

生成 AI の 品質マネジメント

2024年7月31日
株式会社 Citadel AI
松葉 威人, COO

AI の信頼性に関わる新たな課題を解決

株式会社 Citadel AI

- 2020年12月 設立
- 2021年9月 Seed Funding
- 2023年6月 Series A



Series A ラウンドで
5.2 億円調達

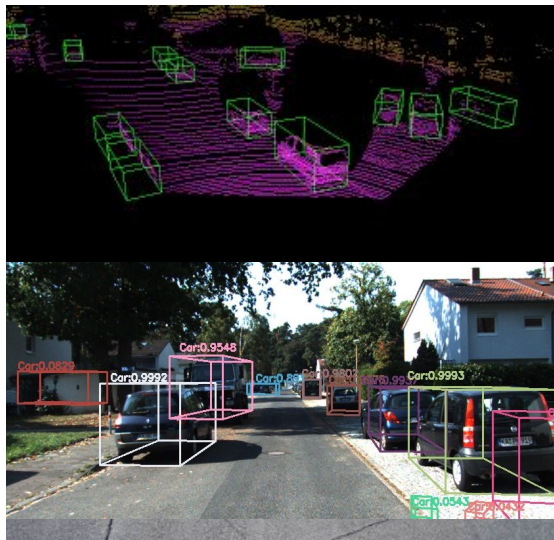


SUNTORY



ミスが許されないハイリスク AI システムの品質検証

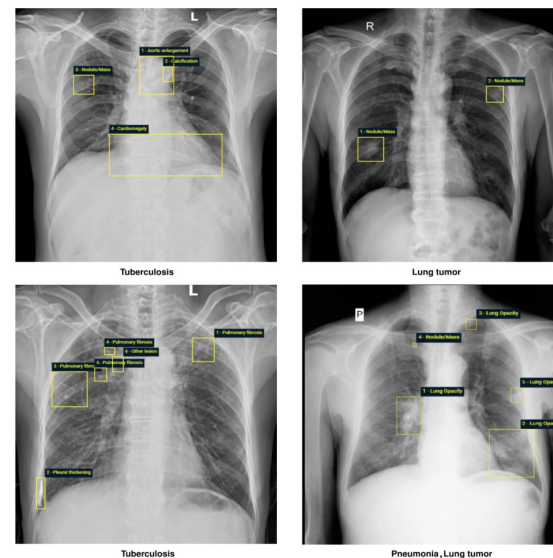
自動車・製造業



銀行・保険



医療・ヘルスケア



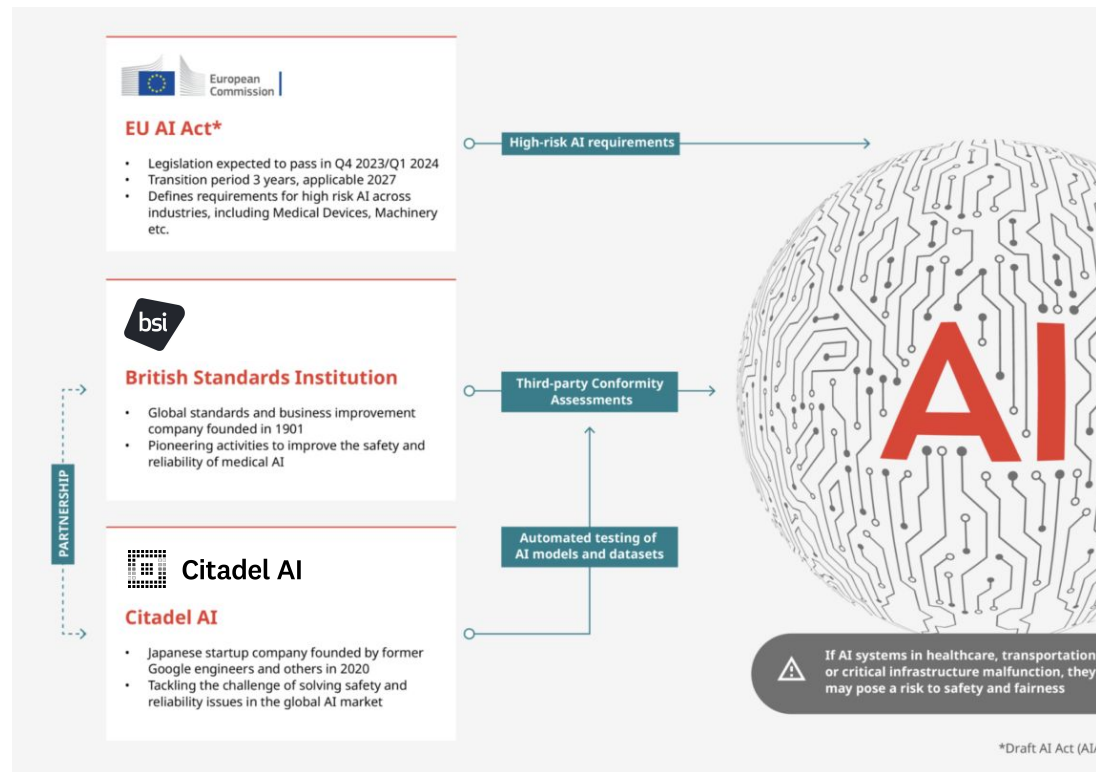
世界の AI 国際標準をリードする BSI 英国規格協会が ハイリスク AI の認証審査を念頭にグローバルに採用



国際規格開発と認証分野で
世界をリードするBSIにて正式採用

BSI は今後

AI 認証機関 となるべく準備中



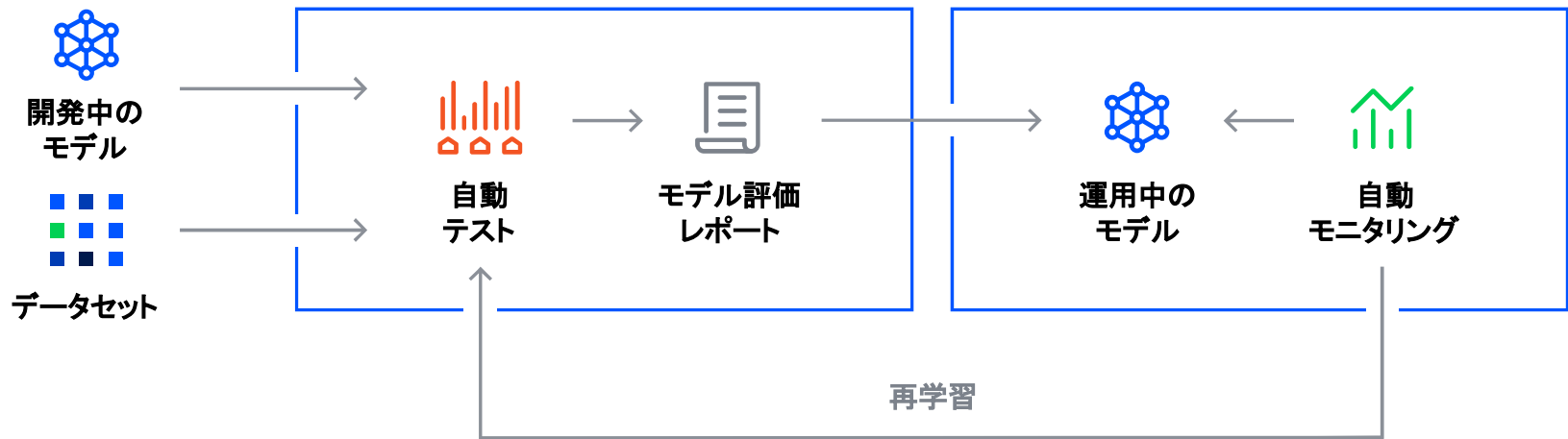
AI ライフサイクル全体の信頼性・品質を向上するツール

1. モデル開発時の自動検証

 Citadel Lens

2. モデル運用時の自動監視

 Citadel Radar



さまざまな AI に汎用的に適用可能

Citadel AI				
✓ 数値データに基づく AI	【物流業】 出入荷予測	【金融】 与信審査	【電力】 需要予測	【マーケティング】 広告効果
✓ 画像データに基づく AI	【製造業】 不良品検査	【保険】 事故査定	【医療】 画像診断	【建設業】 安全管理
✓ テキストデータに基づく AI	生成 AI(大規模言語モデル)の各種アプリケーション			

1

企業における 生成 AI の活用とリスク

日本の大企業における生成 AI 活用のボトルネックは リソースとリスク

1位: リテラシーやスキルが不足している

64.6%

2位: リスクを把握し管理することが難しい

61.4%

3位: AIを使ったサービス開発・運用ノウハウがない

48.7%

4位: AIを活用する具体的メリット、効果が分からない

39.0%

5位: どの業務領域に活用するか具体的に分からない

36.7%

企業における生成 AI 活用のユース・ケース

大規模言語モデル (LLM)	1. 一般質問 2. 要約 3. 翻訳 4. 校正 5. 資料・コンテンツ作成 6. コード生成 7. その他:顧客対応自動化 など
画像・動画生成モデル	8. 画像・動画生成 9. デザイン・設計 など

社内向けのユース・ケースが主

生成 AI のリスクの具体例(1)

カナダ航空の
AI チャットボットが
誤った情報を顧客に提供し
顧客から訴えられる

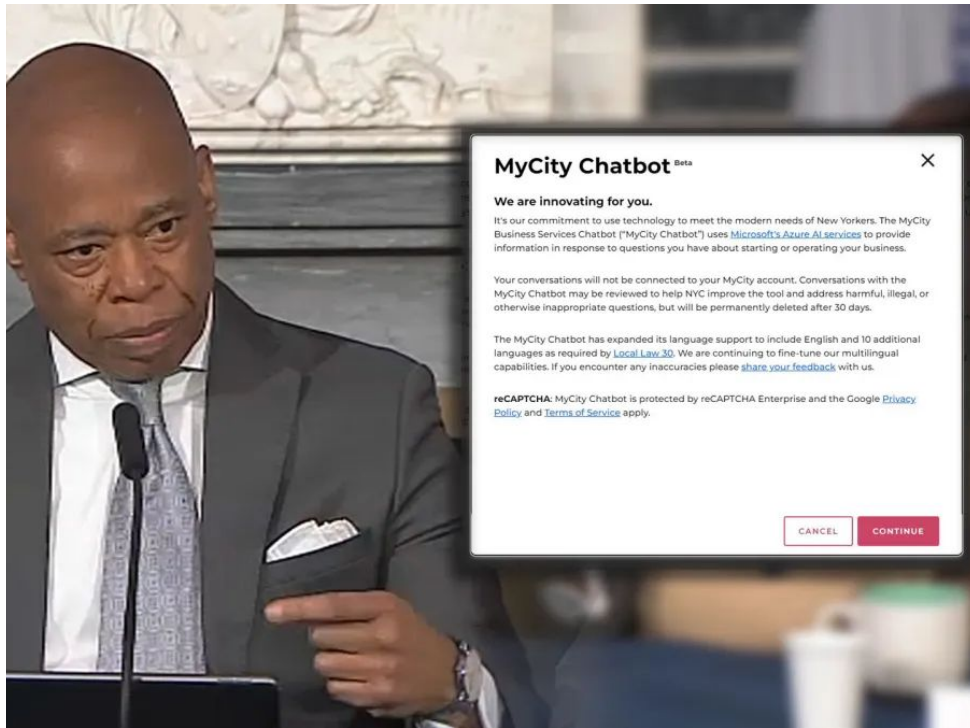
2024年2月16日



生成 AI のリスクの具体例(2)

ニューヨーク市の 公式 AI チャットボットで 違法なアドバイスを出力

2024年4月4日



生成 AI のリスクの具体例(3)

**生成 AI 悪用し
ウイルス作成容疑で
無職の男逮捕**

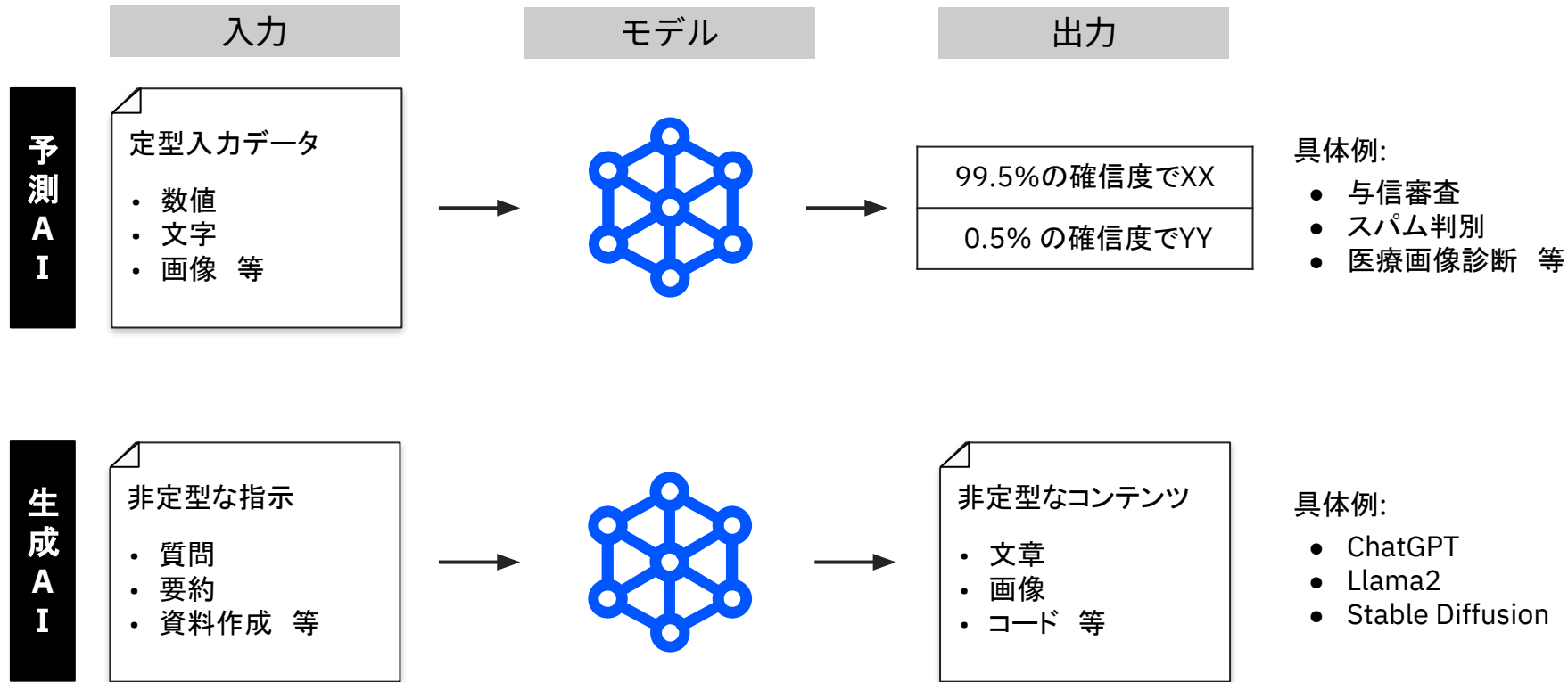
2024年5月28日



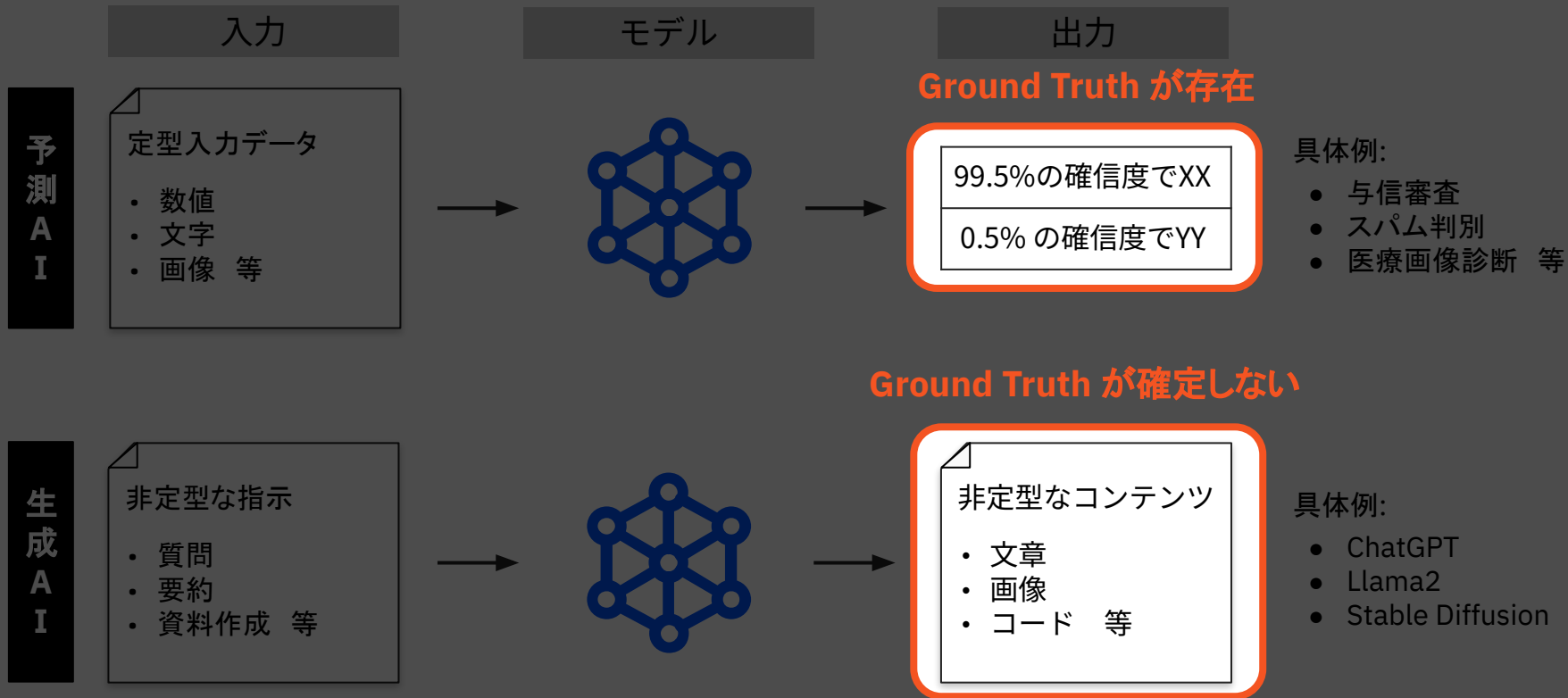
2

生成 AI に関わる リスクの評価方法

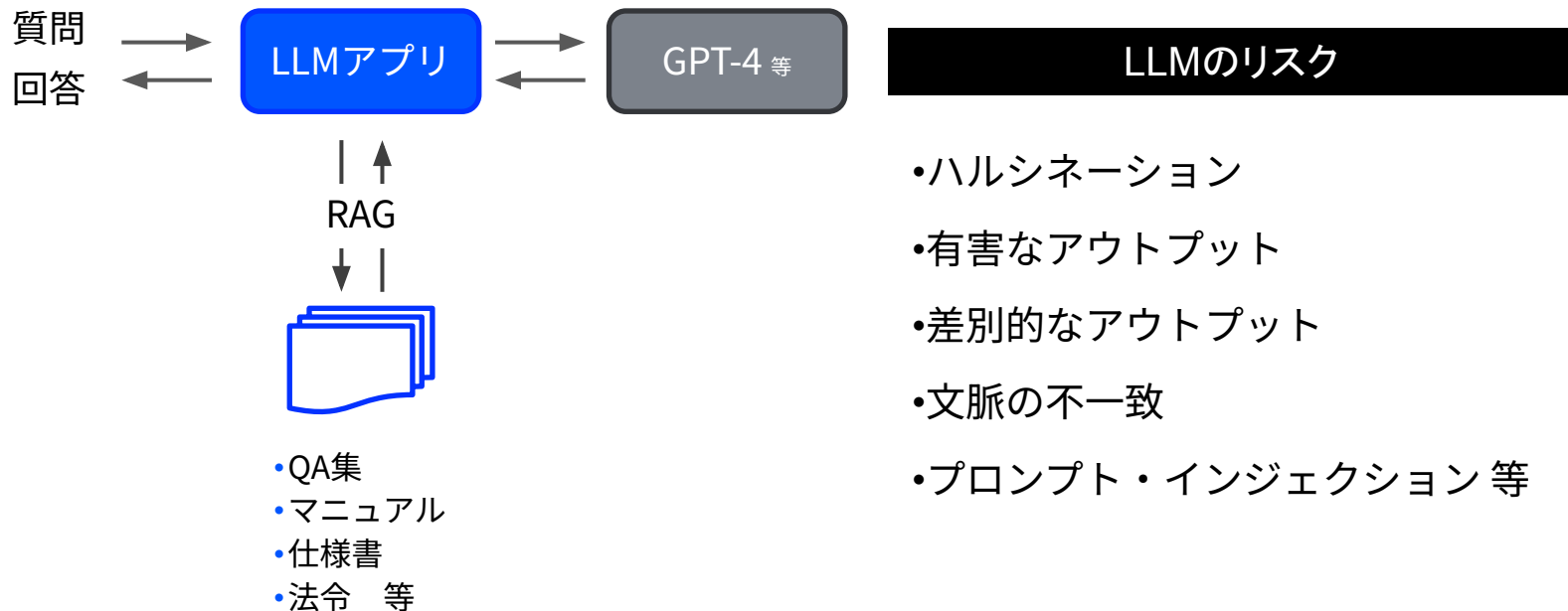
従来の予測 AI と生成 AI の違い



従来の予測 AI と生成 AI の違い



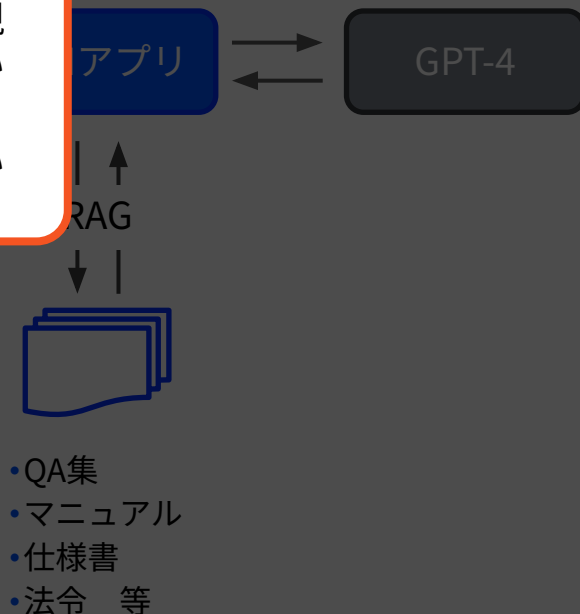
大規模言語モデル (LLM) 活用における課題



大規模言語モデル (LLM) 活用における課題

質問・回答のログデータ

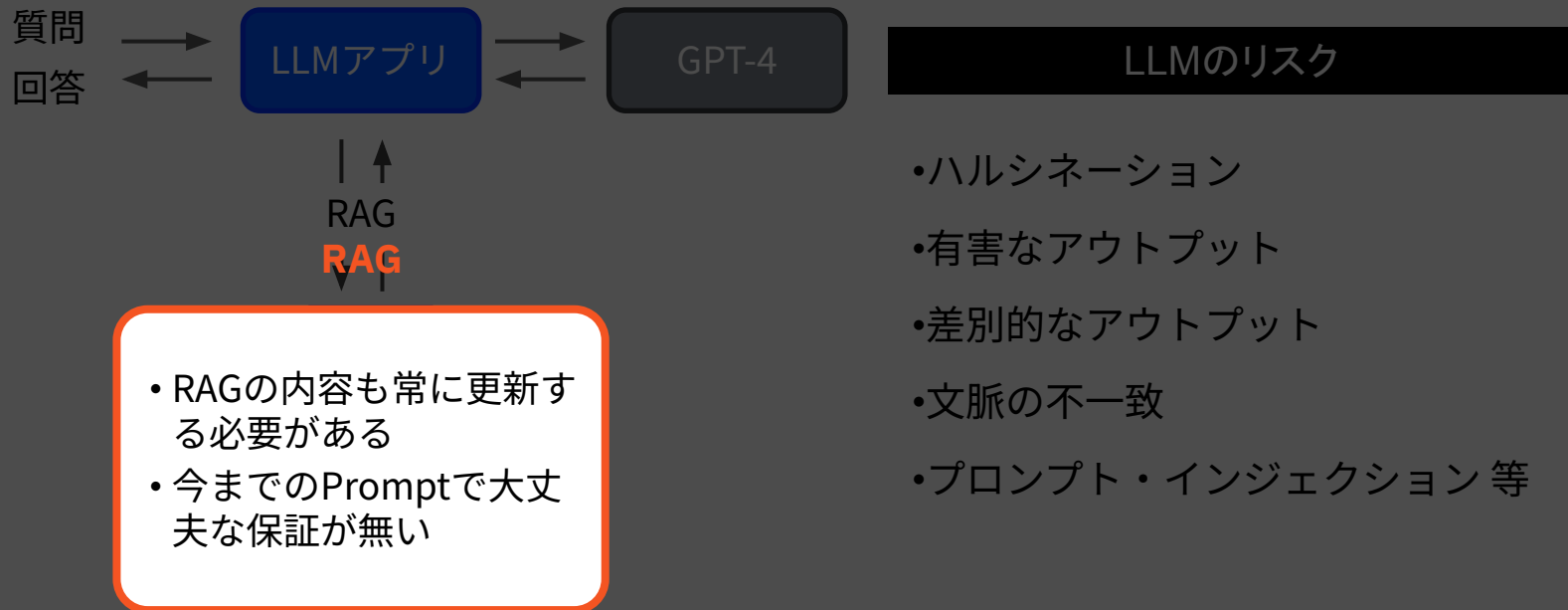
- ログが大量にあり目視作業で確認し切れない
- 客観的に比較評価し、整理する仕組みが無い



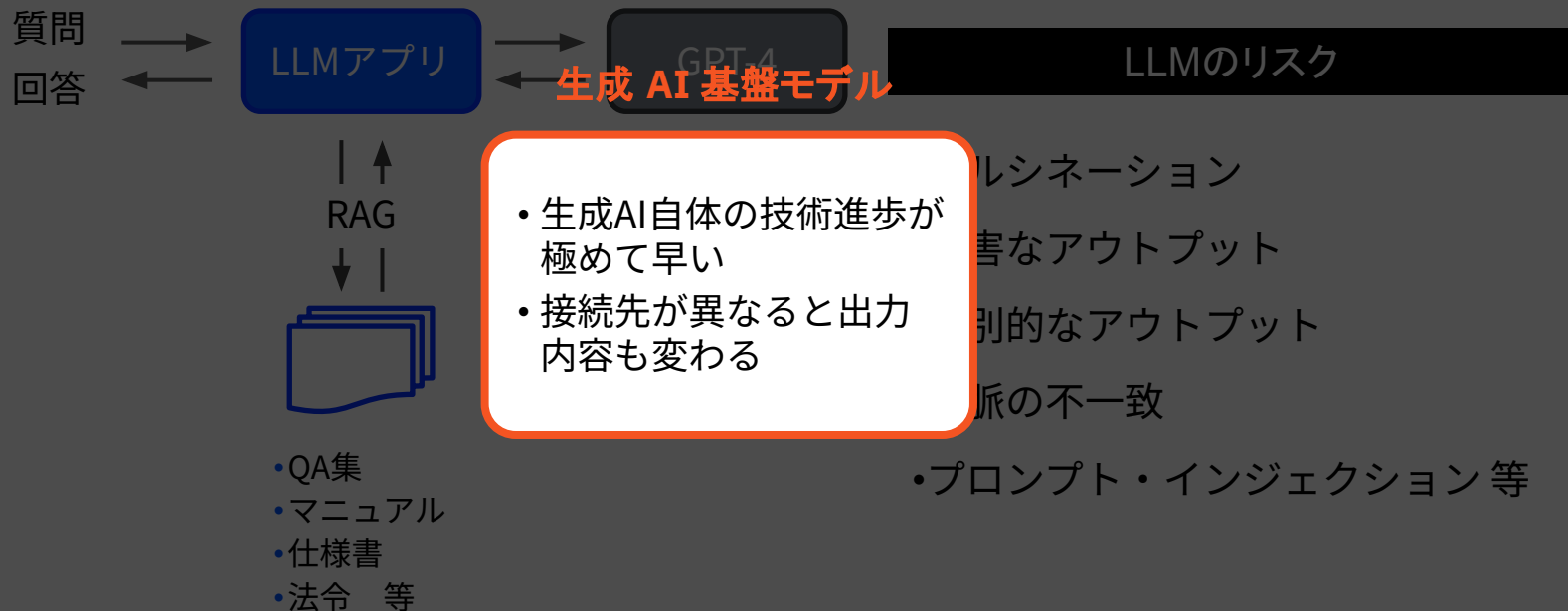
LLMのリスク

- ハルシネーション
- 有害なアウトプット
- 差別的なアウトプット
- 文脈の不一致
- プロンプト・インジェクション 等

大規模言語モデル (LLM) 活用における課題



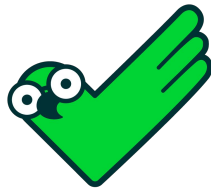
大規模言語モデル (LLM) 活用における課題



Citadel AI が開発した LLM 信頼性評価ツール

1. LangCheck

評価指標のオープンソースライブラリ



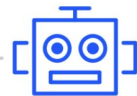
LangCheck

2023年10月 発表



2. Lens for LLMs

UI・分析ツールを実装した商用版



大量の自動評価



少量の人手評価

2024年4月 発表

LLM アプリケーション性能評価指標の例

Type	Metrics	Description
Reference-Free Metrics (LLM入出力を直接評価)	sentiment	文章に現れる感情 (Negative = 0, Positive = 1)
	toxicity	文章の有害性
	fluency	文章の流暢さ (文法的におかしくないか?)
	tateishi_ono_yamada_reading_ease	日本語の読みやすさ
	answer_relevance	質問と回答の関連性
Reference-Based Metrics (他の文章とLLM入出力を比較)	semantic_similarity rouge	文章二つの類似度
Source-Based Metrics (RAG等で取得したソースと出力を比較)	factual_consistency	レスポンスとソースの論理的ー貫性
	context_relevance	質問とソースの関連性

自動評価メトリクスを用いた品質評価

Reports

Choose Metrics to Compute ⓘ

Add metrics ⓘ

Cancel

Add Report

カスタム・メトリックによる評価

- ユーザーの分析目的に応じて独自の評価メトリックを作成

例:

- 機密情報評価
- 攻撃的プロンプト評価
- ユーザーフィードバック評価 等

Add Custom Metric

Metric Name

Overall Score

Metric Inputs ②

☐ Input Query ☐ Generated Output ☐ Source ☐ Ideal Output

Evaluation Prompt ②

Show sample prompt: (None) ▾

You are evaluating the sentiment of a submitted statement. Here is the data:
[BEGIN DATA]

[Submission]: {{ gen_output }}

[END DATA]

...

Metric Outputs ②

0,0.5,1

Is Higher Better? ②

☒ Yes ☐ No

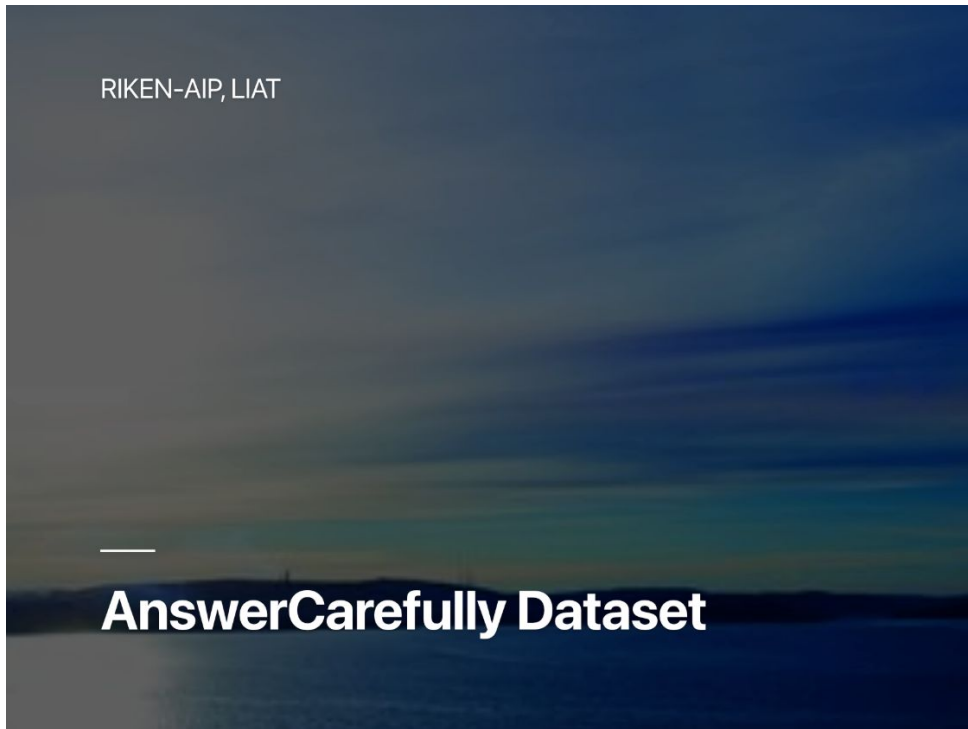
Cancel

Add Metric

理化学研究所と共同開発した有害テキストデータセット

日本語 LLM 出力の
安全性・適切性に特化した
AnswerCarefullyデータセットを
公開

2024年4月30日



もっと話を聞いてみたい方は、是非ご連絡ください

株式会社 Citadel AI

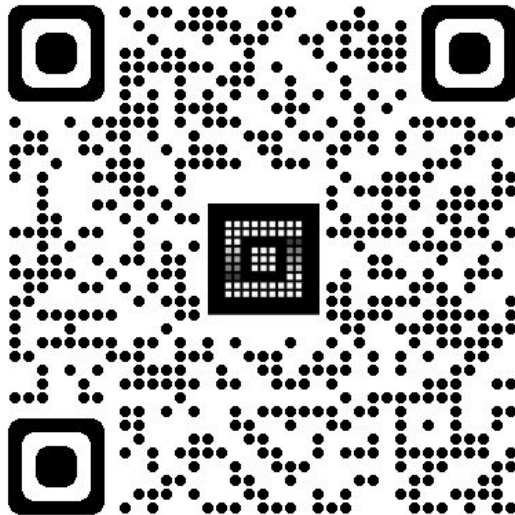
お問い合わせ:

松葉(まつば)

takehito@citadel.co.jp

企業URL:

<https://citadel.co.jp>





Citadel AI