

外部接続ゼロ、情報漏洩リスクもゼロ。金融機関待望の『完全オンプレミス』A I

金融特化型オンプレミス生成A I 商品説明会のご案内

2026 年

1/19
(月)1/26
(月)1/29
(木)14:00 ~ 15:00
オンライン開催

金融機関が直面する「生成A I 導入の3つの壁」を解決します

1. 情報の壁

機密情報や個人情報等を外部サーバ（クラウド）へ出すことができない

2. 精度の壁

汎用A I では金融実務特有のルールに正しく回答できない（ハルシネーション）

3. コストの壁

アカウント課金制では全社展開に膨大なコストがかかる

▼ これらの壁を「完全オンプレミス × 金融実務知」で一掃します
本説明会でご紹介するソリューション～次世代A I パッケージ～

Local-Chat AI

実務知 × RAG

- 銀行研修社が60年に亘り蓄積した「実務知」をデータベース化し、RAG技術で搭載
- 融資推進支援やコンプラチェックなど金融実務に即した正確な回答を導き出します
- 最新モデル gpt-oss-120B を採用。OpenAI o4-mini に匹敵する高い推論性能をオフラインで実現。
- 外部接続なしのセキュアな環境により機密情報を含むプロンプトも安心して入力可能

オンプレミス AI 議事録

外部流出リスクゼロ

- 機密性の高い会議や顧客面談の音声を外部流出リスクを排除した社内ネットワーク内で自動文字起こし・要約します
- 専用GPUサーバにより1台で月間約2,000時間の音声処理を定額で高速に実行可能
- 専門用語登録や自社独自の議事録フォーマットへの柔軟なカスタマイズに対応
- 標準化された高精度な要約により、作業工数の劇的な削減と迅速な意思決定を支援

共通基盤：完全オンプレミス × 定額使い放題

お申込み方法

お申込みは、氏名、メールアドレス、会社名などを下記登録ページにご入力ください。

●1/19（月）開催

https://zoom.us/webinar/register/WN_a2-7TGN6RY-QOULvGWrnMg#/registration



●1/26（月）開催

https://zoom.us/webinar/register/WN_Ap3Z5bnkQpe0s4D-mnU9Bw



●1/29（木）開催

https://zoom.us/webinar/register/WN_tCnMKBvHSrKYZV7mpcfTbw#/registration



■ 主なご説明内容

01 Local-Chat AI（金融特化型 AI チャット）

【金融実務知の活用】

銀行研修社が 60 余年に亘り蓄積してきた知見を基に編集した「業種別融資推進ガイド」「金融取引ルールブック」等の専門コンテンツを RAG 技術でデータベース化し、ハルシネーションを限界まで抑制して正確な回答を導き出すロジックを説明します。

【高度な推論性能】

2025 年最新の OSS モデル「gpt-oss-120B」を採用。OpenAI の o4-mini 級に匹敵する高度な推論性能を、外部接続一切なしの完全オフライン環境で実現する技術的背景をご紹介します。

02 オンプレミス AI 議事録（AI 議事録ソリューション）

【セキュアな構造化技術】

金融特有の専門用語や各金融機関独自のフォーマット（取締役会、経営会議、ALM 委員会等専門会議）に合わせ、外部流出リスクをゼロに保ちながら高精度に要約・構造化を行うプロセスを詳しくご説明します。

【圧倒的な処理能力を経済性】

1 台の専用サーバで月間 2,000 時間分の音声処理を可能にする「定額使い放題」の仕組みと、金融機関への導入実績に基づく効率的な運用ノウハウをご紹介します。

参考 Role Coach-AI（AI ロールプレイ）

【客観的なスキル評価】

若手職員の営業対応を AI が客観分析し、説明力・質問力等のレーダーチャートでスキルを可視化する「AI コーチ」機能の評価アルゴリズムについて解説します。

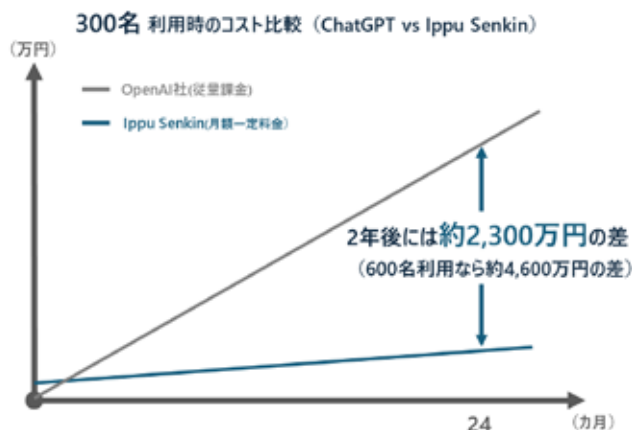
【組織知の承継】

貴社のハイパフォーマー（エース社員）のスク립トを AI に学習させ、属人化したノウハウを組織全体の資産として底上げする、特化型 Agent のオーダーメイド構築手法をご紹介します

■ なぜ今オンプレミスなのか

2025 年、生成 AI の常識が塗り替えられました。

最新のオープンソースモデル（OSS）と GPU の進化により、クラウド型に依存しない「真の自社運用」が実用水準に到達しました。



圧倒的なコスト

従量課金なし全職員で従量課金なし。全職員で 24 時間使い続けても定額。クラウド型と比較で数千万円規模のコストを抑制。

高精度

最新モデル「gpt-oss-120B」を採用。
OpenAI の最新モデル（o4-mini 等）に匹敵する高い推論精度をオフラインで実現。