

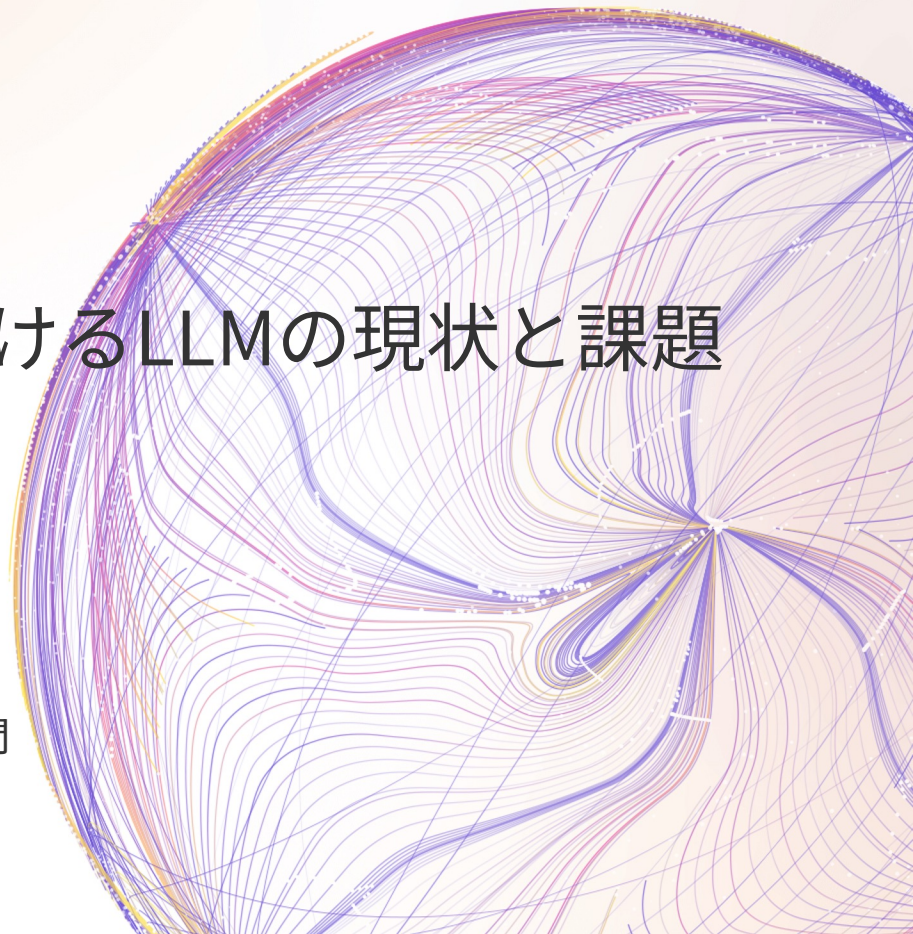
日本のヘルスケア分野におけるLLMの現状と課題

2025年9月11日

SB Intuitions株式会社 事業戦略部

JaDHA(日本デジタルヘルスアライアンス) 特別顧問

碓崎 裕晃



INDEX

目次

1

ヘルスケア分野におけるLLMの現状

2

論文紹介：
業務プロセス全体におけるLLMの性能評価

3

国産LLM（Sarashina）の紹介

- 既にさまざまなベンチマークが提案されている。
- 一般的には大学生～大学院生レベルと同等の性能が出ると言われている。

	名称	概要	引用元
1	MedQA	いわゆる医師国家試験のデータセット	https://github.com/jind11/MedQA
2	PubMedQA	PubMedから作成された質問応答データセット	https://github.com/pubmedqa/pubmedqa
3	JMED-LLM	日本語の医療分野におけるデータセット	https://github.com/sociocom/jmed-llm
4	AnswerCarefully Dataset	安全性・適切性に特化したインストラクションデータ	https://llmc.nii.ac.jp/answercarefully-dataset/

- 2022年のChatGPT登場以降、LLMの性能は上がり続けている。
- ただし、性能評価は主に医師国家試験のような静的な知識テストに依存してきた。

既存のデータセット、
ベンチマーク

モデルの基礎的な能力を評価するためには
極めて重要。

一方で実際の臨床現場における
動的なプロセスには対応できていない。

実際の臨床現場における業務

実際の臨床現場、実際の業務では、
医師や医療従事者は
患者との継続的なやり取りや検査結果に基づ
き、仮説を形成し修正し続ける。

- 臨床現場における全ての業務は繋がっている。
- それぞれのSTEPで医師や医療従事者は、状況に応じて対応を臨機応変に変化させる。



- 静的な性能しか測ることができないという課題感に対しては、GoogleのAMIEやMicrosoftのMAI-DxOに関する先行研究がアプローチしている。
- これらの研究は、静的なタスクから会話的対話や逐次的な情報収集といった動的プロセスへの評価の移行を進めた点で、大きな成果を出している。
- しかし依然として、臨床現場の業務プロセス全体の中で、どの具体的な認知課題がモデルの弱点であるのかを明確には示していない。

	名称	概要	引用元
1	AMIE (Google)	医療面接と診断を支援するための対話型AIシステム	https://research.google/blog/amie-a-research-ai-system-for-diagnostic-medical-reasoning-and-conversations/
2	MAI-DxO (Microsoft)	診断を支援するためのAIツール	https://microsoft.ai/news/the-path-to-medical-superintelligence/

INDEX

目次

1

ヘルスケア分野におけるLLMの現状

2

論文紹介：
業務プロセス全体におけるLLMの性能評価

3

国産LLM（Sarashina）の紹介

- A Novel Framework for Evaluating the Clinical Reasoning Process of Large Language Models: A Comparative Study in Nephrology
 - 引用元 : <https://www.medrxiv.org/content/10.1101/2025.09.04.25334460v1>

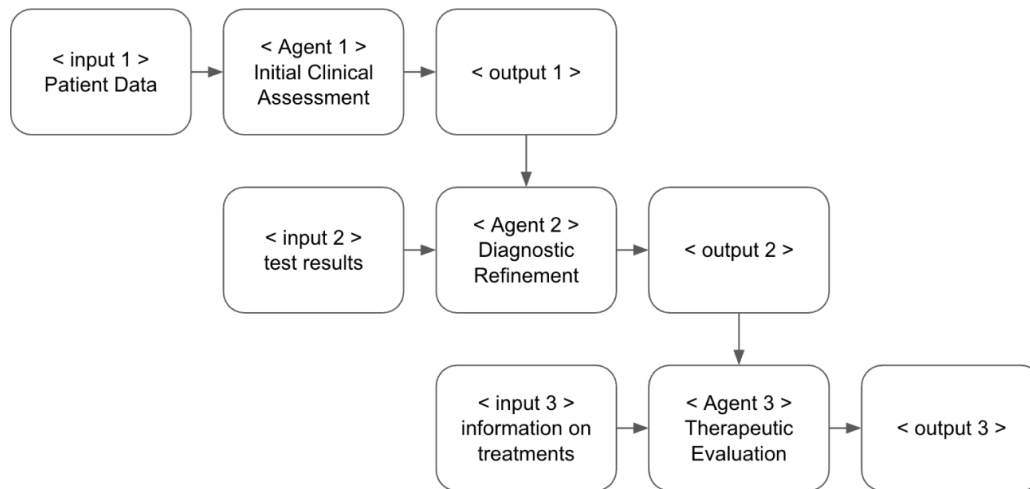
**A Novel Framework for Evaluating the Clinical Reasoning Process of Large
Language Models: A Comparative Study in Nephrology**

Brief title: Evaluating LLM Clinical Reasoning in Nephrology

Yuichiro Yano;^{1,2} Hiroaki Kakizaki;³ Hajime Nagasu;⁴ Seiji Kishi;⁴ Takeo Koshida;⁵
Yoshihito Nihei;⁵ Akira Hirano;⁴ Masaomi Nangaku;⁶ Hirotake Mori;¹ Toshio Naito;¹
Mizuki Ohashi;¹ Shoichi Maruyama;⁷ Isao Matsui;⁸ Yoshitaka Isaka;⁸
Yusuke Suzuki;⁵ and Naoki Kashihara⁴

- AIの最終的な出力だけでなく、「臨床推論そのもののプロセス」を専門医が精査した。
- 各症例を「病歴」「検査結果」「その後の臨床経過」の3段階に分けて評価を実施。

System Overview Diagram of the Multi-Agent Clinical Reasoning Process



- 本研究では、4つのモデルを対象に評価を実施した。
- 業務の3段階に合わせてエージェントを構築し、9つの質問を解かせてその性能を評価した。

Agent 1: Initial Clinical Assessment

1. 患者の医療上の問題の体系的要約
2. 鑑別診断とその根拠
3. 作業仮説を検証するために必要な追加の質問、または診察の一覧
4. 推奨される診断検査（その目的と期待される結果を含む）

Agent 2: Diagnostic Refinement

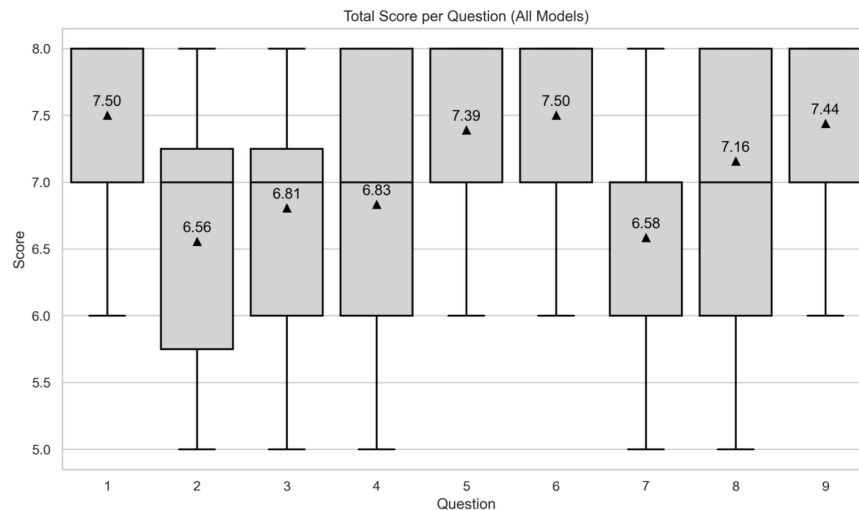
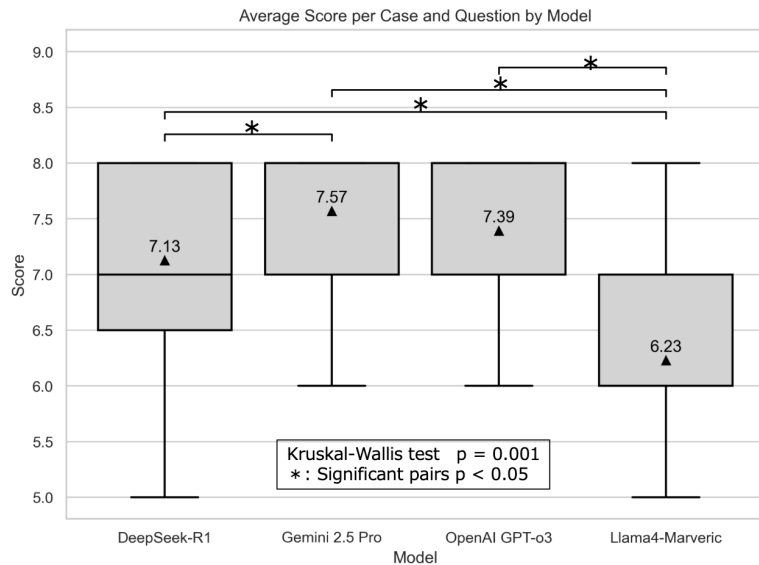
5. 検査結果を解釈し、重要な特徴を強調する
6. この段階でもっとも可能性が高い診断を示す
7. これまでの情報を総合して治療計画を提案する

Agent 3: Therapeutic Evaluation

8. 治療の効果や安全性を見守るための方法をまとめる
9. 患者の状態が改善しない場合や予想外に悪化した場合に備えた対応策を提案する



- 4名の腎臓専門医が評価し、1名の盲検化された研究者によって集計・解析された。
- 各質問は3点尺度で採点（0 = 誤り、1 = 妥当ではあるが最適ではない、2 = 正しい）



論文紹介 モデル別、質問別の結果

- 4つのLLM間で全体的なパフォーマンスに統計的に有意な差が認められた。
- さらにプロセス全体の中でも特に生成AIが弱い部分が明らかになった。

<STEP 1>

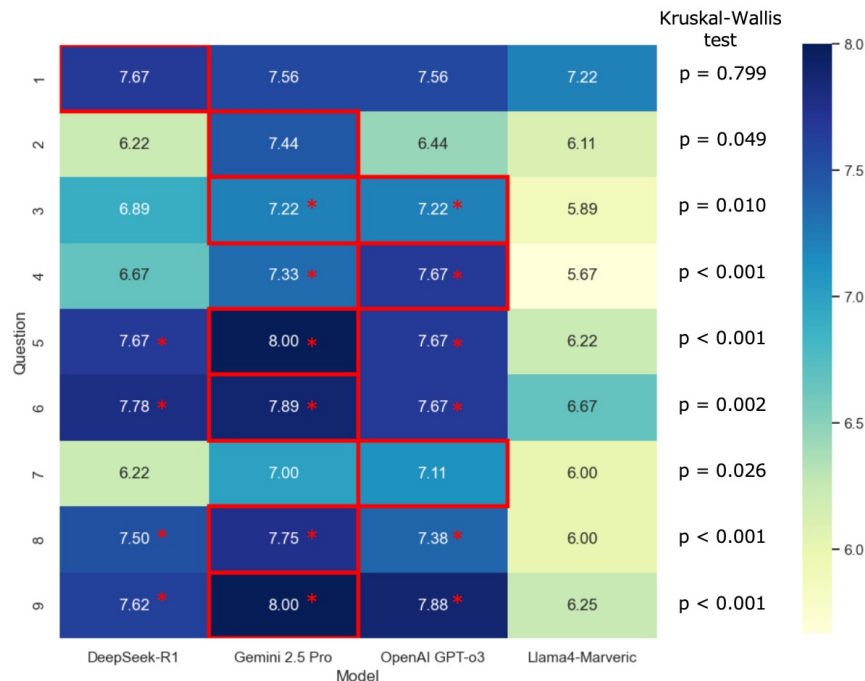
1. 患者の医療上の問題の体系的要約
2. 鑑別診断とその根拠
3. 作業仮説を検証するために必要な追加の質問または診察の一覧
4. 推奨される診断検査（その目的と期待される結果を含む）

<STEP 2>

5. 検査結果を解釈し、重要な特徴を強調する
6. この段階でもっとも可能性が高い診断を示す
7. これまでの情報を総合して治療計画を提案する

<STEP 3>

8. 治療の効果や安全性を見守るための方法をまとめる
9. 患者の状態が改善しない場合や予想外に悪化した場合に備えた対応策を提案する



- 一連のプロセスの中で、特にLLMが弱いのは「質問2：鑑別診断とその根拠」および「質問7：治療計画」だった。
- 最新のLLMは、情報の要約や、検査結果の解釈といった個別タスクには優れていても、最適な介入を計画するような、より高度で統合的なスキルには依然として能力のギャップがあることを示している。
- 本研究ではRAGや、ドメイン特化モデルは利用していない。ヘルスケアなどの専門性が高い分野においては、モデル単体の利用では限界があることが改めて確認できた。

INDEX

目次

1

ヘルスケア分野におけるLLMの現状

2

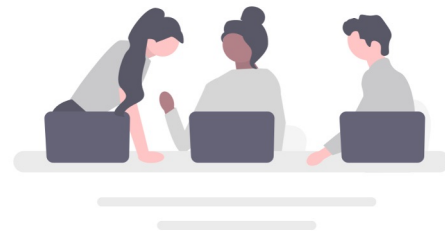
論文紹介：
業務プロセス全体におけるLLMの性能評価

3

国産LLM、Sarashinaのご紹介

SB Intuitions

- 日本語性能の高い大規模基盤モデルの研究開発
- 日本の文化・商習慣にあった安心安全なAIサービスを提供



生成AIの開発経験のあるエンジニアがグループ各社から集結

社名	SB Intuitions株式会社（英文：SB Intuitions Corp.）
----	---

所在地	東京都港区海岸一丁目7番1号
-----	----------------

設立年月日	2023年3月27日（社名変更日：2023年8月1日）
-------	-----------------------------

事業内容	日本語に特化した国産の大規模言語モデルの研究開発、生成AIサービスの開発、販売、提供
------	--

代表者	代表取締役社長 兼 CEO 丹波 廣寅
-----	---------------------

株主	ソフトバンク株式会社 100%
----	-----------------

AI戦略を積極投資

LLMの内製開発への取り組み、ビジネス応用まで広くカバー

ソフトバンクグループ

AGI革命を加速する群戦略
AIへの積極投資を実行

クリスタル・
インテリジェンス



連携

ソフトバンク

AGI革命を加速する群戦略
最先端のAI計算基盤を世界
初導入

国内最大級の
AI計算機版の整備



連携

SB Intuitions

国産LLMの自社開発
国内最大級モデルの実現

日本語に特化した
大規模言語モデルの構築



連携

Gen-AX

生成AI時代における
トランスフォーメーションを
支援

生成AIの社会実装
(SaaS/コンサル)



ソフトバンク社の国内最大級のAI計算基盤構築

SB Intuitions

- 1兆パラメータの基盤モデル開発に向け、国内最大級※のAI計算基盤を構築

FY23 10月 稼働開始

FY24 10月 稼働開始

FY25 7月 稼働開始

世界初導入

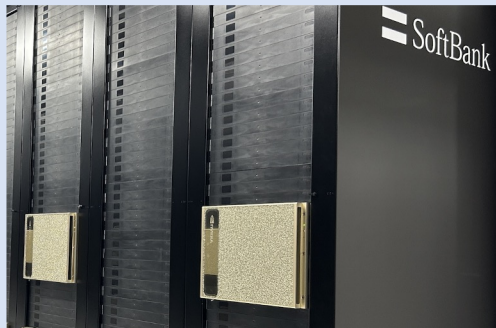
NVIDIA DGX™ A100

約2,000GPU (0.7EFLOPs)



NVIDIA DGX™ H100

約4,000GPU (4EFLOPs)



NVIDIA DGX™ B200

約4,000GPU (9EFLOPs)



累計

約2,000GPU



約6,000GPU



約10,000GPU

※現在稼働している「NVIDIA Blackwell GPU」を搭載した「NVIDIA DGX SuperPOD」（AI計算基盤）として世界最大。ソフトバンク調べ（2025年7月23日時点）

日本語国産LLM(Sarashina)



Sarashina

SB Intuitions株式会社が開発するAIモデルシリーズ。

日本語に特化したモデルとして開発されており、
日本語の理解や生成において高い性能を発揮する。

日本語で生成AIを作る意義

言語特性への対応



『漢字・ひらがな・カタカナ共存』や
『敬語・方言のニュアンス』を学習し、
汎用モデルでは捉えられない
日本語の複雑さを自然に反映できる。

文化・文脈への理解



日本語の表現には、
文化的背景や暗黙の理解が
(「忖度」「空気を読む」など) 密接に関わる。
こうしたニュアンスをより正確に反映できる。

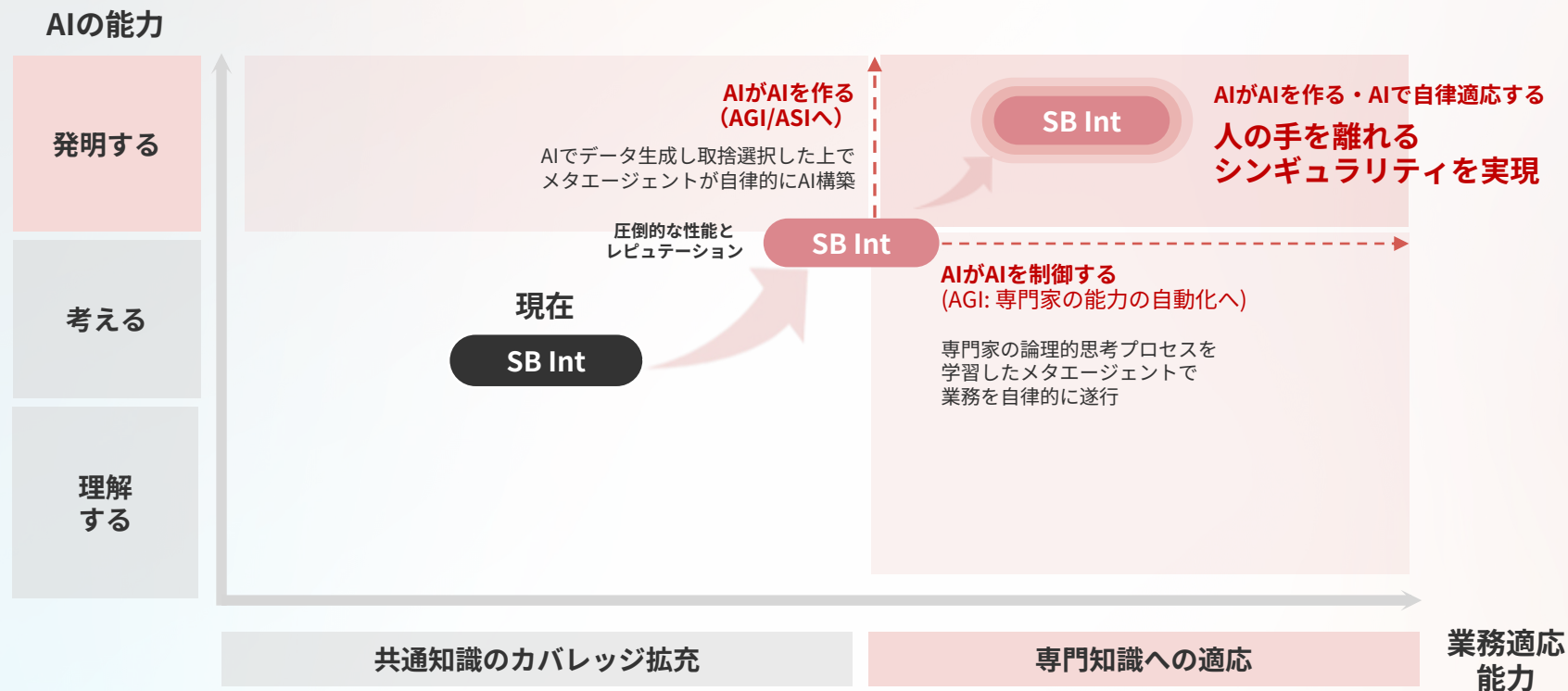
規制・倫理への対応



各国の法規制や倫理観は異なる。
日本語モデルを国内で開発することで、
日本の法律や社会的規範に沿った
運用が可能になる。

実現するポジション

SB Intuitions



あらゆる産業の競争力の源泉となり 世の中を成長させる



医療



法律



金融



製造



製薬



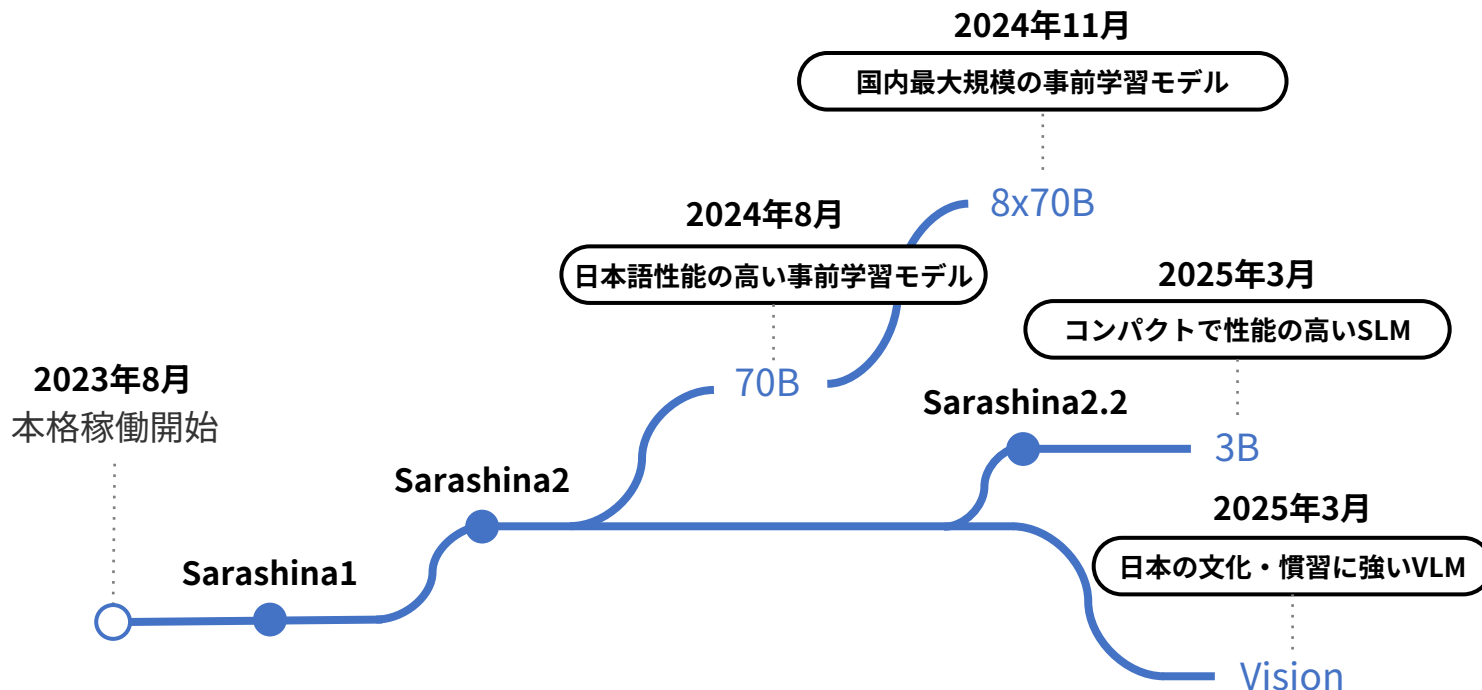
公共

大規模モデル

AIプラットフォーム

国内最大規模のAI計算基盤

- Sarashina1、Sarashina2、Sarashina2.2シリーズを開発し、オープンモデルで公開



与えられた文章の続きを生成

ユーザーの指示に的確に回答

事前学習

事後学習

(Fine-tuning / アラインメント)

日本語フルスクラッチ

SB Intuitions

(NTT、PFNなど)

日本語事前学習モデル

主に
日本語データ

日本語
事後学習モデル

データの安全性を
確保しつつ用途に応じて
性能を制御可能

日本語継続事前学習

国内事業者

(ELYZA、ABEJAなど)

海外製
オープンモデル
Llama、Qwenなど

日本語データ
継続事前学習

日本語
事前学習モデル

主に
日本語データ

日本語
事後学習モデル

性能・技術の外国依存
ライセンスによる制限

英語フルスクラッチ

海外事業者

(OpenAI、Anthropicなど)

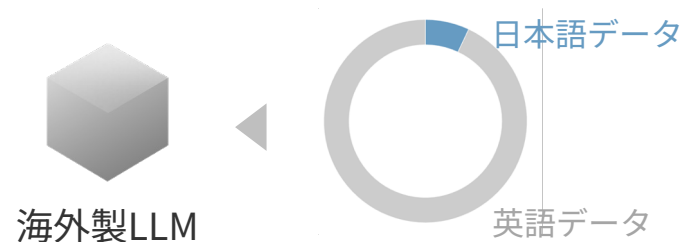
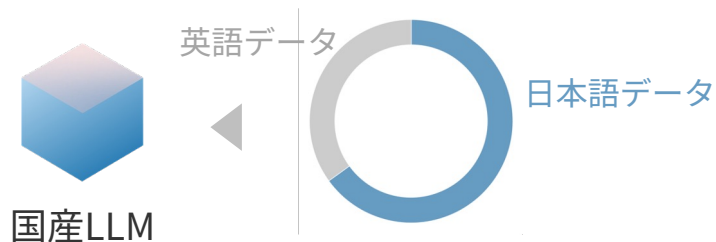
英語(多言語)事前学習モデル

多言語データ

英語(多言語)
事後学習モデル

日本語性能に不足
データの越境リスク

- 海外製モデルは学習データのほとんどが英語データであり、日本語データの割合は数パーセントに留まる
- 学習させるデータ（日本語/英語）によって、モデルの性能に差異が生じる※



※割合は概念図であり、正確な割合を反映するものではありません。

日本語データの学習効果が顕著

▶ 日本の知識に関する質問応答、日英翻訳

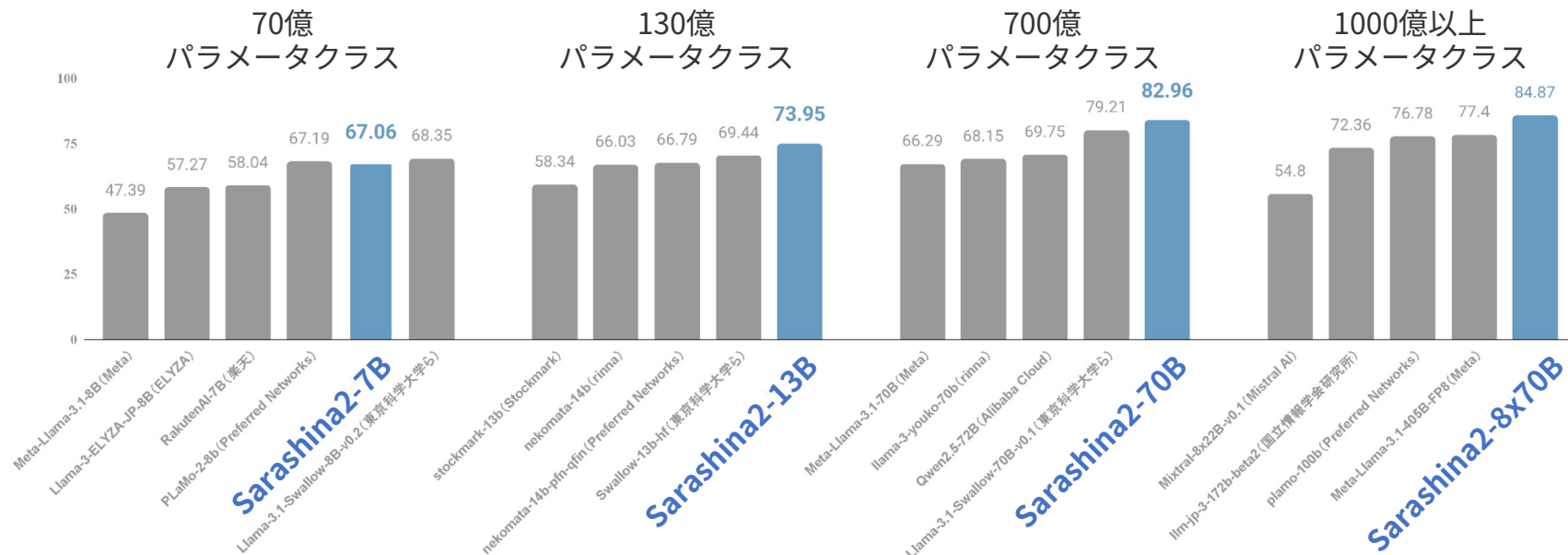
英語データからでも学習可能

▶ 日本語の一般教養、算術推論、コード生成

※第261回自然言語処理研究発表会での東京科学大学と産業技術総合研究所に所属する研究者らが発表した「LLMに日本語テキストを学習させる意義」より

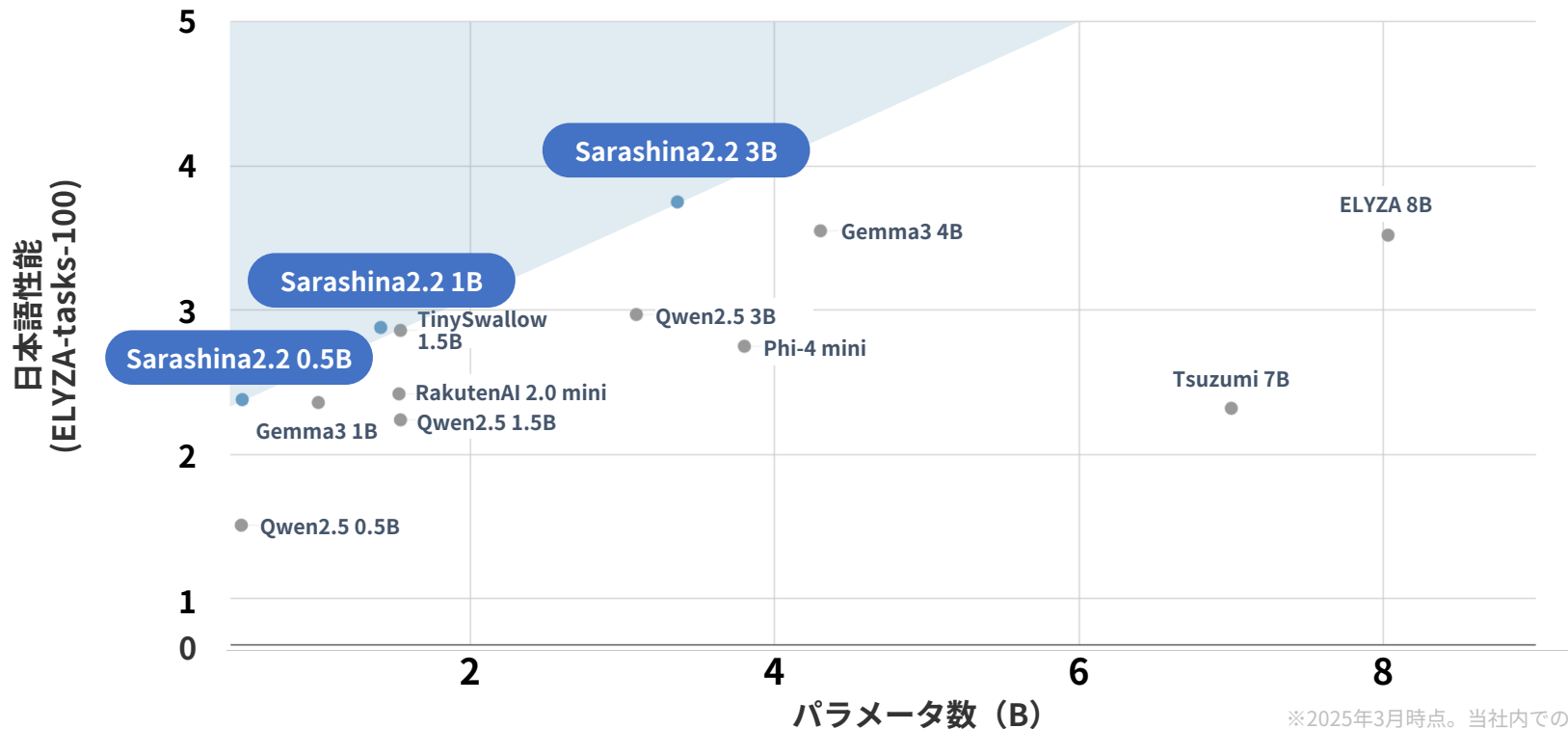
- 「Sarashina2」が各パラメータクラスにおいてトップクラスの日本語性能を記録

日本語質問応答評価セットによる評価（平均スコア）



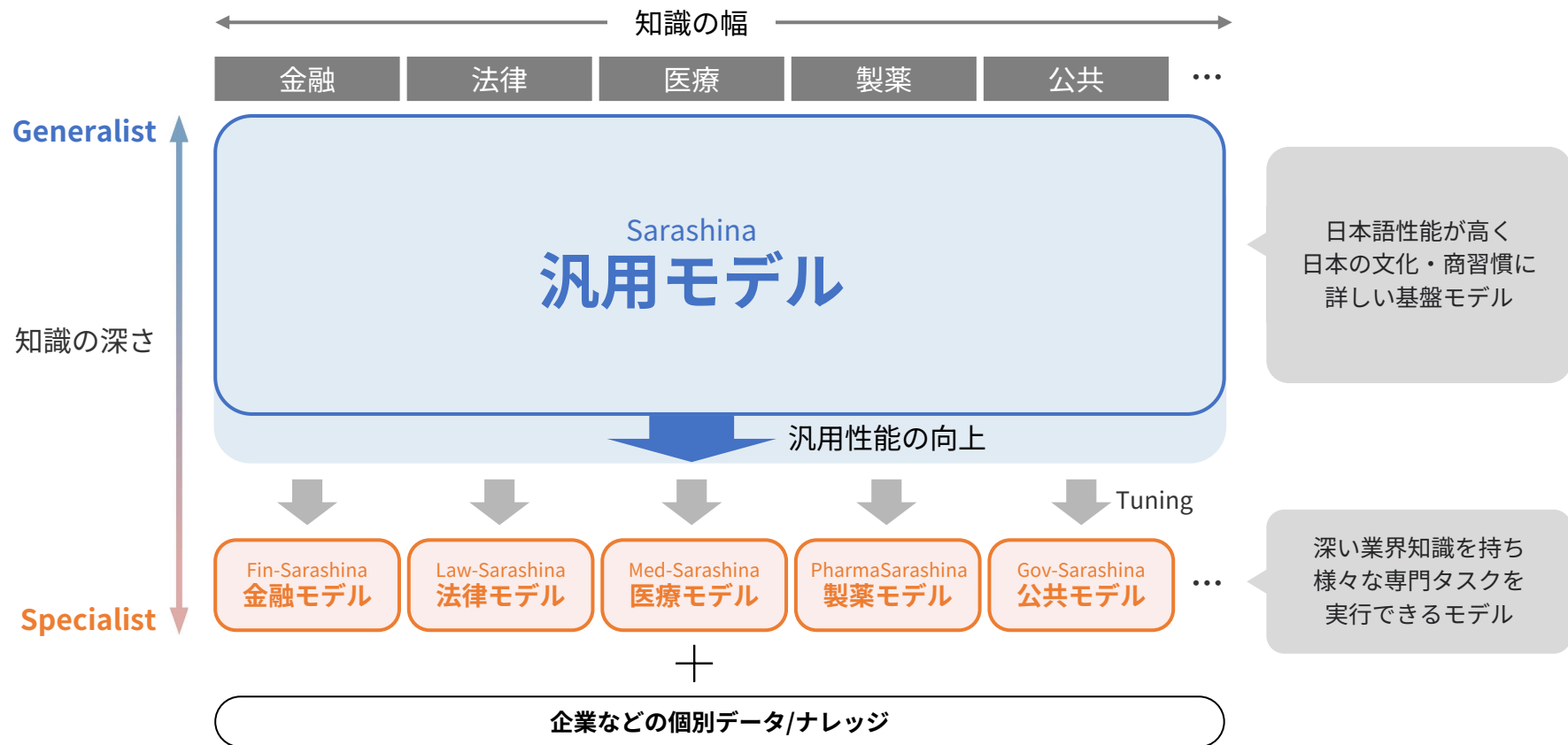
※2025年3月時点。当社内での比較による。

- 「Sarashina2.2」が同じパラメータ規模で最も高い性能をマーク



※2025年3月時点。当社内での比較による。

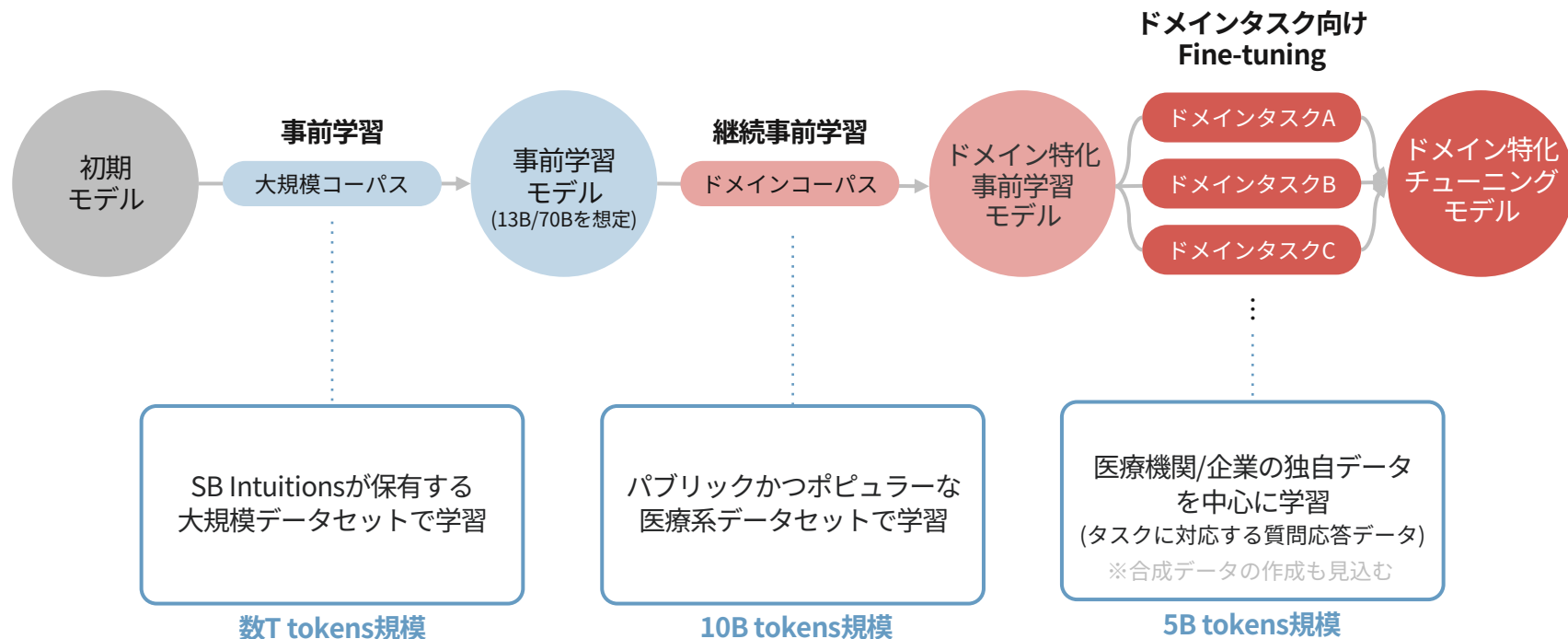
専門領域もカバーできる性能に向けて

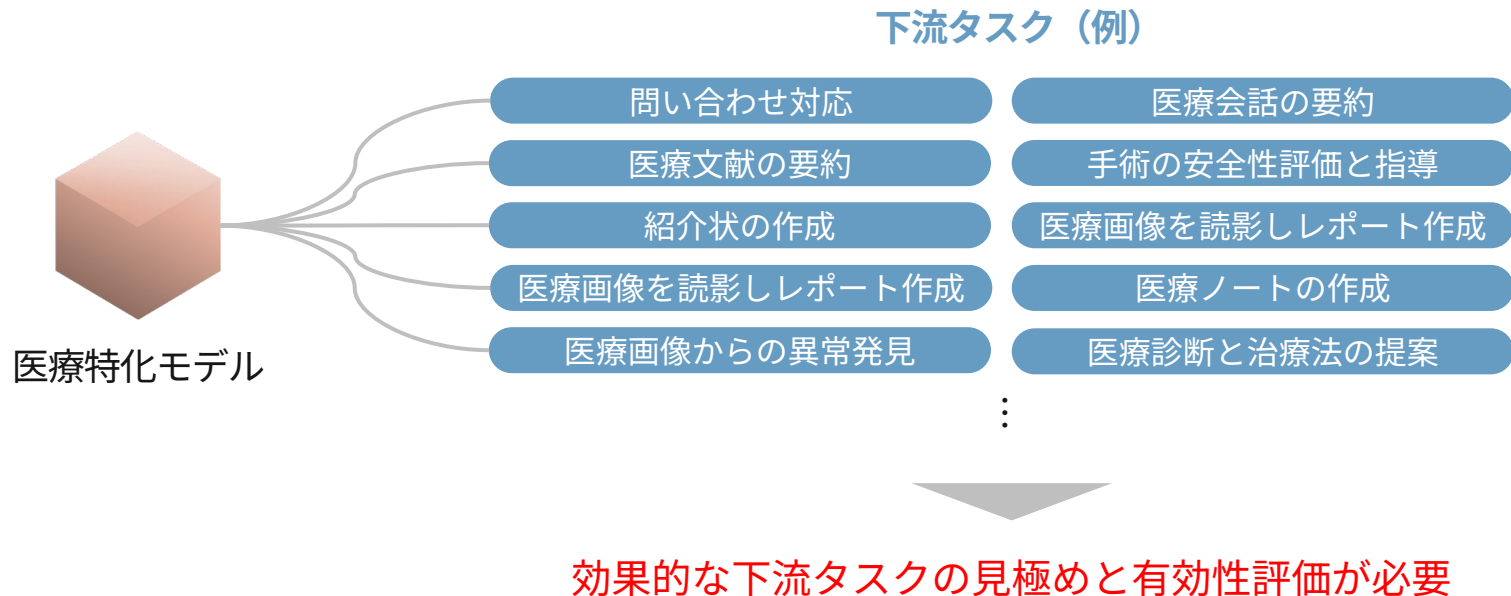


- ドメイン適応には、以下のアプローチをうまく組み合わせることが必要となる
- モデル単体の議論ではなく、RAGやPromptingも視野に入れた総合力で性能を測る

Prompt Engineering	- モデルの出力をより正確に制御するために、ドメイン固有の適切な指示やコンテキストを与える手法。 - 新しいデータの追加やモデルの再学習を必要とせず、ドメイン固有のタスクに対する性能を向上させられる最も簡単な方法。
外部知識の活用 (RAG)	- モデルが保持していない外部知識を活用することで、モデルの回答の精度や適用範囲を広げる方法。 - モデルのトレーニングが最新情報や特定のドメインに限られていない場合、外部ソースからリアルタイムで知識を取得してより正確な回答をすることが可能。
Instruction Fine-tuning	- モデルが特定の命令に従ってより正確に動作するようにFine-tuningする方法。 - 特定のタスクやドメインに関するデータセットを使用して、モデルを学習することで、特定のタスクやドメインに適応した回答を行うことが可能となる。
継続事前学習	- 既存のbaseモデルに対して、特定のドメインデータセットを用いて、さらに事前学習を行う方法。 - 大量の一般的な知識を持つモデルに対して、特定のドメインのテキストデータを用いて再学習を行うことで、そのドメインにおける専門知識や理解を向上させる。

コスト





- Sarashina miniから先行して商用開始へ

2025年3月

2025年6月

2025年 秋

2025年度以降

商用モデル完成
Sarashina mini
(700億パラメータ)



最終確認
(社内トライアル)



商用提供開始



大規模・高性能
モデルの構築



高性能なSarashina miniを
組み合わせてさらに大規模・高性能化



直感を、知性へ

