

CFA 協会ブログ

No.553

2021 年 11 月 26 日

「機械学習：説明せよ、さもなければ吹き飛ばせ」

Machine Learning: Explain It or Bust

[ダン・フィルプス, PhD, CFA](#)

「簡単に説明できないのであれば、それは理解していないということです。」

複雑な機械学習 (ML) においてもこれと同じことが言えます。

機械学習は、今日では ESG リスクを計測し、取引を執行し、株式の銘柄選択やポートフォリオの構築にも使えるようになりましたが、その最も強力なモデルは依然としてブラックボックスのままとなっています。

投資業界にわたる機械学習の加速度的な普及により、透明性の低下と投資決定の説明方法に関して、まったく新しい懸念が生じています。率直に言って、[「説明不可能な機械学習アルゴリズムは \[...\] 許容できない水準の法規制リスクに企業を晒すものなのです。」](#)

平たく言って、あなたがあなたの投資にかかる意思決定を説明できないのであれば、あなたの会社、そしてあなたのステークホルダーはひどく困りますよね。説明は、もしくは、直接的な解釈と言った方が更に良いかもしれませんが、それゆえにとっても重要なのです。

人工知能 (AI) と機械学習をこれまで活用してきた他の主要産業において、優秀な人々がこの課題と格闘してきました。私たちのセクターにおいても、投資のプロフェッショナルよりもコンピューターサイエンティストを重用し、あるいは、単純でただちに何にでも使えるような機械学習アプリケーションを投資の意思決定プロセスに投げ込むような人たちにとって、機械学習はすべてを変えてしまっています。

商用化されている機械学習ソリューションには、現在2つのタイプがあります。

1. 解釈可能な AI は、直接読んだり解釈したりすることができる比較的シンプルな機械学習を使用します。
2. 説明可能な AI (XAI) は、複雑な機械学習を利用して、その説明を試みるものです。

XAI は将来のソリューションとなる可能性があります。しかし、それは遠い将来の話です。私の 20 年間のクオンツ運用と機械学習の研究に基づけば、解釈可能性こそが、現在および予見可能な将来においては、機械学習と人工知能の能力を理解し利用するためにみなさんが目を向けるべきだと信じています。

その理由についてご説明しましょう。

ファイナンスにおける第二の技術革新

機械学習は将来、現代投資マネジメントの重要な部分を形成することになるでしょう。そのことについては広範なコンセンサスがあります。機械学習のおかげで、コストのかかるフロントオフィスの人員を削減でき、古くなったファクターモデルを新たなものに取り替えることができ、巨大で増大し続けるデータプールに梃子入れすることができ、究極的にはより狙いを定めた特注的な方法で、アセットオーナーの目標を達成できるようになるのです。

しかしながら、投資マネジメントにおいて技術がなかなか普及しないことは古い話であり、機械学習もまた例外ではなかったのです。最近まではそうなのです。

直近 18 か月間余りの ESG の興隆と、それを評価するために必要となる巨大なデータプールという外部資源の活用は、機械学習への移行をターボエンジンのように加速させた主要な要因でした。

これらの新しい専門人材とソリューションへの需要は、この 10 年間余りの間、すなわち 1990 年代中盤に金融業界を襲った最後の大きな技術革新以降に私が見てきたものすべてを、既に凌駕してしまいました。

機械学習という兵器の開発競争のペースは心配の種です。一見してわかるように、新たに専門家を名乗りはじめた人たちが取り込まれていることは憂慮すべきです。この革新が、投資ビジネスよりもむしろコンピューターサイエンティストによって取り入れられているかもしれないというところが、すべての可能性の中で最も悩ましいものかもしれません。投資決定の説明とは、いつも投資ビジネスにおける骨の折れる理論的解釈にあるものです。

解釈可能な簡潔性か？説明可能な複雑性か？

解釈可能な AI は、シンボリック AI (SAI) あるいは「古き良き AI」とも呼ばれ、そのルーツは 1960 年代に遡ります。しかし、それは再び AI 研究の最前線に躍り出ています。

解釈可能な AI システムは、ルールベースであることが多く、ほぼ決定木のようなものです。もちろん、過去に何が起きたかを理解するうえで決定木は有用ですが、将来予測のツールとしては厳しく、概してデータを過学習してしまっています。しかしながら、今では解釈可能な AI システムは、ルール学習のためのずっと強力で洗練されたプロセスを有しています。

これらのルールはデータに適用されるべきものです。そのため、まさにベンジャミン・グレアムとデビッド・ドッドが作成した投資ルールのように、直接的に検証、精査、解釈が可能です。これらのルールはおそらく簡潔と言えますが、強力であり、そしてルール学習が正しく行われたならば、安全でもあります。

一方、説明可能な AI、すなわち XAI は、これとはまったく異なります。XAI は、直接解釈することが不可能なブラックボックスモデルの内部動作を説明しようとします。ブラックボックスにおいて、インプットとアウトプットは観察可能ですが、その間にあるプロセスは不明瞭であり、推測することしかできません。

これが、XAI が一般に試みようとすることです。すなわち、ブラックボックスプロセスの説明に向けて推測とテストを行います。XAI は、異なるインプットが結果に及ぼすかもしれない影響のあり方を示すために、可視化という方法を利用しています。

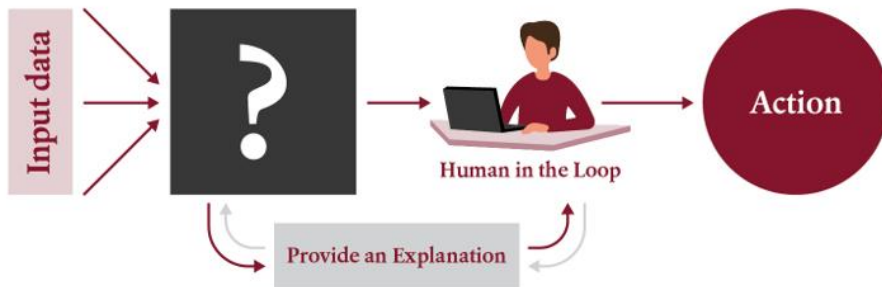
XAI は未だその初期段階にあり、かつ、課題の多い分野であることがわかっています。これらは、機械学習アプリケーションに関する限り、判断を先延しにして解釈可能 AI へと向かうべきである 2 つの強い理由となっています。

解釈か説明か？

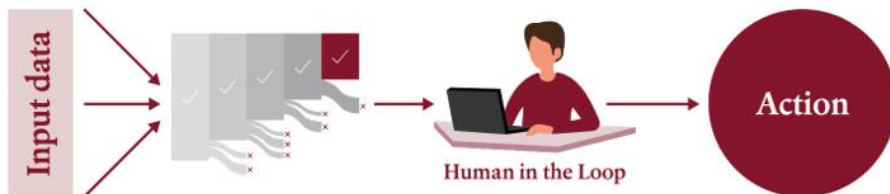
1. Black-Box Learners, Opaque Outcomes



2. Explainable-AI: Black-Box Learner Explained



3. Interpretable-AI

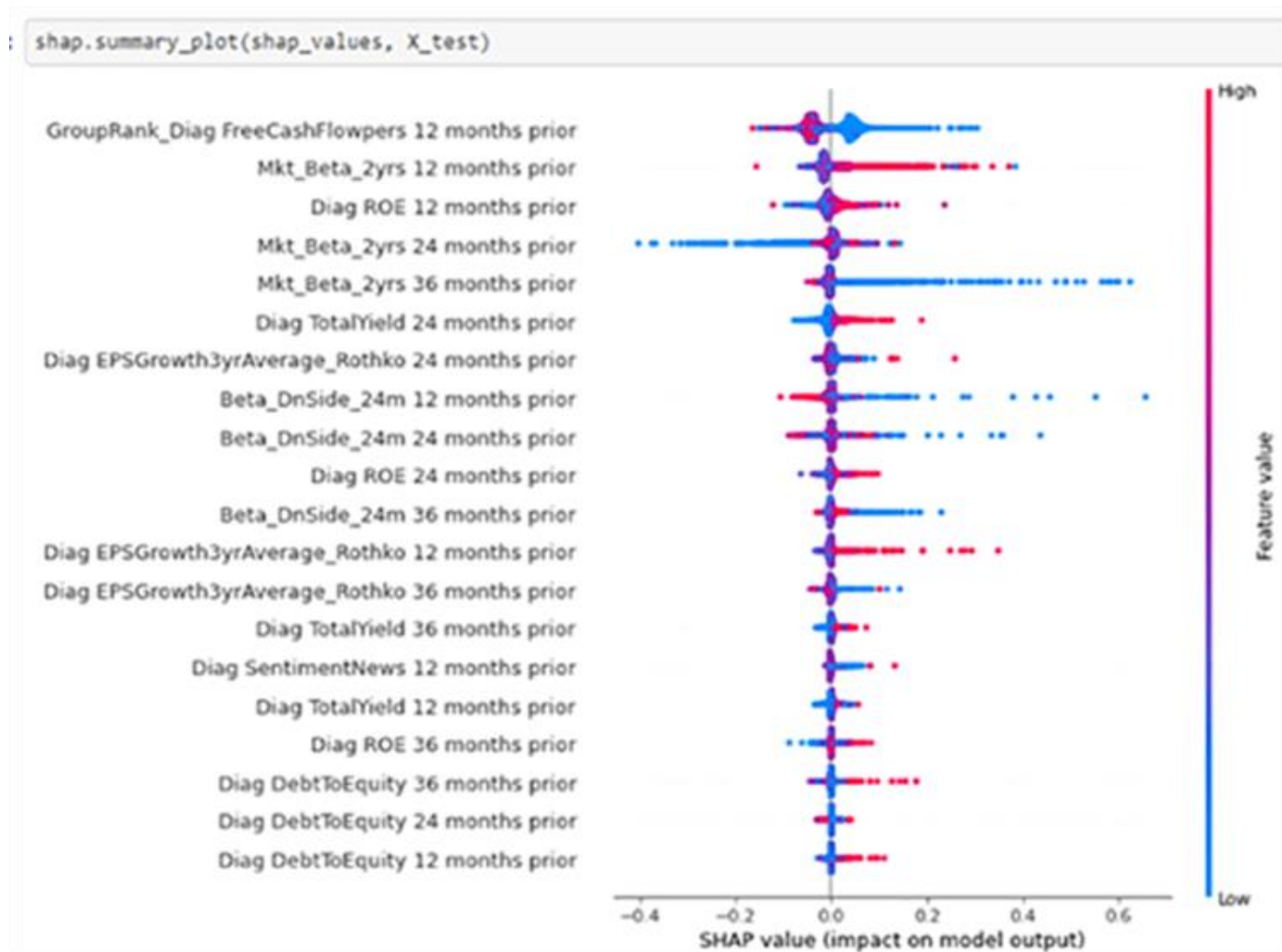


ファイナンス分野でより共通的に用いられている XAI アプリケーションのひとつが SHAP (SHapley Additive exPlanations)です。SHAP は、ゲーム理論のシャープレイ値にその起源を持ち、[ワシントン大学の研究者たちによって、ごく最近になって開発されました。](#)

次のイラストは、わずか数行の Python コードの実行結果である株式銘柄選択モデルの SHAP による説明を示しています。しかし、その説明を理解するためには、それ自体の説明が更に必要となっています。

SHAP は素晴らしいアイデアであり、機械学習システムの開発のためにはとても有用ですが、取引エラーをコンプライアンス担当執行役員に説明するためにそれに頼ろうとするのであれば、優れたプロジェクトマネージャーが必要になるでしょう。

コンプライアンス執行役員向けの機械学習？ニューラルネットワークを説明するためのシャープレイ値の利用



注：これは、ある新興国市場エクイティのユニバースにおいてより高いアルファの株式を選択するために設計されたランダムフォレストモデルの SHAP による説明です。過去のフリーキャッシュフロー、マーケットベータ、ROE 等をインプットとして利用しています。図の右側は、インプットがどのようにアウトプットに影響を及ぼしているかを示しています。

ドローン、核兵器、ガン診断・・・そして株式銘柄選択？

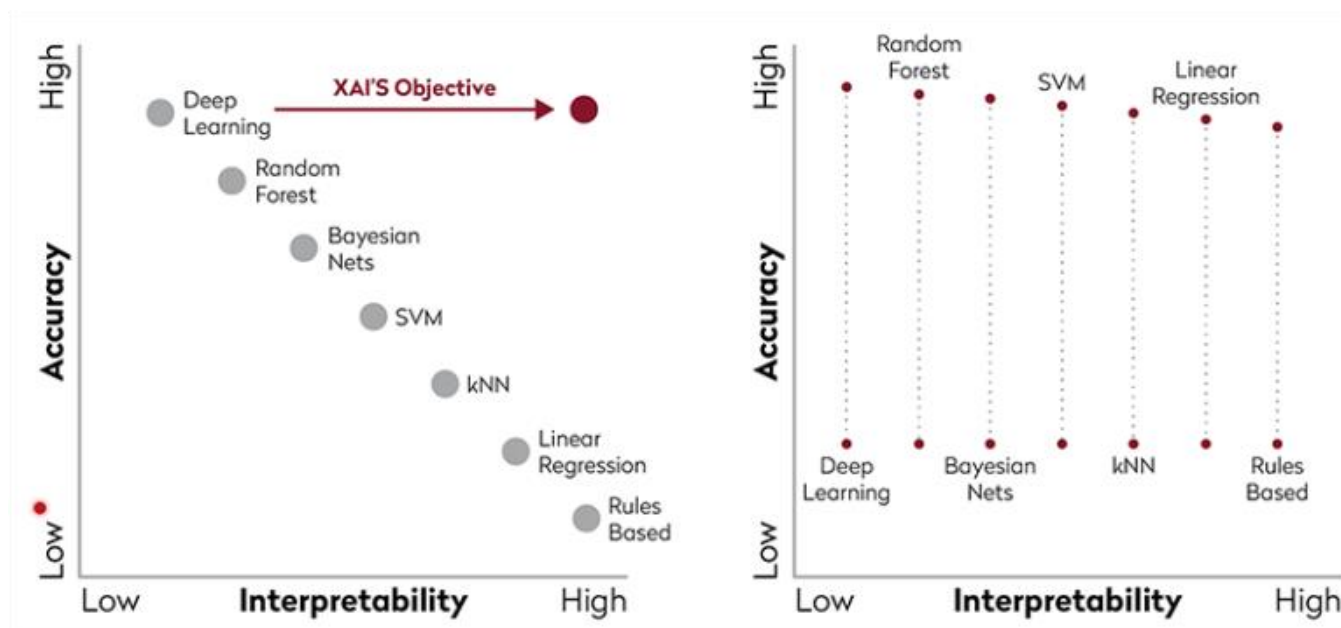
医療研究者と防衛産業は、ファイナンスの分野よりもずっと長い間、説明か解釈かの問いに対して答えを探ってき経緯を有しています。これらの分野は、強力なアプリケーション特定型の解決策を見出してきましたが、未だいかなる一般的な結論にも達していません。

[米国防高等研究計画局\(DARPA\)は、思考開発的な研究を実施して、解釈可能性を機械学習システムのパワーの伝達を妨げるひとつのコストとして識別しました。](#)

下のグラフは、この結論を機械学習の様々なアプローチとともに描いています。同分析においては、より解釈可能なアプローチであればあるほど、その複雑性は低下し、それゆえにその正確性も低下するものとされています。複雑性が正確性と関連しているのであればこれは確かに真

であると言えるでしょうが、オッカムの剃刀の原理(節約の原理)とこの分野の腰の重い研究者たちは、これに同意しかねています。正しい面を示唆するどちらかの図が、現実をよりよく表しているのかもしれませんが。

解釈可能性は本当に正確性を低下させるでしょうか？



注：シンシア・ルーディンは、XAI の支持者が強く主張する（左図）ほどには、正確性は解釈可能性に関連していない（右図）と述べています。

最高責任者たちの複雑性バイアス

正確なブラックボックスとそこまで正確ではないが透明性を有するモデルという誤った二分法は、もはや行き過ぎです。何百人という科学者のリーダーたちと投資会社のエグゼクティブたちがこの二分法によってミスリードされた場合、その他大勢の人々もまた騙されてしまうかもしれないことを想像してみてください。－ [シンシア・ルーディン](#)

複雑性を当然のこととする説明可能性陣営に織り込まれた仮説は、例えば[たんぱく質のフォールディング\(折り畳み\)問題を予測する](#)といったように、機械学習が極めて重要な分野のアプリケーションにおいては正しいことかもしれません。しかし、株式銘柄選択を含むその他のアプリケーションにおいては、それはそれほどなくてはならないものではないかもしれません。

[2018 年の説明可能な機械学習チャレンジにおける狂奔](#)が、このことを実証しました。それは、ニューラルネットワークのためのブラックボックスチャレンジになるであろうと見られていましたが、AI 研究者のスーパースターであるシンシア・ルーディンとそのチームは、これとは違った考えを持っていました。彼女たちは、解釈可能な、より簡単に読むことのできる機械学習モデルを提案しました。それはニューラルネットベースのものではなかったもので、一切の説明を必要としませんでした。それは既に解釈可能だったのです。

ひょっとしたら、ルーディンによる最も印象的なコメントは、「ブラックボックスモデルを信頼することとは、モデルの方程式を信頼することだけではなく、方程式を構築する元となるデータベース全体も信頼するということを意味する」だったかもしれません。

彼女の論点は、行動ファイナンスのバックグラウンドのある方々にはお馴染みでしょう。ルーディンは、さらにもう一つの行動バイアスについて認識しています。つまり、複雑性バイアスです。私たちは、シンプルなものより複雑なものをより魅力的に考えがちです。彼女のアプローチは、[解釈可能な AI vs 説明可能な AI に関する最近の WBS ウェビナー](#)で彼女が解説したとおり、ブラックボックスモデルはベンチマークを提供するためだけに利用し、その後に類似した正確性を持つ解釈可能なモデルを開発するというものです。

人工知能という兵器の開発競争を主導する最高責任者たちは、過剰な複雑性への総力をあげた探究を継続する前に、いったん立ち止まってこのことをよく考えてみた方がいいかもしれません。

株式銘柄選択のための解釈可能で監査可能な機械学習

複雑性を必要とする目的がある一方で、複雑性が害となる目的もあります。

株式銘柄選択もそのひとつです。「解釈可能で、透明で、監査可能な機械学習」(原題: [“Interpretable, Transparent, and Auditable Machine Learning”](#))の中で、デイビット・ティレス、ティモシー・ローと私は、解釈可能な AI を、株式投資マネジメントにおける銘柄選択のための、ファクター投資に対する拡張可能な代替的手法として提案しました。私たちのアプリケーションは、シンプルな機械学習アプローチの非線形指数関数を用いて、シンプルで解釈可能な投資ルールを学習します。

それは、複雑でなく、解釈可能で、拡張可能で、そして、私たちが信じているように、ファクター投資を引き継ぎ、それをはるかに超えているといったところが目新しい点となります。実際のところ、私たちのアプリケーションは、私たちが何年にもわたって試用してきたはるかに複雑なブラックボックスアプローチとほぼ同じようによく機能します。

私たちのアプリケーションにおいて透明性とは、監査可能であって、コンピューターサイエンスにさほど詳しくないステークホルダーにも意思伝達や理解が可能であるということを意味しています。それを説明するために XAI は必要ありません。それは直接解釈することが可能なのです。

株式銘柄選択のために過剰な複雑性は必要ないという長年抱いてきた信念が、私たちがこの研究を公表した動機でした。実際、このような複雑性は、株式銘柄選択にとってはほぼ確実に有害なものとなります。

解釈可能性は、機械学習において最も重要なものです。それに取って代わるものが、あまりに回りくどいたために全ての説明が説明のための説明を際限なく必要とするような複雑性なのです。

どこまでいったらそれは終わるのでしょうか？

人間にとっての機械学習

説明か解釈か、それはどちらなのでしょう？ 論争は激しくなっています。最も将来志向で考える投資会社においては、機械学習の躍進を支える研究に何億ドルもの資金が投じられています。

最先端の技術につきものの、出だしの失敗、突発事象、資本の無駄遣いは避けられません。しかし、今のところ、また予見可能な将来においては、解釈可能な AI がその解決策なのです。

2つの自明の理を考えてみてください。物事が複雑になればなるほど、説明の必要性も高くなります。また、物事が十分に解釈可能であればあるほど、説明の必要性は低くなります。

将来、XAI はより確立され、理解され、ずっと強力なものになるでしょう。今のところ、それは初期段階にありますし、投資マネージャーに彼らの会社とステークホルダーを許容できない水準の法規制リスクの可能性に晒すように頼むということは、やり過ぎと言うものです。

一般的な目的を持つ XAI は、次のようなもののたとえで言うシンプルな説明を、現時点では提供していません。

簡単に説明できないのであれば、それは理解していないということです。

執筆者

Dan Philips, PhD, CFA

(翻訳者: 荒木 謙一、CFA)

英文オリジナル記事はこちらに

<https://blogs.cfainstitute.org/investor/2021/11/26/machine-learning-explain-it-or-bust/>

注) 当記事は CFA 協会 (CFA Institute) のブログ記事を日本 CFA 協会が翻訳したものです。日本語版および英語版で内容に相違が生じている場合には、英語版の内容が優先します。記事内容は執筆者の個人的見解であり、投資助言を意図するものではありません。

また、必ずしも CFA 協会または執筆者の雇用者の見方を反映しているわけではありません。