

2022年度 JUAS研究成果発表 AI研究会

2023年4月13日(木) 14:30—15:00
AI研究会

運営方針

目的： 自社事業に資する成果物を持ち帰って頂くこと

研究会に参加してAIに関する知見を強化して頂きたいのですが、それにとどまらず、皆様の業務に実際に役立つものを持ち帰って頂きたいです

成果物： 事業適用シナリオ、適用への知見、プロトタイプなど

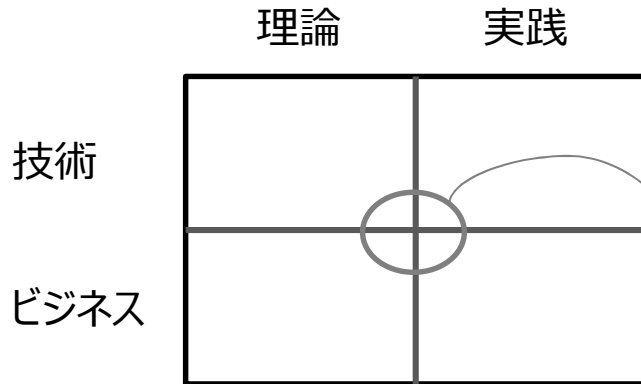
近いリサーチクエストンをお持ちの方で分科会を編成し、個人の興味を重視しつつ、分科会全体としてより大きな成果物をまとめて頂きます（成果物の範囲は上記に限りません）

企業横断： AIに関与する方の企業の境界を越えた連携

JUASという組織を活用して、通常の業務では得難い、他企業でAIに関連する業務に携わる方との関係性を構築頂き、より多面的な貢献につなげて頂くことを期待しております
GAFA等が促す事業の変化への対応にもつながれば何よりです

分科会編成とリサーチクエスト

技術／ビジネス、理論／実践の2軸で参加者の皆様の問題意識をクラスタリング、分科会毎にリサーチクエストを立てて頂き、研究に取り組みました



様々な問題意識を提示頂きました。
横断的に、AI人材に関する意識も出され、
境界領域の重要性を再認識しました

人材×データRQ：
AIプロジェクトを円滑に企画・推進・運用
していくために必要なロール、スキルセット、
教育手段は何か？

技術理論RQ：
AIやデータ分析の結果を納得してもらう
手段は何か。どのように実現できるか？

技術実践RQ：
機械学習の一連のプロセスを実践することで、
なにを学べるか？

ビジネス理論RQ：
定性的な観点から費用対効果が高い、AIビジネスで
共通する特徴にはどのようなものがあるのか？

ビジネス実践RQ：
成功/失敗するAIの特徴とは何なのか？

年間スケジュール

Zoomによるオンライン開催とハイブリッド開催を交えて活動しました

#	日付	会議名	場所	テーマ
1	6月23日(木) 16:00開始	第1回 定例研究会	JUAS	・2021年度活動内容共有 ・AWS紹介、アカウント作成
2	7月28日(木) 16:00開始	第2回 定例研究会	JUAS	・分科会編制（リーダーも決定）
3	9月15日(木) 16:00開始	第3回 定例研究会	JUAS	・分科会活動 ・JMOOC紹介
4	11月24日(木) 16:00開始	第4回 定例研究会	JUAS	・分科会活動
5	12月22日(木) 16:00開始	第5回定例研究会	JUAS	・分科会中間発表
6	1月12日(木) 16:00開始	第6研究会	JUAS	・分科会活動
7	3月2日(木) 16:00開始	第7回定例研究会	JUAS	・分科会成果発表
*	4月13日(木)	2022年度成果発表会	JUAS	・2022年度成果報告

AI研究会 2022

人材×データ分科会

リサーチクエストションが決まるまで

■ 人材、データに関する課題意識

- 人材教育は、何をどのように行えばよいかわからない
- 人材育成に向けて、リテラシーのレベル感をどう定義すればよいかわからない
- プロジェクトの担当者間で、お互い何ができるのかを理解してほしい
- AI開発とITシステム開発の人材の違いがわからない

■ 昨年度までの人材系分科会の研究テーマ

- FY2019:パイロットプロジェクト成功に向けた人材定義
- FY2020:継続的なプロジェクト運営に向けた組織と人材定義
- FY2021:継続的なプロジェクト運営に向けたスキル定義
- FY2022:何にフォーカスするか？



PoCから本運用へ

研究要素の追加

■ 研究テーマの方向性

- FY2021まではAIプロジェクトの人員にフォーカスしていたが、社内外のステークホルダーまでロールの範囲を広げる
- 各ロールが持つべきスキル定義と、それをもとに育成プログラム案作成を目指す
- AIプロジェクトならではの観点を含める

リサーチクエスチョン



AIプロジェクトを円滑に企画・推進・運用していくために
必要なロール、スキルセット、教育手段は何か？

■ 成果物イメージ

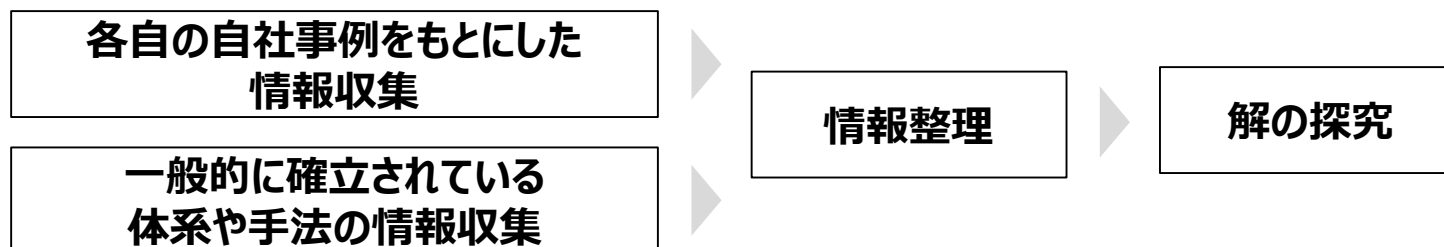
- AIプロジェクトに関わるロール一覧
- 各ロールが担うタスク一覧
- 各ロールが所持すべきスキル一覧
- スキルを獲得するための育成プログラム(案)

■ 成果物の活用イメージ

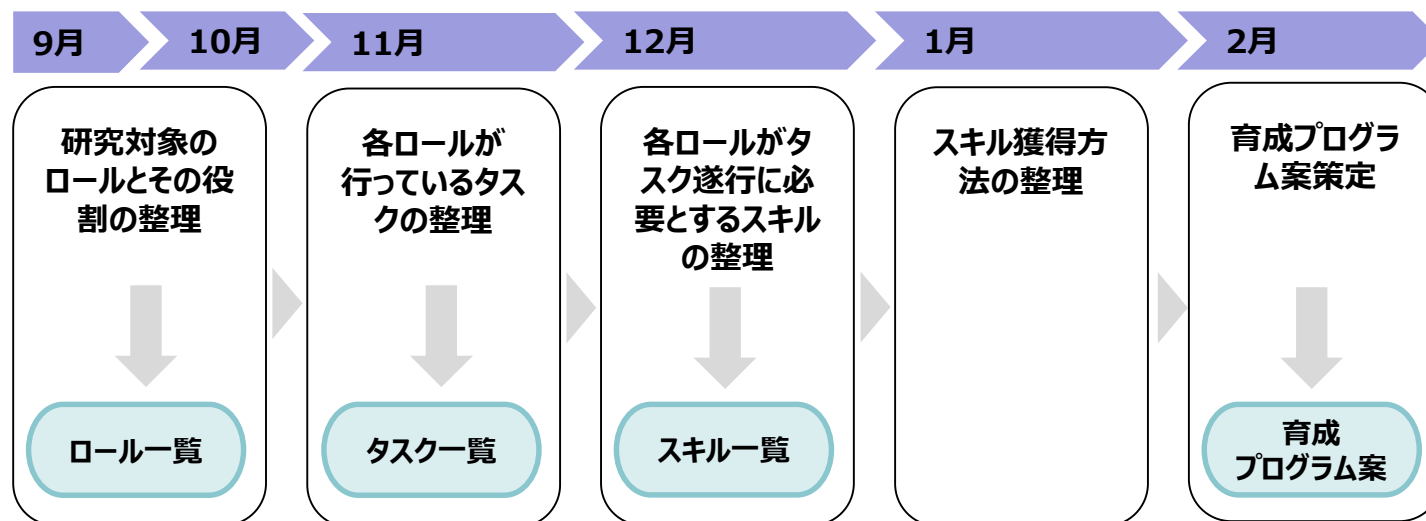
- AIプロジェクト開始時に、必要な人材を不足なく集めるための資料として活用する
- AIプロジェクト体制構築のために周囲を説得する材料として活用する
- AIプロジェクト継続やシステム運用に向けての人材育成プログラム案として活用する

研究方法

■ 基本的なアプローチ

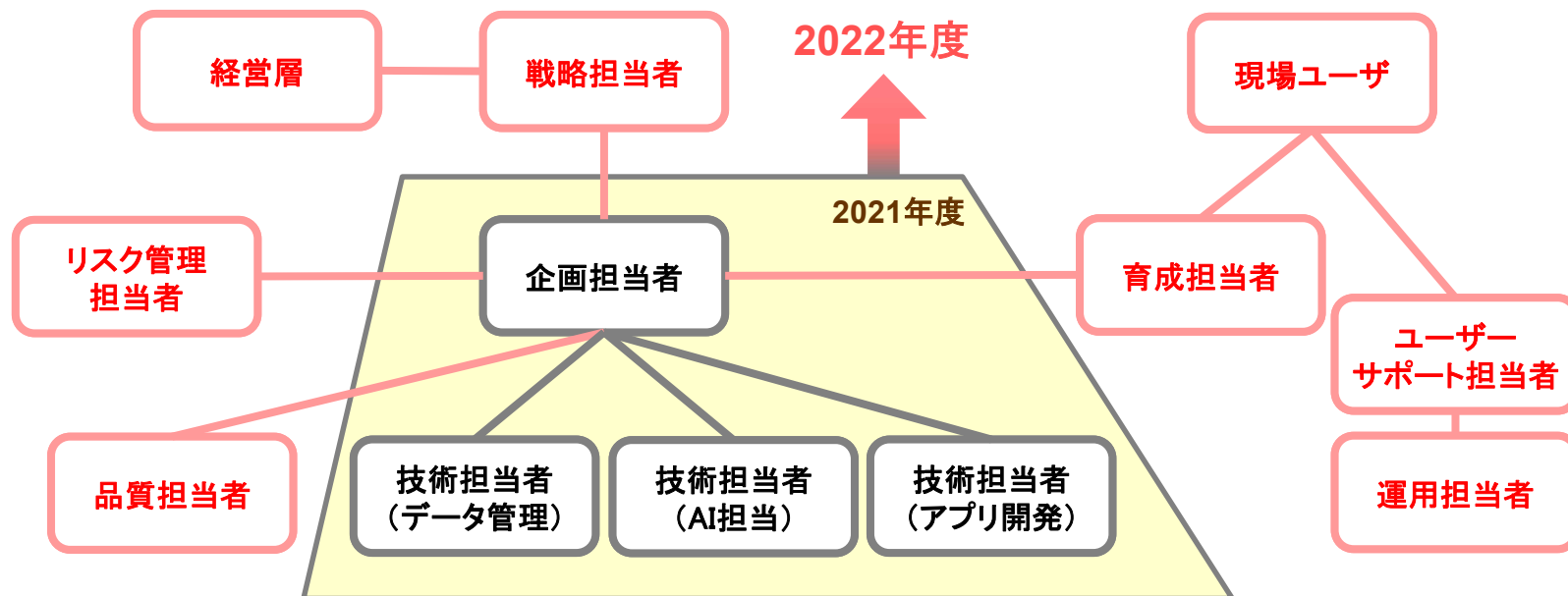


■ ロードマップ



研究対象のロールとその役割を整理

2021年度はAIシステム開発者（PM、リーダー、エンジニア）を対象としていた。
本年度はAIプロジェクトのステークホルダ全体へ対象を広げ、ロールを再整理した。



プロジェクトのフェーズを定義、各ロールの主な役割と各フェーズにおけるタスクを整理した。



各ロールが行うタスクの一覧整理

AIプロジェクトにて必要となる各ロールが担う役割、及び各フェーズで実施すべきタスクを一覧化により整理を行った。

#	ロール	ロールの役割	企画フェーズ	PoCフェーズ			本開発フェーズ		運用フェーズ	
				PoC	PoC評価	本開発企画	本開発	トレーニング	運用	評価
1	経営層	➢事業計画立案 ➢意思決定	-	-	PoC評価	投資・方針判断	-	-	-	システム運用評価 次期計画更新
2	戦略担当者	➢事業計画に基づくプロジェクト選定 ➢プロジェクトの横断的管理、予算管理	戦略立案 投資管理・評価 予算管理	-	PoC評価	本開発方針判断 本開発予算管理	-	-	人材育成計画策定	システム運用評価
3	リスク管理 担当者	➢情報セキュリティ管理 ➢法務確認	セキュリティ観点の評価 開発委託契約確認 データの著作権評価	-	-	セキュリティ要件確認 開発委託契約確認	セキュリティチェック	-	セキュリティ講習	-
4	企画担当者	➢プロジェクトの企画立案 ➢プロジェクト推進、管理	問題分析／要件定義 企画立案、予算、提案 ツール選定 開発委託先選定、契約 データの著作権評価	推進、管理	PoC評価報告 (技術評価 、要件評価)	予算立案 本開発企画立案、提案 開発委託契約	開発推進、管理 テスト リリース計画 リリース判定	トレーニング企画 トレーニング状況管理	利用状況確認	システム運用評価 システム運用評価報告
5	技術担当者 (データ管理)	➢データの選定、分析、管理 ➢運用時の問題解決支援	PoC企画内容の技術検討 データ選定	開発 (データ管理)	PoC結果纏め	技術検討 データ収集 データ管理基盤選定	本開発 (データ管理)	-	システム保守 データ要望対応	システム運用状況纏め
6	技術担当者 (AI担当)	➢AIモデル作成・評価 ➢運用時のモデル改善	PoC企画内容の技術検討	開発 (モデル作成)	PoC結果纏め	AI技術検討	本開発 (モデル調整)	-	モデル改善	システム運用状況纏め
7	技術担当者 (アプリ開発)	➢アプリケーション開発 ➢運用時の問題解決支援	PoC企画内容の技術検討	開発 (アプリ開発)	PoC結果纏め	アプリケーション技術検討	本開発 (アプリ開発)	-	システム保守	システム運用状況纏め
8	品質担当者	➢システム／AIの品質管理 ➢品質向上施策実施	本質評価計画策定	-	品質評価	品質評価計画策定 障害時対応計画策定	品質評価 品質向上施策実施	-	-	システム運用評価
9	育成担当者	➢ユーザトレーニング ➢運用者トレーニング	-	-	-	トレーニング対象者選定	-	トレーニング企画・実施 トレーニング教材作成 マニュアル作成	追加トレーニング実施	-
10	ユーザサポート 担当者	➢システム利用に関する照会対応 ➢一次調査	-	-	-	-	-	マニュアル作成	照会対応・一次調査 マニュアル改訂	システム運用状況纏め
11	運用担当者	➢定常的なシステム運用 ➢運用時の問題解決、障害時対応	-	-	-	-	-	-	システム運用、監視 障害時対応	システム運用状況纏め
12	現場ユーザ	➢要望出し ➢システム利用・評価	問題提起・要望出し AsIs/ToBe提起	-	PoC評価	-	-	トレーニング受講	システム利用	システム運用評価

各ロールがタスク遂行に必要とするスキルの整理

- ✓ 前頁で整理した各ロールのタスク実行に必要なスキルを、大きく4グループに分けて定義した。
- ✓ 次頁からの詳細説明では、ロールにおいて保有がマストなスキルグループを○、ベターもしくはスキルのうち一部が必須であるスキルグループを△、保有していなくても問題ないスキルグループを－で表示する。
- ✓ 詳細スキルのうち、各ロールで必要となるスキルは赤字で明示する。

スキルグループ	スキル概要
AI基礎力	AIを開発・運用するプロジェクトに参画し、成果をビジネスに役立てる上で必要となる知識やマインド
ビジネス力	業務における課題を理解・分析し、AIを利用することで達成する目標を定め、課題解決に繋げるためのスキル
データエンジニア力	非機能要件(保守性・セキュリティ等)を考慮した上で、AIの利用環境構築およびソフトウェアの設計・実装・運用を行うスキル
データサイエンス力	ビジネスの変革・課題解決に向けて、必要なデータの収集・解析の仕組みの設計・実装・運用を行うスキル

参考：データサイエンティストスキルチェックリスト、デジタルスキル標準

スキル獲得方法の整理

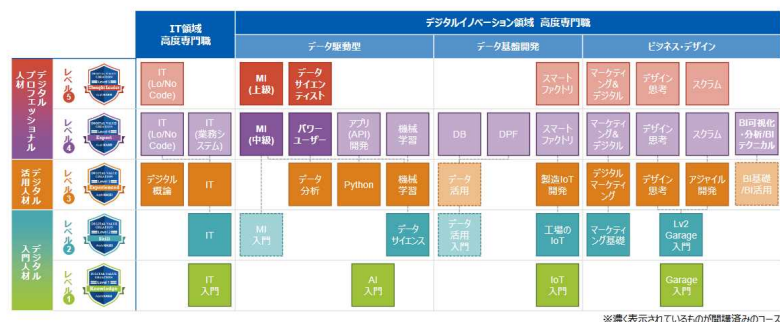
AIの分野は常に技術更新されるため、スキルの測定と、スキルの獲得の循環が必要と考える。
現在のスキル測定およびスキル獲得手段を調査した。

スキル	スキル測定	スキル獲得
AI基礎力	- ITパスポート - Ai人材検定 for business - AI For Everyone	- 書籍 - セミナーやワークショップの活用 - e-learningの活用 - デジタル人材育成支援サービスの活用 - 教育プログラムの内製
ビジネス力	- ITパスポート - ITコーディネータ - Ai人材検定 for business - 人工知能プロジェクトマネージャー試験	
データエンジニアリング力	- 各種スペシャリスト試験 (IPA) - 各種ベンダー資格 - 基本/応用情報技術者試験 - E資格 - AI実装検定 - 画像処理エンジニア検定 - Pythonエンジニア認定データ分析試験 - 統計検定 (DS基礎/発展) - Ai人材検定 for engineer	
データサイエンス力	- G検定 - データサイエンティスト検定 - ITパスポート - AI実装検定 - 統計検定	

育成プログラム案作成

AI人材育成を先進的に行っている企業を例に、育成プログラム作成に際し重要な要素を整理した。

■ 育成プログラム例(旭化成株式会社)



※濃く表示されているものが関連深いコース

育成プログラム作成のステップとポイント

- ✓ 社内に必要なロールの定義
- ✓ ロールごとの育成プログラムの作成
 - ・ 実践経験を積むプログラムを組み、個人の原体験を増やすことが重要
- ✓ 認定制度等による社員スキルの可視化
- ✓ 認定制度と人事制度の連携
- ✓ 時間がかかることを前提に継続的な取り組みとする

■ 近年の人材育成の傾向

今年度以前は、少数精鋭で技術を学び、そのメンバーがエバンジェリストとして有識者を増やすという手段が取られていた。

今年度は、IPAによるとITパスポートの受検傾向としてIT関連企業以外の受検者が増加しており、各企業が有識者のみでなく全社員がリテラシーを持つという考え方に舵を切っていると思われる、今年度以降もその流れが継承されるものと考えられる。

参考：旭化成DX戦略説明会 <https://www.asahi-kasei.com/jp/ir/library/business/pdf/221213.pdf>

結論と今後の課題



AIプロジェクトを円滑に企画・推進・運用していくために 必要なロール、スキルセット、教育手段は何か？

■ 結論

- AIプロジェクトは12のロールを持つメンバーで推進され、企画、PoC、開発、運用のフェーズにおいて、各々のロールが遂行するタスクが存在する。
- AIプロジェクトは、AI基礎力、ビジネス力、データエンジニア力、データサイエンス力のスキルを持つメンバーを必要とし、遂行すべきタスクと照らし合わせて、各ロールが必要なスキルセットを継続的に獲得・強化することにより、円滑な運営を目指すことができる。
- スキル可視化、また、スキル獲得のための教育手段は、オンラインの活用等により増加し、近年はそれらの活用に加え、企業独自の育成プログラム推進の事例も見られる。

■ 今後の課題

- 各スキルセットの獲得・強化のために必要な具体的な教育手段を定義するまで至らなかったため、提言としてまとめることを次の課題とする

AI研究会 2022

技術実践分科会

リサーチクエスト

メンバーで話し合い、検討した結果、この分科会でやりたいこととしては以下の方向性となった

機械学習の一連のプロセスを実践することで、なにを学べるか

なぜ上記をやりたいのか

DX推進やデータドリブンは、全員がデータの収集・分析・議論・意思決定・自動化に貢献することが求められる

技術系の人はもちろん、ビジネスサイドやビジネスサイドに近い立場にいる人間でも、機械学習の一連のプロセスの中で、どこでどのようなことが行われているか俯瞰して知る必要があるのではないか

研究方針

- 実際に手を動かしながら、機械学習を学んでいく
- ビジネス上の課題を疑似設定をしたうえで、下記のプロセス（※）全体を通して実践してみる
- これにより各社でのAI関連の開発や育成施策を模索する上でベースの知識を取得する

※機械学習のプロセス（点線の赤枠を主な対象プロセスとした）



機械学習のプロセスと実施内容の定義

	プロセス	実施内容（一般的に実施すると考えられること）
1	ビジネスゴールの設定	・ビジネス要件を理解する。 ・ビジネス・クエスチョンを形成する。 ・プロジェクトのデータ要件を検討する。 ・この実装から生まれる可能性のある新しいビジネスプロセスを検討する。
2	機械学習の問題設定	・目的変数を決定する。 ・ビジネス・クエスチョンに基づき、機械学習タスクを定義する。（回帰or分類） ・許容誤差を含む主要なパフォーマンス指標を決定する。
3	データ収集	・データソースを確認する。（外部データか自社内データか。） ・収集方法を決定する。（リアルタイムで取得するのかバッチで取得するのか。） ・データ格納先を決定する。
4	EDA&データ前処理	・データを可視化など行い、データの特徴を要約し理解する。 ・データの特徴を損なわないよう、データクリーニングや置換、バランシング、スケーリングなどを行う。
5	特徴量エンジニアリング	・手持ちのデータからドメイン知識などを駆使し、新たな特徴量を生成する
6	学習・チューニング・評価	・問題に適した機械学習アルゴリズムを選択する。 ・過学習にならないよう、ハイパーパラメータを修正したり、評価方法をクロスバリデーションやホールドアウト法などの中から選択し実行する。

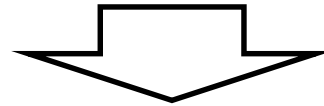
※機械学習のプロセス（点線の赤枠を主な対象プロセスとした）



機械学習の題材

機械学習の題材を決める上での考慮点

- ・メンバー全員が同じ題材に取り組んでみて、それぞれが気づいた点を持ち寄って議論したい
- ・あらかじめ「ビジネスゴール」が設定されている
- ・機械学習に必要なデータが用意されている



データ分析コンペティションサイトのテーマを題材から選択することとし、その対象を**NISHIKA**（国内最大級のデータサイエンスコミュニティ）の「**生鮮野菜の価格予測**」とした



生鮮野菜の価格予測

Nishika株式会社

締切: 2022/08/08 (残り50日)

参加: 513人 投稿: 602件

賞金/賞品: 12万円+有機野菜ギフトカード メダル・スコア付与: あり

- ・生鮮野菜の価格予測がテーマとなっています
- ・目的変数は、大田市場にて相対取引されている**各野菜の日次の卸売価格**です
- ・予測期間は最も多くの種類の野菜が市場に出ている**2022年5月**とします
- ・予測対象の野菜は、以下**16品目**です

きゅうり,こまつな,じゃがいも,そらまめ,だいこん,
なましいたけ,にんじん,ねぎ,はくさい,ほうれんそう,
キャベツ,セルリー,トマト,ピーマン,ミニトマト,レタス

<https://competition.nishika.com/yasai/summary>

※上の画像はNISHIKAのサイトより引用

- ・時系列予測というやや難しいタスクですが、例えば、今回併せて配布する**天候データ**を用いると、予測精度を向上させることができる野菜もあるとのことです

実施手法（1）

- ・メンバーそれぞれが、以下のいずれかの方法で機械学習を実施することとした

1. AutoMLツールを用いて機械学習を行う

- ・ノーコードでデータの可視化や機械学習を行うことが出来る
 - ・前処理機能は乏しいので、事前にデータを加工する必要がある
 - ・出来ることはサービスで用意されていることだけなので、用途は限られる
 - ・ある程度自動で学習できるのは非常に便利
- いくつか似たようなツールがあげられるが、各自の環境で利用しやすいものを選定
今回は「Amazon Sage Maker Canvas」と「DataRobot」を利用することとした

DataRobotについて

・DataRobotとは、「優れた予測を素早く誰でも」という理念のもと開発された機械学習を自動化するAIプラットフォーム(SaaSシステム)

・2012年にトップデータサイエンティストを集めてボストンに設立される。2017年には日本オフィスも立ち上げ

・世界中のデータサイエンティストが集結し分析力を競い合っている「Kaggle」において高成績を収めたデータサイエンティストのデータ分析ノウハウが組み込まれたシステム

・必要なデータを投入するだけで、データ前処理、モデル構築、モデルデプロイ、予測実行までのデータサイエンスプロセスをトータルで自動化してくれる

・各特徴量がどの程度効いているのか？どんな流れで前処理、学習を行っているのか？などのプロセスも可視化してくれるのが強み



参照：DataRobotとは？製品概要・機能一覧・解決できる課題・活用目的や使い方を徹底解説
https://www.ogis-ri.co.jp/column/analysis/data_analysis/c106145.html

SageMaker Canvasについて

・Amazon SageMaker Canvasは、2021年11月30日に機能提供が開始されたいわゆるAutoMLツールに相当するもの。

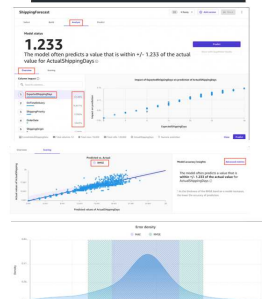
・ビジネスアナリストが、コードを記述したりMLの専門知識を必要とせずにMLモデルを構築して正確な予測を生成できるようにする、新しい視覚的なコード不要の機能を持っている

・「コードを1行も書かずに、不正行為の検出、チャーン(解約率)の削減、在庫の最適化など、ビジネスに不可欠なユースケースに対処できる」と謳っている

・当初は海外の4つのリージョンだけであったが、2022年4月14日に東京(アジアパシフィック)リージョンでも利用できるようになった。

・その後も機能改善や追加がなされている

・内部的にはSagemaker AutopilotやAmazon Forecastの機能が使われているということで、AWSならではの分析が出来るのが強みになっていると思われる

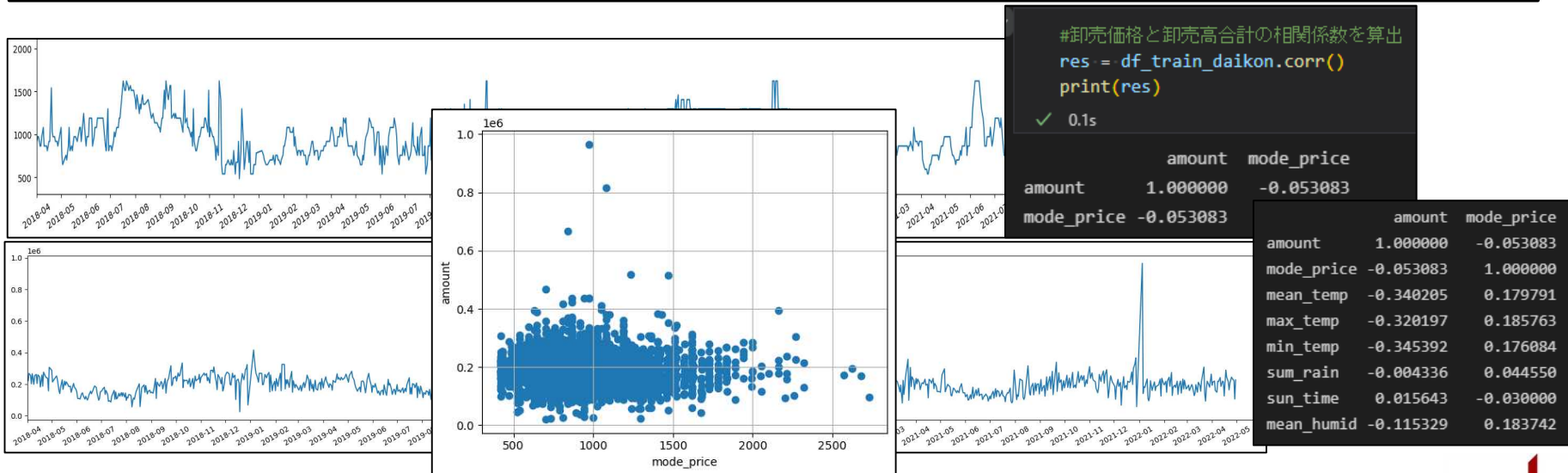


実施手法（2）

- ・メンバーそれぞれが、以下のいずれかの方法で機械学習を実施することとした

2. Pythonで各種ライブラリを用いて機械学習を行う

- ・コードで記載するため、自由に機械学習を行うことが出来る
- ・pandasなどを用いてデータを自在に加工できる
- ・matplotlibなどを用いてグラフ可視化が容易にできる
- ・学習に用いるアルゴリズムに何を使うと精度が上がるのか見極めるのが難しい
- ・前処理や特徴量エンジニアリングなど、工夫の仕方で分析のし易さや予測精度を上げることが出来る



実施過程とポイント（抜粋）

1. 実施計画書を作成し、目的とゴールを明確にする！

- ・何のために機械学習をするのか、どんな課題に対し、どのような改善策や成果を期待しているのかなど、「課題、目的、ゴールを明確にすること」が必要である
- ・どんな工程があるか？どれくらい大変か？スケジュール感はどれくらいになりそうか？を認識合わせでき、関係者に説明できる。

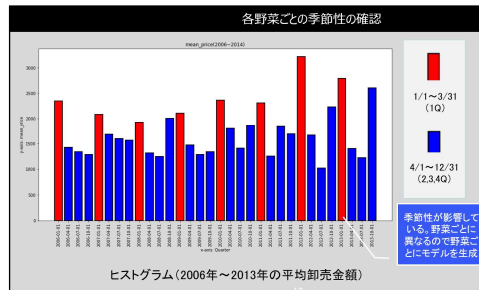
テーマ	生鮮野菜の価格を予測する	依頼組織・依頼者	小売業者●●（仮）
PJT概要 と現状課題	■課題（仮に設定） 生鮮野菜の価格は大きく変動するが、どのような条件で変動しているかが未解明。小売業者の方が予算計画を立てづらい状況となっている。 ■分析の目的（今回は②を採用） ①各野菜が安いタイミングで購入できるよう野菜の金額変動の傾向を明らかにすることで、誰もが上手に（できるだけ費用を抑えて適切なタイミングで）買い物ができるようにする ②野菜の販売価格を予測することで、野菜を小売する人たちが向こう3か月の予算計画の精度を向上させることができるようにする		
利用データ	・東京都の生鮮野菜の卸売価格、及びその産地のデータ ・各産地ごとの天候データ		
アウトプット イメージ	・統計分析による仮説検証結果のレポート（目的変数と各説明変数・説明変数同士の相関、季節性の有無など） ・特徴量を時期や産地のみとした「AI予測モデル」及び「予測&精度評価結果」→時系列 or 回帰 ※金額の予測 ・特徴量に天候データを加えた「AI予測モデル」及びその予測&精度評価結果→時系列 or 回帰 ※金額の予測		
目標とする 成果/達成 条件	・仮説の検証結果から、野菜価格が変動する要因を可能な限り多く特定する。 ・AIモデルを構築し、予測結果を検証できている。実際に活用できると小売業者の担当者が判断できるモデルの精度が出ている状態。		
個人情報の 取扱い	なし	体制案	依頼組織：小売業者●●（仮）→データ提供 分析組織：パーソルホールディングス山本→データ加工・分析

実施過程とポイント（抜粋）

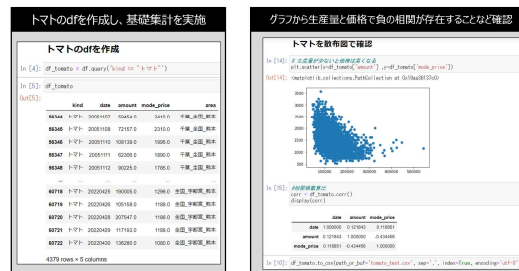
2. 俯瞰的視点からデータを眺め、可視化や基礎集計を行う！

- ・まず全体データを眺め、傾向や特徴はあるか、何らかの相関があるかどうか、季節性はどうかなどある程度つかんでおくことは重要である
- ・いきなりデータ加工をするのではなく、どのような前処理や特徴量エンジニアリングを行うのが効果的なのか、どういうデータに加工すれば精度があげられそうかの「仮説」を立てる

基礎集計 - 季節性の有無確認



基礎集計 - 各野菜ごとの相関分析



データの表記ゆれの補正、下記のような特徴量を生成したりした。
時系列データを扱うにはどうするとよく時系列の特徴を損なわずにこめるかという点で「ラグ特徴量」を生成するというのも実現した。

日別のデータを月単位に集約、かつ基本統計(平均、最大、最小)を特徴量に追加

data	mean_temp	max_temp	min_temp	mean_rain	max_rain	min_rain	mean_sun	max_sun	min_sun
9003	20120107	-1.7	1.1	11.2560	16.4	4.4	15.1960	19.3	7.5
9004	20120108	-2.9	1.4	5.42258	13.3	-1.2	9.21290	21.6	0.4

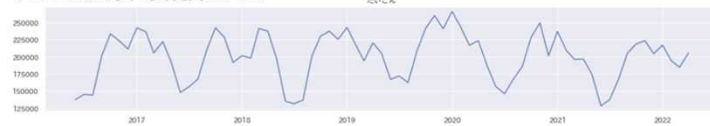
ラグ特徴量を生成

天候データ (その月の平均気温) 等を1,3,6,9,12カ月前のデータをシフトさせラグ特徴量を作った

kind	date	amount	mode_price	area	year	month	mean_mean_temp_1pre	max_mean_temp_1pre	min_mean_temp_1pre	mean_sun_1pre	max_sun_1pre	min_sun_1pre	mean_sun_12pre	max_sun_12pre	min_sun_12pre
だいこん	20051107	201445.0	735.0	千葉	2005	11	18.8387	24.7000	13.4000	6.4000	5.84000	87.0000	0.0	0.0	0.0

1カ月前データ ← 12カ月前データ →

◆だいこんの月平均卸売高のグラフ



◆だいこんの月平均価格のグラフ



- ・価格と卸売高は反比例し、卸売高が増加すると価格はやや下がる傾向があるように見える
- ・また、だいこんは旬が冬であり、夏場に収穫量が減ることもわかる
- ・2017～2018年は冬にもかかわらず、だいこんの卸売高が減り、価格が高騰したことが分かる
- ・天候データは毎年似たようなデータになっており、特に価格との関係性に傾向があるとは言えない

可視化②: 参考情報の分析

天候データのすべての項目を使うべき？

	amount	mode_price	mean_temp	max_temp	min_temp	sum_rain
amount	1.000000	-0.053083	-0.340205	-0.320197	-0.345392	-0.004336
mode_price	-0.053083	1.000000	0.179791	0.185763	0.176084	0.044550
mean_temp	-0.340205	0.179791	1.000000	0.984190	0.985423	0.029609
max_temp	-0.320197	0.185763	0.984190	1.000000	0.946722	-0.008894
min_temp	-0.345392	0.176084	0.985423	0.946722	1.000000	0.066956
sum_rain	-0.004336	0.044550	0.029609	-0.008894	0.066956	1.000000
sun_time	0.015643	-0.030000	0.008379	0.121960	-0.097054	-0.347164
mean_humid	-0.115329	0.183742	0.487886	0.427230	0.537627	0.391975

- ⇒説明変数同士で相関が高い項目がある。
- ⇒気温は「平均気温だけにする」

実施過程とポイント（抜粋）

3. 仮説を元に、「前処理/データ加工」、「特徴量エンジニアリング」を行う！

- ・前工程での「仮説」を元にデータを加工する
- ・欠損値の補完を行い、必要な説明変数（列）を追加・生成する
- ・機械学習に用いるために、データを1つのファイルに結合する
- ・AutoMLツールによっては、そういった作業をある程度自動で行ってくれるものもある

DataRobotへデータ投入

The screenshot shows the DataRobot web interface. On the left, under 'データ品質' (Data Quality), there's a section for '特徴量' (Features) with a list of variables: 'date (E)', 'date (I)', 'date (M)', 'date (F)', 'amount', 'mode_price', and 'area'. A blue box highlights the 'date' variables and a callout says '特徴量を自動生成' (Automatically generate features). On the right, under 'データポイント' (Data Points), there's a section for '特徴量' (Features) with a list of variables: 'mean_temp', 'max_temp', 'min_temp', 'sun_time', 'mean_humid', and 'area'. A blue box highlights the 'mean_temp' variable and a callout says 'カラムごとの集計結果を算出' (Calculate aggregation results for each column).

データの事前処理

欠損・異常データの確認

date	210336 non-null	int64	count	210336 non-null	209969 non-null	209969 non-null	209969 non-null
mean_temp	209969 non-null	float64	mean	2.013413e+07	15.980279	20.117185	11.773782
max_temp	209969 non-null	float64	std	5.198143e+04	8.817219	8.984162	9.196040
min_temp	209969 non-null	float64	min	2.054111e+07	-16.000000	-8.000000	-24.500000
sun_time	209969 non-null	float64	25%	2.059051e+07	8.300000	12.900000	4.000000
mean_humid	209969 non-null	float64	50%	2.013110e+07	16.300000	20.900000	12.100000
area	209934 non-null	float64	75%	2.010851e+07	22.900000	27.400000	19.700000
	210336 non-null	object	max	2.022103e+07	33.900000	41.100000	30.500000

trainデータ・testデータの野菜を確認

```
df_price['kind'].value_counts().keys()
Index(['キャベツ', 'レタス', 'ほうきい', 'ねぎ', 'きゅうり', 'トマト', 'だいこん', 'ピーマン', 'にんじん', 'たまねぎ', 'ほうれんそう', 'かぼちゃ', 'なましいたけ', 'かぶ', 'さといも', 'セリリー', 'ミニトマト', 'いんげん', 'とうもろこし', 'えだまめ', 'そらまめ', 'ふか', 'さやえんどう', 'レインにがり', 'まつたけ', 'アスパラガス', 'みずな', 'ななめな', 'うめ', 'アスパ', 'なつめ', 'ピーズ', 'なな', 'えのきだけ', 'うど', 'しめじ'], dtype='object')

df_price['area'].value_counts().keys()
Index(['全国', '北海道', '東北', '関東', '中部', '関西', '四国', '九州', '沖縄', '北海道', '青森', '岩手', '秋田', '山形', '福島', '茨城', '栃木', '群馬', '埼玉県', '千葉県', '東京都', '神奈川県', '新潟県', '富山県', '石川県', '福井県', '山梨県', '長野県', '岐阜県', '静岡県', '愛知県', '三重県', '滋賀県', '京都府', '大阪府', '兵庫県', '奈良県', '和歌山県', '徳島県', '香川県', '愛媛県', '高知県', '福岡県', '佐賀県', '大分県', '熊本県', '鹿児島県', '沖縄県'], dtype='object')
```

特徴量エンジニアリング

```
df_price['area'].value_counts().keys()
Index(['全国', '北海道', '東北', '関東', '中部', '関西', '四国', '九州', '沖縄', '北海道', '青森', '岩手', '秋田', '山形', '福島', '茨城', '栃木', '群馬', '埼玉県', '千葉県', '東京都', '神奈川県', '新潟県', '富山県', '石川県', '福井県', '山梨県', '長野県', '岐阜県', '静岡県', '愛知県', '三重県', '滋賀県', '京都府', '大阪府', '兵庫県', '奈良県', '和歌山県', '徳島県', '香川県', '愛媛県', '高知県', '福岡県', '佐賀県', '大分県', '熊本県', '鹿児島県', '沖縄県'], dtype='object')

df_price['area'] = df_price['area'].replace(japan_dic)
df_price['area'].value_counts().keys()
Index(['全国', '北海道', '東北', '関東', '中部', '関西', '四国', '九州', '沖縄', '北海道', '青森', '岩手', '秋田', '山形', '福島', '茨城', '栃木', '群馬', '埼玉県', '千葉県', '東京都', '神奈川県', '新潟県', '富山県', '石川県', '福井県', '山梨県', '長野県', '岐阜県', '静岡県', '愛知県', '三重県', '滋賀県', '京都府', '大阪府', '兵庫県', '奈良県', '和歌山県', '徳島県', '香川県', '愛媛県', '高知県', '福岡県', '佐賀県', '大分県', '熊本県', '鹿児島県', '沖縄県'], dtype='object')
```

重要度の可視化

feature	importance
10 mean_humid	0.142831
2 year	0.135940
9 sun_time	0.128305
7 min_temp	0.111918
8 sum_rain	0.108939
6 max_temp	0.106145

feature	importance
4 day	0.093855
3 month	0.085475
5 mean_temp	0.082495
1 area_x	0.004097
0 kind	0.000000

- Pythonを用いて機械学習を行う（LGBMを用いる人が多かった）
- AutoMLツールは、最適なアルゴリズムを用いてモデルを自動生成してくれるほか、どの変数がモデルに影響を与えたのかを確認することが出来るものもある

The screenshot shows the 'Mod Manager' application interface. At the top, there's a header with 'Mod Manager' and a description: 'モジュールのインストール・管理・更新が可能なツールです。' Below this, there are two red circular icons with white symbols. The main section is titled 'インストール済み' (Installed) and lists several mods. Each mod entry includes a green checkmark icon, the mod name, version, and a brief description. The mods listed are:

- Electric Net Expressor (v1.0)**: 電気のネットワークを拡張し、電気のネットワークを拡張します。
- Electric Net Expressor (v1.0)**: 電気のネットワークを拡張し、電気のネットワークを拡張します。
- Generalized Addressable 1.0**: 電気のネットワークを拡張し、電気のネットワークを拡張します。
- Light Gradient Boosting (v1.0)**: 電気のネットワークを拡張し、電気のネットワークを拡張します。
- Light Gradient Boosting (v1.0)**: 電気のネットワークを拡張し、電気のネットワークを拡張します。

[illegible]

2022/4美国、字测推移

凡例

- model02_p50
- model01_p50
- model_genso

date

機械学習のワークフローを自動化する
Python製のオープンソース、ローコード機械
学習ライブラリ。

Model	MAE	MSE	RMSE	MAPE	RMSE	MAPE	17.5mm
1st Bayesian Ridge	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
2nd Random Forest Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
3rd Gradient Boosting Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
4th Extra Trees Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
5th Least Angle Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
6th Ridge Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
7th Lasso Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
8th Light Gradient Boosting Machine	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
9th XGBoost	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
10th CatBoost	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
11th Stacking Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
12th Boosting Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
13th Ensemble Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
14th Linear Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
15th Ridge Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
16th Lasso Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
17th Elastic Net Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
18th Support Vector Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
19th Kernel Ridge Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
20th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
21th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
22th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
23th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
24th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
25th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
26th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
27th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
28th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
29th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70
30th Gaussian Process Regression	164.1504	36322.0000	187.8408	3.2785	3.1875	3.4375	0.70



実施過程とポイント（抜粋）

5. 生成したモデルで予測し、結果を検証する！

- ・予測結果と実績値を比較し、予測精度を見る
- ・精度が悪い場合は、何が影響したのか、データの加工で改善できそうかなどを検討したり、違うアルゴリズムを用いて学習させてみたりする



RMSPEの値：平均二乗パーセント誤差の平方根
(16品目の平均値)
※参考値： ルールベース（前月平均をそのまま用いたもの）
は、0.328012（でした）

	CASE1: 2016年5月～2022年3月の データで学習	CASE2: 2016年～2022年の3月 ～5月のデータで学習	CASE3: 2022年1月～2022年3月 のデータで学習
SageMaker Canvas（天候データを結合した場合）	0.369154	0.238671	学習不能
SageMaker Canvas（天候データを結合しない場合）	0.233046	0.224064	学習不能
LightGBM（天候データを結合した場合）	0.093813	0.079617	0.351693
LightGBM（天候データを結合しない場合）	0.150142	0.113623	0.351693

予測値と実績値の確認

予測結果

	A	B	C	D	E	F	G	H
row_id	date	date (月)	date (日)	amount	area		prediction	mode_price
1	0	2013	3	19	111763	全国_宇都宮_熊本	1486.720801	1470
2	1	2021	8	19	152024	東京_全国_福島	1299.471493	1296
3	2	2021	6	24	182109	全国_宇都宮_熊本	1154.505931	1188
4	3	2020	8	27	251022	東京_全国_青森	1643.25711	1620
5	4	2012	1	28	101583	全国_名古屋_熊本	1504.898404	1579
6	5	2014	1	31	88094	全国_名古屋_熊本	1422.965318	1365
7	6	2007	4	21	116165	千葉_全国_宇都宮	1329.830349	1365
8	7	2018	11	21	88463	千葉_全国_熊本	1661.734575	1620
9	8	2012	6	23	169756	全国_宇都宮_水戸	1218.43079	1365
10	9	2017	10	23	113791	千葉_全国_水戸	1599.636886	1512
11	10	2015	7	11	138231	東京_全国_青森	1253.80595	1188
12	11	2009	4	30	169987	全国_宇都宮_熊本	1452.96007	1470
13	12	2016	7	26	127311	東京_全国_青森	1223.740274	1188
14	13	2021	11	18	127678	全国_名古屋_熊本	1899.809293	1944
15	14	2009	7	10	93812	全国_青森	978.7013647	997.5
16	15	2015	3	31	102092	全国_宇都宮_熊本	1842.288866	1836
17	16	2017	7	12	117615	東京_全国_青森	1158.353016	1188
18	17	2012	4	19	127266	全国_宇都宮_熊本	1886.949722	1785
19	18	2006	8	2	130703	全国_前橋_青森	1277.630688	1365
20	19	2021	9	27	163828	東京_全国_福島	2083.070884	2052
21	20	2016	7	8	153898	東京_全国_青森	1240.323086	1296
22	21	2009	7	3	145760	全国_宇都宮_熊本	975.5267013	945
23	22	2012	4	21	114815	全国_宇都宮_熊本	1877.685603	1785

Mode_priceは手で入力し、モデル精度を確認

予測モデル作成

予測結果:モデル①とモデル②の日別の誤差率

date	mode_price	mdl1-error-ratio	mdl2-error-ratio
2022/04/01	864	11.8%	-1.4%
2022/04/02	864	3.8%	3.6%
2022/04/04	864	-7.6%	-4.0%
2022/04/05	918	-6.9%	-8.1%
2022/04/07	864	9.4%	7.3%
2022/04/08	864	23.7%	9.3%
2022/04/09	864	-2.2%	8.0%
2022/04/11	972	0.9%	-4.0%
2022/04/12	1,080	-6.4%	-6.0%
2022/04/14	1,080	-10.9%	-17.8%
2022/04/15	1,188	-10.0%	-18.6%
2022/04/16	1,512	-37.0%	-40.3%
2022/04/18	1,080	-25.1%	-6.9%
2022/04/19	1,080	0.4%	-15.4%

⇒モデル①②ともに大まかな推移は実績と大きく外れていない。
ただ実績で大きく価格が変動した日(4/16など)を拾いきれない。
また、説明変数の多いモデル②の方が常に精度が良いわけではない。

各メンバーの気づき

プロセス	気づき
ビジネスゴール・目的設定	関係者にAIプロジェクトを理解いただくために機械学習プロセスの説明、PJT計画書を作成する必要がある →どんな工程があるか？どれくらい大変か？スケジュール感はどれくらいになりそうか？を認識合わせできる
データ前処理・特徴量エンジニアリング	データ分析するアプローチに絶対的な正解は存在しない。とにかく色々なことを試してみても多角的な視点を持てる材料をそろえることが重要と感じた
	ノーコードと言いつつ、そのままのデータではうまく学習に使えなかったため、結局Pythonによるデータの加工(前処理)の実施が必要となった
	データを眺めてみるだけである程度傾向や特徴がわかる。とにかくデータを投入するのではなく、データを観察・分析し、取捨選択を行うことが予測精度の向上につながる
	AIモデルを作る前段で仮説を立てるために基本的な統計知識、ドメイン知識が必要 これらの知識レベルによりモデルの精度に大きく影響する。
学習・チューニング・評価	AutoML系のツールでは、いくつかの項目を指定するだけで最適なアルゴリズムで学習してくれるので、学習時はこういったツールを活用も検討していくべきである

総括

「気づき」が多かったのは？

データ前処理プロセス、特徴量エンジニアリングプロセス

その理由はなんだろうか？

学習&予測に必要なデータを揃えるには、まず予測対象/予測材料のデータの分析を行い、仮説を立て、その上で必要なデータの収集、結合、加工が必要になる。

そのため、前処理プロセスに時間をかけた人が多かったのではないかな。

このことから言えることは？

- 収集データをどのように整形するか？が予測精度に大きく影響するため、最も人が考え、手を動かして取り組むべき重要なプロセスである。
- またこのプロセスを進めるには、データ加工のスキル、統計などデータ分析の基本的な知識が必要になる。
- データを収集・加工・分析してみないと、「要件を実現できるものか？」を判断しづらいが、実際に手を動かしてデータを見れば、ある程度「実現可能性」が見えてくる

今回、機械学習のプロセスを一通り実践し、各フェーズで重要な点を「気づき」として理解・整理することができた。

また、特にデータ分析のプロセスにおいて、統計などの知識や収集したデータを加工・可視化する手法の習得が必要だと体感することができた。

今後AI・機械学習を活用する際には、関係者にこのプロセスを説明し理解いただきつつ、おさえるべきポイントに注意しながら我々が主体となって推進していきたい。

AI研究会 2022

ビジネス実践分科会

Research Question

チーム

ビジネス実践

議論
テーマ

Q. 成功/失敗するAIの特徴とは何なのか？

～過去から現在、未来のAIトレンドから学ぶ、成功/失敗するAIの特徴（ROIを含む実装のポイント等）整理
とチームメンバーの事業領域を検証フィールドとしたビジネス実証～

検討
プロセス

最先端のAI事例や過去の事例等を基にAIトレンドを整理、実フィールドでの事例仮説検証までの実施を目指す！

STEP①

最先端/これから流行る未来のAI調査

- 最先端応用事例
- 先端ファウンデーションモデル
(GPT-3, DALL-E的なモデル)

事例・文献調査

特徴整理(トレンド)

STEP②

過去/現在のAI調査

- AI導入効果
- AI導入における課題
- AI導入事例

事例・文献調査

効果・課題整理

STEP④

OUTPUT整理

- AI技術トレンドの遷移
- 過去と未来のAI事例/特徴
- 実FieldでのPoC実施結果

STEP③

実ビジネス領域でのAI実践

- チームメンバの事業領域にて、調査を通じて得られた仮説を検証

OUTPUT
イメージ

AIトレンド、事例

過去
AI現在
AI未来
AI

AI技術と実用化事例の変遷を整理

- 事例・文献調査
- 効果・課題整理
- 事例・文献調査
- 効果・課題整理
- 事例(予想含)
- 特徴(トレンド)

実ビジネス領域でのAI実践結果

事例調査を通じて得られた
AI適応ビジネス仮説

チームメンバの事業領域で
実際にAIを使い仮説を検証

今回取組の纏め

(1) 過去/現在から最先端、これからAI進む方向性

□ 今回活動の結果纏めは以下の通り。

調査結果サマリ

過去/現在のAI

- AI導入にて**成功/失敗を左右するポイント**は以下3点、ビッグデータ/データ活用時代から言われてきたものと同様
 - ①**AI活用の人材・体制**
 - ②**AI導入におけるROI**
 - ③**情報セキュリティ**
- 課題としては新規性はなく、AIに限らず、ビッグデータ/データ活用の時代から言われていたものと感じる。この壁を乗り越えた企業が、コア事業/業務に当たり前のようAIを活用して、競争力を増すことになるだろう。
- 既に、既存AI活用成功のコツは世の中でフレームカされつつあるため、調査し特徴を学ぶことが非常に重要

最先端/未来のAI

- 画像生成AI(Dall-E2、Stable Diffusion、Muse等)と自然言語生成AI(ex.ChatGPT、Bard等)を含み、**生成系AIの進化が凄い**
- 100%AIにお任せはまだ早い、アイデア発想や文章/コードドラフト作成など、**人間の作業支援には十分活用可能**
- タスク特化型AIの精度向上が進展する一方で、**汎用型AI(Gato等)の研究も進む**
- 少し先の未来だが**量子技術活用で学習効率が大幅にUPするかも？**

実ビジネスでのAI実践

仮説：生成系AIを業務に活用出来ないか？

<検証シナリオ>

- ①PythonプログラミングへのChatGPT活用
- ②SQLプログラミングへのChatGPT活用
- ③オフィスデザインへのDall-E活用

<検証結果>

- プログラミング業務やデザイン業務へは**人間支援の位置づけで十分に活用可能なレベル**
- 出力の信頼性は必ずしも高くなく、**AIを扱う側の人間の知識やAIを扱うためのテクニックが必要**

これからAIが進む方向性(チーム内予想)

- 特定領域の識別、予測、会話、自動化から、**従来困難と考えられていた創造性が必要となる業務へAI活用のスコープが拡張**
- 例えば、**デザイン業務はAIが自動生成、システム開発は要件だけ伝えればAIが自動的に開発してくれる時代**が来るかもしれない
- また、量子技術活用により大量データの高速学習も可能となり、**物理解析のような分野へもAIがより普及していくもの**と考える
- 一方で、**AIを利用する人間側の知識もより一層重要**となり、AIによる事故やAI活用の倫理にはより一層留意が必要
- 民間企業では、新しいAI技術を業務にうまく活用できるか否かにて、**企業間競争力の差がより一層顕著になるもの**と考える

(2) 成功/失敗するAIとは？

□ ビジネス実践チームが考える特徴は以下4点。

成功するAIの特徴とは？（考察）

- ① 以下の①～③を十分に検討出来ていること
①AI活用の人材・体制、②AI導入におけるROI、③情報セキュリティ
- ② 世の中に溢れる過去事例やフレームワークから学ぶこと
- ③ 技術進展速度が速い中、トレンドをいち早く捉え、柔軟に吸収できること
- ④ AIの技術力だけでなく、使う側の人間に知識があること

(2) 成功/失敗するAIとは？ ～ChatGPT版～

- ChatGPTに対して、ビジネス実践チームのResearch Questionに対する回答を聞いてみた。

KA

～過去から現在、未来のAIトレンドから学ぶ、成功/失敗するAIの特徴（ROIを含む実装のポイント等）整理とチームメンバーの事業領域を検証フィールドとしたビジネス実証～というテーマのもとで、成功/失敗するAIの特徴とは何なのかを教えてください。



成功するAIの特徴：



1. 具体的な目的：AIには明確な目的が設定されていることが重要です。
2. 質の良いデータ：質の良いデータがあることが重要です。これにより、モデルの正確さが向上します。
3. 評価手法：適切な評価手法が選ばれていることが重要です。
4. 実際の利用者の要件：実際の利用者の要件に基づいて設計されていることが重要です。
5. 適切なリソース：適切なリソースが投入されていることが重要です。

失敗するAIの特徴：

1. 模糊な目的：AIには明確な目的が設定されていない場合、失敗する可能性が高いです。
2. 悪質なデータ：データが汚染されている場合、モデルの正確性が低下する可能性があります。
3. 不適切な評価手法：適切な評価手法が選ばれていない場合、モデルの正確性が低下する可能性があります。
4. 利用者の要件が考慮されていない：実際の利用者の要件が考慮されていない場合、モデルの使用性が低下する可能性があります。
5. リソース不足：リソースが不足している場合、モデルの性能が低下する可能性があります。

AI研究会 2022

技術理論分科会

リサーチクエスト

AIやデータ分析の結果を納得してもらう
手段は何か。どのように実現できるか。

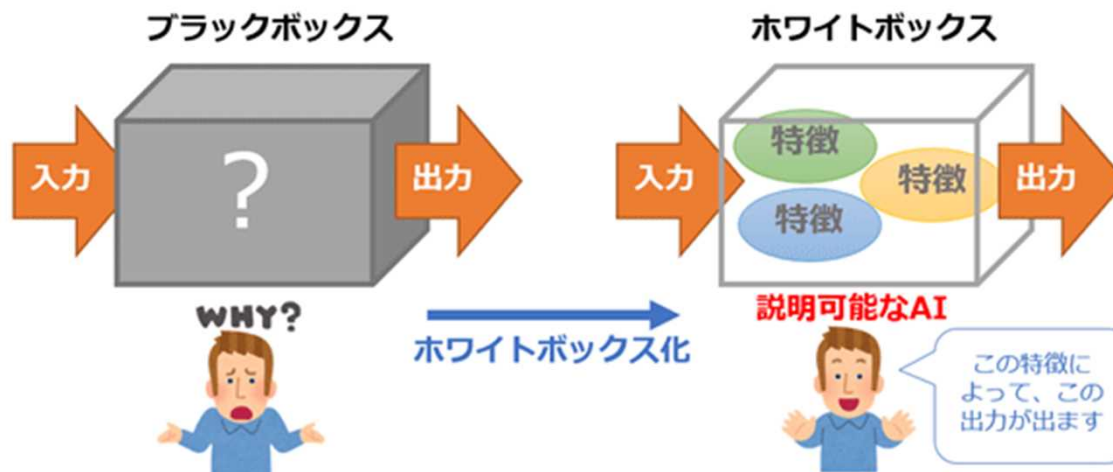


- ✓ XAI（Explainable AI：説明可能なAI）は、どのようなことが実現できるか調査を行う。（どの分野でどんな手法が使われているか）
 - 相手に説明する・納得してもらえる方法は？
 - 予測結果の説明は、精度向上に寄与する可能性もある
- ✓ 各人で実際に手を動かして、チャレンジしてみる。

XAIとは

説明可能なAI（人工知能）（Explainable AI : XAI）とは、言葉通り、予測結果や推定結果に至るプロセスが人間によって説明可能になっている機械学習のモデル（つまりAIの本体）のこと、あるいはそれに関する技術や研究分野のことを指す。ちなみに「XAI」は、米国のDARPA（Defense Advanced Research Projects Agency：国防高等研究計画局）が主導する研究プロジェクトが発端で、社会的に広く使われるようになった用語である。

類義語として、解釈可能性（Interpretability、もしくは解釈性）という用語もある。これも言葉通り、機械学習モデルが予測、推定を行ったプロセスを、人間が解釈可能であるかどうかを指す。



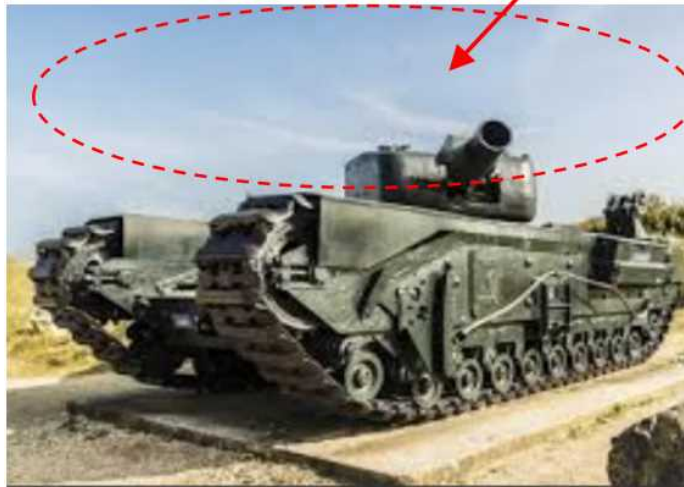
出所) ITmedia <https://atmarkit.itmedia.co.jp/ait/articles/1908/19/news022.html>

説明可能性が求められる例

アメリカ陸軍、敵・味方の戦車を識別するAIの例

- AIは高精度である一方、評価や判断根拠を明確に説明できないことが多い
(**ブラックボックス問題**)
- この事例ではテストにおいて高い精度を上げていたが、本番導入したところ非常に精度が低かった。
- 調査により**訓練データにバラツキがあった**ことが判明。
→ 味方の戦車は“晴れ”の日に撮影されたものが多く、敵の戦車は“曇り”の日に撮影されていた。

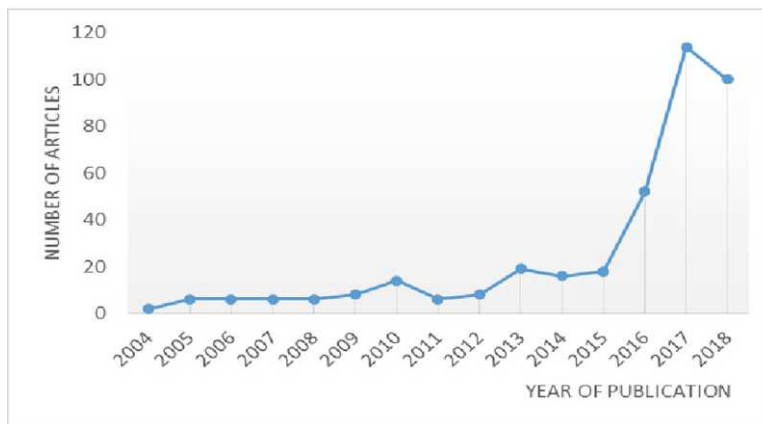
AIは、戦車ではなく、こうした空の様子を見て判断していた



XAI研究の進展



出所) <https://www.slideshare.net/SatoshiHara3/ss-238926593>



有名どころは、「LIME」、「SHAP」、「Grad-CAM」。左図は、XAI関連の論文数であり、爆発的に増加している。

出所) <https://ieeexplore.ieee.org/document/8466590>

今年度の取り組みテーマ（取り組みテーマ）

画像データ

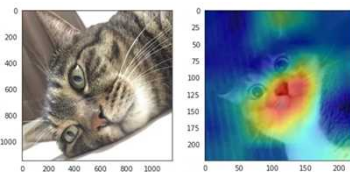
学習済みモデルVGG19を利用し、入力画像に対して推論結果を出力。Grad-CAMにより推論の根拠となるヒートマップ作成。



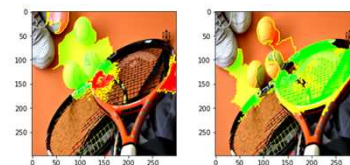
EfficientNetを利用し、Grad-CAMでヒートマップ作成により予測結果に影響を与えている箇所の可視化。



ImageNetで事前学習済みのResNet50にGrad-CAMを実装し、予測に影響を与えた個所の可視化。

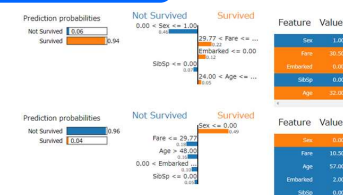


画像認識モデルにLIMEを利用し寄与する領域の可視化。



テーブルデータ

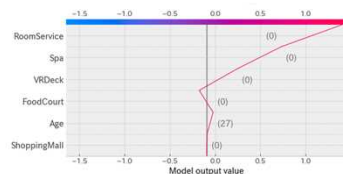
タイタニック号の搭乗者データ (kaggle)を利用し、LIME・SHAPにより生存に影響のある変数の確認。



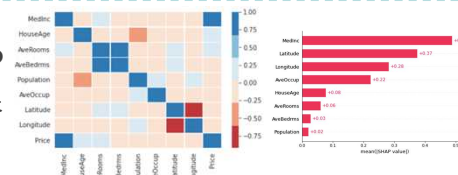
東京都の不動産取引価格の予測モデルを利用し、重要な特徴量の可視化。



SpaceTitanicのデータを利用し、SHAPにより寄与度の可視化。

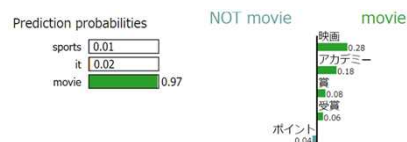


住宅価格データを利用し、SHAPにより相関係数ヒートマップ、結果に影響を与える項目の可視化。



テキストデータ

テキスト分類モデルにLIMEを適用し、影響度の高い単語の可視化。



次ページ以降で、ピックアップし取り組みをご紹介します

今年度の取り組みテーマ（画像データ）

学習済みモデルVGG19を利用し、入力画像に対して推論結果を出力。Grad-CAMにより推論の根拠となるヒートマップ作成。

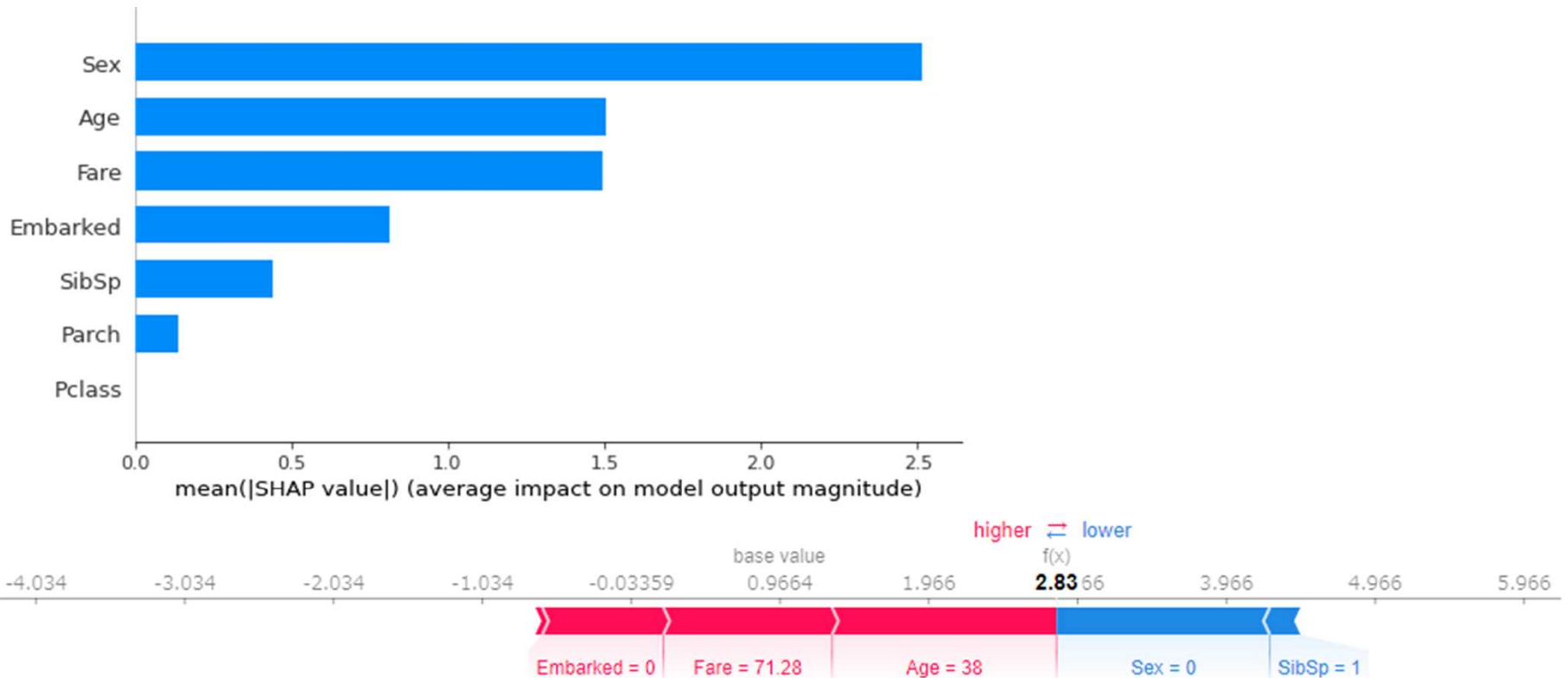
学習済みモデルのVGG19（※辞書形式で1000個のラベルの分類を行うモデル）を利用し、入力画像に対して、予測を行い、推論結果を出力する。併せて、推論の根拠となるヒートマップを可視化し、推論結果及び推論結果に近いクラス4種の併せて、5種の推論根拠となるヒートマップをプロットとして画像保管する。



入力画像はエジプシャン・マウのため、推論結果は適切な判断である。一方でスコアの近いクラス分類ではオオヤマネコや、トラ猫、ジャガー猫などがあり、ヒートマップの参照先では主に“顔”と、体の模様を推論に参照して、猫と判断している。

今年度の取り組みテーマ（テーブルデータ）

タイタニック号の搭乗者データ(kaggle)を利用し、SHAPにより生存に影響のある変数の確認。



上グラフより、生存者の傾向として性別（Sex）と年齢(Age)、運賃(Fare)が大きく影響していることが分かる。

下グラフでは個別のインスタンスに対しても直感的に変数の影響を確認できる。

今年度の取り組みテーマ（テキストデータ）

テキスト分類モデルに対してXAI技術（LIME）を適用し、どのような情報が得られるか確認

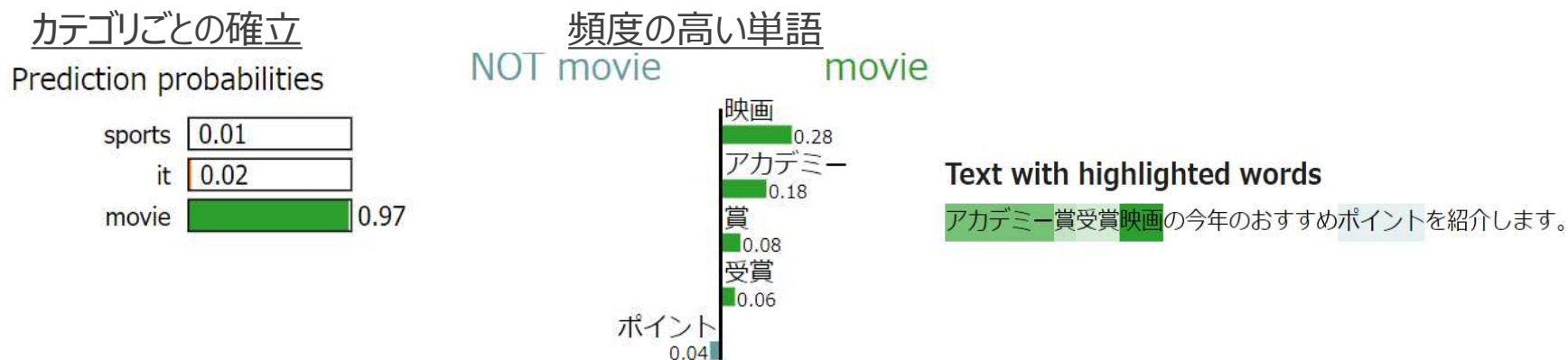
■テキスト分類モデルの概要

BERTをファインチューニングして、入力文章を「sports / it / movie」のカテゴリに分類するモデルを作成。
※ファインチューニングについては、Livedoorニュースコーパスを使用。

入力

アカデミー賞受賞映画の今年のおすすめポイントを紹介します。

入力



入力文章は「 movie 」カテゴリに属すると判定（左図）

その根拠は、映画やアカデミーの単語が大きく影響していることが確認できる（中央・右図）

結論

リサーチクエスト

AIやデータ分析の結果を納得してもらう
手段は何か。どのように実現できるか。



- ✓ **XAI(Explainable artificial intelligence)の技術を利用**することにより、映像や自然言語、数値など様々なデータに適用でき、なぜそうなったのかを確認できるようになる。
- ✓ AIでの処理結果を納得してもらうためには、**何を説明できれば良いか**を目星をつけて、それに**適したXAI手法**を用いる必要がある。
→ その為には、様々なXAI手法を理解り利用できる必要がある。
- ✓ 物体認識や検出ではGrad-CAMを使用すれば、モデルが判断した個所を可視化して確認・説明ができる。また、誤った判断については、モデルの精度向上につなげることができる。

AI研究会 2022

ビジネス理論分科会

リサーチクエスト

定性的な観点から費用対効果が高い、
AIビジネスで共通する特徴にはどのようなものがあるのか？

<前提事項>

- ・「費用」とは、AI導入にかかる初期導入コストと運用コストを指す
- ・「効果」とは、事業会社でAIを導入した場合の利益増大／費用削減の効果に着目する。また、公共の利益も効果とみなす

<リサーチクエスト設定の背景>

分科会メンバーが共通に感じている課題意識として以下がある

- ・AIのPoCは成功事例ばかり喧伝されているが、現実とは乖離を感じている
- ・PoCで終了することも多い
- ・成功例に共通の特徴があれば、自社のビジネスに活かしたい

結論

定性的な観点から費用対効果が高い、
AIビジネスで共通する特徴にはどのようなものがあるのか？



- ✓ DB（データ）が揃っている（アノテーションが不要な状態）
- ✓ 求められる精度が高くない（AIに完璧を求めない）

これらの特徴があれば、費用対効果が高い傾向があると考え
一方で、上記特徴が整っていないとAIビジネスが成功しない
とは断定できない（事例が限定的なため）

研究方法

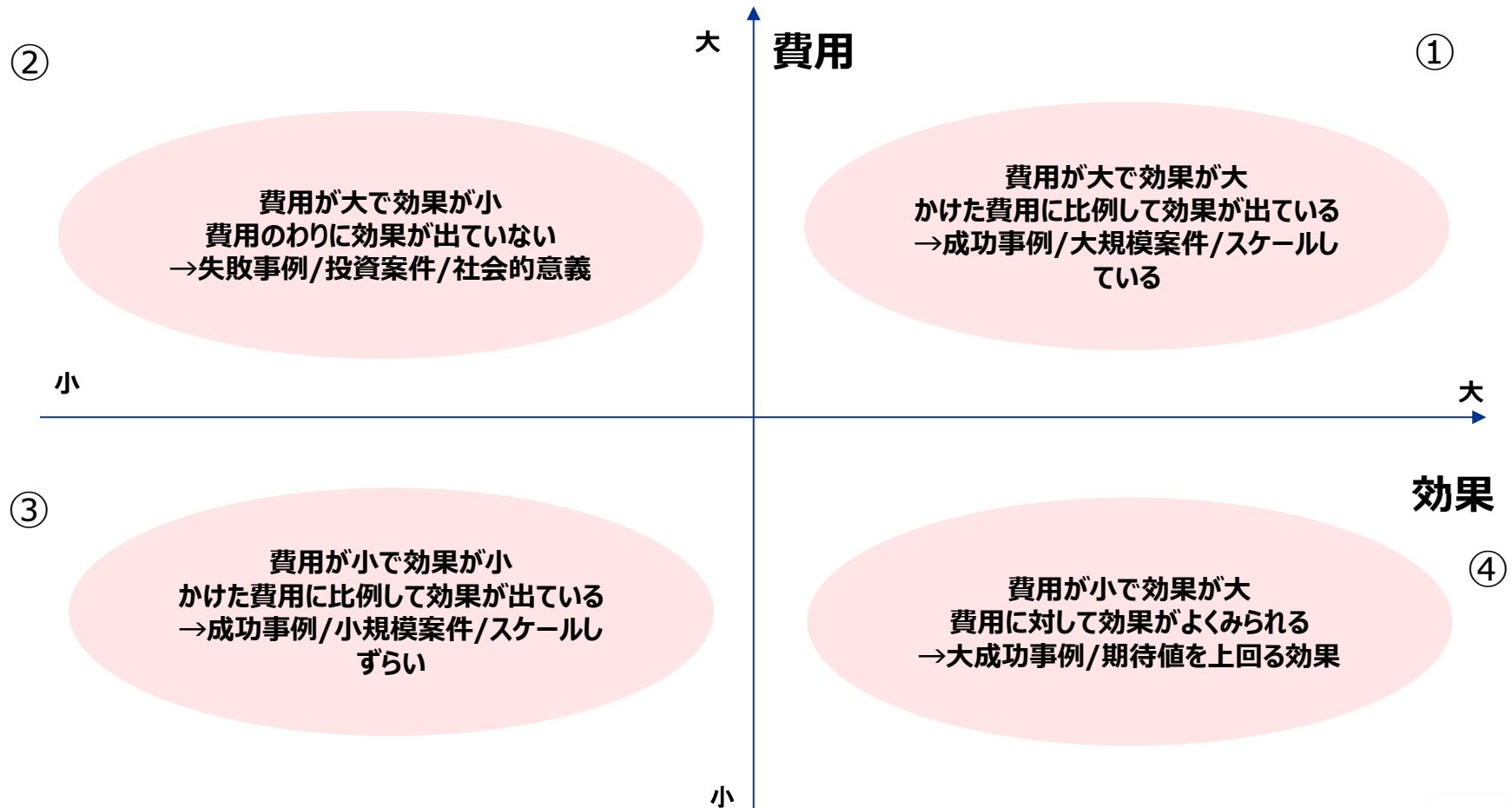
AIビジネス（プロジェクト）の事例を調査し、**定性的**に、費用と効果を推定
費用の大小、効果の大小で4象限に分類して、特徴を抽出した

各象限の特徴を分析するにあたり以下の軸を検討した

軸	分類	分類	分類
新規・代替	新規: 7、9、14、16、19	代替: 1、3、4、5、6、8、10、11、12、 13、15、17、18	新規+代替: 2
業界	物流、水産業(2)、報道機関等 教育、金融・保険(2)、ヘルプデスク 旅行、製造(3)、広告(3)、小売(3) 農業、スポーツ		
技術(データ)	画像: 1、3、5、6、10、13、14、17 動画: 18 自然言語処理: 4、7、11、12、15、 16 音声解析: 2 テーブルデータ: 8、9、19		
海外・国内事例	国内事例: 1-19	海外事例:	
構造化データ・非構造化データ	非構造化: 1、3、5、6、10、13、14、 17、18、4、7、11、12、 15、16、2	構造化: 8、9、19	
回帰(時系列)か・分類か	分類: 1、2、3、4、5、6、7、9、1 1、 12、13、15、16、17、18	回帰: 8、10(金額推計の場合)、14、 19	

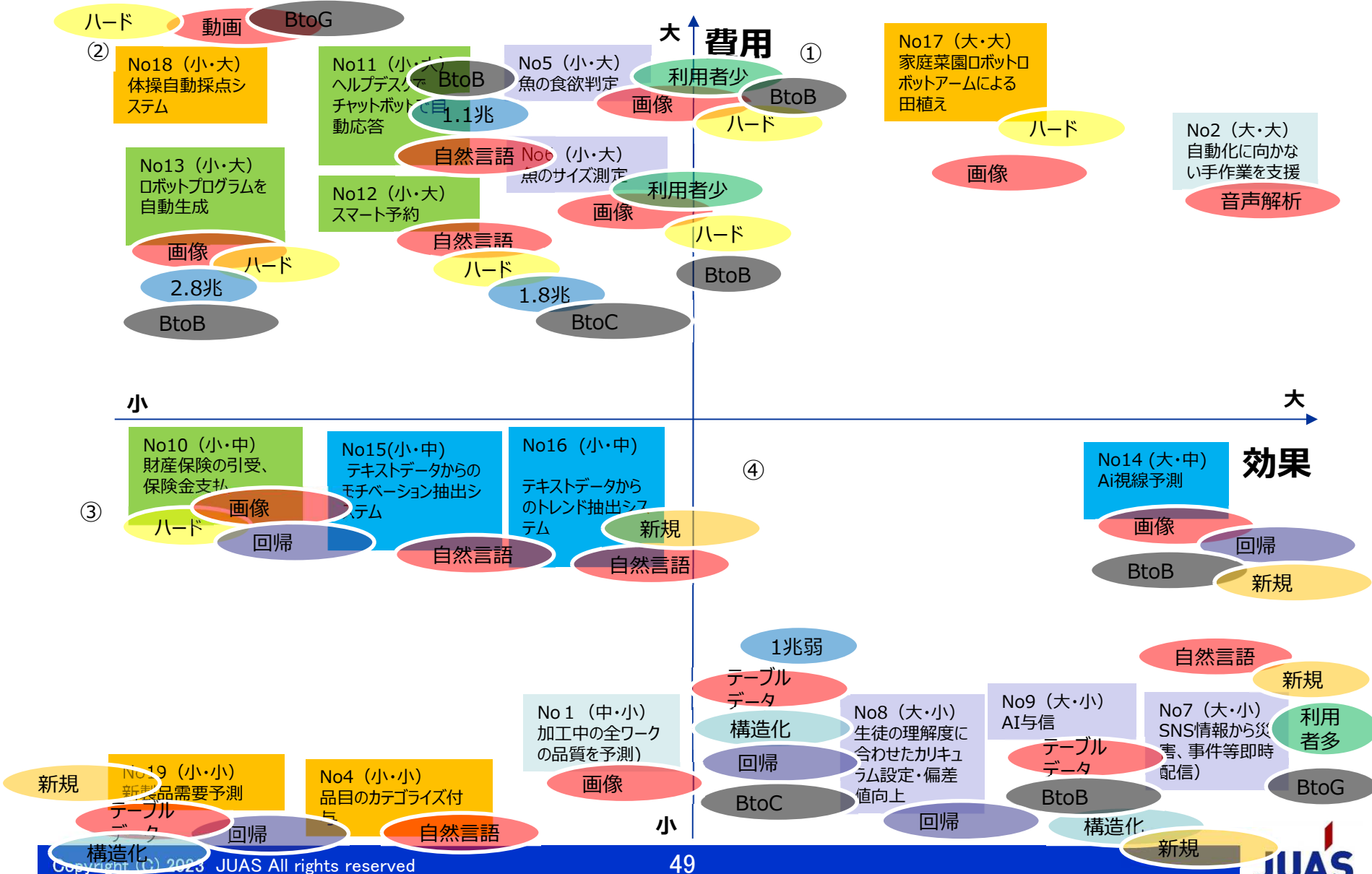
マトリックスの解釈

事例を定性的な観点から効果と費用の大小によって分類した
各象限を ④ > ① > ③ > ② の順に成功度合を順位付けし、特徴を考察した



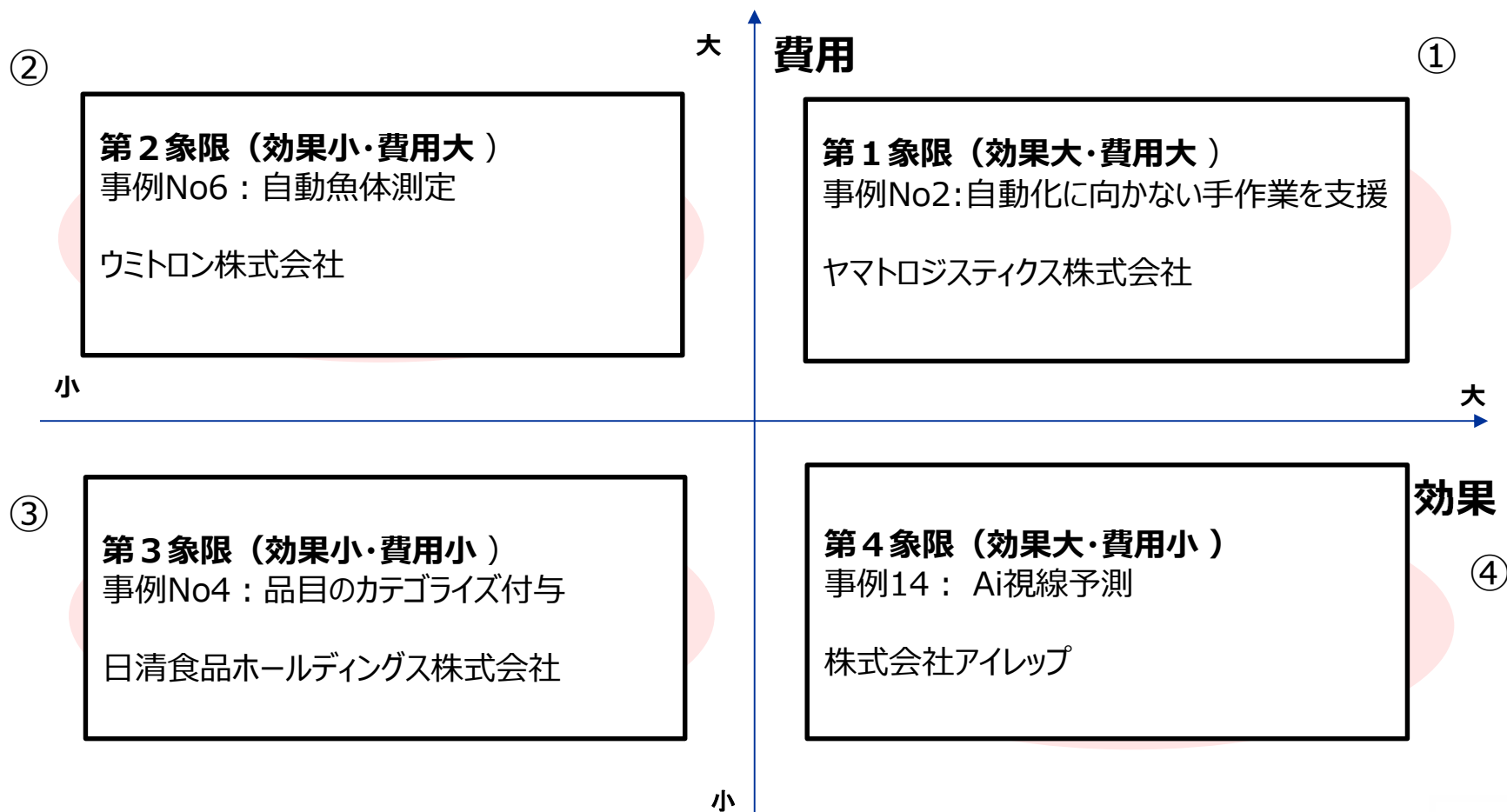
分析マップ（タグあり）

事例No〇（効果・費用）



典型事例（象限毎）

各象限毎の典型事例をピックアップ



費用に関する検討事項（考察）

どのような要素がAIビジネスの費用を下げるのかについて議論した。

「ハードの有無」

「データベースの有無」

上記2点がAIビジネスの費用に大きく寄与すると考えられる。

「当該AI導入のためのハードの有無」

当該AIを導入するためにハードが必要となるかどうかで、AIビジネス自体の費用の大きさが強く影響を受けることが想定される。

「既存のデータベースの有無」

既存のデータを利用できることで、データ収集・データベース整備・アノテーションに必要なコスト（時にハードが必要となる）が削減されるので、既存のデータベース有無はAIビジネスの費用に大きく影響する。

データベースの有無はアノテーションの費用の有無とも考えられる。回帰分析、教師なし学習の場合、アノテーションが不要となるため、それらのアルゴリズムを用いるかどうかAIビジネスの費用に影響をもたらす。

効果に関する検討事項（考察）

どのような要素がAIビジネスの効果を高めるのかについて議論した。

「アウトプットに求められる精度が低いAI：技術代替的なAI」

「（導入企業目線での）伸び代の有無」

上記2点がAIビジネスの効果に大きく寄与すると考えられる。

「アウトプットに求められる精度が低いAI：技術代替的なAI」

人の部分代替ができる技術であれば、導入したAIの精度を人が判断できるため、精度が低いままでもビジネスが業務フローにのりやすく、導入ハードルが低いといえる。

そのためAIビジネスの導入が比較的開発の早い段階から進み、受け入れ後の効果も人との協業により、結果として効果が高く得られると想定できる。

「（導入企業目線での）伸び代の有無」

導入企業側から見て、伸び代が有るAIビジネスほど効果が高いと認識できる。伸び代とは、「市場の大きさ」、「汎用性の高さ」などが挙げられる。

「AIビジネス導入先の市場の大きさ」：

導入先市場の規模が大きく大規模事業者の比率が高い場合、1件のビジネス導入で得られる効果が大きいと判断できる。

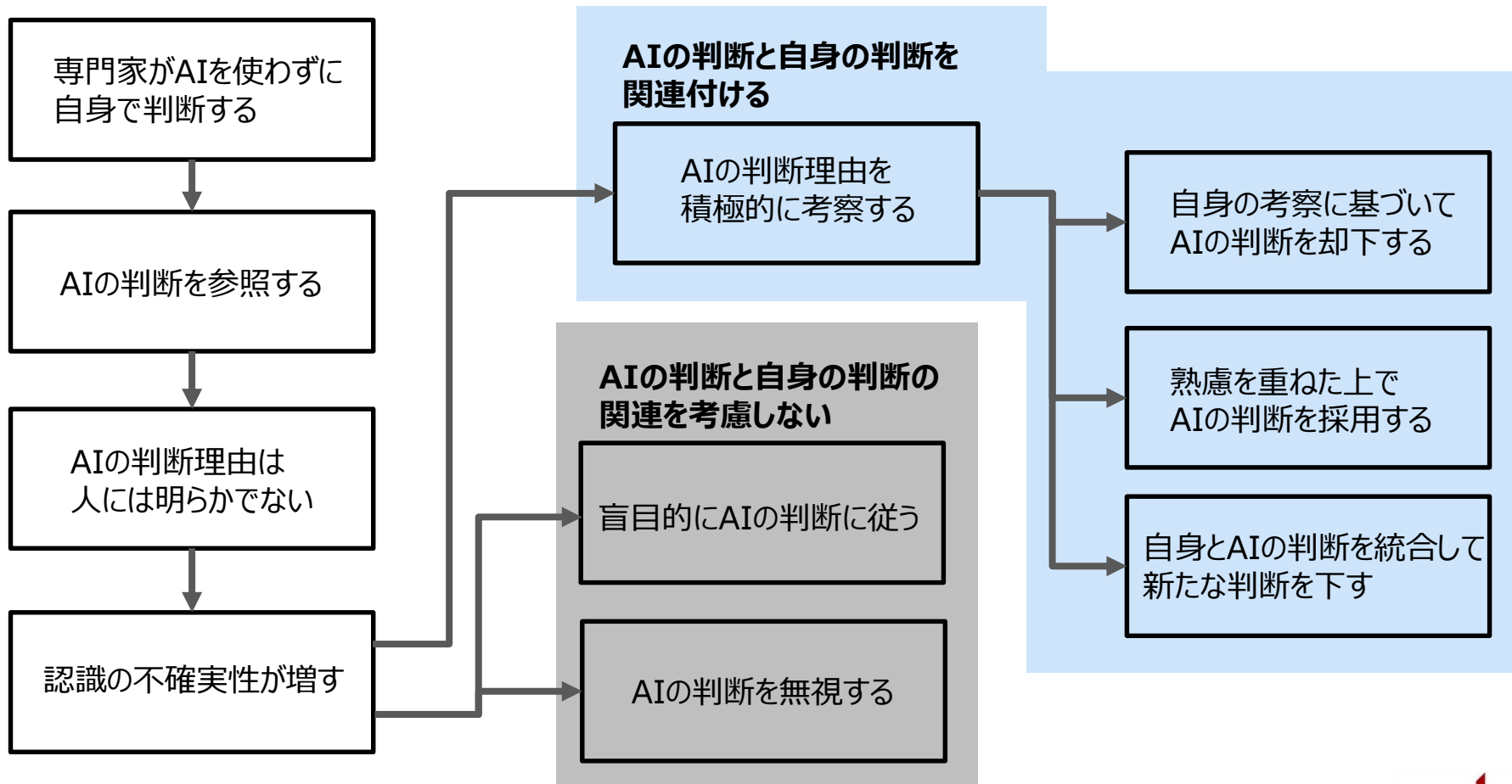
「AIビジネス適用シーンの汎用性の高さ」：

導入した技術の汎用性が高い場合、AIビジネス自体の汎用性も高くなり、1件のビジネス導入で得られる効果が大きくなると判断できる。

AI導入の効果は、ユーザーの使い方にも大きく依存

Sarah Lebovitz, Hila Lifshitz-Assaf, Natalia Levina (2022) , Organization Science 33(1):126-148

レントゲン診断のような、専門的かつ重要度が高い判断にAIツールを導入する場合、その効果はユーザーの使い方にも大きく依存するため、



2023年度AI研究会 みなさまのご参加をお待ちしております！