研究および GPIF 現場での試験活用を行った。現場での試験活用の結果、従来の運用スタイルといった指標では捕捉し得ない運用行動の変化を検知可能であることを確認した。さらに、Resembling を応用して GPIF のマネジャーストラクチャーの運用行動の類似性に基づいた分散度合いの評価および運用機関の選定における既存委託先と候補ファンドの類似性評価の適応可能性を示唆した。本内容は 2019 年 12 月に中間報告書 [4, 5] としてリリースされた。

- 第二フェーズ後半 本誌内容（2020 年 1 月-6 月）

第二フェーズ前半の調査研究結果である Resembling を次の 2 つに発展した：

- “Mutual Resemblance”：各運用機関の類似性を評価。主に運用機関の選定で使用。
- “Self Resemblance”：自身の「クセ」変化を検知。運用機関の選定および管理で使用。

本誌（要約版）では次の Case study を紹介する：

- 運用機関の選定
 - 「既存委託先 A と選定候補先 B の違いの決め手は何か？（国内株式）」（Mutual Resemblance の試験活用）
 - 「選定候補先 C の PM 変更（国内株式）」（Self Resemblance の試験活用）
- 既存委託先の管理
 - 「なぜ既存委託先 D はコロナショック相場で運用行動を変えたか？（外国株式）」（同上）

なお、本誌は調査研究の一部について言及する最終報告書フルレポートの要約版である。

2 運用機関の選定

Case study：「既存委託先 A と選定候補先 B の違いの決め手は何か？（国内株式）」
（Mutual Resemblance の試験活用）

運用機関の選定では、GPIF 担当官が運用機関に関する定量的な実績を勘案した多岐にわたる定性項目（投資方針、運用プロセス、組織・人材、等）を分析の上その結果を総合判断する [6]。扱う情報は複雑で分析結果の統合に多くの時間を要するため、統合された情報を直感的に捉えやすい形に表現できれば負荷軽減が見込める。また、定性項目による評価は外部より「客観性に欠けるのではないか」等の課題について GPIF 経営委員会で議論されている [7]。これらの問題や懸念を一定程度解決しうることを目指し AI 活用を調査研究した。図 1 は国内株式アクティブ運用の既存委託先と選定候補先 7 社をマップ上に銘柄保有ウエイトデータを入力データとして AI が出力したものである。全体的に次のことが言える：

- 既存委託先（灰色の点群）は運用スタイルごとにグループ分けされている（赤・青・緑の点線で囲まれた領域）。
- 総じて選定候補先 7 社（桃色の点群、ただし選定候補先 B は橙色）は、赤点線で囲われた領域を中心に分布しており、この領域内のファンドと同じ運用スタイルをとっている。
- その中で、選定候補先の内 4 社は赤点線で示されたスタイルに近接しており、残り 3 社は相対的に青や緑の領域に近いという違いがある。

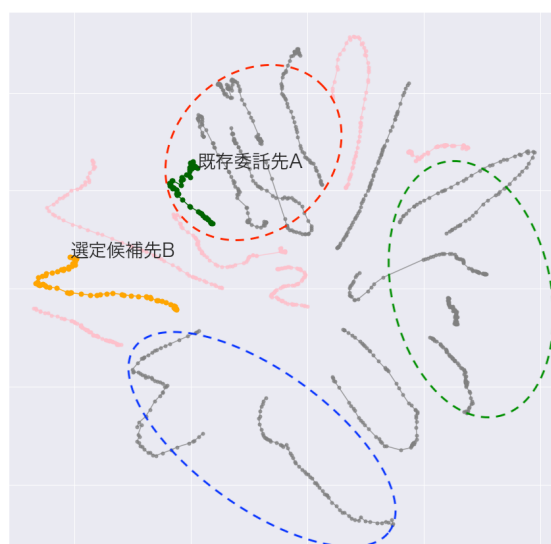


図 1: 既存委託先と選定候補先の類似度マップ

既存委託先 A: 緑色、その他既存委託先: 灰色、選定候補先 B: 橙色、その他選定候補先: 桃色

次に、1つのサンプルとして既存委託先Aと選定候補先Bに焦点を当て、AIは何を違いの決め手と見ているかを解明する。図2、図3、図4は銘柄が各運用機関のマップ上の位置配列にどの程度影響を与えているかを可視化したものである。探索の結果、銘柄①と銘柄②が選定候補先Bを特徴づけていることがわかる。図5は銘柄①と銘柄②が属する業種③で同様に可視化したものである。その結果、選定候補先Bは例外的に業種③選好が強いことがわかる。

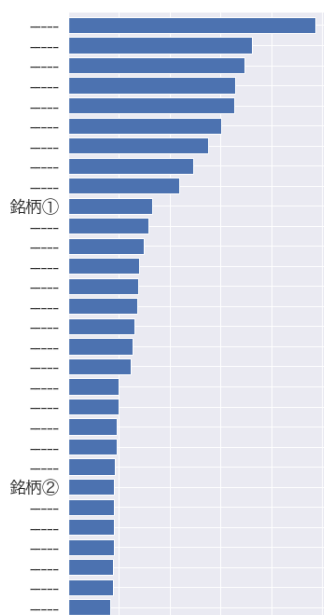


図2: 個別銘柄の重要度



図3: 銘柄①に対する選好度（青色の領域は強い選好、橙色の領域は弱い選好をあらわす）



図4: 銘柄②に対する選好度（青:強/橙:弱）

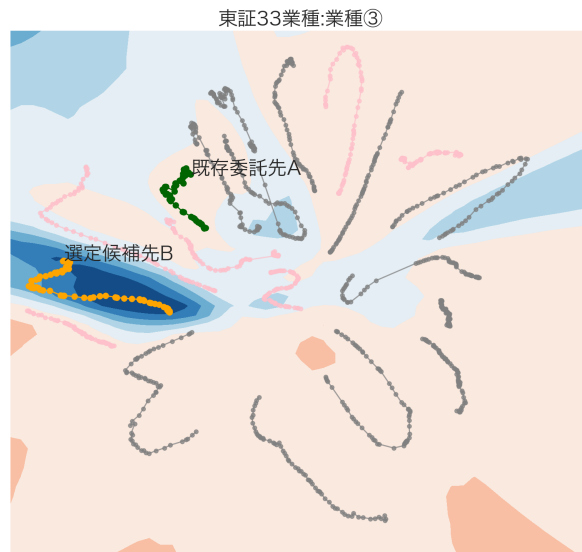


図 5: 業種③に対する選好度（青:強/橙:弱）

なお、図 6 および図 7 より既存委託先 A は銘柄④と業種⑤を選好し、選定候補先 B は弱いことがわかるが、銘柄①、銘柄②および業種③ほど決め手にはならない。その他分析結果含め、AI は既存委託先 A と選定候補先 B の違いの決め手を業種③と見ていることがわかった。この結果は、業種ウェイトの傾向から示唆される内容と一致している。次のステップは、選定候補先 B はなぜ当該業種に特に足下着目しているか分析の上、業種選定理由と見通しおよび Valuation 等を確認することが次のアクションとなる。一般的に、選定において選定候補先の何に焦点を当てるかは個人の経験や能力に依る傾向がある。この AI を活用することにより、一般的な分析ツールによる結論を、平易に可視化した形で理解することができる。

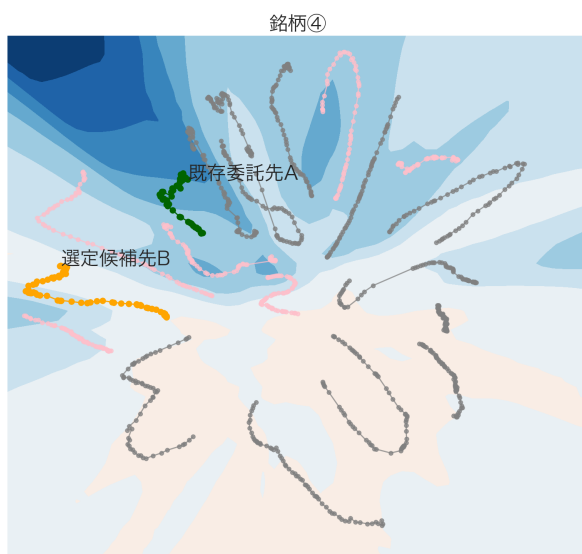


図 6: 銘柄④に対する選好度（青:強/橙:弱）

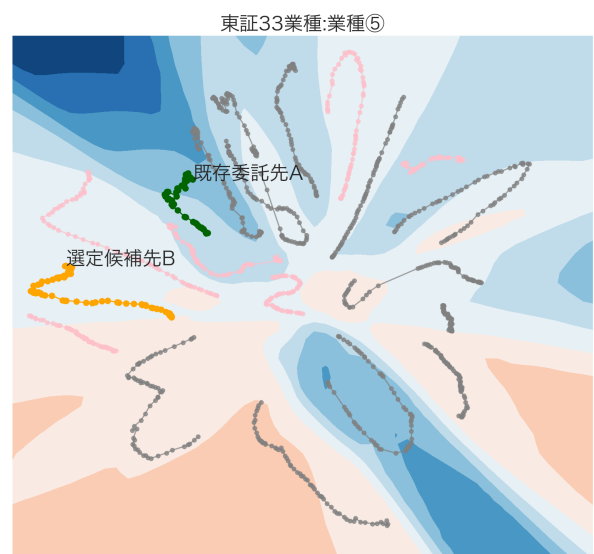


図 7: 業種⑤に対する選好度（青:強/橙:弱）

方法論

運用機関の選定にあたり、運用機関間の運用行動の類似性を評価する“Mutual Resemblance”を開発した。その Mutual Resemblance のモデルには変分オートエンコーダー (Variational Auto Encoder、以下 VAE [8, 9, 10, 11, 12]) という高次元の情報を低次元に圧縮する Encoder と、圧縮された情報を復元する Decoder から構成される深層学習技術が用いられている。合わせて、XAI [13, 14, 15, 16, 17, 18, 19, 20, 21] の一手法である SHapley Additive exPlanation Value (以下 SHAP) を組み合わせ、AI が注目している特徴を説明可能なものとした (図 8 参照)。

VAE を用いた理由は次のとおりである：

- Encoder と Decoder でデータを統合・可視化するために必要な機能を有する。
- 敵対的生成ネットワーク (GAN) [22] 等の最新のアルゴリズムと比較するとシンプルなモデルだが、それゆえにモデルの挙動が把握しやすく、学習に必要なコストも低廉である。
- ネットワーク構造を複雑にすることで広範な問題に適用できることから、将来的なデータ量拡大に対応する拡張性も備えている。
- 近年問題になっている AI 説明性について、VAE 機能に SHAP を組み合わせることで一定解決が可能である。
- 今回の分析では月次レベルの比較的粒度の粗いデータを入力データとしているが、VAE は安定した結果を出すことができる。

VAE に GPIF が所有する銘柄保有ウエイトのデータを入力すると、Encoder による圧縮の過程で情報が統合される。この圧縮された情報を取り出し平面上に配置することで、運用機関間の大まかな関係を直感的に理解するためのマップが作られる。さらに、このマップ上に Decoder で復元した情報を Projector を通して銘柄、業種、ファクター等集約し投影することで、マップの領域が表す特徴を可視化することができる。言い換えると、Projector はマップを様々な視点で見るためのレンズに相当し、業種、銘柄、ファクターなどのレンズを切り替えながらマップを見ることで、各運用機関のより詳細な情報を直感的に捉えることができる。レンズの種類は多岐にわたるが、SHAP が示す重要度を参考にすることで、見るべきレンズを絞り込むことができ、AI が注目する特徴を効率的に把握できる。このように、画像認識の場合は SHAP のみで解釈可能であるが、運用機関の数値データの場合は数値や個別銘柄に基づく解釈には限界があり、SHAP に加えて VAE の出力も活用するという工夫を施した。

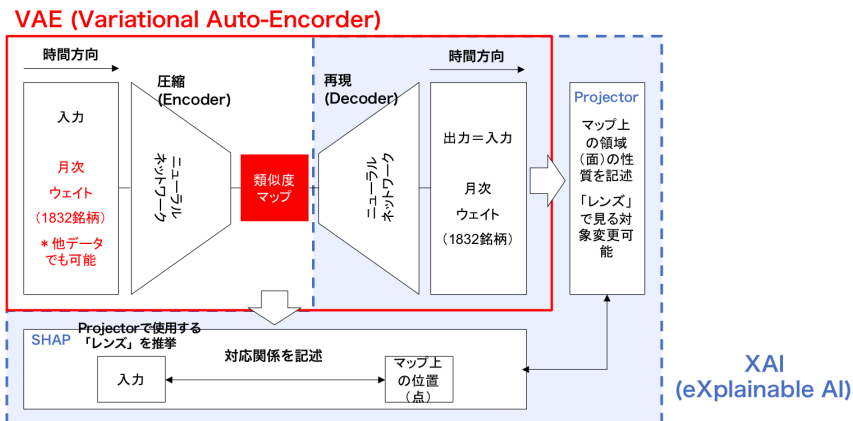


図 8: VAE と SHAP を用いた運用機関分析ツールの概念図

Case Study : 「選定候補先 C の PM 変更（国内株式）」（Self Resemblance の試験活用）

図 9 は選定候補先 C の運用行動が 2016 年から 2017 年にかけて変化したことを検知したものである^{*2}。確認の結果、当該期間に PM の変更があった。AI により、PM 変更による運用スタイルの変化の大きさ、トランジションの期間などを理解することができたことが分かった。選定の過程において、その変化の大きさ・期間についてより深掘りした質問をすることで、選定候補先の組織・人材の活用に関する考え方の 1 サンプルを理解することが可能となろう。

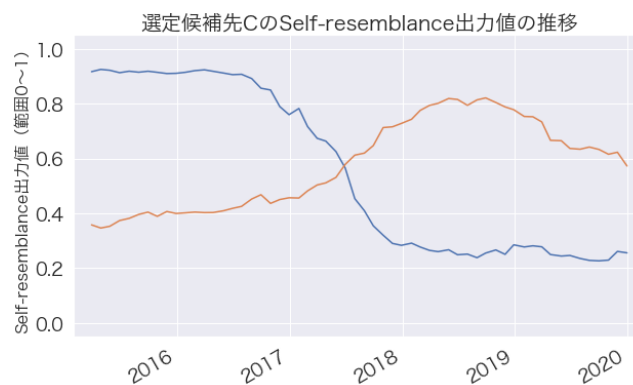


図 9: 選定候補先 C の Self Resemblance の推移

青色線：学習期間 2015 年 1 月～2017 年 4 月、橙色線：学習期間 2018 年 1 月～2019 年 6 月

^{*2} 折れ線が高い位置にあれば「らしさ」の合致度が高いことを示し、これが大きく変化した場合は「らしさ」からの乖離を示唆する。

方法論

運用会社自身の「クセ」の変化を評価する Self Resemblance を用いて検知を行った。この Self Resemblance は、2019 年 12 月リリースした「人工知能（AI）によるファンド行動学習についての委託研究業務中間報告書」で論じた Resembler と同一のものである（図 10 参照）。方法論の主なポイントは次のとおりである：

- 月次という比較的粒度の粗いデータを使用
- 時間軸の訓練データにて出力
- 分析対象とする全既存委託先と選定候補先合わせた計 20 ファンドに関する GPIF 所有の売買行動データおよび保有銘柄データ、ならびに外部データとしてファクトセットおよび ICE Data Services から取得したマーケットデータを深層ニューラルネットワークで学習し、ファンドらしさ（クセ）を抽出

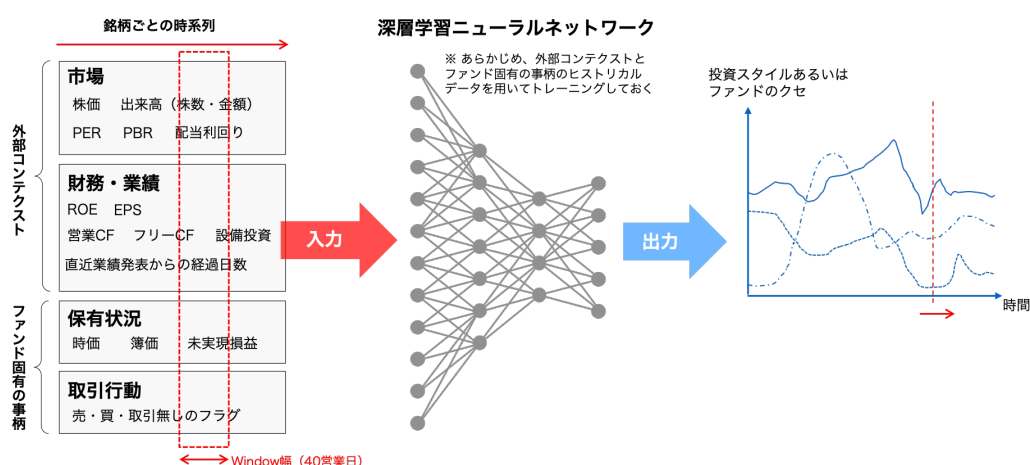


図 10: Self Resemblance の概念図

3 既存委託先の管理

Case study: 「なぜ既存委託先Dはコロナショック相場で運用行動を変えたか? (外国株式)」 (Self Resemblance の試験活用)

2020年2月-4月のいわゆるコロナショック相場での GPIF 既存委託先の運用行動を確認した結果、多くの既存委託先は従来と変わらぬ一貫した運用を継続するも、一部既存委託先で運用行動の変化を検知した。図 11 は外国株式アクティブ運用の既存委託先Dの Self Resemblance がコロナショックの時期に変化していることを検知したものである。分析結果、一部の運用行動が変化していた。当初は、コロナショック相場を踏まえた緊急的対応または予定していた運用方針の変更を予定より早期に開始したといった原因を想定していたが、既存委託先に確認した結果、過去より継続していたリスク管理モードを切り替えた、ということであった。リスク管理モード切り替えは運用ガイドライン内の変更で重大事案には該当せず、かつコロナショックが運用方針に多大な影響を与えていないことが確認できた。AI がコロナショックをどう解釈したかはより長期のデータを踏まえて分析を行う必要がある。

過去、稀ではあるが、GPIF において既存委託先の重大事案は起きており、その都度対策を講じている。一方、GPIF は少数のリソースで数多くの業務をこなしており、森羅万象見渡すのは限界がある。AI によって重大事案のみを検知することは現段階では困難ではあるが、既存委託先の行動変化の検知を繰り返し検知事例を継続的に蓄積すれば、重大事案の捕捉も可能となるかもしれない。

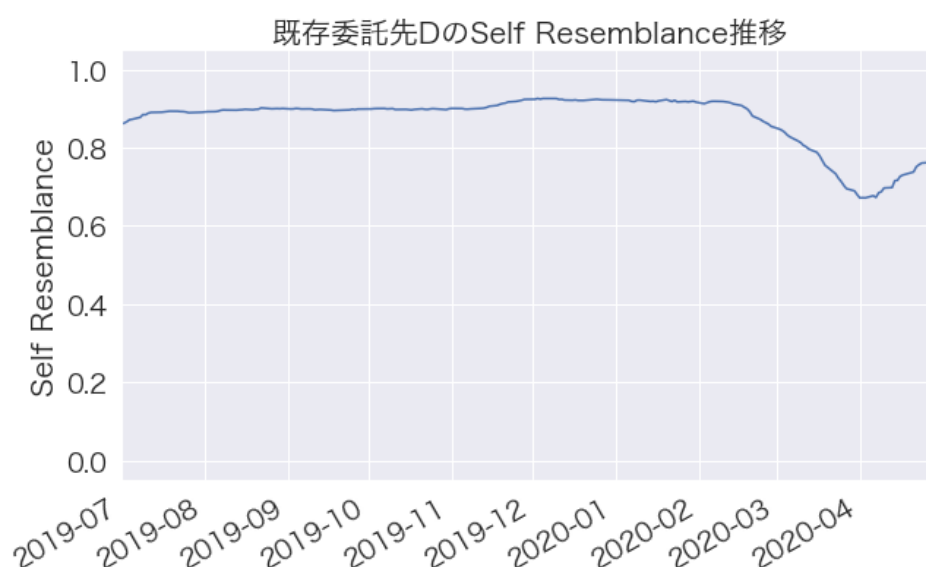


図 11: 委託先Dの Self Resemblance

4 今後の研究

GPIF は第 4 期中期計画の効率的な業務運営体制の確立の中で「業務運営の高度化・効率化のために、AI、RPA 等の先端技術を積極的に活用する」ことを掲げ、AI 活用を計画に定めている。今後も研究を継続するとした場合は、残された課題の中で次の 3 点は深掘りして調査を進める必要があると考える：

- AI 活用には一定年数の事例蓄積が必要。一方、本誌で論じた運用機関の選定での試験活用は約半年で一事例を扱ったのみで、高度化に向け事例蓄積を継続する必要がある。
- 本誌内容は選定プロセス後半での試験活用に関するものであったが、GPIF 経営委員会で議論された「(候補)ユニバース全体に対していいファンドを選べ」[23] られるよう第一次審査でより多くの選定候補先を評価できるように試験活用の範囲を広げる必要がある。
- 選定プロセスのどこに AI を位置づけ、何を最終目標とするか整理する必要がある。なお、選定プロセスにおける AI の位置づけと目標次第ではあるが、GPIF 内に蓄積されている運用機関の選定に関するノウハウも AI に入力し、人間の過去の判断や知見も組合わせたハイブリッド出力とすることにより、選定能力の均一化（属人性をなくした汎用化）を狙うという将来展望もありうると考えている。

公的アセットオーナーで AI 活用を試行しているのは一部にとどまっている。その中でも GPIF は運用機関の選定や既存委託先の管理で AI を活用するという視点の独創性が評価され、アセットオーナーおよびヘッジファンド向け金融専門誌である EQDerivatives 誌より 2 年連続受賞し、アジア地域の機関投資家向け金融専門誌である Asianinvestor 誌からも受賞、かつ多くのアセットオーナー、運用会社や海外メディアより注目された。今回、AI 活用において「何が AI の予測に寄与しているか分からない」というブラックボックス問題を解決するため、アセットオーナーのデータ特性を踏まえた XAI を搭載するなど、技術面でさらに進化を遂げた。

本調査研究は、基礎研究から試験活用まで行ってきた。基礎研究の活用では様々な課題に直面し、長い時間を要することが多い。このような基礎研究の活用の重要性を理解して時間的および資金的投資ができるのは、長期的観点で取り組むことができる公的アセットオーナーである GPIF ならではのことである。当社にこの取り組みの機会を与えてくださった GPIF には改めて多謝申し上げる。

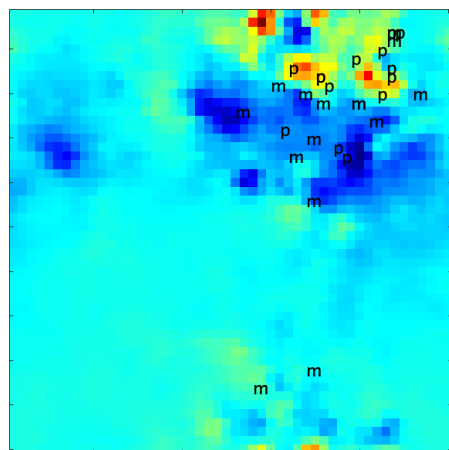
運用機関の行動予測（実験的取り組み）

様々な市場環境に対してどの様にファンドが行動するかを予測する“Replicator”に関して実験的取り組みを行なっている。次の様な利用シーンを想定している：

- リスクシナリオ分析の精緻化：
 - － 現ポートフォリオを前提に市場リスク発生時の損失を推定する一般的なシナリオ分析に対し、Replicator は市場リスクに対するファンドの反応を考慮した上で損失を推定でき、より精緻なシナリオ分析が可能となる。
 - － 市場環境変化に対し全既存委託先が類似行動をとるのか、Replicator により事前にマネジャーストラクチャーが収斂する可能性を評価することが可能となる。
- 運用機関の選定：
 - － 既存委託先と比べて新規運用機関の情報は限定的である。Replicator が新規運用機関の投資行動を学習し、様々な市場環境での行動を事前把握できる。

かつてのファンドリターンベースのポートフォリオ複製は学習において価格変化ノイズの影響を受けやすく精度の問題があったが、本手法はファンドの行動を予測することに着目することで本問題を回避している。

Replicator のアルゴリズムに多様体学習法のひとつである自己組織化写像（SOM） [24, 25] を利用した。ファンド属性値として、GPIF 所有のアクティブウエイトおよびファクトセットならびに ICE Data Services から市場シナリオとして取得したファクターリターン [26, 27] およびファクター値を入力すると、各市場環境での行動予測のアクティブウエイトが出力される。下図は行動予測結果の一例である。赤・黄は Replicator によってファンドが保有を増加させると予測された銘柄群を、水色・青の領域は保有を減少させると予測された銘柄群を表す。一部ずれがあるものの、実際に保有を増やした銘柄（p）が赤色部分に、実際に保有を減らした銘柄（m）が青色部分にそれぞれ分布しており、相応の予測結果が得られている。一方、現時点では、ファンド種別により精度にばらつきがあり、モデル改良やデータ拡張を継続している。



Replicator による行動予測結果

Appendix：データ出力（一部を抜粋して掲載）

運用機関の選定

Case study：「既存委託先 A と選定候補先 B の違いの決め手は何か？（国内株式）」

Projector の出力

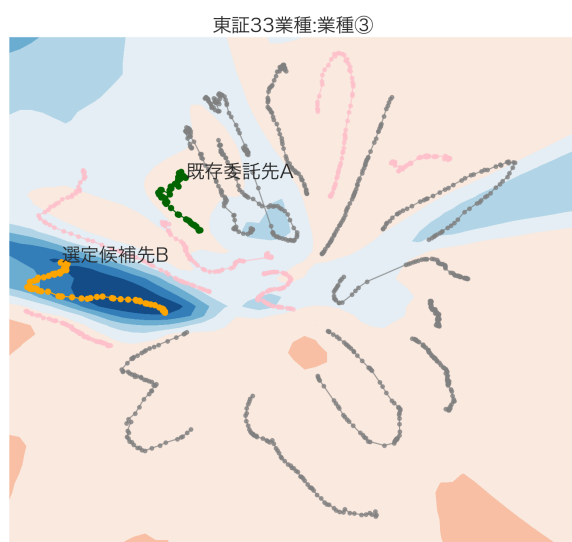


図 12: 業種③に対する選好度（青:強/橙:弱）

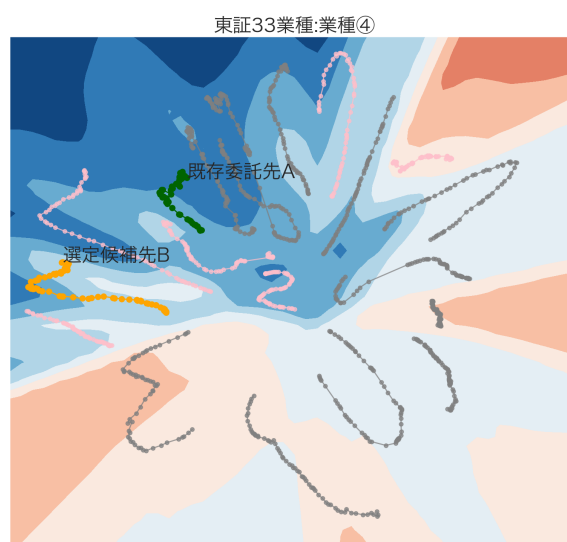


図 13: 業種④に対する選好度（青:強/橙:弱）

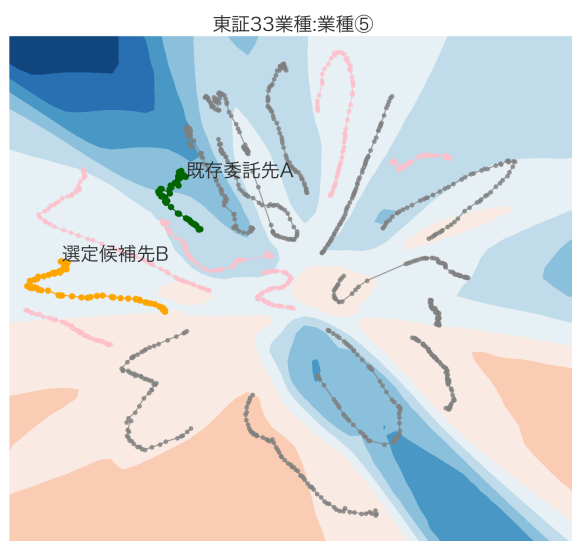


図 14: 業種⑤に対する選好度（青:強/橙:弱）

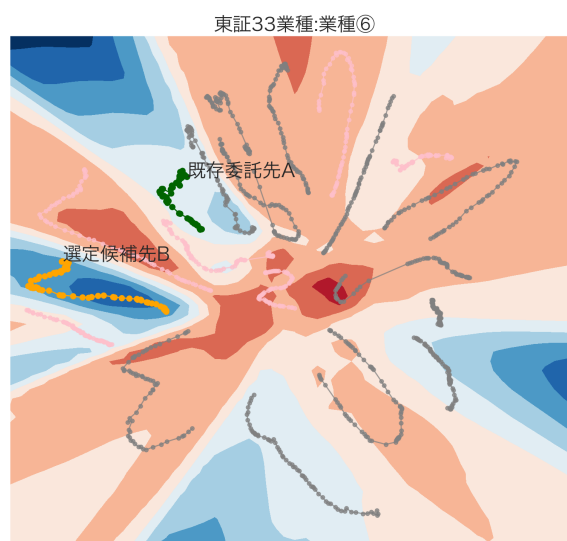


図 15: 業種⑥に対する選好度（青:強/橙:弱）

Case Study：「選定候補先 C の PM 変更（国内株式）」

月次粒度データでの Self Resemblance の出力

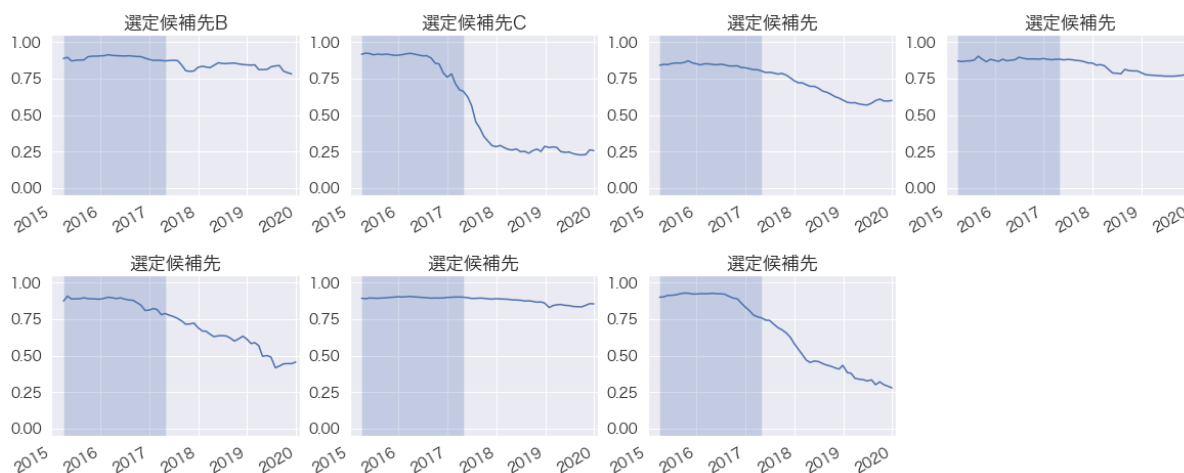


図 16: 選定候補先の Self resemblance の推移（学習期間 2015 年 1 月～2017 年 4 月）

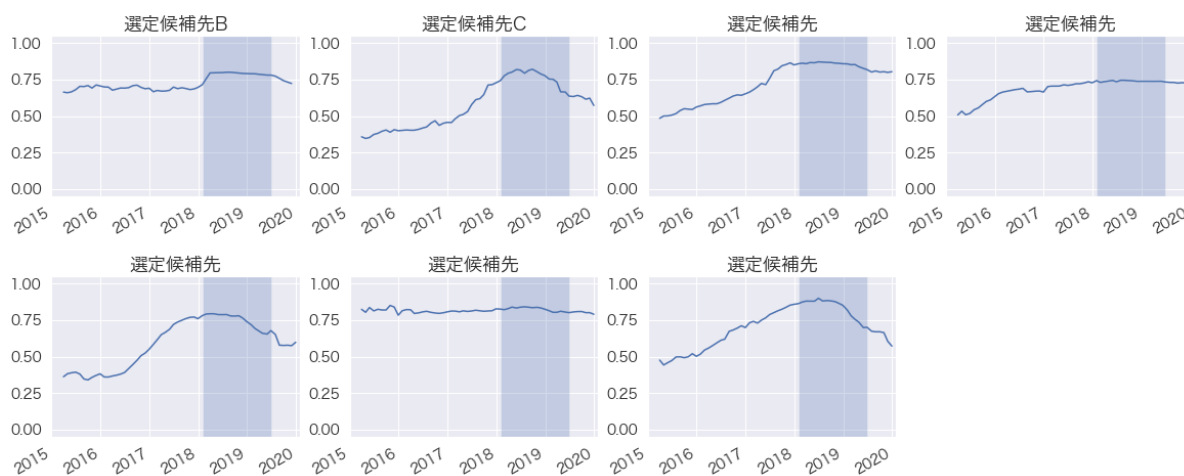


図 17: 選定候補先の Self resemblance の推移（学習期間 2018 年 1 月～2019 年 6 月）

Case study : 「なぜ既存委託先 D はコロナショック相場で運用行動を変えたか（外国株式）」
 コロナショック相場時の国内株式ファンドの Self Resemblance の出力

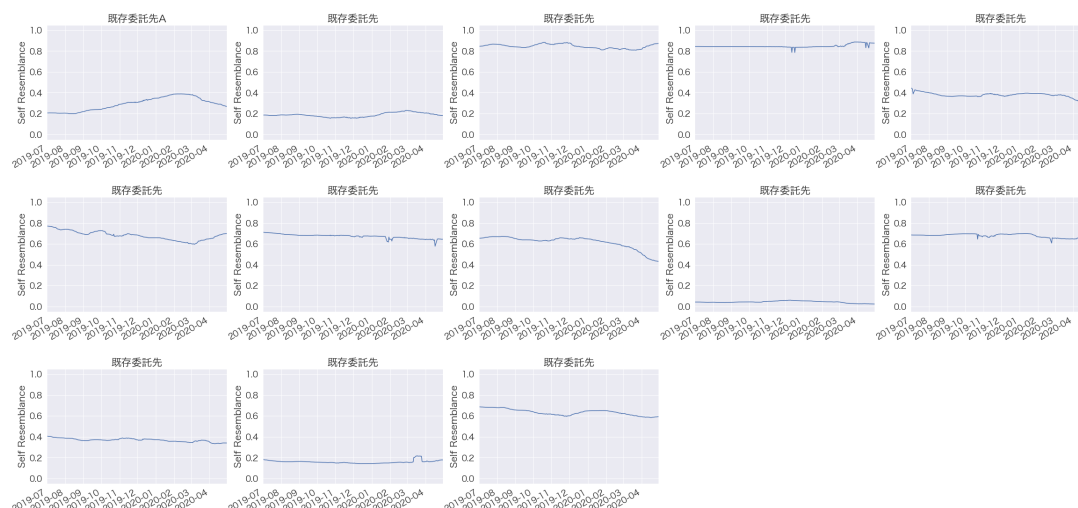


図 18: 既存委託先の Self resemblance の推移（学習期間 2015 年 1 月～2017 年 4 月）

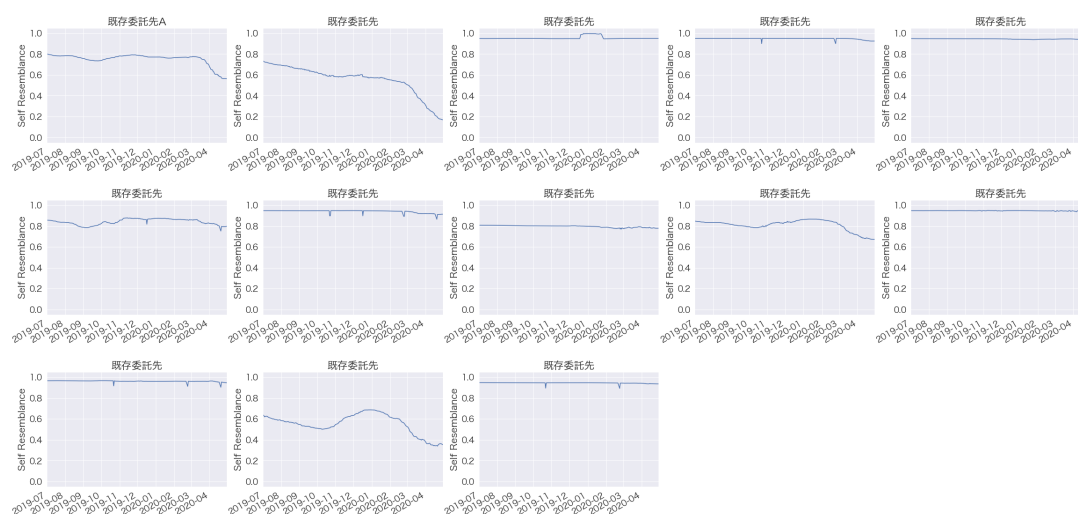


図 19: 既存委託先の Self resemblance の推移（学習期間 2018 年 1 月～2019 年 6 月）

コロナショック相場時の外国株式ファンドの Self Resemblance の出力

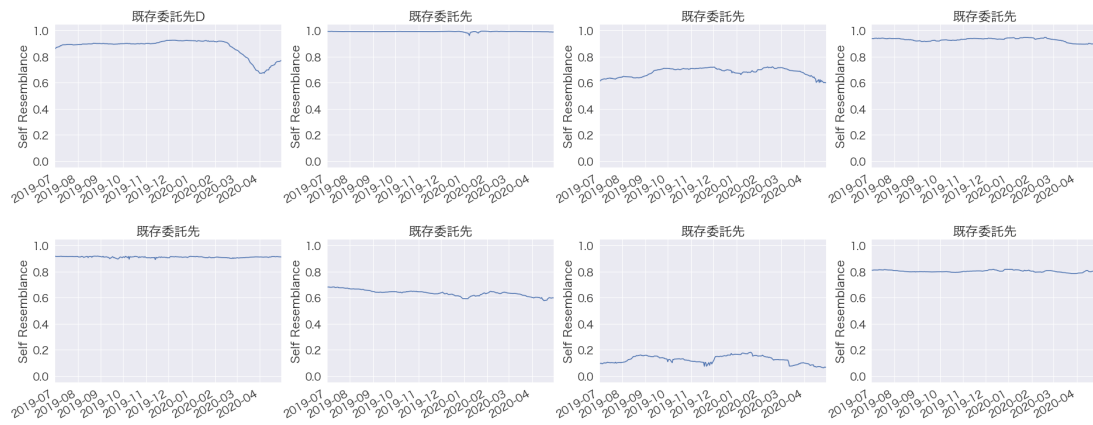


図 20: 既存委託先の Self resemblance の推移 (学習期間 2015 年 1 月～2017 年 4 月)

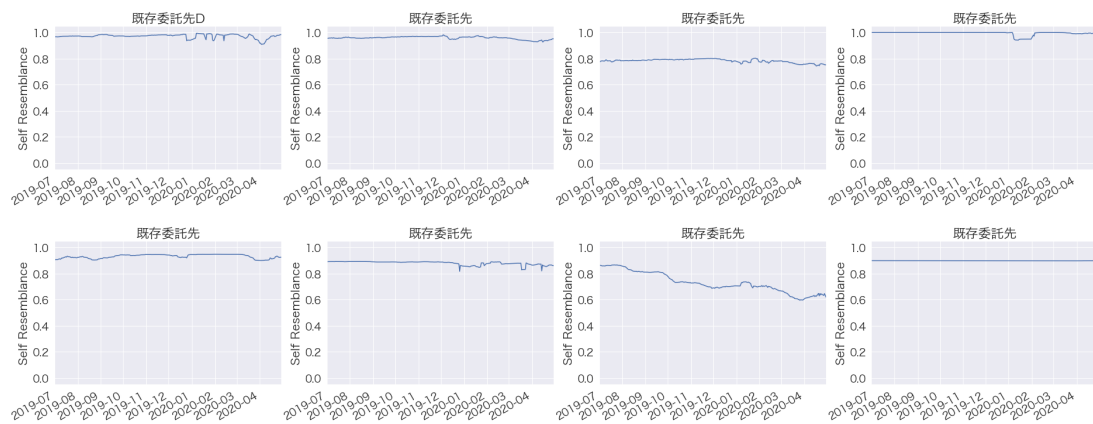


図 21: 既存委託先の Self resemblance の推移 (学習期間 2018 年 1 月～2019 年 6 月)

参考文献

- [1] Takahiro Sasaki, Hiroo Koizumi, Takao Tajiri, and Hiroaki Kitano. A study on the use of artificial intelligence within government pension investment fund' s investment management practices (summary report). Tokyo, Japan: Government Pension Investment Fund, Mar 2018.
- [2] 人工知能（A I）が運用に与える影響についての調査研究業務報告書（要約版）. 株式会社ソニーコンピュータサイエンス研究所, 2018 年 3 月.
- [3] G. E. Hinton and R. R. Salakhutdinov. Reducing the dimensionality of data with neural networks. *Science*, Vol. 313, No. 5786, pp. 504–507, 2006.
- [4] Takao Tajiri, Takahiro Sasaki, and Hiroaki Kitano. A step toward a cybernetic whale: An interim report on “a study on the use of artificial intelligence for learning characteristics of funds' behavior” . Tokyo, Japan: Government Pension Investment Fund, Dec 2019.
- [5] 人工知能（A I）によるファンド行動学習についての委託研究業務中間報告書. 株式会社ソニーコンピュータサイエンス研究所, 2019 年 12 月.
- [6] 年金積立金管理運用独立行政法人. 2019 年度業務概況書.
- [7] 年金積立金管理運用独立行政法人. 第 3 回経営委員会議事概要.
- [8] Diederik P Kingma and Max Welling. Auto-encoding variational bayes, 2013.
- [9] Diederik P. Kingma, Danilo J. Rezende, Shakir Mohamed, and Max Welling. Semi-supervised learning with deep generative models, 2014.
- [10] Carl Doersch. Tutorial on variational autoencoders, 2016.
- [11] Danilo Jimenez Rezende, Shakir Mohamed, and Daan Wierstra. Stochastic backpropagation and approximate inference in deep generative models. In Eric P. Xing and Tony Jebara, editors, *Proceedings of the 31st International Conference on Machine Learning*, Vol. 32 of *Proceedings of Machine Learning Research*, pp. 1278–1286, Beijing, China, 22–24 Jun 2014. PMLR.
- [12] Yoshua Bengio, Aaron Courville, and Pascal Vincent. Representation learning: A review and new perspectives. *IEEE transactions on pattern analysis and machine intelligence*, Vol. 35, No. 8, pp. 1798–1828, 2013.
- [13] Scott M Lundberg and Su-In Lee. A unified approach to interpreting model predictions. 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*, pp. 4765–4774. Curran Associates, Inc., 2017.
- [14] Marco Tulio Ribeiro, Sameer Singh, and Carlos Guestrin. ”why should i trust you?”: Explaining the predictions of any classifier, 2016.
- [15] Scott M Lundberg, Gabriel Erion, Hugh Chen, Alex DeGrave, Jordan M Prutkin, Bala

- Nair, Ronit Katz, Jonathan Himmelfarb, Nisha Bansal, and Su-In Lee. From local explanations to global understanding with explainable ai for trees. *Nature machine intelligence*, Vol. 2, No. 1, pp. 2522–5839, 2020.
- [16] Scott M Lundberg, Bala Nair, Monica S Vavilala, Mayumi Horibe, Michael J Eisses, Trevor Adams, David E Liston, Daniel King-Wai Low, Shu-Fang Newman, Jerry Kim, et al. Explainable machine-learning predictions for the prevention of hypoxaemia during surgery. *Nature biomedical engineering*, Vol. 2, No. 10, pp. 749–760, 2018.
- [17] Erik trumbelj and Igor Kononenko. Explaining prediction models and individual predictions with feature contributions. *Knowledge and Information Systems*, Vol. 41, pp. 647–665, 2013.
- [18] Avanti Shrikumar, Peyton Greenside, and Anshul Kundaje. Learning important features through propagating activation differences. 2017.
- [19] A. Datta, S. Sen, and Y. Zick. Algorithmic transparency via quantitative input influence: Theory and experiments with learning systems. In *2016 IEEE Symposium on Security and Privacy (SP)*, pp. 598–617, 2016.
- [20] Sebastian Bach, Alexander Binder, Grégoire Montavon, Frederick Klauschen, Klaus-Robert Müller, and Wojciech Samek. On pixel-wise explanations for non-linear classifier decisions by layer-wise relevance propagation. *PloS one*, Vol. 10, No. 7, p. e0130140, 2015.
- [21] Stan Lipovetsky and Michael Conklin. Analysis of regression in game theory approach. *Applied Stochastic Models in Business and Industry*, Vol. 17, No. 4, pp. 319–330, 2001.
- [22] Ian J. Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. Generative adversarial networks, 2014.
- [23] 年金積立金管理運用独立行政法人. 第4回経営委員会議事概要.
- [24] T. Kohonen. "Self-Organization and Associative Memory", Vol. 8. 1984.
- [25] T. Kohonen. The self-organizing map. *Proceedings of the IEEE*, Vol. 78, No. 9, pp. 1464–1480, 1990.
- [26] Eugene F. Fama and Kenneth R. French. Common risk factors in the returns on stocks and bonds. *Financial Economics*, Vol. 33, No. 1, pp. 3–56, Feb 1993.
- [27] Mark M. Carhart. On persistence in mutual fund performance. *The Journal of Finance*, Vol. 52, No. 1, pp. 57–82, Mar 1997.

著作:

年金積立金管理運用独立行政法人

制作:

田尻 貴夫

村上 朝成

橋戸 公則

北野 宏明

株式会社ソニーコンピュータサイエンス研究所

〒141-0022

東京都品川区東五反田 3-14-13 高輪ミュージズビル 3F

Tel: 03-5448-4380