

業務 AI エージェントの 実行境界

権限・承認・監査から設計する安全な自動化

CRM・メール・社内データを AI と接続する前に
組織で合意すべき実装原則と判断手順

1. エグゼクティブ・サマリー

「CRM を見て次のアクションを提案してほしい」「問い合わせメールの下書きを作してほしい」「社内文書を調べて回答してほしい」。生成 AI の活用は、文章を生成する段階から、CRM、メール、文書保管、チケット管理といった業務システムを操作する段階へ進んでいる。Function Calling や Model Context Protocol (MCP) は外部システムとの接続を容易にするが、接続できることは、その操作を自動実行してよいことを意味しない。

業務エージェントのリスクは、モデルが誤った文章を生成することだけではない。権限外の文書を検索する、誤った商談レコードを更新する、承認されていない宛先へメールを送る、同じ処理を重複して実行する、といった副作用が生じうる。特に、検索結果やメール本文に含まれる指示をモデルが信頼してしまう間接的なプロンプトインジェクションは、ツールの不適切な呼び出しや情報流出へつながりうる。NIST と OWASP は、生成 AI・エージェントのリスクを、技術性能だけでなく、ガバナンス、アクセス制御、人による監督、運用監視を含むシステムの問題として整理している¹²³。

本稿でいう「実行境界」とは、モデルの提案と、外部システムへの副作用との間に置く、決定論的な認可・検証の境界である。LLM が「このメールを送るべきだ」と出力しても、メール送信が実行されるとは限らない。送信先、添付、送信元、承認の有無、有効期限、利用者の権限、業務ルールをアプリケーション側で確認し、条件を満たす場合にだけ実行する。

本稿の主張は明確である。業務 AI エージェントの導入要件は、モデル名ではなく、次の六つで決めるべきである。

- ・ 誰の指示で、誰の権限として実行するか。
- ・ どの操作・対象・項目までを許可するか。
- ・ 何を人が確認し、承認はどの実行内容へ結び付くか。
- ・ 誤り、重複、途中失敗をどう止め、復旧するか。
- ・ 実行の経緯を後から調査できるか。
- ・ どの失敗を検証してから、権限を広げるか。

2. 接続・実行・主体を分けて考える

2.1. LLM は認可の根拠ではない

業務エージェントは、利用者の依頼を LLM が解釈し、ツール呼び出しとして業務システムへ渡す構造を取ることが多い。ここで、モデル出力をそのまま API 呼び出しへ渡す設計は危険である。MCP や Function Calling は、ツールを呼び出すための仕組みであり、業務上の認可・承認・監査を代替するものではない。

実行境界の基本構造

利用者の指示 → LLM による提案 → 入力・権限・業務ルールの検証 → 必要時の承認 → ツール実行 → 監査ログ

モデルは提案を担い、認可・検証・実行・記録は決定論的なアプリケーション層で担う。

たとえば「取引先 A へ見積もりを送って」という依頼に対し、モデルは候補の宛先、件名、本文、添付を作成できる。しかし実行境界は、取引先 A の定義、宛先ドメイン、添付ファイルの機密区分、送信元、送信が許可された時刻、承認対象と実行内容の一致を別途確認する。モデルの自信度は、これらの条件を置き換えない。

2.2. 主体を混同しない

設計時には、少なくとも次の四つを分けて記録する。

主体	確認すべきこと
利用者	誰が依頼したか。対象データや操作を業務上許可されているか。
実行主体	どのサービスアカウントまたは委任トークンが API を呼ぶか。利用者より強い権限を持っていないか。
承認者	誰が、どの対象・操作・差分・有効期限を確認して承認するか。
管理者	接続アプリ、OAuth スコープ、ツール一覧、ログ保存、停止権限を誰が管理するか。

「開発者の API キーで動くエージェント」は、利用者の権限より広いデータを読んだり更新したりするおそれがある。OAuth のスコープも、業務上の適切さを保証しない。OAuth Security Best Current Practice は、トークンの漏洩や誤用を前提に、権限を最小化し、対象リソースを制限することを推奨している⁴。委任では、誰のために行う操作か（subject）と、どの実行主体が行ったか（actor）を区別する。RFC 8693 はこの区別をトークン交換において定義し、RFC 8707 はトークンの利用先を特定リソースへ制限する仕組みを定めている^{5 6}。

NIST のゼロトラストアーキテクチャは、ネットワーク上の位置や資産の所有だけを根拠に暗黙の信頼を与えず、リソースへのアクセス前に主体と条件を検証する考え方を示している⁷。この文書は AI エージェント専用ではないが、利用者の指示、実行主体、対象リソース、操作条件を要求ごとに確認する設計原則として参考になる。

2.3. 参照操作も無害ではない

「読むだけなら安全」とは限らない。横断検索の対象に人事情報、顧客契約、個人メールボックスが含まれれば、回答画面やログ、外部 LLM への送信を経由して情報が流出する可能性がある。まずは利用者の既存権限を引き継ぐか、公開範囲が明確な文書コーパスに限定する。大量取得、権限外の検索結果、外部 URL へのアクセス、機密データを含むログにも上限と検査を設ける。

3. 実行権限を 4 段階で設計する

段階は業界標準ではない。本稿では、導入範囲を組織で合意するための実務上の分類として、次の四段階を用いる。重要なのは、すべての操作を最終段階へ進めることではない。業務への影響と復旧可能性に合う段階で止めることである。

段階	典型操作	人の関与	実行前の条件	監査・評価
1. 参照	公開文書の検索、CRM の検索、要約	通常は不要。ただし権限外・大量取得は停止。	利用者権限、検索範囲、取得件数、送信先を検証。	検索対象・結果・引用を記録。権限外検索と根拠欠落を評価。
2. 下書き	メール案、CRM 更新案、タスク案	利用者が保存・送信・更新を実行。	モデルに送信・更新ツールを渡さない。	提案内容と最終変更を分けて記録。誤紐付け・宛先誤りを評価。
3. 承認付き実行	CRM メモ追記、社外メール送信、チケット更新	指定された承認者が実行ごとに承認。	対象、操作、引数または差分、実行主体、有効期限を照合。	承認・実行・結果を関連 ID で記録。承認後差替え、却下、期限切れを評価。

4. 条件付き自動実行	定型の社内タスク作成、限定した通知	通常不要。ただし停止・例外対応の担当を置く。	許可リスト、固定テンプレート、上限、罫等キー、例外条件を満たす。	継続監視。重複実行、上限超過、例外、ログ欠落を定期確認。
-------------	-------------------	------------------------	----------------------------------	------------------------------

3.1. 承認は「要約への同意」ではない

有効な承認は、「この AI の提案を承認する」という曖昧な同意ではない。承認画面には、少なくとも対象レコードまたは宛先、実行する操作、更新前後の差分または送信本文・添付、実行主体、有効期限を表示する。承認後に本文や引数が変われば、再承認を要求する。これにより、承認疲れや、承認済み内容を実行前に差し替える問題を減らせる。

3.2. 条件付き自動実行は、検知できる条件で止める

「AI が自信を持てないとき」は、運用上の停止条件として曖昧である。止める条件は、機械的に検知できる形にする。たとえば、社外ドメインへの送信、金額・契約・権限フィールドの変更、添付ファイルの追加、必須項目の欠落、許可リスト外の操作、予算・件数上限の超過、参照根拠の欠落を停止条件にする。

高影響または不可逆な操作、具体的にはメール送信、情報開示、削除、支払、契約・権限変更は、初期導入で承認を残す。自動化を検討する場合も、対象を定型通知などに限定し、実行件数と影響範囲の上限を設ける。

4. 業務別の設計例

4.1. CRM：広い更新 API ではなく、限定ツールを渡す

CRM で「商談を更新できる」ことは、金額、確度、担当者、ステージ、契約条件などをすべて更新できることを意味してはならない。エージェントに汎用の update_record を渡す代わりに、用途ごとの限定ツールを設ける。

業務操作	初期段階	実行境界の例
商談・顧客の検索と要約	参照	利用者のレコード共有範囲に限定し、結果件数と機密項目を制限する。
商談メモ・次回アクション案の作成	下書き	対象商談を利用者が確認して保存する。モデルに更新権限は渡さない。
特定メモ欄への追記	承認付き実行	更新可能なオブジェクト・項目・文字数を限定し、差分を承認する。
金額、確度、ステージ、担当者、削除	原則禁止または個別承認	業務ルールと責任者による別フローを優先する。

Salesforce 等では OAuth スコープに加えて、オブジェクト、項目、レコード共有の制御がある。どれか一つだけでは不十分である。エージェント用の資格情報は、利用者の権限を超えないように設計し、必要なら更新可能な項目だけを受け付ける中間 API を置く。NIST SP 800-53 は、最小権限、職務分離、監査可能なイベント記録をアクセス制御の基本としている⁸。

4.2. メール：下書きと送信を分ける

メールでは、下書きを作ることと外部へ送信することの影響が大きく異なる。まずは、会話や CRM メモを基にした下書き作成に留める。利用者が宛先、件名、本文、添付、BCC、送信元を確認してから、メールクライアントまたは承認済み送信フローで送信する。

Microsoft Graph では下書きの作成・更新と送信に異なる権限を設定できる。一方、メールサービスによっては、下書き作成と送信を同じスコープで許可する場合がある。その場合、OAuth スコープだけに依存せず、エージェントへ送信ツールを公開しない。承認済みリクエストだけを検証して送信 API へ渡す中間サービスを設ける。

送信を自動化する場合にも、送信元、宛先ドメイン、テンプレート、添付の可否、送信時刻、1 日当たりの上限、禁止語や機密情報の検査を固定する。未知の宛先、添付追加、社外送信、受信メール本文が直接指示として使われた場合は停止し、人へ渡す。

4.3. 文書検索と申請・チケット

領域	開始しやすい操作	保留・承認が必要な操作	確認事項
社内文書・RAG	公開 Wiki・現行マニュアルの検索、出典付き要約	個人メール、人事、顧客契約、未公開経営資料の横断検索	利用者権限の継承、文書の公開範囲・有効期限、引用、取得上限
申請・チケット	起票案、分類、担当候補、社内メモ案	優先度・SLA・担当・完了状態の変更、顧客向けコメント	対象チケット、状態遷移、担当範囲、顧客可視性、重複起票

社内検索では、回答品質だけでなく、誰がどの文書を参照できたかを確認する。申請・チケットでは、作成、割当、優先度変更、完了、外部公開コメントを同じ「更新」と見なさない。状態遷移と可視性を操作ごとに分け、初期は下書きまたは承認付き実行から始める。

4.4. 設計ケース：自動化の範囲を段階的に決める

以下は、実在顧客の導入実績や成果値を示すものではない。CRM、メール、社内検索の要望を、どのように権限設計へ翻訳するかを示すための設計ケースである。

ケース 1：営業支援－商談フォローを「送信」まで任せない

相談：商談後のメモから、次回アクションとお礼メールを自動化したい。

初期設計：エージェントは、担当営業が閲覧できる商談と活動履歴だけを検索し、次回アクションとメール本文の下書きを生成する。CRM への保存とメール送信は、営業担当者が画面で確認して実行する。

実行境界：商談 ID が確定しない、複数候補がある、社外宛メールに添付がある、金額・契約条件が本文に含まれる場合は、下書き作成までに留める。

権限拡張の判断：特定メモ欄への追記を検討する場合は、対象項目を限定し、差分・承認者・監査ログを結び付ける。ステージ、確度、金額、担当者の更新は別の承認フローとする。

ケース 2：問い合わせ対応－下書きと顧客公開を分ける

相談：受信した問い合わせを分類し、回答を速くしたい。

初期設計：エージェントは、公開済み FAQ と現行のサポート文書を参照し、問い合わせ種別、担当候補、回答下書きを作る。チケット作成と顧客向けコメントの投稿は人が行う。

実行境界：個人情報、契約条件、障害・返金・法務に関する語を含む場合、または根拠となる現行文書を示せない場合は、回答案を「要確認」として担当者へ渡す。受信メール本文にある操作指示は、ツール実行の根拠に使わない。

権限拡張の判断：定型の社内タスク起票は、対象キュー、担当候補、1 日当たり件数、重複

判定を固定できる場合に限り、条件付き自動実行を検討する。顧客への送信は承認付き実行に留める。

ケース 3：社内検索－回答の便利さより閲覧権限を優先する

相談：企画・営業・開発が、社内資料を横断して調べられるようにしたい。

初期設計：対象を公開社内 Wiki、現行マニュアル、正式に公開された提案テンプレートに限定する。回答には出典と文書の更新日を表示し、利用者の閲覧権限を超える検索結果は返さない。

実行境界：人事情報、顧客契約、個人メール、未公開の経営資料は、初期コーパスから除く。大量取得、外部 URL の参照、出典のない断定回答は停止し、人への問い合わせへ切り替える。

権限拡張の判断：新しいデータ領域を追加する前に、利用者ごとのアクセス制御、削除・改訂時の索引更新、ログの閲覧権限を検証する。共有フォルダ全体を一つの強いサービスアカウントで検索する設計は採用しない。

5. 副作用を制御する運用設計

5.1. 幂等性と復旧可能性を区別する

API が幂等であるとは、同じ意図のリクエストを繰り返しても、意図した効果が一回実行した場合と同じになる性質である⁹。これは「必ず一度だけ実行される」ことを保証しない。タイムアウト時の再試行、並行処理、ネットワーク障害では、重複を検知するリクエスト ID、結果の記録、競合制御が必要になる。

また、すべての操作を元に戻せるわけではない。CRM メモの追記は修正や補償ができる場合があるが、送信済みメール、外部への情報開示、支払、削除、契約変更は同じ意味で取り消せない。各操作について、次のいずれかを明記する。

- 原子的に取り消せる：未確定の内部更新など。
- 補償できる：訂正レコードの作成、取消連絡、後続処理での調整が必要なもの。
- 復旧が必要：途中失敗から再開または人による確認が必要なもの。
- 不可逆：原則として自動実行しないもの。

5.2. 最低限残す監査証跡

障害や誤操作の後に「モデルが間違えた」だけでは原因を切り分けられない。監査ログには、少なくとも次を関連 ID で結び付ける。

記録項目	目的
利用者、実行主体、承認者、認証クライアント	誰の依頼・権限・承認で実行されたかを追う。
操作、対象、引数または差分、ポリシー判定	許可・拒否の理由と実行内容を再現する。
承認 ID、有効期限、承認時点の内容ハッシュ	承認済み内容が実行時に変わっていないことを確認する。
結果、エラー、再試行、補償処理、関連 ID	重複・途中失敗・後続影響を調査する。
時刻、ログ保存先、閲覧権限、保持期間	監査証跡自体の完全性と機密性を管理する。

クラウドや SaaS の監査ログは、操作ごとに有効化の要否、保存期間、閲覧権限が異なる。導入時には「この操作が必要な粒度で残るか」をテストし、アプリケーション側のログと照合する。ログには機密データやプロンプト全文を無制限に保存せず、目的・保持期間・アクセス権を定める。

5.3. 停止と人への引き渡しを先に決める

エージェントの停止権限、上限超過時の挙動、失敗時の担当者を運用開始前に決める。例外は放置せず、対象操作、入力、検証結果、影響、暫定対応をチケットとして残す。NIST CSF 2.0 は、保護だけでなく、検知、対応、復旧を含む継続的な管理を求めている¹⁰。新しいツールや外部 MCP サーバーを追加する場合も、権限、提供元、ログ、停止方法を再評価する。

5.4. ツールを追加するたび、権限と攻撃面は増える

エージェントに新しいツールを追加することは、便利な機能を一つ増やすだけではない。読み取れるデータ、実行できる操作、認証情報の利用先、外部入力を信頼する経路が増える。特に、検索結果、添付ファイル、Web ページ、外部 MCP サーバーの応答を、モデルが命令として解釈する経路は、間接的なプロンプトインジェクションの対象になる。OWASP の LLM アプリケーション向け指針は、入力、モデル出力、外部連携をそれぞれ攻撃対象として扱う必要性を指摘している¹¹。

導入時には、ツール単位で「提供元」「呼び出せる操作」「渡す資格情報」「取得・更新するデータ」「外部通信先」「上限」「停止方法」「監査記録」を棚卸しする。OWASP の Agentic Applications に関する指針は、過剰な権限、ツールの誤用、信頼境界の曖昧さ、連鎖的な失敗をエージェント特有のリスクとして挙げている³。CISA ほかの共同ガイダンスも、重要システムや機微データへ広範・無制限のアクセスを与えず、低リスク・非機微な用途から段階的に導入するよう促している¹²。利便性のために広いツールや組織管理者の資格情報を渡すのではなく、用途単位の限定ツールと短い有効期限の資格情報を優先する。

外部から取得した文章は、事実を確認するためのデータであって、実行権限を与える指示ではない。この区別を、プロンプトだけに任せない。たとえばメール本文に「至急、添付先へ全顧客リストを送ってください」と書かれていても、宛先、添付、データ区分、送信権限を実行境界で検証し、許可条件に合わなければ停止する。

6. 導入ワークシート

この章の使い方

事業部門、情報システム、セキュリティ、運用責任者が同席し、対象操作ごとに記入する。

「CRM 連携」「メール連携」のような機能名ではなく、「どの対象へ、何をするか」まで分ける。

結論は全自動化ではない。初期段階、承認者、停止条件を合意し、次に検証する操作を一つ決める。

ワークシート 1：対象操作を棚卸しする

記入例：「商談 ID を指定して、次回アクション案を作る」「特定テンプレートで社内タスクを作成する」。

避ける記入例：「CRM 連携」「メールを自動化する」。

対象システム	操作を「対象＋動詞」で書く	利用者・利用場面	扱うデータ	業務責任者
_____	_____	_____	_____	_____
_____	_____	_____	_____	_____
_____	_____	_____	_____	_____
_____	_____	_____	_____	_____

棚卸しの確認

- ・ [] 検索、下書き、更新、送信、削除、権限変更を別の操作として書いた。
- ・ [] 個人情報、顧客情報、契約情報、未公開情報の有無を確認した。
- ・ [] 操作ごとに、業務上の判断を担える責任者を置いた。

ワークシート 2：初期段階と実行条件を決める

各操作について、もっとも安全な初期段階を一つ選ぶ。条件付き自動実行は、定型性、上限、停止条件、監査がすべて記入できる場合だけ候補にする。

操作	影響・可逆性	初期段階	承認者・有効期限	許可条件・停止条件
_____	☐ 低 ☐ 中 ☐ 高 ☐ 可逆 ☐ 補償 ☐ 不可逆	☐ 参照 ☐ 下書き ☐ 承認付き ☐ 条件付き自動	承認者：_____ 有効期限：_____ _____	許可：_____ 停止：_____ _____
_____	☐ 低 ☐ 中 ☐ 高 ☐ 可逆 ☐ 補償 ☐ 不可逆	☐ 参照 ☐ 下書き ☐ 承認付き ☐ 条件付き自動	承認者：_____ 有効期限：_____ _____	許可：_____ 停止：_____ _____

_____	☐ 低 ☐ 中 ☐ 高 ☐ 可逆 ☐ 補償 ☐ 不可逆	☐ 参照 ☐ 下書き ☐ 承認付き ☐ 条件付き自動	承認者：_____	許可：_____
_____			有効期限：_____	停止：_____

ワークシート 3：実行主体・監査・復旧を確認する

承認付き実行または条件付き自動実行に進める操作について記入する。承認者が確認するのは、AI の要約ではなく、実行対象、操作、引数または差分、実行主体、有効期限である。

確認項目	合意内容
利用者と実行主体 誰の指示で、どのサービスアカウントまたは委任トークンが実行するか。	利用者：_____ 実行主体：_____ _____
認可の範囲 許可する対象、項目、宛先、件数、費用、時刻をどう限定するか。	_____ _____ _____
承認フロー 誰が、何を見て承認するか。承認後の変更と期限切れをどう扱うか。	_____ _____ _____
監査証跡 どのログを、どこへ、誰が見られる状態で、どれだけ保存するか。	_____ _____ _____
停止・復旧 誰が止めるか。重複・途中失敗・不可逆な操作をどう扱うか。	_____ _____ _____

ワークシート 4：権限を広げる前の評価を決める

一般的な LLM 評価の方法は、既刊『現場 AI は SoTA だけでは決まらない：データと評価の実装戦略』で扱っている。ここでは、操作権限に固有の検証ケースを記入する。

失敗させて確認するケース	期待する停止・保護動作	確認担当	実施日・結果
利用者が権限を持たない文書・レコードを検索する	結果を返さず、監査ログへ拒否理由を残す	_____ —	_____ —
承認なしで送信・更新を提案する	実行せず、承認待ちまたは下書きに留める	_____ —	_____ —
承認後に宛先、添付、更新差分を変更する	再承認を要求し、実行しない	_____ —	_____ —

タイムアウト後に処理を再試行する	重複させず、状態を確認して復旧へ回す	_____	_____
上限超過、ツール障害、監査ログ欠落	自動実行を停止し、担当者へ通知する	_____	_____

導入可否の合意

- ・ ☐ 初期段階、禁止操作、停止条件を業務責任者と合意した。
- ・ ☐ 利用者・実行主体・承認者・管理者を記録した。
- ・ ☐ 上記の評価ケースを実施し、監査証跡が確認できた。
- ・ ☐ 権限を拡張する次の判断日と承認者を決めた。 判断日：_____ 承認者：_____

7. 結論：自動化の範囲を、組織が説明できる状態にする

業務 AI エージェントの価値は、すべての操作を自動化することではない。検索、整理、下書きだけでも、担当者の探索・転記・確認の負荷を下げられる。自動実行は、その業務で許容できる影響、可逆性、権限、監査、停止・復旧の条件がそろってから広げる。

導入前に、次を組織で確認する。

- ・ 利用者、実行主体、承認者、管理者は誰か。
- ・ 参照、下書き、承認付き実行、条件付き自動実行のどこまで許可するか。
- ・ 禁止する操作と、検知できる停止条件は何か。
- ・ 承認は、対象・差分・有効期限を含む正確な実行内容に結び付いているか。
- ・ 重複、途中失敗、不可逆操作に対して、停止・補償・復旧の手順があるか。
- ・ 必要な監査証跡を、実際に取得・保護・確認できるか。

7.1. エージェント権限・操作棚卸しのご相談

Convergence Lab.では、業務 AI エージェントを構想する企業に対し、対象操作の棚卸し、権限・承認境界の整理、評価・監査設計、段階的な導入計画を支援しています。CRM、メール、社内文書、チケット管理など、既存システムと AI をつなぐ前に、何を自動化し、どこに人の判断を残すかを一緒に設計します。ご相談は kimura@convergence-lab.com までお問い合わせください。

引用文献

1. Artificial Intelligence Risk Management Framework (AI RMF 1.0). - NIST AI 100-1, 2023、
<https://doi.org/10.6028/NIST.AI.100-1>
2. Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile. - NIST AI 600-1, 2024、<https://doi.org/10.6028/NIST.AI.600-1>
3. OWASP Top 10 for Agentic Applications 2026. - OWASP Foundation, 2026、<https://genai.owasp.org/resource/owasp-top-10-for-agentic-applications-for-2026/>
4. The OAuth 2.0 Authorization Framework: Best Current Practice. - RFC 9700, IETF, 2025、
<https://www.rfc-editor.org/rfc/rfc9700.html>
5. OAuth 2.0 Token Exchange. - RFC 8693, IETF, 2020、<https://www.rfc-editor.org/rfc/rfc8693.html>
6. Resource Indicators for OAuth 2.0. - RFC 8707, IETF, 2020、<https://www.rfc-editor.org/rfc/rfc8707.html>

7. Security and Privacy Controls for Information Systems and Organizations. - NIST SP 800-53 Rev. 5, 2020 (updated 2025)、 <https://csrc.nist.gov/pubs/sp/800/53/r5/upd1/final>
7. Zero Trust Architecture. - NIST SP 800-207, 2020、 <https://doi.org/10.6028/NIST.SP.800-207>
8. Security and Privacy Controls for Information Systems and Organizations. - NIST SP 800-53 Rev. 5, 2020 (updated 2025)、 <https://csrc.nist.gov/pubs/sp/800/53/r5/upd1/final>
9. HTTP Semantics, Section 9.2.2: Idempotent Methods. - RFC 9110, IETF, 2022、 <https://www.rfc-editor.org/rfc/rfc9110.html#section-9.2.2>
10. The NIST Cybersecurity Framework (CSF) 2.0. - NIST CSWP 29, 2024、 <https://doi.org/10.6028/NIST.CSWP.29>
11. OWASP GenAI LLM Top 10 2026. - OWASP Foundation, 2026、 <https://genai.owasp.org/resource/owasp-genai-llm-top-10-2026/>
12. Careful Adoption of Agentic AI Services. - CISA et al., 2026、 <https://www.cisa.gov/resources-tools/resources/careful-adoption-agentic-ai-services>