

# AIエージェント 人材基準

AIエージェント時代に組織の成果を最大化する人材の定義  
— 「作る人」と「監督する人」

人とAIが共存する社会をつくる。その実現を、いちばん近くで支えるのが「人」である。本書は、AIエージェント時代に組織の成果を最大化する人材を、現場の役割として定義した育成の共通言語である。

DRAFT

発行日

2026年6月

バージョン

Version 0.5

発行元

一般社団法人AICX協会

執筆・監修：小澤 健祐（おざけん）

本書は現場での運用を通じて継続的に改訂される草案です。

## — EXECUTIVE SUMMARY

# 要旨

AIが「作業」を引き受ける時代において、人材の価値は作業の巧拙から、設計と監督へと移る。本基準は、その移行を「作る人（アーキテクト）」と「監督する人（オペレーター）」という現場の役割で捉え直し、研修・等級・採用に使える共通言語として定義する。

これまでの生成AIは、ChatGPTに代表されるように、人が指示（プロンプト）を工夫して「使う」ものだった。ところが新しいAIエージェントは、目的（ゴール）を渡せば、自分で段取りを決めて働く。人の関わりは、操作の反復から、任せて・見て・必要なら介入する「監督」へと変わる。「**使う**」から「**任せる**」へ——この変化が、求められる人材像を根本から書き換える。

そこで本基準は、これからの現場の人材を「土台＋3つの役割」で整理する。全員が持つべき土台が共通基礎スキル（AIエージェントの特徴・限界・リスクの理解）であり、その上に、何を任せるかを決める**ストラテジスト**、エージェントを作る**アーキテクト**、動くエージェントを監督する**オペレーター**の3つの役割が立つ。各役割は5つの管掌領域に分解し、習熟度は基礎／実践／高度の3レベルで、可否の行動基準と成果物（証拠）によって測る。

さらに本基準は、現場への導入で最も大切な「業務の棚卸し」と「現場の暗黙知の循環（AI-SECI）」を扱い、また自律エージェント特有の最大のリスクである「監督の形骸化（自動化バイアス）」への備えを、監督人材の核心能力に位置づける。運用にあたっては、自社の現在地（多くの企業はまだ初期段階にある）から逆算することを推奨する。**本基準は、いまの「使う」段階を否定するものではない。**むしろ現在地を出発点として尊重し、そこから次の段階へ無理なく進むための指針である。国のデジタルスキル標準（DSS）を置き換えるものではなく、AIエージェント実装という一点に解像度を加える民間リファレンスである。

## — CONTENTS

## 目次

|      |                                  |
|------|----------------------------------|
| 第1章  | 背景と目的 —— なぜ「使えるスキル」では足りないのか      |
| 第2章  | 本基準の全体像 —— 4つの軸                  |
| 第3章  | 実現ステップ（点→線→面→立体）と現在地             |
| 第4章  | 人材構成の本質 —— 作る人と監督する人             |
| 第5章  | 役割定義 —— 共通基礎スキルと3つの役割            |
| 第6章  | 各役割の5つの管掌領域                      |
| 第7章  | 現場への導入 —— 業務の棚卸しと暗黙知の循環（AI-SECI） |
| 第8章  | 監督設計 —— 自動化バイアスへの備え              |
| 第9章  | 導入段階別の人材像 —— 誰を、どの順で育てるか         |
| 第10章 | 習熟レベルと評価                         |
| 第11章 | 業務シーン別の適用例                       |
| 第12章 | 役割別スキル一覧 —— 3つの役割が取得すべきスキル       |
| 第13章 | 会社として、何から始めるか —— 導入アクションプラン      |
| 認定資格 | 本基準に対応した認定資格のご案内                 |
| おわりに | 代表理事メッセージ                        |
| 付録A  | DSS（デジタルスキル標準）との整合性              |
| 付録B  | 前提・出典について                        |

## CHAPTER 01

# 背景と目的 —— なぜ「使えるスキル」では足りないのか

生成AIの普及とともに、「AIを使いこなす人材になろう」というスキル論が広く語られるようになった。プロンプト（指示文）の書き方を学び、上手に質問できる人が重宝される、という考え方である。これは入口としては正しい。しかし、自律的に動くAIエージェントの時代を見据えると、この問いの立て方だけでは足りない。

理由は単純である。AIが賢くなるほど、上手な指示を出さなくても、そこそこの答えが返ってくるようになる。つまり「操作の巧拙」が生む差は、年々小さくなっていく。一方で、価値が大きくなっていくのは「何を・どこまでAIに任せ、その仕事をどう見届けるか」を決める力だ。これは操作のスキルではなく、**設計と監督のスキル**である。AIが優秀になればなるほど、こちらの重要性は増していく。

本書は、この変化を抽象的な能力論ではなく、現場で実際に分かれる役割として捉え直す。なぜなら「設計できる人材を育てよう」と言うだけでは、誰が・何を・どこまでできればよいのかが曖昧で、研修にも評価にも落とせないからだ。役割と、その中身と、習熟の度合いを具体的に定義してはじめて、人材は「育てる・測る・採る」の対象になる。

## 本基準の立場 —— 現在地を否定しない

本基準は、いまの「使う」段階を否定するものではない。多くの企業がまだ「使う」段階にいるという現在地を、当然の出発点として尊重する。そのうえで、そこから次の段階へ無理なく進むための地図として設計している。いま立っている場所を否定せず、そこから逆算して育てる——これが本書を貫く一貫した立場である。

## 本書が提供するもの・想定読者

本書が提供するものは、AIエージェント人材を育成・評価・採用するための共通言語である。経営者にとっては投資と組織設計の地図として、人事にとっては研修カリキュラムや等級・職務要件の土台として、現場のリーダーにとっては「自分のチームの誰が、どの役割を、どのレベルで担っているか」を見立てる物差しとして使えることを狙った。

## 本書のスコープ

本書が対象とするのは「AIエージェントを構築し、運用（監督）する」という一領域に限られる。デジタル人材育成の全体像は国のデジタルスキル標準（DSS）が広くカバーしており、本書はそれを置き換えるものではない。あくまで、AIエージェント実装という狭く深い現場に、民間の実装知で解像度を加える補助的なリファレンスである（DSSとの対応関係は付録Aを参照）。

## 本基準の全体像——4つの軸

本書を読み進める前に、全体像を押さえておきたい。本基準は「ステップ」「役割」「管掌領域」「レベル」という4つの言葉で構成される。これらはよく似て聞こえるが、実はまったく別の問いに答えるものである。混同すると一気に分かりにくくなるため、最初に4つの意味をはっきり分けておく。

図1－本基準の全体像・読み方の地図

- ① **実現ステップ**：「使う」から「任せる」へ——点→線→面→立体（第3章）
- ② **役割**：「任せる」時代の人材を、土台+3つの役割で定義（第4・5章）  
共通基礎スキル＝土台／ストラテジスト＝戦略／アーキテクト＝構築／オペレーター＝監督
- ③ **管掌領域**：各役割を「5つの管掌領域」に分ける（第6章）
- ④ **レベル**：3つの習熟レベル+役割別スキル一覧で育て・測る（第10・12章）

言い換えると、それぞれの軸は次の問いに対応する。① 実現ステップは「AIはどこまで自律しているか」、② 役割は「誰が関わるか」、③ 管掌領域は「その役割は何を担うか」、④ レベルは「その人はどれだけできるか」。①はAI側の話、②～④は人側の話である。

表1－本基準を構成する4つの軸

| 軸        | 答える問い          | 内容                                                       |
|----------|----------------|----------------------------------------------------------|
| ① 実現ステップ | AIはどこまで自律しているか | 点→線→面→立体。エージェントがどこまで自律して動くかの段階（第3章）。AI側の話で、人の役割そのものではない。 |
| ② 役割     | 誰が関わるか         | 共通基礎スキルを土台に、戦略（ストラテジスト）・構築（アーキテクト）・監督（オペレーター）の3役割（第5章）。  |
| ③ 管掌領域   | その役割は何を担うか     | 各役割が責任を持つ範囲を5つに分けたもの（第6章）。例：監督なら観測・評価・介入・調整・循環。          |
| ④ レベル    | どれだけできるか       | その人の習熟度。基礎／実践／高度の3段階（第10章）。役割・管掌領域とは別の軸。                 |

**読み方のコツ**。〔ステップ〕はAIがどこまで自律しているか、〔役割〕は誰が関わるか、〔管掌領域〕はその役割が担う範囲、〔レベル〕はその人の習熟度。たとえば「面の段階の会社で、オペレーターを担う人が、その観測の領域を実践レベルでこなす」というように、4つは掛け合わせて使う。縦と横が違う軸なので、混ぜずに読むのがコツである。

## 実現ステップ（点→線→面→立体）と現在地

AIエージェントは、ある日いきなり完成形が現れるわけではない。企業の現場では、ごく簡単な使い方から始まり、段階的に「任せられる範囲」が広がっていく。本基準では、その広がりを「点を線に、面を立体に」という4つの段階で表す。

これは、数学でいう0次元から3次元になぞらえた比喻である。点はそれ単体（1回きりの利用）、線は点がつながった流れ（決まった手順の自動化）、面は線が広がった状態（複数のツールやデータを横断）、立体は面に「時間という奥行き」が加わった状態——いま目の前の処理をこなすだけでなく、AIが未来に向けた目標を自ら立て、計画して動き続ける。次元が一つ上がるごとに、AIに任せられる範囲と、AIが自分で決める割合が増えていく。軸は一貫して「ルート（段取り）と目標を、誰がどこまで決めるか」である。

図2－AIエージェントの実現ステップ（点→線→面→立体）



使う — プロンプトが主役（特に点）

任せる — コンテキスト・ハーネス設計+監督

### 身近なツールでいうと

抽象的なままだと分かりにくいので、各段階を、多くの企業が実際に使うツールになぞらえて説明する。

- **点（単発チャット）** —— ChatGPTやGoogleのGemini、Microsoft Copilotなどに、その都度チャットで質問して使う段階。便利だが、毎回人が指示を出し、結果も人がそのまま使う。最も多くの人が今いる地点である。
- **線①（チャットボット）** —— よく使う指示を一度「型」にして、繰り返し使えるボットにする段階。ChatGPTの「GPTs」、Geminiの「Gems」、MicrosoftのCopilot Studioで作る業務特化ボットがこれにあたる。一度作れば、あとは呼び出すだけ——ここから「任せる」の入口になる。
- **線②（ワークフロー）** —— チャットの枠を超え、社内システムやデータベースとつないで、複数の処理を自動で流す段階。Power Automate、Zapier、Difyのようなツールで、「申請が来たら内容を判定して登録し、担当に通知する」といった一連の流れを組む。
- **面（半自律AIエージェント）** —— 1本の決まった流れを超えて、複数のツールやデータを横断しながら、AIが判断の一部を担い始める段階。Copilot CoworkやMicrosoft Scout、Claude Code、Gemini Sparkなどがこれにあたり、人が要所で確認・指示しつつ、AIがある程度自分で段取りを進める。
- **立体（完全自律AIエージェント）** —— 人は細かなゴールすら与えず、方向性を示すだけ。AIが時間軸を持って、未来に向けた目標を自ら設定し、計画を立て、継続的に動き続ける段階。ただし、人手をまったく介さずにここまで到達した実例は、現時点ではほぼ存在しない。各社が目指している到達点であり、現状はまだ将来像に近い。

### 「使う」から「任せる」へ——質的な変化

ここで大切なのは、人とAIの関係が「使う」から「任せる」へ移っていくことだ。純粋に「使う」のは、単発チャット（点）の段階だけ。人がその場で指示を出して結果を受け取る、プロンプトエンジニアリングが主役の世界である。だが、チャットボット（GPTsやGems）まで進むと様子が変わる。一度ボットを作ってしまうと、あとは特定の

言葉を入れるだけで処理が進む。プロンプトを工夫して「使う」というより、作ったポットに「任せる」段階に入っているのだ。ワークフロー・面・立体と進むほど、その「任せる」の度合いは深まっていく。だからこそ、プロンプトの巧拙が効くのは主に「点」であり、それ以降は「何を・どこまで任せ、どう監督するか」という設計と委任が主役になる。

表2ー「使う」と「任せる」・どこで切り替わるか

|        | 「使う」が中心           | 「任せる」が中心                              |
|--------|-------------------|---------------------------------------|
| 関わり方   | 人が指示（プロンプト）を出して使う | AI・システムに任せて自走させる                      |
| 中心となる力 | プロンプト（とくに「点」で重要）  | コンテキスト・ハーネス設計／委任／監督                   |
| 該当ステップ | 点（単発チャット）         | チャットボット／ワークフロー／面／立体                   |
| イメージ   | ChatGPTにその都度質問する  | GPTsで作ったポットに任せる、処理が自動で流れる、AIが目標を立てて自走 |

## プロンプトから、コンテキスト・ハーネスへ

「使う」から「任せる」へ移ると、効くスキルそのものが変わる。プロンプトエンジニアリングが主役なのは「点」まで。任せる段階に進むと、主役は次の2つへ移っていく。ひとつは**コンテキストエンジニアリング**——AIに「何を・どんな情報を渡すか」を設計することだ。指示文だけでなく、参照すべきデータ（RAG）、過去のやりとり（メモリ）、役割や制約の設定までを含めて、AIが良い判断をくだせる文脈（コンテキスト）を整える。もうひとつは**ハーネスエンジニアリング**——AIが実際に動くための「土台（ハーネス）」を設計することだ。使えるツール、与える権限、越えてはいけないガードレール、処理を繰り返すループなど、エージェントが自律して安全に動ける周辺の仕組みを組み上げる。どちらもアーキテクト（作る人）の中核であり、プロンプトの巧拙に代わって、これらの設計力がものを言うようになる。

### ワンポイント

現在地を見誤らないこと。いま大半の企業は、単発チャット（「使う」）から、チャットボットやワークフロー（「任せる」の入口）あたりにいる。面・立体はゴールであって、多くの現場の「今」ではない。とくに完全自律の「立体」は**実例がほぼなく**、将来像に近い。運用時は「自社は今の段階か」から逆算してほしい。

小澤健祐（おざけん）

ツール名は2026年時点の例であり、各社の名称・機能は頻繁に更新される。ワークフローは決めた流れどおりに動くが、面以降はAIが自分で判断するぶん、監督という役割が本格的に効いてくる（第9章で詳述）。

## 人材構成の本質——作る人と監督する人

AIに仕事を「任せる」段階に進むと、人とAIの関係そのものが変わる。ここに本基準の中心となる考え方がある。

従来の道具——PCやSaaS、そしてAIへの単発プロンプト（点）——は、人が手を動かして「使う」ものだった。表計算ソフトは人がセルに数式を入れて初めて動くし、単発プロンプトもまた、人が毎回その場で指示を打ち込んで初めて答えが返る。だがAIエージェントは違う。目的を渡せば、自分でルートを決めて働き続ける。人は一手ずつ操作しない。つまり、人がやるのは「使う」ことではなく、任せて・出力を見て・必要なら介入する「**監督**」である。

この変化を踏まえると、これからの現場の人材は、大きく2つに分かれる。エージェントを**作る人**と、動くエージェントを**監督する人**である。前者はエージェントが動き出す前の世界をつくり、後者は動き出した後のエージェントを見届ける。役割の性質がまったく違う。

表3－作る人と監督する人

|        | 作る人 —— アーキテクト                                | 監督する人 —— オペレーター                     |
|--------|----------------------------------------------|-------------------------------------|
| 担うこと   | エージェントそのものを設計・構築する。動き出す前の世界をつくる。             | 動いているエージェントを監督する。任せて、見て、異常や例外で介入する。 |
| 役割の中心  | モデル・文脈・ツール・権限・評価の仕組みを組み上げ、「自律して動ける状態」を成立させる。 | 出力の良し悪しを見極め、止める線を超えたら介入し、改善を回す。     |
| 現場での人数 | 監督する人より少数。大企業はCoEを核に現場でも育てる。                 | 各職種に広がり、最も人数が多くなる。                  |

### なぜ多くの人が「監督する人」になるのか

たとえるなら、作る人は工場のラインを設計する技術者であり、監督する人は、自律稼働する設備を見守り、異常があれば手を打つオペレーターである。エージェントを一から作れる人は、専門性が高く、どうしても少数になる。一方、出来上がったエージェントを自分の業務で監督する人は、あらゆる部署に広がる。だから人数の上では、監督する人が最も多くなる。「**AIを使う人**」を「**AIを監督する人**」と捉え直すこと——ここが本基準の出発点である。ただし、その監督は放っておくと形骸化する（第8章）。だからこそ、監督にも設計が要る。

## 役割定義 —— 共通基礎スキルと3つの役割

「作る人」と「監督する人」を中心に据えると、現場の人材構成は3つの役割で整理できる。全員が持つ共通基礎スキルを土台に、上流で何を任せるか決める戦略、エージェントを作る構築、動くエージェントを監督する運用が立つ。順に見ていく。

図3－AIエージェント人材の構成・土台と3つの役割



### 土台：共通基礎スキル（AIエージェント・リテラシー）

3つの役割の前提として、全職種・全等級が共通して持つべき「理解」の層がある。何かを実務として行うより前に、AIエージェントとは何か、何ができて何ができないか、どんなリスクがあるかを、ひととおり理解している状態だ。これは特別な技術ではなく、AIエージェントと安全に付き合うための土台となる知識である。具体的には次を押さえている。

表4－共通基礎スキルで理解しておくこと

| 項目       | 理解しておくこと                              |
|----------|---------------------------------------|
| エージェントとは | AIエージェントの定義と、点→線→面→立体の自律度の段階          |
| できること・限界 | 得意なこと・苦手なこと・できないこと、出力のばらつきや再現性の限界     |
| 主なリスク    | 誤出力（ハルシネーション）・情報漏洩・権限の暴走・バイアスなどの典型リスク |
| 安全な扱い方   | ガードレール・権限・人による確認の必要性、機密とコンプライアンスの基本   |
| 人とAIの関係  | 「使う」から「任せる・監督する」への変化、最終責任は人にあるという原則   |

共通基礎スキルは「実務」ではなく「理解」である。だから、監督を実際に行うオペレーターとは別物として、土台に置く。全員がここを押さえたうえで、実際にエージェントを監督し始めた人が、オペレーターの役割（基礎レベル）に入っていく。なお本基礎は、生成AIの一般リテラシー（プロンプトの基礎や一般的なAIリスク。既存の研修・資格で代替可）の上に乗る、「自律して動くエージェント特有」の理解に絞っている。

### その上に立つ3つの専門役割

土台の上に、性質の異なる3つの役割が立つ。「何を任せるか決める」ストラテジスト、「作る」アーキテクト、「監督する」オペレーターである。

表5－3つの専門役割

| 役割                           | 位置づけ                                            | 主担当     |
|------------------------------|-------------------------------------------------|---------|
| <b>ストラテジスト</b><br>何を任せるか決める人 | どの業務を・どこまでエージェントに委譲するかを決め、費用対効果（ROI）と組織設計を担う上流。 | 戦略・委任設計 |
| <b>アーキテクト</b><br>作る人         | エージェントを設計・構築する。文脈・ツール・権限・評価の仕組みを組み上げる。          | 構築・実装   |
| <b>オペレーター</b><br>監督する人       | 動くエージェントを監督し、評価・介入・改善を回す現場の主役。人数が最も多い。          | 運用・監督   |

#### 既存の職種を「置き換える」のではなく「重ねる」

この3役割は、既存の職種を置き換えるものではない。営業・人事・法務・経理といった今の職種はそのままに、その上へ「この人はいま、どの役割としてエージェントに関わっているか」という視点を重ねるものである。一人が複数の役割を兼ねることも、既存の職務と両立することも当然に起こる。国のデジタルスキル標準（DSS）との関係も“置き換え”ではなく“重ね合わせ”であり、既存の人材定義を壊さずに解像度だけを上乘せできる（付録A）。

この3役割は「分業の体制」であって、「3人必要」という意味ではない。規模や段階によっては1人が兼ねることも多い。誰が何人必要かは第9章「導入段階別」で扱う。

## 各役割の5つの管掌領域

この章の要点は一つ——役割は「5つの担当範囲（管掌領域）」に分解すると、研修や評価にそのまま落とせる、ということだ。

役割の名前だけでは、その人が実際に何を担うのかが見えない。そこで本基準では、3つの役割を、それぞれ5つの管掌領域（その役割が責任を持つ5つの範囲）に分けて定義する。役割の「中身」が具体的に見えることで、研修の到達目標や等級・職務要件を組む土台になる。3つの役割が、いずれも5つの管掌領域できれいに並ぶのが、本基準の骨格である。

表6－3つの役割 × 5つの管掌領域

| # | ストラテジスト（戦略）                          | アーキテクト（構築）                                     | オペレーター（監督）                            |
|---|--------------------------------------|------------------------------------------------|---------------------------------------|
| 1 | 可視化／業務とプロセスの棚卸し、現状とボトルネックの把握         | 基盤（モデル設定）／どのAIモデルを使うか、コスト・精度・速度のバランスを設計        | 観測／何を・どの頻度で見えるか決め、出力と挙動をモニタリング        |
| 2 | 選定／「任せる／人が担う／協働する」の切り分け              | 知能（コンテキスト整備）／プロンプト・RAG・メモリで、AIに渡す文脈（コンテキスト）を設計 | 評価／基準に照らした出力の良し悪しの判断（自動化バイアスに流されない確認） |
| 3 | 委任設計／委任範囲・リスク許容度・成功条件の定義             | 行動（ツール権限整備）／使えるツールの接続、権限、ガードレールを設計             | 介入／例外・逸脱の見極めと、止める線を越えた時の介入・エスカレーション   |
| 4 | 投資／ROIの試算と、着手の優先順位・投資配分              | 統合／ワークフロー・計画・マルチエージェント連携の構造化                   | 調整／結果をもとにした委任範囲・止める線・コンテキストの更新        |
| 5 | 変革／会社全体のAI変革（AX）の進み具合に合わせた、役割・組織の再設計 | 運用／評価・観測の仕組みづくりと本番運用への作り込み                     | 循環／知見の形式知化と、組織での共有・再利用                |

ストラテジストは「可視化→選定→委任設計→投資→変革」と、業務を見渡してから委ね方を決め、投資と組織設計へ進む。アーキテクトは「基盤→知能→行動→統合→運用」と、AIの土台から組み上げて運用に耐える形にする。オペレーターは「観測→評価→介入→調整→循環」と、見て・判断して・手を打って・直して・組織に広げる。これは一方向の工程というより、繰り返すマネジメントサイクルである。

3つの役割は独立ではなく連結する。アーキテクトの管掌領域5「運用（評価・観測の仕組みづくり）」が、オペレーターの監督を成立させる土台になる。ストラテジストの委任設計が甘ければ、下流の構築も監督も空回りする。上流・構築・監督がかみ合って、初めて「効く自律化」になる。なお、アーキテクトの「知能」はコンテキストエンジニアリング、「行動・統合」はハーネスエンジニアリングにあたる（第3章）。

# 現場への導入—— 業務の棚卸しと暗黙知の循環（AI-SECI）

AIエージェントを現場に入れるとき、最大の壁は技術ではない。「現場をどれだけ分かっているか」である。現場の実態を知らないまま上から導入すると、業務に合わず、結局使われずに終わる。本章では、現場目線でエージェントを根づかせるための要点を整理する。

## 出発点は「業務の棚卸し」

導入の第一歩は、現場の業務を一つずつ洗い出すことである。どこが定型作業で、どこに人の判断が要り、どこがボトルネックになっているか。そのうえで「AIに任せる／人が担う／協働する」を切り分ける。これはストラテジストの管掌領域<sup>1</sup>「可視化」にあたる（第6章）。この棚卸しを飛ばすと、効果の薄いところにAIを入れてしまい、「導入したのに変わらない」という結果になりやすい。逆に、現場の業務がきちんと見えていれば、最も効くポイントから着手できる。

## 現場の「暗黙知」という壁

現場には、ベテランの勘や経験で回っている「暗黙知」がある。「この条件なら念のため確認する」「この取引先は例外扱い」といった判断基準が、言葉になっていないまま共有されている。エージェントに任せるには、この暗黙知を、ルール・手順・プロンプト・評価基準といった「形式知」に変換しなければならない。ここが導入の最難所であり、現場を知る人（多くはオペレーター）こそが、その「翻訳者」になる。現場の暗黙知を汲み取れるかどうか、エージェントが使い物になるかどうかを分ける。

## 暗黙知を循環させる——AI-SECIモデル

暗黙知と形式知の変換は、一度きりではなく、ぐるぐると循環させることが肝心だ。理論に見えるかもしれないが、やることはシンプルである。たとえば問い合わせ対応なら——①ベテランの「こう見分けて、こう返す」という勘を聞き出し、②まずはチャットボットや簡単なワークフローとして形にして任せ、③FAQや過去履歴を組み込んでコンテキストを強化し、より高度なワークフローへ育て、④現場で使って外れを直しながら、人もAIも学んでいく。この①～④を回し続けるのがAI-SECIだ。経営学者・野中郁次郎のSECIモデルを、AIエージェント時代に応用したものである。

図4－AI-SECIモデル・現場の知を、AIと一緒に育てる4ステップ

**① 共有・認識する**

SOCIALIZATION－共同化（暗黙知→暗黙知）

暗黙知を暗黙知のまま、共有し「認識する」段階。ワークショップでの共有や、ベテランへのインタビュー・観察を通じて、現場に眠る勘・判断（暗黙知）を掘り起こし、まず存在をつかむ。

**② 書き出す**

EXTERNALIZATION－表出化（暗黙知→形式知）

暗黙知を、ルール・プロンプトで指示し、チャットボット～半自律エージェントとして、まず形にして任せる。まず動かす。

**③ つなげる**

COMBINATION－結合化（形式知→形式知）

ナレッジ・ツールを組み合わせ、より高度なエージェントに育てる。コンテキストを強化（RAG）し、ツール・データを統合、ハーネスを整備する。

**④ 使って磨く**

INTERNALIZATION－内面化（形式知→暗黙知）

現場で使い、観測・評価する中で新しい勘・コツを得る。得た暗黙知（新しい勘）が①に戻り、循環する。

①→②→③→④→① 現場とAIで、知を回し続ける

この4ステップは、本基準の用語とも対応している。①業務の棚卸し・暗黙知の共有（共同化）→②チャットボット～半自律エージェントとして、まず形にする（表出化）→③コンテキスト強化・高度なワークフロー化で育てる（結合化）→④使って観測・評価し、新たな勘を得る（内面化）、という流れだ。この循環を回すのが、オペレーターの管掌領域5「循環」（第6章）である。

**人事業務で考えてみる ―― 入社手続きの例**

たとえば、入社手続きをAIエージェントに任せるケースを、一つの流れで追ってみる。①**聞き出す（共同化）**――ベテラン人事担当者に「雇用形態によって必要書類が変わる」「外国籍社員は在留資格の確認が要る」「社会保険の加入条件には例外がある」といった経験・判断のポイント（暗黙知）をヒアリングし、観察・対話しながら把握する。②**書き出す（表出化）**――聞き出した暗黙知を、チェックリスト・業務フロー・プロンプトとして整理し、まずAIが実行できる形（形式知）にする。③**つなげる（結合化）**――就業規則・社会保険制度・雇用契約ルール・過去の対応事例を組み合わせ、RAGなどでコンテキストを強化し、雇用形態や勤務条件に応じて適切に案内できるエージェントへ育てる。④**使って磨く（内面化）**――実際に入社手続きで使いながら観測・評価し、例外対応や新しい判断ポイントを蓄積する。得た知見（新しい暗黙知）は再び①へ戻り、AIと現場と一緒に学び続ける循環になる。

現場の暗黙知をAIへ移し、AIの働きから得た新しい知見を現場へ戻す。つまりAIエージェントの導入は、一度作って終わりではなく、現場と一緒に知を育て続ける営みなのだ。この循環が回り始めると、エージェントは現場に馴染み、組織の財産として蓄積されていく。

## 監督設計——自動化バイアスへの備え

はじめに言葉を整理しておく。本書でいう監督とは、動くAIエージェントを観測・評価し、必要に応じて介入・改善していく継続的な活動を指す。その活動を担う役割がオペレーターであり、活動を個人任せにせず仕組みとして実現することを監督設計と呼ぶ。

「監督する人」を中心に置かなら、避けて通れない現実がある。それは、**監督は放っておくと形骸化する**ということだ。

エージェントが正しい確率が上がるほど、人は確認をやめていく。最初はきちんと見ていても、「どうせ合っている」が続くと、やがて中身を見ずにハンコを押すだけになる。これは人間の自然な傾向で、**自動化バイアス**と呼ばれる。さらに監督疲れも重なる。皮肉なことに、AIの精度が上がって安心したところが、いちばん事故が起きやすい状態に近い。残りの数%の誤りが、誰のチェックも通らずにすり抜けてしまうからである。

### 設計思想——人ではなく、仕組みで守る

したがって監督は、担当者の気合や注意力に頼ってはならない。個人の注意力に依存せず、**仕組みで守る——人ではなく仕組みで担保する**必要がある。これは、確認という行為を「気づいたらやる」個人の努力から、「決めた通りに必ず回る」再現性のある仕組み（メカニズム）へと変えることを意味する。本基準は、この「監督が形骸化する前提で、仕組みとして備えを設計できること」を、オペレーター（監督）の核心能力に位置づける。具体的には次の3つである。

表7ー 監督を形骸化させないための3つの備え

| 備え                   | 内容                                                                                      |
|----------------------|-----------------------------------------------------------------------------------------|
| ① サンプル確認<br>(抜き取り確認) | 全件は見ない。そのかわり、リスクの高い出力と、一定割合とを、あらかじめ決めたルールで抽出し、必ず人が確認する仕組みにする。「気づいたら見る」を「決めて見る」に変えるのが要点。 |
| ② 止める線<br>(停止条件)     | 「ここを超えたら必ず人に回す」という境界を、金額・対象・例外パターンなどで事前に決めておく。判断の難所をAIに丸投げせず、人に戻す設計。                    |
| ③ 確認の点検              | 見逃し率・介入率・確認にかけた時間を計測し、「確認そのものが形骸化していないか」を見張る。仕組みが効いているかを定点観測する。                         |

この3つは、第7章で述べた「現場の暗黙知」とも深く結びつく。どの出力が危ないか、どこに「止める線」を引くべきかは、現場の判断基準そのものだからである。監督設計とは、現場の暗黙知を、エージェントを見届ける仕組みへと翻訳する作業でもある。

## 導入段階別の人材像——誰を、どの順で育てるか

必要な人材は、会社が今どの段階（点→線→面→立体）にいるかで変わる。「これからはAI監督の時代だから、全社員を監督できるようにしよう」と一足飛びに走っても、そもそも監督すべき自律エージェントがまだ社内にはない、ということが起きる。大切なのは、自社の現在地を見定め、育てる順番を間違えないことだ。段階ごとに、必要な人材の重心を見ていく。

### 点・線の段階——大半の企業はここにいる

単発でChatGPTを使ったり、一部の部署がボットや簡単な自動化を試したりしている段階である。自律して動くエージェントはまだ少ない。ここで最優先すべきは、全員の共通基礎スキルの底上げと、少数の作り手（アーキテクト）の確保だ。監督（オペレーター）の出番はまだ小さいので、無理に監督人材を量産する必要はない。まずは基礎の理解を全社に広げ、業務の棚卸しを進めながら、小さな構築から経験を積むのがよい。

### 面の段階——監督が本格的に必要なになる

エージェントが複数のツールを横断し、判断の一部を自分で担い始める段階である。ここで初めて、監督するオペレーターが本格的に必要なになる。同時に、「どの業務を・どこまで任せるか」を決めるストラテジストの役割が重みを増す。現場の担当者の多くが、自分の業務を担うエージェントの監督役へと回り始める。第8章の監督設計が効いてくるのも、この段階からである。

### 立体の段階——組織的な運用体制が前提になる

ゴールを渡せばAIが自走する段階では、監督・運用・ガバナンスを組織として回す必要がある。専任のアーキテクトとオペレーターの体制に加え、責任の所在や権限、ルールの設計が欠かせない。ここまで到達する企業は、今はまだごく一部である。

表8－導入段階別に必要な人材の重心

| 段階                          | 育成の重心                                      | 状況                       |
|-----------------------------|--------------------------------------------|--------------------------|
| <b>点・線</b><br>単発チャット～ワークフロー | 全員の共通基礎スキル＋少数の作り手（兼任可）。監督の出番はまだ小さい。        | 大半の企業はここ。基礎の底上げと小さな構築から。 |
| <b>面</b><br>半自律AIエージェント     | 監督（オペレーター）が本格化。ストラテジストが委任範囲を決め、現場担当が監督に回る。 | 監督設計（第8章）が効き始める。         |
| <b>立体</b><br>完全自律AIエージェント   | 専任のアーキテクト／オペレーター体制と、責任・権限・ガバナンスの設計。        | 今はごく一部。組織的な運用体制が前提。      |

### 規模によって、役割の持ち方は変わる

役割を3つに分けるのは、あくまで「分業の体制」であって、「3人雇わなければならない」という意味ではない。中小企業や非IT部門では、1人～少数がこの3役割を兼ねるのが現実的である。その場合、エージェントの構築そのものはノーコードツールやSaaS、外部パートナーに任せ、社内は「何を任せるか決める（戦略）」と「ちゃんと監督する」に集中するとよい。一方、大企業では、エージェントの構築を情報システム部門に集約するのではなく、CoE（Center of Excellence）のような横断組織を核に、各現場が自らエージェントを構築できる人材を育てる「分散型」が有効である。CoEが標準・ガードレール・ナレッジを整え、現場のアーキテクト人材を支援することで、構築のス

ピードと現場への適合を両立できる。どちらの形であっても、「いま自分がどの役割として、このエージェントに関わっているのか」を意識することに意味がある。

## 習熟レベルと評価

「基準」と名乗る以上、レベルは測れなければならない。誰が見ても同じように「この人はこのレベル」と判断できなければ、人材育成や評価、等級設計などに活かせないからだ。本基準は習熟度を3段階で定義し、各レベルに「合格の目安（できる）」を主軸に、行動で線を引く。抽象的な形容詞ではなく、行動で測るのが要点である。

レベルは「役割」ではなく「管掌領域ごと」に測る

このレベルは役割そのものにつくのではなく、第6章で示した「5つの管掌領域」のそれぞれについて測る。同じ人でも、ある管掌領域は実践、別の領域は基礎、ということが起こりうる。各管掌領域×各レベルの具体的なスキルは、第12章の役割別スキル一覧で展開する。

図5－習熟の3レベル（基礎 → 実践 → 高度）



習熟度が上がる ↗

表9－習熟レベルの定義（できること）

| レベル    | できる（合格の目安）                                              | まだの目安                          |
|--------|---------------------------------------------------------|--------------------------------|
| Lv1 基礎 | 決められた基準でエージェントの出力を確認し、おかしいと気づいたら止めて人に上げられる。             | 出力をそのまま通してしまい、何を確認すべきか分かっていない。 |
| Lv2 実践 | 自分の担当業務のエージェントを、手順書がなくても一人で運用・監督でき、例外やイレギュラーにも自分で対処できる。 | 用意された手順の範囲でしか動けず、外れると手が止まる。    |
| Lv3 高度 | エージェントの評価・監督・運用の仕組みを整え、他の人がそれで回せる状態に標準化できる。             | 自分一人なら回せるが、仕組み化・標準化ができず属人的。    |

さらに、レベルを言葉だけで判定すると評価がぶれる。そこで本基準は、各役割が各レベルで実際に残す成果物を「証拠」として用いる。下表は、その成果物の例である。

表10－各役割・各レベルで残す成果物（証拠）

| 役割      | 基礎                             | 実践                                       | 高度                                    |
|---------|--------------------------------|------------------------------------------|---------------------------------------|
| ストラテジスト | 自部署の業務棚卸しリスト（任せる／担う／協働の仕分け）    | 委任業務の優先順位マップ＋ROI試算、PoC企画書                | 全社の委譲ロードマップ、AX投資の配分計画、役割・組織設計案        |
| アーキテクト  | ガイドに沿った単機能エージェント（定型業務の自動化）     | 業務要件から設計したエージェント（文脈・ツール・権限・ガードレール込み）＋設計書 | マルチエージェント構成、評価・観測基盤、本番運用に耐える改善体制      |
| オペレーター  | 日次の出力チェック記録、基準に沿った異常時のエスカレーション | 監督手順書（サンプリング基準・止める線・例外対応フロー）、改善ログ        | 監督ダッシュボードと介入基準の設計、見逃し率の管理、監督チームの体制づくり |

レベルの本当の証明は、資格や研修の修了ではない。担当したエージェント導入が、設定した成果（時間削減・品質・コスト等）を実際に出し、事故なく回り続けたか。成果から人材レベルへ評価を返すループを持って、初めて「できる」が裏づけられる。

## 業務シーン別の適用例

役割は、現場では必ずセットで現れる。同じ業務にエージェントを入れるとき、誰かが「何を任せるか」を決め（戦略）、誰かが「エージェントを作り」（構築）、誰かが「動くエージェントを監督する」（監督）。この3つは、業務の種類が変わっても共通して必要になる。代表的な9つのシーンで、3つの役割がそれぞれ何をするかを並べた。ここでは特定の役職名ではなく、3役割の働きそのものを示している。

### これは「標準」ではなく、適用イメージの一例

表11に示すのはあくまで適用イメージの一例である。どこまで任せ、どう監督するか（たとえば人事での「バイアスの点検の仕方」）は、業界・企業・体制によって大きく異なる。ここでの記述を「唯一の標準」と読まず、3つの役割の働き方をつかむ手がかりとして、自社の状況に合わせて読み替えてほしい。

表11－業務シーン別・3つの役割の働き（適用イメージの一例）

| シーン                       | ストラテジスト（戦略）                       | アーキテクト（構築）                          | オペレーター（監督）                          |
|---------------------------|-----------------------------------|-------------------------------------|-------------------------------------|
| <b>営業</b><br>提案・見積もり      | 提案業務のどこを任せるか、見積りの承認ラインと費用対効果を決める  | 案件情報の収集・提案書ドラフト・見積り計算を行うエージェントを構築する | 数字と提案の筋を確認し顧客提出前に差し戻す。金額の上限超は必ず人が承認 |
| <b>マーケティング</b><br>コンテンツ制作 | どの制作物をどこまで任せるか、ブランド基準とチェック体制を決める  | 記事・SNS・広告コピーの草案を作るエージェントを構築する       | 事実・トーン・著作権を確認し、公開前に必ずレビューする         |
| <b>カスタマーサポート</b><br>一次対応  | 一次対応のどこまで任せるか、品質基準と引き継ぎ方針を決める     | FAQ連携・問い合わせ分類・回答ドラフトのエージェントを構築する    | 回答品質を監視し、難しい案件を引き取り、改善を回す           |
| <b>経理・財務</b><br>請求・経費処理   | どの処理を任せるか、金額のしきい値と監査要件を決める        | 請求書の読み取り・仕訳・経費チェックのエージェントを構築する      | 仕訳と証憑を照合し、しきい値超や例外を承認・差し戻す          |
| <b>人事</b><br>採用・社内問い合わせ   | どこを任せるか、公平性と機密保持の基準を決める           | 書類スクリーニング・候補者連絡・社内Q&Aのエージェントを構築する   | 可否に関わる判断は必ず人が確認し、バイアスを点検する          |
| <b>法務</b><br>契約レビュー       | どの契約類型を任せるか、リスク許容度とエスカレーション条件を決める | 契約書のリスク抽出・条項比較・要約のエージェントを構築する       | 抽出されたリスクを精査し、重要条項は必ず人が最終判断する        |
| <b>製造現場</b><br>設備保全・検査記録  | どの記録・判断を任せるか、安全と品質の線を決める          | 設備データ連携・異常検知・検査記録のエージェントを構築する       | AIの判定を現物と照合し、異常を見極めて止める             |
| <b>物流・調達</b><br>在庫・発注     | どの発注・補充を任せるか、発注上限と欠品リスクの線を決める     | 需要予測・在庫監視・発注ドラフトのエージェントを構築する        | 発注量と在庫を確認し、上限超や異常な需要変動に介入する         |
| <b>自治体・公共</b><br>住民問い合わせ  | どの問い合わせを任せるか、正確性と公平性の基準を決める       | 制度情報を参照する一次対応エージェントを構築する            | 回答の正確さを監視し、制度の例外や苦情を引き取る            |

役割は肩書きで固定されない。同じ人が、ある業務ではオペレーター（監督）、別の業務では基礎を理解して使う側、というように関わり方は変わる。大切なのは「いま自分は、どの役割としてこのエージェントに関わっているか」を意識できることである。

## CHAPTER 12

# 役割別スキル一覧—— 3つの役割が取得すべきスキル

第6章で示した各役割の「5つの管掌領域」を、現場で取得すべき具体的なスキルへと展開する。**研修カリキュラム・資格・職務要件を設計する際のベースとして活用できる**よう、3つの役割ごとに、第6章の「5つの管掌領域」で束ね、各スキルに習熟レベル（Lv1＝基礎／Lv2＝実践／Lv3＝高度）を付した。ただし、とくに職務要件は、企業ごとの組織体制・役割分担・業界特性によって変わる。そのまま当てはめるのではなく、自社の段階や職種に合わせてカスタマイズして用いることを前提としてほしい。なお、各スキルの習熟は最終的に、表10の「成果物（証拠）」と、その成果物が実際に成果を出したか（第10章）で裏づけられる。スキル一覧は入口であり、評価の軸は「何を残せたか」に置く。

レベルは第10章の習熟3段階に対応する。Lv1＝基礎（理解と定型対応）、Lv2＝実践（一人で設計・運用）、Lv3＝高度（仕組み化・標準化・全社設計）。

表12－ストラテジスト（戦略・委任設計）

| レベル                | スキル項目            | 概要説明                                 |
|--------------------|------------------|--------------------------------------|
| ① 可視化——業務を見える化する   |                  |                                      |
| Lv1                | 業務フロー可視化         | 業務をAs-Isで洗い出し、入出力・判断点・例外をフロー図で可視化する  |
| Lv1                | 業務タスク分解          | 業務を判断・作業・確認・連絡に分解し、属人性・定型性・頻度を評価する   |
| Lv1                | 業務量・工数の把握        | 各タスクの所要時間・頻度・件数を定量的に把握する             |
| Lv1                | ボトルネック分析         | 時間・コスト・品質のボトルネックを特定し、改善の優先度を判断する     |
| ② 選定——何を任せるかを切り分ける |                  |                                      |
| Lv1                | 定型・非定型の判定        | ルールで回る業務と、判断を要する業務を切り分ける             |
| Lv1                | AI活用適性の理解        | どんな業務がAIに向き・向かないかを判断する               |
| Lv1                | 委任・協働の仕分け基礎      | タスクごとに「任せる／人が担う／協働する」を切り分ける          |
| Lv2                | 優先順位付け・PoC選定     | どの業務から着手するかを費用対効果で決める                |
| ③ 委任設計——任せ方を設計する   |                  |                                      |
| Lv2                | 業務プロセス再設計（BPR実践） | AIを組み込んだTo-Be業務フローを設計し、人とAIの役割分担を定める |
| Lv2                | 業務分解と再構成         | タスクを分解し、AI担当・人担当・協働へ組み替える            |
| Lv2                | 標準業務手順（SOP）設計    | 協働前提のSOPを設計（判断分岐・品質チェック）             |
| Lv2                | 例外・エスカレーション設計    | 「ここを超えたら人に戻す」条件を設計する                 |
| Lv2                | 委任範囲の設計          | どこまで任せ、どこで人が承認するかの線を引く               |
| ④ 投資——効果と配分を見極める   |                  |                                      |
| Lv2                | 効果測定・ROI算出       | 工数・コスト・品質の観点で効果を測定し、ROIを算出・報告する      |
| Lv3                | 全社AIエージェント戦略立案   | 経営ビジョンと整合した全社戦略（優先順位・投資・ロードマップ）を策定する |
| Lv3                | AX投資ポートフォリオ設計    | 投資配分と段階計画を設計し、効果を最大化する               |
| Lv3                | AX KPI設計・経営報告    | 変革の進捗を経営指標で可視化し、経営に報告する              |
| ⑤ 変革——組織を動かす       |                  |                                      |
| Lv1                | 自社MVVの理解と言語化     | ミッション・ビジョン・バリューを判断軸として使える形に言語化する     |
| Lv1                | 組織構造・業務分掌の把握     | 各部門の役割・分掌を把握し、導入の影響範囲を捉える            |
| Lv1                | ステークホルダーの特定      | 誰が影響を受け、誰の合意が要るかを把握する                |
| Lv2                | 導入プロジェクト推進       | PoC計画・タスク分解・スケジュール／リスク管理で小さく始める      |
| Lv2                | 現場チェンジマネジメント     | 現場の不安・抵抗に向き合い、対話とサポートで受容を促す          |
| Lv2                | AIリテラシー教育・育成     | メンバーが活用できるよう教育プログラムを設計・実施する          |
| Lv3                | AIガバナンス体制設計      | 品質・セキュリティ・コンプライアンス・倫理を包括する体制を設計する    |
| Lv3                | 組織カルチャー変革推進      | AIとの協働を前提とした働き方・組織文化を醸成する            |
| Lv3                | クロスファンクショナル連携    | 多様な部門・ステークホルダーを巻き込み、組織横断で推進する        |
| Lv3                | 人材ポートフォリオ再設計     | 業務変化を見据え、人材要件の再定義・リスクリングを計画する        |

表13－アーキテクト（構築・実装）

| レベル                     | スキル項目               | 概要説明                                      |
|-------------------------|---------------------|-------------------------------------------|
| ① 基盤 —— 土台をつくる          |                     |                                           |
| Lv1                     | AIエージェントの基本理解       | 定義・種類（シングル／マルチ、自律／半自律）と主要フレームワークを理解する     |
| Lv1                     | LLM・モデルの基礎理解        | モデルの特性・得手不得手・選定の基礎を理解する                   |
| Lv1                     | ワークフロー自動化基礎         | Zapier・Make・Power Automate等の自動化の基本概念を理解する |
| Lv1                     | ノーコード／ローコード構築       | GPTs・Dify等でエージェントやボットを試作できる               |
| Lv1                     | データ・API基礎           | データ形式・API・認証など連携の基礎を理解する                  |
| ② 知能 —— コンテキストを設計する     |                     |                                           |
| Lv1                     | プロンプト設計基礎           | 役割設定・制約条件・出力形式などプロンプト設計の基本を押さえる           |
| Lv1                     | コンテキスト情報の整理         | 業務ルール・判断基準・ドメイン知識を洗い出し、構造化する              |
| Lv1                     | RAG・知識参照の基礎         | 外部知識を参照させる仕組み（RAG）の基礎を理解する                |
| Lv2                     | コンテキストエンジニアリング実践    | システムプロンプト・参照データ・ツール定義・会話履歴を最適に設計する        |
| Lv2                     | メモリ・状態管理設計          | 会話履歴・長期記憶・状態の持ち方を設計する                     |
| Lv3                     | エンタープライズコンテキスト基盤設計  | 全社で参照する知識・ルール・ブランド・MVVを統合管理する基盤を設計する      |
| Lv3                     | 動的コンテキスト最適化戦略       | 状況・ユーザー・タスクに応じ、RAG・メモリで文脈を動的に最適化する        |
| ③ 行動 —— ツールと土台（ハーネス）を組む |                     |                                           |
| Lv2                     | AIエージェントワークフロー設計    | 業務要件を入力トリガー・処理・判断・出力・例外処理に落とす             |
| Lv2                     | ツール連携・API設計         | 外部ツール・API連携を設計（MCP／Function Calling等）     |
| Lv2                     | ガードレール・権限設計         | 使えるツール・権限・禁止事項など、安全に動く土台を設計する             |
| Lv3                     | ハーネスエンジニアリング（高度）    | ツール群・実行ループ・自律動作の土台を高度に設計する                |
| ④ 統合 —— つなぎ、広げる         |                     |                                           |
| Lv2                     | システム統合設計            | 既存システム・データベースとの統合を設計する                    |
| Lv2                     | セキュリティ・データ保護実装      | 権限・機密・ログなど安全要件を実装する                       |
| Lv3                     | マルチエージェント・アーキテクチャ設計 | 複数エージェントが連携・協調する全体構成を設計する                 |
| Lv3                     | エージェント間連携プロトコル設計    | 役割分担・協調・調停の仕組み（A2A等）を設計する                 |
| ⑤ 運用 —— 支える基盤をつくる       |                     |                                           |
| Lv2                     | エージェント評価・改善         | 出力品質・処理速度・コストを定量評価し、改善を回す                 |
| Lv2                     | 監視・ログ基盤の構築          | 監督を支える観測・ログ基盤を構築する                        |
| Lv3                     | 可観測性・評価パイプライン設計     | 継続評価・回帰テストを自動化する仕組みを設計する                  |
| Lv3                     | コスト・性能最適化           | 推論コスト・レイテンシ・精度のバランスを最適化する                 |

表14－オペレーター（運用・監督）

| レベル               | スキル項目             | 概要説明                             |
|-------------------|-------------------|----------------------------------|
| ① 観測 ―― 動きを見る     |                   |                                  |
| Lv1               | 出力レビュー基礎          | 決められた基準で出力を確認し、誤り・逸脱を見つけて止める     |
| Lv1               | ログ・実行履歴の確認        | エージェントが何をしたかを履歴で追えるようにする         |
| Lv2               | モニタリング・サンプリング確認設計 | 何を・どの頻度で見えるか決め、サンプリング確認の仕組みを運用する |
| Lv2               | 異常検知・品質監視         | 誤り・品質劣化の兆候を早期に見つける               |
| ② 評価 ―― 良し悪しを判断する |                   |                                  |
| Lv1               | AIの限界・リスク理解       | ハルシネーション・情報漏洩・権限暴走などのリスクを理解する    |
| Lv1               | 機密・個人情報の取り扱い      | 入力してはいけない情報と、安全な使い方を理解する         |
| Lv2               | 評価基準の運用           | 品質基準に照らし、自動化バイアスに流されず良し悪しを判断する   |
| ③ 介入 ―― 止める・直す    |                   |                                  |
| Lv1               | エスカレーション基礎        | 自分で判断できない出力・例外を、適切な担当へ正しく上げる     |
| Lv2               | 止める線・例外対応         | 停止条件（金額・対象・例外）を運用し、逸脱に介入する       |
| Lv2               | インシデント対応          | 問題発生時に切り分け・暫定対応・報告を行う            |
| ④ 調整 ―― 配分とチューニング |                   |                                  |
| Lv2               | 業務分担の調整           | AIと人の担当配分を、運用しながら見直し最適化する        |
| Lv2               | プロンプト・設定の微調整      | 現場で軽微なチューニングを行い、精度を保つ            |
| Lv3               | 監督体制の設計・標準化       | 監督ダッシュボード・介入基準・見逃し率管理を設計し、標準化する  |
| Lv3               | 自動化バイアス対策の設計      | 人が見なくなる問題への構造的な歯止めを設計する          |
| ⑤ 循環 ―― 知を回す      |                   |                                  |
| Lv2               | 暗黙知の形式知化          | 熟練者の判断基準・勘所を引き出し、参照可能なルール・データにする |
| Lv2               | 改善フィードバック         | 監督で得た知見を構築側に渡し、改善を回す             |
| Lv3               | 監督の監督             | 確認そのものの形骸化を計測・監査する仕組みを設計する       |
| Lv3               | 組織ナレッジマネジメント運用    | 現場知を形式知化し、組織で蓄積・更新・共有する仕組みを回す    |
| Lv3               | 継続的改善サイクルの構築      | AI-SECIを回し、現場とAIで知を育て続ける仕組みを作る   |
| Lv3               | 監督KPI設計・報告        | 監督の効き目を指標化し、経営・現場に報告する           |

#### 個人のスキルと、組織の仕組みを取り違えない

オペレーター（監督）のLv3（監督体制の設計・標準化、自動化バイアス対策の設計、監督の監督、監督KPI設計・報告など）は、担当者個人の作業スキルというより、組織としての運用設計・仕組みづくりの能力である。第10章のとおりLv3は「仕組み化・標準化・全社設計」の段階であり、これを個人が担うか、専任チームやCoEが担うかは組織の規模による。表14は“個人のスキル”と“組織の仕組み”をレベルの高さで橋渡ししている、と読むと役割・スキル・仕組みの混在感が解ける。

本一覧は出発点であり、完成形ではない。業種・職種により必要なスキルは異なる。とくにオペレーター（監督）の領域は、自律エージェントの普及とともに今まさに体系化が進む新しい領域であり、現場での実践を通じて更新していくことを前提とする。

# 会社として、何から始めるか—— 導入アクションプラン

本基準は、読んで終わりの整理ではない。明日からの動きに落とせる。会社としてまず取り組むべきことを、4つのステップにまとめた。規模や段階を問わず、型は同じである——**小さく始め、監督の仕組みを整え、成果で育てる。**

図6－導入の4ステップ（ロードマップ）



大切なのは順番だ。多くの失敗は、ツール選びから入る・監督を後回しにする・一気に全社へ広げる、という順番の誤りから起きる。まずは1業務で型をつくり、監督の仕組みごと回してから横に広げる。各ステップで「誰が担い、何を出すか」を次の表に示す。

表15－導入アクションの4ステップ・何を・誰が・何を出すか

| ステップ                | 主な動き                                                                    | 中心となる役割        | アウトプット                          |
|---------------------|-------------------------------------------------------------------------|----------------|---------------------------------|
| STEP 1<br>現在地を知る    | 主要業務を棚卸しし、点→線→面→立体のどこにいるかを把握。定型性・頻度・インパクトで「任せる候補」を3つほどに絞る（第3・6章）        | ストラテジスト        | 業務マップ／優先業務リスト／役割アサイン案           |
| STEP 2<br>小さく始める    | インパクトが高くリスクの低い1業務でPoC。委任範囲・標準業務手順（SOP）・「止める線」を設計し、エージェントを構築する（第5・8・12章） | ストラテジスト＋アーキテクト | 動くエージェント1つ／運用ルール・SOP／停止条件       |
| STEP 3<br>監督の仕組みを作る | 広げる前に、観測・評価・介入の仕組みを用意。サンプリング確認・エスカレーション・自動化バイアス対策を回す（第8章）               | オペレーター         | 監督の運用ルール／介入・停止基準／インシデント手順       |
| STEP 4<br>育てて標準化    | 成果（時間・品質・コスト）で評価し人材レベルへ返す。スキルを研修・等級・採用へ。暗黙知を形式知化し横展開・ガバナンスへ（第7・10・12章）  | 3役割＋経営         | 成果レポート／育成・採用計画／全社ロードマップ・ガバナンス方針 |

## 実行チェックリスト

自社の動きを、そのまま確認できるようにした。1業務でひと回しするごとに、上から順にチェックしていけばよい。

### STEP 1 現在地を知る

- 主要業務を5～10個書き出し、入出力・判断点・例外を可視化した
- 各業務を「定型／非定型・頻度・ミスの影響」で評価した
- 任せる候補業務を3つ程度に絞った
- 3つの役割（戦略・構築・監督）の担い手を仮置きした

### STEP 2 小さく始める

- インパクトが高くリスクの低い1業務を選んだ
- 委任範囲（どこまで任せ、どこで人が承認するか）を決めた
- 標準業務手順（SOP：手順・判断分岐・品質チェック）を用意した
- 「ここを超えたら人に戻す」停止条件を設計した
- エージェントを構築し、小さく動かした

### STEP 3 監督の仕組みを作る

- 何を・どの頻度で確認するか（観測）を決めた
- サンプル確認の基準を決めた
- エスカレーション先と手順を定めた
- インシデント時の対応手順を用意した
- 「人が見なくなる」自動化バイアスへの歯止めを入れた

### STEP 4 育てて標準化

- 成果（時間・品質・コスト）を測り、記録した
- 役割別スキルを研修・等級・採用に反映した
- 現場の暗黙知を形式知化して共有した
- 横展開のロードマップとガバナンス方針を定めた

よくあるつまづき。① ツール選びから入る（先に業務の棚卸しを）。② 監督を後回しにする（広げる前に仕組みを整える）。③ いきなり全社展開する（まず1業務で型をつくる）。④ 資格や研修の修了をゴールにする（成果で測る）。この4ステップを1業務で回しきれば、同じ型で次の業務へ広げられる。全社の変革は、特別な号令からではなく、この小さなサイクルの積み重ねから始まる。

# 認定資格のご案内—— 学びを、証明できる力に

本基準は、AIエージェント時代に必要な人材像を示した地図である。この人材像を体系的に学んで身につけ、第三者に証明する具体的な方法の一つが、AICX協会の認定資格である。本基準が定義する人材像を、実際に「測り、証明できる」ものにする。学んで終わりにせず、現場で成果を出せる力を、共通のものさしで可視化する。受験前の準備に役立つ公式ガイドブックも、無料で配布している。

## 見える化

必要な人材像を、共通のことばで可視化

## 活用できる

研修・採用・等級づくりにそのまま使える

## 証明できる

第三者の認定で、社内外にスキルを示せる

## 本基準に対応した認定資格

### AI AGENT STRATEGIST

#### AIエージェント・ストラテジスト

こんな方へ：経営企画・事業開発・コンサルタントなど、ビジネス側で構想を描く方。AIエージェント・業務変革・組織変革の3つを柱に、生成AIを生かした企画構想から要件定義までを担う人材のスキルを可視化する。

[aicx.jp/certifications/strategist](http://aicx.jp/certifications/strategist)

### AI AGENT ARCHITECT

#### AIエージェント・アーキテクト

こんな方へ：エンジニア・開発者・SIerなど、実装・運用を担う方。ストラテジストが描いた構想を、実際のAIエージェント構築ツールで効率的かつ安全に実装・運用する人材のスキルを可視化する。

実装 / 運用 / セキュリティ

— 公式ガイドブックを無料配布中

資格の全体像・出題範囲・学習の進め方を1冊に。受験を検討する前に、まずは無料でダウンロードを。

## おわりに——代表理事メッセージ



一般社団法人AICX協会 代表理事

**小澤 健祐（おざけん）**

「使う」から「任せる」へ。これは、言葉の違いではありません。本質がまったく異なります。これまで、仕事の主体はつねに人間でした。道具がどれほど進化しても、最後に手を動かし、判断するのは人だった。AIエージェントの時代は、その主体が人間でなくなる、史上初めての転換です。

私は年間300回を超える講演や、多くの大企業のアドバイザーとして数多くの現場に立ち会ってきましたが、断言します。この転換は、並大抵の努力で成し遂げられるものではありません。これまでのDXの延長線上で考えてはいけません。業務の進め方も、組織の形も、責任の所在も、そして人材の育て方も、すべてを根底から疑い直す必要があります。

いま求められているのは、既存の仕組みの改良ではありません。「まったく新しいシステムになる」という前提に立つことです。とはいえ、それは今いる場所を否定することではありません。現在地から一歩ずつ、任せる範囲を広げていく。本基準は、その道のりで人とAIの関係を捉え直すための共通言語であり、現場で叩きながら育てていく草案です。忌憚のないご意見をいただければ幸いです。

人とAIが共存し、人がより創造的な仕事に向かっていける社会へ。その実現は、特別な誰かではなく、それぞれの現場でAIと向き合う一人ひとりの積み重ねから始まります。本基準が、その確かな一歩になることを願っています。



一般社団法人AICX協会 代表理事

**小栗 伸**

AIエージェントに仕事を「任せる」ことは、導入した瞬間に完成するものではありません。現場で使い、結果を見て、違和感を拾い、少しずつ直していく。その日々の改善を引き受ける人がいて、初めてAIは本当に業務の力になります。

一方で、その改善を担う人は、決して楽ではありません。AI/DX推進部門、事業部のリーダー、現場の担当者、情報システム、人事・育成、企画部門など、それぞれの持ち場で変化を前に進めようとしている方々は、新しいやり方への戸惑いや、これまで通りの進め方を望む声にも向き合いながら、何度も説明し、巻き込み、試行錯誤を重ねています。私は、そうしたすべての推進者、現場で奮闘されている一人ひとりの挑戦に、心から敬意を表します。

本基準は、そうした推進者だけのものではありません。これから担い手となる方々が「自分もAIエージェント時代の人材になる」と目標を持ち、周囲に宣言し、少しずつ役割を引き受けていくための共通言語です。業務を棚卸しし、小さく試し、監督の仕組みを整え、成果で人材を育てる。そのサイクルを、各社の現場で回していくことが重要です。

AIを使える人から、AIとともに仕事を育て、組織を前に進める人へ。AICX協会は、その挑戦を後押しし、現場で変化を引き受ける人が一人でも多く生まれる社会を、皆さまとともにつくりていきたいと考えています。

## DSS（デジタルスキル標準）との整合性

本基準は単独で使えるが、国のデジタルスキル標準（DSS、経済産業省・IPA ver.2.0／2026年4月公表）と矛盾しないよう設計している。企業がDSSベースの人材定義と併用できるよう、対応関係をここに整理する。本基準はDSSを置き換えるものではなく、AIエージェント実装という一点に解像度を加える民間リファレンスである。

表16－本基準の役割とDSS類型の対応（重ね合わせ）

| 本基準の役割  | 対応するDSSの位置                                  | 補足                                                                                 |
|---------|---------------------------------------------|------------------------------------------------------------------------------------|
| 共通基礎スキル | DXリテラシー標準／共通スキル「AI実装・運用」「AIガバナンス」の理解部分      | DXリテラシー標準の延長で、AIエージェント特有の特徴・限界・リスクの理解に絞る。理解の層なので最もきれいに重なる。                         |
| ストラテジスト | ビジネスアーキテクト類型（＋デザイナー類型の構想・合意形成）              | その射程のうち「どの業務を委譲し、どうROIを設計するか」に焦点を絞る位置づけ。                                           |
| アーキテクト  | ソフトウェアエンジニア類型＋データマネジメント類型（＋データサイエンティスト等の一部） | 技術領域別の類型を横断し、エージェントを成立させる5つの管掌領域を統合する視点を担う。                                        |
| オペレーター  | 明確な対応類型なし（「AI実装・運用」の運用／「AIガバナンス」に部分的に重なる）   | 自律エージェントを日々監督する専任役割はDSSの6類型に明確な対応がない。標準の不備ではなく“いま生まれつつある”新役割であり、現場視点で先に言語化する価値が高い。 |

対応は“主従”ではなく“重ね合わせ”である。ある人材がDSS上はビジネスアーキテクトでありながら、本基準上はストラテジストとして読める、といった多重対応を想定する。企業は既存のDSSベースの人材定義を壊さずに、AIエージェント観点の解像度だけを上乘せできる。

# 前提・出典について

## 本基準の前提と今後の更新

本基準は「人がAIを監督する」時代を前提にしている。将来、AIが自己監督・自己修正を高度化すれば、「監督する人」の比重も同じ理屈で変わりうる。本基準は固定の正解ではなく、現場で叩いて更新し続ける土台（草案）として位置づける。記載の役割・管掌領域・レベルは、自社の段階や業種に合わせて取捨選択・カスタマイズして用いることを想定している。

## 出典・免責

**出典：**経済産業省・独立行政法人情報処理推進機構（IPA）「デジタルスキル標準 ver.2.0」（2026年4月16日公表）。SECIモデルは野中郁次郎・竹内弘高による知識創造理論に基づく。本書に記載のDSSの構成・人材類型は公表資料に基づく要約である。本文中のツール名・製品名は各社の登録商標であり、2026年時点の例として挙げたもので、各社の名称・機能は変更されることがある。

**免責：**本書は一般社団法人AICX協会による民間リファレンスであり、国の標準を代替・改変するものではない。本書の内容は発行時点の見解であり、将来予告なく改訂されることがある。特定の製品・サービスの導入や成果を保証するものではない。

**位置づけ：**本書はAIエージェント実装に関わる人材育成の共通言語を提供することを目的とした草案である。