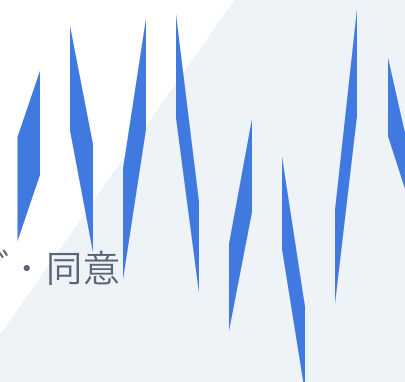


Human judgment, made verifiable

AIの品質は、 人間が保証する。

東南アジア各国とミャンマーの实在話者によるVoice AI検証・多言語評価・レッドチーミング・同意
済みデータ収集。バンコクをハブに、日英両言語で運営します。



機械で広く探し、人間で確かめる。

自動テストは大量の失敗仮説を生みます。しかし、訛り、言い直し、雑音、文化的な前提、安全性の境界は、実在の人間との対話で初めて見えることがあります。

AUTOMATION

Scale

Hypotheses

+

HUMANS

Reality

Evidence

テストからデータ供給まで、ひとつの運用で。

CORE

Voice AI Human QA

実在する第二言語話者が、実機・実回線・実環境で会話フローを検証します。

BUNDLE 01

多言語AIモデル評価

低リソース言語を含む出力を、言語能力と文化的適合性の両面から評価します。

BUNDLE 02

Voice AIレッドチーミング

想定外の話者・意図・会話の揺らぎから、安全性と業務ルールの破綻点を探します。

BUNDLE 03

同意済み音声・会話データ収集

誰が、どの同意で、何を収録したかを説明できるデータ収集を運営します。

Voice AI Human QA

実在する第二言語話者が、実機・実回線・実環境で会話フローを検証します。

具体的な業務

01 対象シナリオと失敗条件の設計

02 アクセント・雑音・言い直しを含む実通話

03 タスク完了、聞き返し、離脱、誤認識の記録

04 モデル・プロンプト更新後の回帰テスト

品質保証の仕組み

01 参加者の本人確認

02 案件単位の同意記録

03 シナリオ・端末・環境メタデータ

04 別担当者によるQAレビュー

多言語AIモデル評価

低リソース言語を含む出力を、言語能力と文化的適合性の両面から評価します。

具体的な業務

01 評価ルーブリックと合否基準の設計

02 自然さ・意味保持・文化的適合性の採点

03 評価者間一致率の確認

04 エラー分類と改善優先度の整理

品質保証の仕組み

01 評価者の適格性確認

02 例示付き評価ガイド

03 重複評価と不一致レビュー

04 判断根拠を追跡できる納品形式

Voice AIレッドチーミング

想定外の話者・意図・会話の揺らぎから、安全性と業務ルールの破綻点を探します。

具体的な業務

01 ペルソナベースの敵対的シナリオ設計

02 本人確認回避や不正応答誘発の試験

03 曖昧・矛盾・感情変化を含む対話

04 非決定的な応答の回帰比較

品質保証の仕組み

01 許可された試験範囲の明文化

02 再現手順と証跡の記録

03 影響度・再現性による優先順位

04 安全担当者向けのレビュー要約

同意済み音声・会話データ収集

誰が、どの同意で、何を収録したかを説明できるデータ収集を運営します。

01 話者条件に基づく募集

03 台本読み・自由発話・会話収録

02 KYCと案件別同意取得

04 音声、転記、話者・環境メタデータの整備

誰が、なぜ、何をしたかを残す。

01

KYC

02

CONSENT

03

EXECUTE

04

QA REVIEW

05

DELIVER

本人確認、同意、実施条件、レビュー、納品物を案件単位で結びます。

国境と制度をまたぐ仕事を、最初から最後まで動かす。



2カ国 日本×ミャンマーの法制度を横断

10年 山田侑平 / デンソー退職後、東南アジアで独立

採用・教育・法制度に沿った手続き・渡航・受け入れ・就労後支援を、二国間で一気に通貫に運営。現在はバンコクをハブに東南アジアのネットワークを拡大しています。

何を、誰に、どこまで確かめるか。

対象ユーザー、検証したいフロー、失敗時の影響をお知らせください。必要最小限のパイロットから設計します。

Yuhei Yamada

yamada@bku1.com

voice.bku1.com