

AI-Ready 数据组织与平台建设指南

从诊断评估到落地实施的企业实践框架

亚马逊云科技

2026 年

文档信息与修订记录

文档基本信息

文档名称	AI-Ready 数据组织与平台建设指南
英文名称	AI-Ready Data Organization and Platform Guide
当前版本	1.0.1
文档状态	公开
首次发布日期	2026 年 9 月 9 日
最近修订日期	2026 年 9 月 9 日
适用范围	企业数据平台、生成式 AI、RAG、Data Agent、数据治理与组织协同建设
保密级别	公开

作者与审稿人员

类别	姓名
主要作者	Xin Zhang, Junxiong Ke, Yewei Li, Xin Zhao, Huilin Wu
审稿	Xiaoyong Han, Lili Ma

修订记录

版本	日期	状态	主要修改
1.1.0	2026-09-09	完稿	完成总体框架、能力模型、技术平台、组织流程、实施路线图和附录初稿。

1 引言：企业数据组织面临的 AI 转型焦虑与行业驱动力

1.1 本白皮书要解决的企业问题

当前，企业正面临前所未有的 AI 转型压力与机遇。据咨询机构公开发布的调研，多数企业尚未具备支撑规模化高级 AI 应用所需的可信数据质量与标准化治理规范，仅有少数领先企业已经建立起可复用、可规模化的数据就绪能力，相当比例的组织因数据相关的质量、治理或风险问题而延迟、限制或调整了原定的 AI 落地计划^[1]。这一现象背后是一个核心矛盾：企业对 AI 的投资热情持续攀升，但数据基础设施的就绪程度普遍滞后于业务期待。

这一矛盾在企业内部外部的驱动力迫使下被同时放大。基础模型的推理能力、上下文长度和多模态处理能力持续提升，单次调用成本快速下降，[MCP](#)、[A2A](#) 等工具调用协议逐步标准化，使企业接入 AI 的技术门槛大幅降低——这意味着企业不再受限于“技术做不到”，瓶颈已经转移到数据和治理侧。与此同时，率先完成数据能力建设的企业，已经在客户响应速度、运营效率和决策质量上取得可观察的优势，且这种优势会随着场景复用不断放大，落后企业面临的差距正在非线性拉大。客户和员工的期待也已改变：客户期待更快速、更个性化的服务响应，员工期待日常工作中获得 AI 辅助，这种体验基准倒逼企业内部系统加快跟上。围绕 AI 治理、数据隐私和内容安全的监管要求在多个司法辖区日趋明确和收紧，企业若不提前建立数据分类分级、访问审计等能力，后续合规改造的成本将远高于提前投入。传统数据团队的技能结构也面临语义工程、[AI Agent](#) 治理等新技能需求的调整压力，人才结构的重塑本身也在倒逼企业系统性推进，而不是停留在零散试点。技术、市场、监管和人才这几方面的变化几乎同步发生，使“数据是否 AI-Ready”从一个技术团队内部的工程问题，上升为关系企业竞争力和风险敞口的战略问题。

亚马逊云科技在协助企业客户推进 AI 转型的实践中观察到，这一差距正是导致多数 AI 试点项目难以规模化落地的首要瓶颈。企业面临的核心挑战，是数据尚未具备被 AI 可靠、持续、规模化使用的能力而数据规模不足只是其中一种表象。企业需要解决的并非单纯的数据存储或数据汇聚问题，而是如何让分散、异构、持续变化的数据资

产，转化为可被 AI 安全、可信、及时调用的生产要素。这要求企业不仅具备统一的数据发现、访问和共享机制，还需要建立覆盖数据质量、元数据、血缘、权限、隐私保护与合规审计的治理能力。尤其在检索增强生成（RAG）、企业知识助手和智能体等场景中，数据的准确性、时效性和授权边界将直接决定 AI 输出是否可靠、是否可控，以及是否能够真正服务于业务决策和流程执行。

因此，本白皮书聚焦企业从“拥有数据”走向“数据 AI 就绪”过程中**最关键的共性问题**：如何识别现有数据基础与 AI 规模化应用之间的差距；如何构建兼顾敏捷创新与安全治理的数据平台和运营机制；以及如何通过分阶段、可度量的路径，将零散的 AI 试点沉淀为可复用、可扩展、可持续运营的企业级能力。白皮书旨在为企业技术领导者、数据负责人和业务决策者提供一套系统性参考框架，帮助其以可信数据为基础，加速释放生成式 AI 的生产级业务价值。

1.2 目标读者和适用范围

目标读者包括既要业务结果、数据风险和合规要求承担责任的决策者，也包括负责将治理要求、数据能力和 AI 应用落地为工程实践的建设者：

- **CEO、业务负责人及数字化转型负责人**：通过自查工具评估 AI 应用的业务优先级、数据准备度与风险承受能力，明确哪些场景适合优先试点、哪些场景具备规模化推广条件，以及数据与 AI 投入应如何与可衡量的业务价值相结合。
- **CIO、CDO、CAIO 及数据与 AI 平台负责人**：识别企业在数据资产管理、数据质量、数据访问、知识管理、模型与应用运营等方面的能力短板，明确共享数据平台、领域数据产品、企业知识库和 AI 应用之间的责任边界，并制定分阶段的能力建设路线图。
- **CISO、法务、风险与合规团队**：评估敏感数据识别、访问控制、隐私保护、内容安全、审计留痕和责任追溯等能力是否能够覆盖 AI 应用全生命周期，确保创新速度与安全、合规要求之间取得可控平衡。
- **企业架构师、数据架构师、AI 架构师及工程团队**：将自查结果及落地指南转化为具体的架构改进与工程任务，包括数据标准化、元数据与血缘管理、数据质

量监控、权限治理、知识检索、版本管理、可观测性、评估机制和回滚能力建设。

本白皮书及自查工具适用于处于不同 **AI** 转型阶段的企业：无论是仍在探索生成式 **AI** 用例、正在开展局部 **PoC**，还是已上线知识助手、智能客服、开发助手或智能体应用并希望实现跨部门复制和规模化运营的组织，均可据此建立共同的评估语言和改进基线。

在数据范围上，白皮书覆盖结构化、半结构化、非结构化及多模态数据，包括业务交易数据、客户与运营数据、日志数据、文档、图片、音视频及企业知识资产；同时覆盖这些源数据经过采集、清洗、转换、分块、标注、嵌入、索引、检索、上下文组装和模型生成后形成的数据产品与派生制品。重点不在于某一种存储、数据库或模型技术，而在于确保这些数据及其派生资产在全生命周期中具备可发现、可理解、可访问、可治理、可追溯和可持续优化的能力。

在应用范围上，白皮书关注由可信数据支撑的企业 **AI** 场景，包括内容生成、企业知识问答、检索增强生成（**RAG**）、分析辅助、运营洞察、决策建议、智能客服、开发与办公助手、受监督的工具调用，以及具备明确权限边界和人工监督机制的智能体执行。白皮书不以模型训练算法、基础模型能力排名、提示词技巧或单一云服务与厂商选型为重点；这些议题可以作为具体实施方案的一部分，但不构成本白皮书及自查工具的核心评估对象。

1.3 如何阅读本白皮书

本白皮书篇幅较长，涵盖概念定义、评估工具、平台架构、技术细节、组织流程、实施路线图和参考资料等多个层次。以下按章节说明各部分具体讲什么、写给谁看，帮助你根据自身角色和时间选择阅读路径，而不必从头逐字读到尾。

第一章引言篇幅短，介绍企业推进 **AI** 转型时面临的数据就绪度差距、这一问题背后的内外部原因，以及白皮书的目标读者和覆盖范围。无论你是业务决策者还是技术建设者，都建议完整阅读本章，用于建立后续章节共用的背景和语言基础。

第二章什么是 **AI-Ready Data** 分四个部分。“为什么传统数据体系不再充分”分析 **AI Agent** 的出现如何改变数据消费方式、数据从静态输入变为动态组装的上下文、治理

对象从结构化数据扩展到多模态和派生资产、权限管理从静态角色扩展到跨域动态运行时控制，这部分适合所有角色阅读，用于理解后续维度设计的出发点。“**AI-Ready Data** 的概念及场景分类”给出概念边界和 **L1** 至 **L4** 的场景风险分级，业务负责人和架构师都建议阅读，这是贯穿全书的分级语言。“能力模型及六大可衡量维度”逐一展开成本与业务价值、数据基础与质量、平台架构与运维、组织与人才、治理与信任、安全与合规六个维度，管理层可以重点看每个维度开头的“核心问题”和“评估重点”，技术团队则需要关注“关键指标”部分的具体度量方式。“成熟度模型与评估路径”说明四级成熟度模型、如何将企业整体成熟度与场景风险等级结合使用、**45** 题自检工具的打分方法，管理层建议看“如何使用成熟度模型”一节，具体的打分公式和动态分母计算方法可以留给负责自检工作的团队精读。

第三章 **AI-Ready Data** 平台技术含量最高，面向数据架构师、平台工程师和 **AI/Agent** 工程师，建议这几类读者完整阅读。“平台的目标状态与建设原则”给出六项建设原则，管理层可以只读这一节了解总体方向。“能力分层与责任划分”三类典型消费场景技术架构与数据管控数据产品的建设目标分析、**AI** 与 **Agent** 如何使用数据平台运维、变更与持续改进这几节偏工程实现细节，是架构师和工程团队的重点阅读对象，非技术读者可以选择跳过。本章最后一节平台投入与业务价值的度量面向管理层，讲如何分阶段设置代理指标和测算 **ROI**，即使跳过本章其他技术内容，也建议管理层单独阅读这一节。

第四章 **AI-Ready Data** 的组织和流程回到管理和组织层面，适合 **CIO**、**CDO**、业务负责人和负责组织变革的团队完整阅读。“数据组织与流程面临的新挑战”对比集中式、分散式和联邦式三种组织形态各自的适用阶段和缺口。“**AI-Ready** 组织与流程的目标状态”是本章核心，提出责任三角框架，说明数据生产者、领域实践者、业务消费者三种当责职能如何分工，以及 **Agent** 作为第四参与者带来的角色转型，技术团队也建议查看角色转型部分了解自己岗位的转型方向。“支撑目标状态落地的文化与流程”篇幅不长，适合所有人阅读。“让责任可执行的跨团队协作机制”介绍数据契约、**SLA** 和变更通知协议三类具体机制，架构师和数据产品经理需要重点关注。

第五章迈向规模化的组织能力建设把前四章的评估框架、平台能力和组织分工，收敛成一套可落地的分阶段实施路线图，按 **0-30** 天、**30-60** 天、**60-90** 天三个阶段说明每个阶段应该完成什么、如何把单点 **PoC** 逐步扩展为跨部门、可复制的生产能力，并

呼应第二章的 L1 至 L4 场景分级和六大维度自检结果来确定每个阶段的优先级。这一章适合负责实际推进 AI-Ready 建设、需要制定具体项目计划的人阅读，例如项目负责人、数据平台负责人和参与立项的业务负责人；管理层如果只了解节奏和里程碑，可以只看每个阶段的目标产出，不必深入每一步的具体任务清单。

附录 A（亚马逊云科技参考架构）把前面几章的能力分层——数据供应、处理产品化、语义知识、分析 AI 服务、消费激活——逐一映射到具体的亚马逊云科技服务。这一章只对已经决定或倾向使用亚马逊云科技平台作为落地平台的架构师和工程团队有直接参考价值，管理层和其他云平台的读者可以跳过，只需知道白皮书提供了一份对应的技术选型参考。

附录 B（术语表）收录了全书出现的关键术语解释，建议作为随查工具使用，遇到正文中不确定的术语时回来查阅，不需要单独完整读一遍。

如果时间有限，建议至少完整阅读第一章和第二章的前两节，掌握评估语言和风险分级逻辑；再根据自身角色选择第三章深入了解技术建设路径，或第四章深入了解组织落地方式；如果需要马上启动项目，第五章的分阶段路线图是最直接的行动参考；技术选型阶段再查附录 A，遇到术语随时查附录 B。

2 什么是 AI-Ready Data：概念及六大可衡量维度

2.1 为什么传统数据体系不再充分

传统分析体系隐含了一个重要前提：数据最终由人解释。分析师知道字段命名可能不完整，业务负责人会对异常数字提出质疑，操作人员能够在执行前结合上下文复核。许多企业数据体系的“正确性”实际上由数据、文档和人的隐性知识共同维持。人工经验经常会在企业生产过程中将部分数据的语义缺口、数据质量问题暂时掩盖过去。

AI Agent 的出现改变了这一前提。它能够在单次任务中跨越数据库、文档库、搜索索引和外部系统获取信息，动态形成任务上下文，调用模型进行推理和决策，并通过 API、MCP 或其他工具直接触发后续动作。这意味着 AI 不再只是为人提供分析结果或建议，而是开始参与甚至承担部分业务流程的执行。当系统能够从“理解信息”进一步走向“发起行动”时，数据治理的重点也必须随之变化：企业不仅要关注数据是否可

用，还必须确保数据的来源、语义、时效性、权限边界和使用过程能够被验证、约束和追溯。对于具备更高自主性的 **Agent** 而言，身份认证、授权机制、操作范围、人工介入点与审计控制，都需要在其进入业务流程之前被系统化设计。变化的关键不在于“机器是否比人更容易犯错”，而在于错误的传播机制不同。人类消费通常是低频、可解释且存在组织缓冲；**Agent** 消费可以持续、并发、跨系统和自动执行。同一个语义错误可能在多个应用中被复用，同一个过期片段可能被数千次检索，同一个过宽权限可能被工具链放大。企业因此必须把过去依赖个人经验的判断转化为机器可读取、策略可执行和运行可验证的显式约束。

2.1.1 数据从静态输入变为动态组装的上下文

生成式 **AI** 并不总是消费一个稳定、预先验证的数据集。非结构化内容会被解析为文本、表格和图像对象，再被切分、嵌入和索引；运行时检索只选择其中一部分，并与提示词、对话历史、工具结果和记忆组合。行业研究普遍指出，每一次抽取、分块和嵌入都会改变数据点的上下文，正确的完整文档并不意味着被检索出的片段能够支持正确答案。

企业需要治理的最小对象，不再止于源表或源文档，而是所有能够影响决策或行动的派生制品。向量、索引、提示词模板、检索配置、工具返回、记忆和生成输出都需要所有者、版本、来源、适用范围、刷新策略和退役条件。否则企业能够证明“源文档受控”，却不能证明“这一次决策使用了什么”。

W3C PROV-O 提供了跨系统表示和交换来源信息的基础模型，通过 [Entity](#)（对象及实体）、[Activity](#)（活动事件）和 [Agent](#)（代理服务）及其关系描述对象如何被使用、生成、派生和归责。白皮书不建议所有企业直接采用 **PROV-O** 实现，但任何架构都应回答同样的问题：哪个来源、哪个版本、经过什么活动、由哪个主体或软件 **Agent** 处理、形成了哪个派生对象，并最终影响了哪次输出和行动。

2.1.2 从结构化数据治理走向多模态与派生资产治理

传统数据体系主要围绕结构化数据表、指标口径和报表链路进行管理，治理对象通常边界清晰、结构稳定。但在 **AI** 应用中，企业需要同时处理文档、邮件、网页、图片、音视频、代码、日志、会话记录等多种形态的数据。更重要的是，**Agent** 实际使用的

往往不是原始数据本身，而是数据经过清洗、分块、标注、嵌入、索引、检索和上下文组装后形成的派生资产。

这要求治理能力不能只停留在源数据层面。企业还需要了解某个文档片段、向量索引或模型上下文来自哪里，是否仍然有效，包含何种敏感信息，适用哪些使用范围，以及在源数据更新、删除或权限变化后是否能够被同步修正。换言之，治理对象已经从“数据集”扩展为贯穿源数据、知识资产、AI 上下文和生成结果的完整链路。

2.1.3 从静态权限管理走向跨域、动态的运行时治理

传统权限管理通常以用户、角色和系统边界为基础：用户在某个系统中被授予相对固定的访问权限。但 **Agent** 的任务往往跨越数据湖、数据库、文档库、搜索索引、SaaS 应用和外部工具，并会根据当前任务动态检索信息、组合上下文和调用服务。仅依赖静态角色授权，难以判断某次访问是否真正符合任务目的、数据敏感级别和当前业务风险。

因此，企业需要能够兼容不同数据模态、存储平台和访问方式的统一治理能力：无论数据通过 **SQL** 查询、文件读取、向量检索、**API** 调用、**MCP** 工具或应用接口被访问，都应能够基于 **Agent** 身份、代表的用户、任务目的、数据分类和操作风险执行一致的策略。高风险数据访问、跨域数据组合和可能改变业务状态的工具调用，还应在运行时触发脱敏、限制、审批、阻断或人工复核。面向 **Agent** 的治理重点不再只是“谁有权限”，而是“谁在何种上下文下，为完成什么任务，可以访问哪些数据并采取哪些行动”。动态策略执行、最小权限、任务级访问范围和可追溯审计，是这一转变的核心。

综上，传统数据体系虽能够支撑面向人的分析与决策，却难以满足 **Agent** 在多模态、跨域和动态环境中自主理解、调用并执行行动的要求；企业需要从“数据可用”迈向“数据与治理能力全面 **AI-Ready**”，以可信、可控、可追溯的数据基础支撑 **AI** 的规模化生产应用。

2.1.4 AI Agent 既是消费者也是生产者：数据治理对象的根本扩展

AI Agent 在消费数据的同时，几乎总是在生产新的数据——它不是数据管道的终点，而是数据管道中的一个新节点。这一双重身份，是 **AI** 时代数据治理对象发生根本扩展的深层原因。

Agent 在运行中持续产出需要被治理的新资产。一次 **Agent** 任务的执行，会沿途留下大量派生制品：检索命中的 **Chunk** 组合、动态拼装的 **Context**、写入的 **Memory**、工具调用的返回结果、中间推理链、最终生成的结论，以及由此触发的业务动作记录。这些产出并非一次性的临时数据——它们会被写回知识库、沉淀为用户或任务 **Memory**、作为下一次任务的输入被再次检索，甚至被其他 **Agent** 消费。换言之，**Agent** 的输出会回流成为下一轮的输入，如果不加治理，就会形成“AI 引用自身生成结果”的错误闭环：一个未经核验的结论被写回知识库，随后被检索、被引用、被再次生成，错误在无人察觉中自我强化。

生产者身份意味着治理必须前移到“写入”环节，而不仅是“读取”环节。传统治理的重心在于控制“谁能读什么数据”；但当 **Agent** 成为生产者，治理必须同时回答“哪些产出可以被持久化、以什么身份写入、进入哪一层资产、是否需要人工审批”。这要求企业对 **Agent** 的每一类产出明确其所有者、版本、来源、适用范围、刷新策略和退役条件——正如 2.1.1 所述，治理的最小对象已从源文档扩展到所有能影响后续决策或行动的派生制品。特别地，**Agent** 写入共享业务 **Memory** 或知识库的内容，应被视为对企业可信数据核心的注入，须经过与人工数据生产者同等甚至更严格的验证与 **Owner** 审批，而不能因其“由 AI 自动产生”就默认可信。

这一双重身份还改变了责任归属的边界。当 **Agent** 既读又写、既消费又派生时，它在数据血缘中同时扮演消费端和生产端两个角色。**W3C PROV-O** 描述的“哪个来源、经过什么活动、由哪个主体处理、形成哪个派生对象”——在 **Agent** 场景中，这个“主体”第一次不再必然是人，而可能是一个软件 **Agent**。因此血缘记录必须能够区分人类身份与 **Agent** 身份，明确某个派生资产究竟由谁、代表哪个用户、在什么任务下产生，才能在出现错误输出或错误行动时完成追溯与归责。这也呼应了后文责任三角中的关键提醒：**Agent** 是执行者和生产者，但不是责任主体——它产出的数据和触发的动作，其最终当责必须显式落位到人或治理职能上。

2.2 AI-Ready Data 的概念与场景分类

AI-Ready Data，指企业的数据、平台、组织与流程共同具备的一种系统性能力——使高质量、可信、具备业务语义的数据，能够在治理和安全边界内，被模型训练、生成式 AI 应用、Agent 决策与人机协作等各类 AI 工作负载稳定消费，并随着场景扩展而持续复用、演进。

数据能力由数据基础、技术平台、组织分工与治理机制共同决定；任一环节滞后，都可能限制 AI 价值兑现。

2.2.1 相关概念的边界

AI-Ready Data 与以下概念存在区别与联系：

- **Data Quality（数据质量）**：是 AI-Ready 的必要条件之一，但仅有数据质量不足以支撑 AI 规模化。AI-Ready 还需要语义层、权限模型、变更通知和组织能力的配合。
- **Data Governance（数据治理）**：传统治理的细粒度控制本已成熟，AI-Ready 的差异在于主体变成代表用户行动的 Agent、对象扩展到 Chunk 与 Memory 等派生资产、动作从读写数据扩展到改变业务状态，使授权必须从静态绑定转向运行时动态校验。
- **Data Platform（数据平台）**：AI-Ready 平台不仅是存储和计算引擎，还必须提供语义层、检索服务、身份管理和行为审计能力。
- **MLOps / LLMOps**：关注模型生命周期管理。AI-Ready 关注数据侧如何为模型提供可靠输入，两者是互补关系。

2.2.2 场景分类和风险分级

AI 场景按照自主性和影响范围可分为四级：

风险等级	场景类型	自动化边界	典型例子	最低 AI-Ready 要求
L1 辅助参考	内容生成、知识问答	AI 输出仅供人参考，不直接触发业务动作	文档摘要、FAQ 回答、会议纪要、营销文案初稿	数据质量基本合格；知识来源可标注；内容时效性可识别；敏感信息得到基础保护；用户能够核验或反馈结果
L2 分析建议	分析辅助、决策建议	AI 参与判断和推荐，但由人作出最终决策并执行	财务分析、风险评估、销售推荐、运营诊断	数据语义完整；指标与业务规则定义明确；来源、版本和推理依据可追溯；具备准确性与偏差评估；关键结论有人复核

L3 受监督执行	受限工具调用、流程自动化	AI 可发起或执行预定义动作，但必须经过人工审批或规则闸门	创建订单草稿、提交审批、工单流转、受限系统配置变更	细粒度权限模型；Agent 独立身份；任务级授权；工具调用与数据访问全程审计；关键动作可审批、可拒绝、可撤销；具备回滚与异常处置机制
L4 受控自主执行	自主决策与闭环执行	AI 可在预设范围内自主完成多步骤行动，仅在超出边界或发生异常时升级人工处理	自动付款、自动交易、自动资源调度、自动运维修复	全链路可观测；实时策略验证；最小权限与行动限额；数据、上下文、模型和工具调用可追溯；持续效果与风险监控；异常熔断、降级与回滚；明确责任归属和定期独立复核

2.3 能力模型及六大可衡量维度

AI-Ready 并不等同于企业已经部署了大模型、上线了知识助手，或拥有统一的数据平台。其核心在于：企业是否具备用可信数据支撑 AI 理解、推理、调用工具和执行行动的能力，并能够在这一过程中持续控制成本、保障质量、约束风险、追溯责任和验证业务价值。

基于企业 AI 转型中常见的数据、平台、治理和运营挑战，本白皮书提出六大可衡量维度，配套自检工具按这六个维度逐一展开。六个维度共同覆盖从业务投入到数据供给、从平台运行到组织协同、从治理机制到安全防护的完整能力闭环。任何单一维度的短板，都可能限制企业将 AI 应用从局部试点扩展至生产级、跨部门和高自主性的场景。

2.3.1 成本与业务价值（Cost & Business Value）

本维度回答的核心问题是：企业为 AI 场景建设和运营数据能力的投入，是否看得见、算得清、说得明白，以及是否真正形成了可验证的业务产出。

企业应避免将 AI 数据相关成本分散在云资源、软件许可、外部服务、人员工时和项目预算中，导致无法识别实际投入，也无法判断数据平台、知识库、检索能力和治理能力被哪些场景使用、产生了多少价值。特别是在多个业务部门复用同一数据资产、数据服务或 AI 平台能力时，缺乏合理的成本归集与分摊机制，容易造成“建设团队承担成本、使用团队难以量化收益”的失衡。

该维度不仅关注 ROI，也关注企业是否在项目启动前明确业务假设、预期收益、投入规模和可接受的失败条件；是否在项目过程中预设检查点，并基于实际数据决定继续投入、调整方向或暂缓建设。对于 L3 和 L4 场景，建议企业将自动化带来的效率收益、错误处置成本、人工监督成本、潜在风险暴露和恢复成本纳入整体价值评估，而不能只以模型调用量或上线数量衡量成功。

- **关键指标：** AI 数据侧投入可追踪比例、共享资产成本分摊覆盖率、业务价值实现周期、数据准备工时变化、返工次数、场景活跃度、业务采纳率、关键业务指标改善情况。
- **评估重点：** 是否建立 AI 相关数据投入的标签或核算机制；是否定义并实际使用价值衡量方法；重要投入是否有书面评估记录；立项是否明确业务目标、投入预算和退出条件；项目中是否存在基于数据的继续、调整或暂缓决策机制。

2.3.2 数据基础与质量（Data Foundation & Quality）

本维度回答的核心问题是：支撑 AI 场景的数据是否具有足够的可用性、准确性、时效性、语义完整性和可追溯性，能够让 AI 在缺少人工实时纠错的情况下稳定使用。

传统分析场景中，数据质量缺陷可以部分由分析师、业务人员和操作人员的经验吸收；但当 AI 或 Agent 直接基于数据生成回答、提出建议、调用工具甚至执行动作时，数据缺失、口径冲突、更新滞后、来源不明和语义模糊都会直接转化为错误输出或错误行动。因此，AI-Ready 的数据基础不只是“有数据”或“数据已汇聚”，而是数据本身及其面向 AI 的派生资产都具备明确标准和持续控制。对于 L3 和 L4 场景，派生资产的版本切换、重新计算和回滚不仅需要在设计上具备，还应能够在真实事故处置中被实际执行和验证。

本维度同时覆盖结构化与非结构化数据，以及从源数据到 AI 上下文的完整数据链路。对于结构化数据，重点在于完整性、一致性、唯一性、及时性、指标口径和字段语义；对于文档、图片、音视频、代码和其他非结构化资产，重点在于内容质量、来源、版本、时效性、敏感性和适用范围。对于分块文本、嵌入向量、特征表、搜索索引和知识库等派生资产，则需要具备版本管理、变更控制和回滚能力。

- **关键指标：**数据完整率、一致性达标率、数据刷新延迟、质量规则覆盖率、质量告警响应与关闭时长、机器可读元数据覆盖率、端到端血缘覆盖率、统一指标口径覆盖率、派生资产版本化覆盖率。
- **评估重点：**关键 AI 场景的数据是否定义了时效 SLO/SLA；是否针对不同数据模态建立相适配的质量标准；是否将自动化质量检查嵌入数据管道；是否为分块、嵌入和索引等派生资产建立版本与回滚机制；是否为 AI 消费的数据提供机器可读的语义、Owner 和血缘信息；质量事件是否有记录、根因分析、改进措施和回归验证闭环。

2.3.3 平台架构与运维（Platform Architecture & Operations）

本维度回答的核心问题是：企业的数据平台是否已经从服务传统报表和离线分析的基础设施，演进为能够稳定、弹性、可观测地服务大模型、RAG、知识库、Agent 和多工具编排的数据供给平台。

AI 场景会对数据平台提出新的要求。一方面，平台需要能够支持多种数据模态、检索方式和访问路径，包括特征查询、SQL 查询、文档访问、向量检索、知识库召回、API 调用和工具调用；另一方面，平台需要统一管理 AI 应用和 Agent 的数据访问边界，避免每个项目重复建设数据副本、检索链路、权限逻辑和监控体系。

平台能力不应只被理解为“部署一个向量数据库”或“提供一个模型接口”。真正的 AI-Ready 平台需要将数据服务、身份权限、策略执行、检索质量、成本监控、容量规划、变更控制、故障恢复和持续运维结合起来。尤其对于 L3 和 L4 场景，企业应能够在发生数据延迟、索引失效、召回质量下降、权限异常或系统故障时，及时发现问题、限制影响范围并恢复至稳定状态。

- **关键指标：**数据服务与检索接口复用率、数据获取延迟、检索成功率、召回质量、索引更新时效、P95/P99 响应时间、平台可用性、自动扩缩容覆盖率、变更失败率、故障平均恢复时间（MTTR）、数据侧单位调用成本。
- **评估重点：**是否通过统一数据服务向 AI 场景提供特征、知识库、向量检索和上下文召回能力；是否建立 AI 应用和 Agent 的统一身份与细粒度数据访问控制；是否可在统一视图中观测数据延迟、成功率、检索质量和成本；重大变更

是否评估下游 AI 场景影响并具备回滚方案；关键数据管道、索引和服务故障时是否能够自动切换或回退至稳定版本。

2.3.4 组织与人才（Organization & Talent）

本维度回答的核心问题是：企业是否已形成与 AI 数据能力相匹配的职责分工、跨团队协作机制、关键技术与业务能力，以及在发生争议或事件时可执行的决策和响应路径。

AI-Ready 不是单一技术团队可以独立完成的项目。业务团队定义价值和可接受风险边界；数据团队负责数据产品、质量与供给；AI 团队负责模型、检索和应用效果；安全、法务与合规团队负责控制要求和审计边界；平台与运维团队负责稳定运行。若职责主要依赖临时协调、个人经验或项目负责人推动，企业往往能够完成 PoC，却难以让多个 AI 场景持续稳定地运行和复制。对于 L3 和 L4 场景，建议企业明确谁有权批准扩大自动化边界、谁在异常时有权中止执行，避免这一责任在业务、数据和平台团队之间处于真空状态。

本维度不应狭义理解为“拥有多少 AI 工程师”或“是否会 Prompt Engineering”。更重要的是，企业是否建立了把业务、数据、平台、AI、安全和运营工作衔接起来的组织机制；是否有明确的责任归属和升级路径；是否能让知识和经验从单个项目沉淀为可复用的流程与能力。

- **关键指标：**数据产品 Owner 覆盖率、关键 AI 场景责任人覆盖率、RACI 覆盖率、跨团队协作响应时间、重大问题升级时效、联合评审覆盖率、事件响应演练完成率、关键能力培训与认证覆盖率。
- **评估重点：**是否以 RACI 或等效机制明确端到端职责；是否建立 AI 场景上线、下线和重大变更的决策流程；跨部门发生业务、数据或风险分歧时是否有预定义的升级机制；是否对关键数据资产和高风险 AI 场景开展周期性联合评审；数据或 AI 场景发生重大问题时，是否具备经过演练的事件响应 SOP。

2.3.5 治理与信任（Governance & Trust）

本维度回答的核心问题是：企业能否让 AI 在正确的业务语义、可信的数据来源和可控的变更机制之上运行，并在出现争议、错误或影响时追溯责任链条、解释决策依据和完成纠正。

治理与信任的重点不只是制定数据标准或建立数据目录，而是使统一的数据定义、权威来源、质量要求和变更规则能够真正进入 AI 的消费链路。若同一个客户、订单、合同、库存或财务指标在不同系统和团队中具有不同定义，AI 就可能在不同场景中产生相互矛盾的结论；若 AI 可以绕开治理后的权威数据源，直接访问临时副本、原始系统或未经验证的知识库，企业即使拥有治理制度，也无法保证 AI 的实际行为受其约束。对于 L3 和 L4 场景，建议企业能够针对某一次已执行的行动，追溯其依据的数据来源、版本和指标口径，而不仅是证明源数据受控。本维度与安全合规的分工边界见下一维度开头的说明。

随着 Agent 进入业务流程，治理对象还需要从源数据延伸至动态形成的 AI 上下文和派生资产。企业需要能够识别模型回答或行动所使用的数据来自哪里、采用了哪个版本、遵循了什么口径、是否发生过变更，以及当权威数据、治理标准或权限边界变化时，下游 AI 应用是否能及时同步调整。

- **关键指标：**核心受治理数据覆盖率、权威数据源明确率、统一指标口径覆盖率、AI 场景从治理链路取数覆盖率、治理标准版本化覆盖率、变更通知触达率、下游适配完成率、关键输出可追溯率。
- **评估重点：**是否明确核心受治理数据及其权威来源；关键数据定义、命名和计算口径是否统一发布；AI 场景是否被要求从治理后的权威数据源获取核心数据；内部 AI 场景和对外数据服务是否使用一致的数据与指标口径；治理标准变更时，是否具备版本管理、通知、适配验证和问题处置机制。

2.3.6 安全与合规（Security & Compliance）

本维度回答的核心问题是：企业是否已将 AI 应用、Agent 身份、数据处理链路和工具调用纳入既有安全架构与合规体系，使其受到与传统应用及人类用户相当或更严格的身份、权限、审计、隐私和风险控制。

治理与信任主要解决“数据是否以正确、统一、受控的方式被使用”；安全与合规则主要解决“系统和数据是否可能遭受未授权访问、恶意操纵、泄露或违反监管要求”。两者必须协同，但不可互相替代。即使数据定义正确、血缘完整，若 **Agent** 使用高权限共享账号、敏感数据未经处理进入模型上下文、外部工具调用缺乏审计，企业仍会面临严重的安全与合规风险。

AI 和 **Agent** 还引入了传统应用安全之外的攻击面，例如 **Prompt** 注入、知识库投毒、恶意文档诱导、敏感信息泄露、工具滥用、越权访问和 **Shadow AI**。特别是在 **L3** 和 **L4** 场景中，企业需要将 **Agent** 视为具有独立身份和行动能力的非人类主体，为其设置最小权限、任务边界、调用限制、实时策略校验和异常处置机制。

- **关键指标：** **AI** 消费数据分类分级覆盖率、敏感数据自动识别与脱敏覆盖率、数据加密覆盖率、**Agent** 独立身份覆盖率、最小权限策略覆盖率、异常访问检测率、审计日志完整性、**AI** 安全评估或红队演练覆盖率、合规审计通过率。
- **评估重点：** 敏感数据是否在进入训练、嵌入、**RAG** 或 **Agent** 链路前完成自动化分类、标记和脱敏；数据分类分级标签是否真正驱动访问控制与脱敏策略；**Agent** 是否使用独立身份、最小权限和细粒度授权；**AI/Agent** 的身份、会话、数据访问和工具调用是否纳入安全事件响应流程；是否定期评估 **Prompt** 注入、数据投毒、工具滥用等 **AI** 特有攻击面，并验证防护措施的有效性。

六大维度并非并列的检查清单，而是具有依赖关系的能力结构。数据基础与质量、平台架构与运维决定 **AI** 能否稳定获得可用、可信的数据，是其余维度发挥作用的前提；治理与信任、安全与合规决定企业可以把自动化边界开到多大，构成扩大 **Agent** 自主性的硬约束；成本与业务价值、组织与人才决定这套能力能否被长期运营和复制，而不是随项目结束而流失。因此，评估的目的不是让六个维度同时达到高分，而是判断在企业计划推进的具体场景上，哪一个维度当前构成实际约束，下一节将详细解释如何使用成熟度模型与评估路径。

2.4 成熟度模型与评估路径

六大维度描述“企业需要具备哪些能力”，而成熟度模型用于回答“企业当前处于什么位置，以及下一步应该优先建设什么”。成熟度并非对企业整体数字化水平的笼统评价，

而是反映其数据能力、平台能力、治理能力和运行能力，是否足以支撑目标 AI 场景的自动化边界。

企业不应因为拥有某项先进技术、某个成功 PoC 或少量上线应用，就认定自身达到较高成熟度。特别是对于具备工具调用和自主行动能力的 Agent，企业必须同时具备足够的数据可信度、运行时控制、审计能力、组织责任和安全保障，才能逐步扩大自动化范围。需要说明的是，多数企业首次评估的结果落在 L1 至 L2 区间属于正常情况。首次评估的目的是建立企业自身的能力基线，作为后续复评的比较基准，而不是通过某项考核或与其他企业比较。

2.4.1 四级成熟度模型

阶段	名称	核心特征	典型表现	适合的 AI 场景
L1	未起步	相关能力基本空白，实践主要依赖个人经验和临时项目	数据孤岛明显；数据语义、质量和血缘不完整；AI 以零散试验为主；缺乏统一平台、治理与责任机制	低风险探索、内部内容生成、有限范围的知识问答
L2	初步实践	已有局部能力和工具，但覆盖有限、标准不统一，尚未形成体系	建立部分数据整合、元数据或质量监控能力；出现单点 RAG 或助手应用；权限、数据接入和运维仍常按项目单独建设	L1 场景可逐步生产化；低至中风险的 L2 分析建议场景
L3	基本达标	核心机制已建立并覆盖主要场景，但自动化、验证和闭环能力仍需加强	关键数据具备质量标准、语义和血缘；统一数据服务和治理链路开始复用；跨团队职责清晰；平台、检索和安全能力可持续监控	多场景 RAG、分析建议、部分 L3 受监督工具调用
L4	成熟领先	能力体系化、自动化、可追溯，并能够基于运行反馈持续优化	数据、派生资产、模型上下文和工具调用形成端到端可观测链路；策略可运行时执行；变更、故障与安全事件具备验证和闭环；跨部门运营机制成熟	在严格边界内推进 L3；对满足行业监管、业务影响与控制要求的场景，可审慎开放 L4 受控自主执行

需要强调的是，成熟度等级描述的是企业的**整体能力基础**，并不自动决定任何单个 AI 场景的风险等级。企业即使整体达到 L3，也不意味着所有业务都适合进入 L4 自主执行；反之，处于 L2 的企业也可能在范围有限、数据敏感度较低、人工监督充分的条件下安全运行某些 L3 受监督场景。

因此，建议将企业整体成熟度与场景风险等级结合使用：

- **L1 辅助参考场景：**企业至少需要具备基础的数据质量、来源标注、敏感数据保护和人工核验机制。

- **L2 分析建议场景：**除基础能力外，还需要确保业务语义、指标口径、数据来源和关键输出可追溯。
- **L3 受监督执行场景：**需要具备独立身份、最小权限、任务级授权、工具调用审计、人工审批或规则闸门，以及可撤销、可回滚的执行机制。
- **L4 受控自主执行场景：**应在 L3 能力基础上，进一步具备全链路可观测、运行时策略验证、异常检测、自动降级、熔断、回滚、责任归属和持续复核能力。

简言之，场景风险等级决定控制强度，企业成熟度决定其可安全承载的自动化上限。

2.4.2 如何使用成熟度模型

成熟度模型的目的不是帮助企业“证明自己已经足够先进”，而是建立一个可重复使用的评估框架，使不同角色能够围绕同一套事实作出投入、优先级和风险决策。

定位当前位置

企业可以通过六维得分及各维度的具体题目，识别自身处于 L1、L2、L3 或 L4 的整体位置，并避免两种常见误判：

- 因为已上线少数 AI 应用，而高估整体 AI-Ready 能力。
- 因为某一技术能力仍不完善，而低估已经具备的业务价值和可推进场景。

总体成熟度应与维度成熟度一起解读。例如，某企业可能在平台架构与运维上达到 L3，但在治理与信任、安全与合规方面仍处于 L1 或 L2。这意味着其可以继续优化低风险知识问答或内部辅助场景，但不宜贸然将 Agent 接入订单、付款、客户权益或生产系统等高影响流程。

明确进化路径

成熟度评估的关键输出不应只是一个总分，而应是从当前阶段进入下一阶段所需补齐的最短路径。企业可按“业务目标—场景风险—能力短板”的顺序制定建设优先级：

1. 明确计划规模化的业务场景及其目标风险等级。
2. 确认该类场景的最低 AI-Ready 准入要求。
3. 对照六维能力定位无法满足准入要求的短板。

4. 优先建设阻碍业务价值实现或限制自动化边界的能力。
5. 在完成改进后重新评估，并根据结果扩大场景范围或自动化程度。

例如，企业希望从 L2 的“风险分析建议”升级到 L3 的“受监督审批流自动化”。如果评估发现数据质量已基本达标，但 Agent 身份、任务级授权、工具调用审计和回滚机制不足，那么优先事项不应是继续扩大模型规模，而应是完善安全、治理和平台运行控制。

建立共同语言

六大维度和成熟度分级为业务、数据、AI、平台、安全及合规团队提供共同语言。业务团队可以用其说明目标价值与风险承受度；数据团队可以说明数据质量、语义和供给能力；平台团队可以说明服务稳定性与运维能力；安全与合规团队可以明确风险控制和审计要求；管理层则能够基于统一的能力视图作出预算、优先级和责任分配决策。

支撑资源分配

评估结果可直接用作预算、组织和工具建设的依据。企业不宜平均分配资源，而应优先投入对目标场景形成“硬约束”的能力短板。例如：

- 若成本与业务价值得分较低，应先建立成本归集、价值假设和阶段性决策机制，避免持续投入却无法证明收益。
- 若数据基础与质量得分较低，应优先补齐关键数据的质量标准、语义说明、血缘和派生资产管理，避免 AI 在不可信的数据上规模化运行。
- 若平台架构与运维得分较低，应优先建设统一数据服务、可观测性、检索质量监控和故障恢复能力。
- 若治理与信任、安全与合规得分较低，则应限制 Agent 的数据范围和行动边界，优先建立权威数据源、身份权限、审计与安全防护机制。

2.4.3 评估方法与场景准入判定

为避免成熟度定位停留在主观判断层面，配套自检工具将六大维度拆解为 **45** 个可验证问题。每道题应以企业当前实际运行状态为依据作答，并以可出示的证据支撑，例如监控视图、审计记录、变更单、演练报告或事件复盘记录；已立项、已编写制度但尚未在生产运行中执行的能力，不应按已具备作答。具体的选项分值、维度得分率与总体得分率的计算方法，以及成熟度映射区间，见附录 A。

单题评分规则

每道题提供 **A、B、C、D** 和 **N/A** 五种选项，对应如下分值与含义：

选项	分值	含义
A	0 分	未起步：完全不具备相关机制，或仅停留在意识和讨论阶段
B	1 分	初步实践：已有局部尝试，但范围有限、依赖人工或未形成统一机制
C	2 分	基本达标：机制已经建立并覆盖主要场景，但自动化、验证或闭环仍存在缺口
D	4 分	成熟领先：能力体系化、自动化、可追溯，并经过持续验证或实际运行检验
N/A	不计分	当前没有对应场景或能力对象，不纳入该题及所在维度的计分分母

该计分方法将 **D** 档设置为 **4** 分，而非等距的 **3** 分，目的是突出“基本具备”与“体系化、自动化、可验证闭环”之间的实质差异。企业从 **C** 档进入 **D** 档，通常不是多部署一个工具或多编写一份制度，而是需要将能力真正嵌入业务流程和生产运行中，并通过监控、演练、验证和持续改进证明其有效性。

维度得分计算

对每一维度，工具将所有非 **N/A** 题目的得分汇总，并按该维度实际作答题数计算得分率：

$$\text{维度得分率} = \frac{\sum_{i=1}^n s_i}{n \times 4} \times 100\%$$

其中， s_i 为单题得分， n 为该维度实际作答且非 **N/A** 的题目数量。

用动态分母的原因是：不同企业所处阶段不同，并非所有问题均适用。例如，尚未部署具备工具调用能力的 **Agent** 的企业，某些 **Agent** 身份和工具调用相关问题可以标记为 **N/A**；这表示该能力尚不适用于当前评估对象，而不应被误判为“零分”或“安全能力缺失”。同时，**N/A** 也不意味着该能力已经具备——若某项能力属于业务已在开展但尚未建设，应按实际状态作答，而不宜标记为 **N/A**。由于 **N/A** 表示该能力不在本次评估范围内，当某一类场景所依赖的问题被整体标记为 **N/A** 时，本次结论也就不适用于该

类场景的准入判断，企业计划开展此类场景时应先补答相应问题并重新评估。因此建议报告一并披露各维度的实际作答题数与 N/A 分布，使阅读者能够判断结论的覆盖范围。

总体得分与成熟度映射

默认情况下，工具根据所有有效作答题目的总得分计算总体得分率：

$$\text{总体得分率} = \frac{\sum_{d=1}^6 \text{维度}_d \text{得分}}{\sum_{d=1}^6 (n_d \times 4)} \times 100\%$$

其中， n_d 为第 d 个维度中实际作答的题目数量。

总体得分率和各维度得分率均使用同一套成熟度映射规则：

成熟度	得分率区间	说明
L1 未起步	0%–25%	AI 相关数据能力基本空白，多数问题处于 A 或 B 档，缺乏基础机制
L2 初步实践	26%–50%	已有局部实践和工具，但能力未形成体系，较多依赖个人经验和项目推动
L3 基本达标	51%–75%	核心机制已建立并覆盖主要场景，但自动化、验证和闭环能力仍需加强
L4 成熟领先	76%–100%	能力体系化、自动化、可追溯，并具备持续监控、验证和改进闭环

除总体成熟度外，报告还应单独展示每个维度的得分率和成熟度标签。总体得分可以说明企业的整体位置，但无法替代对结构性短板的判断：平台能力领先而安全能力不足，或业务价值机制成熟而数据质量不足，都会限制企业安全扩大 AI 自动化边界。

差距识别与复评机制

建议企业将六维得分率以雷达图或热力图形式呈现，并结合目标场景所需的最低准入条件识别优先级，而不是仅与总分或行业平均值比较。短板维度并不一定都需要立即补齐；优先事项应由企业希望推进的 AI 场景、自主性等级、数据敏感性和潜在业务影响共同决定。

建议至少每季度复评一次；在以下情形发生后，应触发专项复评：

- 新增高影响或高自主性的 **Agent** 场景。
- **Agent** 获得新的数据访问权限、工具调用能力或外部系统连接能力。
- 核心数据源、数据治理标准、检索索引或平台架构发生重大变更。
- 出现数据质量、安全、合规或 **Agent** 误操作事件。

- 计划扩大自动化范围，例如从人工审核升级为规则闸门，或从受监督执行升级为受控自主执行。

2.4.4 自检工具说明

本白皮书配套提供 **AI-Ready Data** 企业自检工具，用于将六大能力维度转化为可操作、可重复、可量化的评估过程。工具包含 **45** 道问题，覆盖成本与业务价值、数据基础与质量、平台架构与运维、组织与人才、治理与信任，以及安全与合规六个维度。

该工具的定位不是替代企业已有的安全审计、数据治理评估、架构评审或合规检查，而是提供一个面向 **AI** 规模化落地的综合视角，帮助企业将分散的能力建设活动与具体 **AI** 场景的准入要求联系起来。

工具支持以下使用方式：

- **多维能力评估：**从业务价值、数据、平台、组织、治理和安全六个维度识别现状，避免仅从模型能力或单一技术平台判断 **AI** 就绪程度。
- **场景准入评估：**将企业整体成熟度与 **L1-L4** 场景风险等级结合，判断某个场景适合停留在辅助参考、进入分析建议、开展受监督执行，还是具备探索受控自主执行的基础。
- **标准化报告输出：**形成场景准入结论、六维雷达图、各维度分项结果与评估范围披露、主要短板、优先行动建议及复评事项，为管理层提供统一的决策依据。总体成熟度可作为整体位置的参考一并呈现，但报告的核心结论是准入判断与短板清单。
- **跨团队沟通：**帮助业务、数据、平台、**AI**、安全和合规团队围绕同一套问题讨论能力差距、责任边界和建设优先级。
- **阶段性跟踪：**通过季度或专项复评，观察能力建设是否真正从局部实践进入体系化、自动化和可验证的运行状态。

通过该工具，企业可以将“是否 **AI-Ready**”从抽象判断转化为一套可讨论、可验证、可追踪的能力建设机制。最终目标不是追求更高的分数本身，而是确保企业在扩大 **AI** 使

用范围和 **Agent** 自主性之前，已经具备与其业务价值和风险水平相匹配的数据基础、治理能力和运行控制。

3 AI-Ready Data 平台：从架构设计到技术建设

3.1 平台的目标状态与建设原则

AI-Ready 数据平台的目标不是增加一套独立于企业数据平台的 AI 技术栈，而是使可信数据、业务语义和智能能力能够以稳定、可复用、可治理的方式服务分析、机器学习、生成式 AI 和 **Agent** 场景。平台既要支持人和应用持续获取一致的数据，也要支持模型和 **Agent** 在缺少人工即时纠错的情况下正确理解、受限使用并留下可验证记录。

首先企业要明确哪些数据属于关键数据，哪个系统或认证数据产品是特定用途下的权威交付对象，以及谁对质量、语义、时效和变更负责。权威并不意味着所有消费者直接查询生产系统；经过验证、具有来源和服务承诺的数据产品、认证视图或快照同样可以成为可信交付对象。**关键在于来源、转换、质量、版本和责任能够被说明和验证。**

平台需要将质量与时效从经验判断转化为可管理的服务承诺。关键数据产品应定义更新频率、允许延迟、质量底线和服务等级，并通过监控持续验证实际达标情况。结构化数据需要检查完整性、一致性、唯一性和时效；文档、图像、音视频等非结构化内容还需要检查解析保真、噪声、敏感信息和适用范围。分块文本、向量、索引、特征和评估集等派生制品同样需要版本、来源和失效条件。

平台的语义和控制能力必须能够被机器执行。字段、指标、实体和关系应具有机器可读定义；人、应用、模型和 **Agent** 应通过可识别身份访问数据；分类、用途、权限和策略应在实际访问链路中生效，而不仅存在于治理文档中。对外数据服务、内部分析和 AI 场景可以采用不同接口和数据粒度，但核心业务概念和指标口径应保持一致。

企业还需要观察数据延迟、质量、查询与检索成功率、容量、成本和异常，并将数据、语义、模型、索引、策略和下游场景通过版本与血缘关联起来。发生质量问题、平台故障或错误输出时，应能够定位来源、限制影响、恢复稳定服务，并通过根因分

析和回归验证确认改进有效。持续监测、事件处理和安全退役是 AI 生产能力的一部分，而不是上线后的附加工作。

归纳为六项建设原则：

1. **权威来源明确。** 关键数据和数据产品具有 **Owner**、来源、用途和责任边界。
2. **质量与时效可承诺。** 关键对象具有质量标准、**SLO/SLA**、监控和处置流程。
3. **语义与来源可理解。** 数据、指标和派生制品具有机器可读语义、血缘和版本。
4. **访问与行动可控制。** 人、应用、模型和 **Agent** 受到身份、最小权限和用途策略约束。
5. **运行结果可验证。** 质量、性能、检索、成本和业务结果能够通过指标与证据检查。
6. **变更与故障可恢复。** 平台具备影响分析、稳定版本、降级、重算、回滚或业务补偿机制。

3.2 三类典型消费场景

不同业务场景对上述五层能力的调用方式可归纳为三种典型模式，企业应先明确场景归属，再评估相应的能力缺口。

3.2.1 Knowledge：经治理的权威知识

用于内容生成和知识问答场景，通常以检索增强生成方式落地：非结构化内容经过解析、分块和向量化后建立索引，运行时按查询检索相关片段并组装进生成过程。这类场景的核心要求是，只有经过验证、具备权威来源标注的内容才能进入被检索的知识库，且 AI 生成的内容不应未经人工审核就直接写回知识库，避免“AI 引用自身生成结果”的错误闭环。

3.2.2 Context 与 Memory：Agent 场景的临时上下文与跨轮次记忆

当 AI 从内容生成走向自主执行工具调用，场景所使用的数据从 **Knowledge** 进一步扩展为 **Context**（某次任务从 **Knowledge**、当前输入、用户身份、会话状态、工具结果和环境中动态选择并组装的临时信息，需要明确的选择理由、**Token** 预算控制和敏感性过滤）和 **Memory**（跨步骤或跨交互保留的状态、偏好、历史或学习结果）。

Memory 应区分会话 **Memory**（随会话结束清除）、用户 **Memory**（需用户同意与访

问控制）、任务 **Memory**（任务完成后归档或清除）和共享业务 **Memory**（进入企业知识前须经验证与 **Owner** 审批）四个层次分别管理生命周期。把所有对话历史无差别当作 **Memory** 保留，是这类场景最常见的质量事故来源。

3.2.3 Semantic Query: 受治理的语义化查询

指 **Agent** 或人类用户面向语义层中已定义的指标和维度发起结构化查询，而不是直接对物理表编写 **SQL**，由语义层统一完成查询编译、权限校验与执行下推。语义层由此成为唯一的指标计算入口，业务口径的所有权归属数据团队，使用方仅拥有组合和调用的权限。

Semantic Query 的价值体现在三个方面。口径一致：指标的聚合方式、过滤条件、去重规则和时间粒度在语义层单点定义并纳入版本管理，所有使用入口共用同一份定义，口径变更可追溯、可回滚、可通知下游，避免“同名不同算”引发的跨部门争议。计算正确：语义层自动处理连接路径解析、粒度对齐、扇出重复计数和比率类非可加指标的聚合顺序，这类错误在手写或模型生成的 **SQL** 中通常不会触发报错，只表现为数值的静默偏差，是最难被发现的质量风险。可治理可审计：行列级权限、脱敏策略、并发和成本约束在语义层统一执行，不依赖各使用端自行实现，同时完整留存查询主体、时间、所引用的指标版本和实际下发的物理查询，形成端到端可审计的记录。

对 **Agent** 而言，这一转变把开放式的 **SQL** 生成，收敛为受限语义空间内的参数填充——查询意图可枚举、可校验，也可以在执行前被拒绝。这使 **Agent** 的数据访问从“结果事后核对”变为“请求事前约束”，显著降低了幻觉和结果不可复现的风险，是 **Agent** 数据使用具备可靠性的前提条件。

语义层保证的是计算逻辑一致，不保证跨权限主体的数值相同——行级权限生效时，不同主体查询同一指标本应得到不同结果。若语义层内并存多个语义相近的指标或多个口径版本，仍然存在选择错误的可能，需要通过命名规范、别名映射和消歧机制加以约束。因此，语义层是 **Semantic Query** 成立的必要条件，而非充分条件，实际收益取决于语义资产本身的建模质量和治理成熟度。

3.3 能力分层与责任划分

逻辑架构描述企业需要具备哪些长期稳定的能力、这些能力由谁负责以及如何被不同场景复用，不直接规定产品和部署方式。如下图所示，**AI-Ready** 数据平台可抽象为五层能力链，并由治理、安全与合规，以及平台工程、评估与运行保障两条横向控制面贯穿。采用五层能力链与两条横向控制面，覆盖从数据进入平台到形成业务结果的完整过程。

AI-Ready 平台逻辑能力架构

以可信数据为基础，构建贯穿数据、语义、AI 与业务的端到端能力链，支撑企业 AI 场景的安全、可控、规模化落地。



第一层是数据与内容供应，管理数据库、事件、文件、文档、图像、音视频和外部数据，并形成具有来源和使用边界的候选资产。

第二层是数据处理与产品化，负责批量、增量和流式处理，执行质量验证、版本管理、数据契约和产品发布。

第三层是语义与知识，管理术语、指标、实体、关系、认证查询、本体和知识，使人和机器能够使用一致业务含义。

第四层是分析与 AI 服务，提供 SQL、特征、模型、检索、Context 和评估能力。

第五层是消费与业务激活，通过 BI、应用、模型推理、RAG、Data Agent、API 和事件把能力交付给消费者，并在受控条件下影响业务状态。

其中，治理、安全与合规控制面贯穿五层，管理数据分类、用途、隐私、权限、策略、例外和批准。平台工程、评估与运行保障控制面负责环境、版本、容量、性能、成本、可观测性、变更、事故和恢复。前者回答什么行为被允许以及在哪里执行控制，后者回答真实发生了什么、是否符合预期以及如何介入和恢复。

3.3.1 统一平台的含义

“统一平台的核心是统一标准与控制接口，而非将所有场景只能使用一个服务端点。

企业可以同时使用数据湖、数仓、湖仓、操作型数据库、搜索、向量和图存储，也可以为不同领域提供专用的数据和 AI 服务。只要这些组件能够接入共同目录、身份、契约、策略、观测和证据接口，就可以形成一个可治理的组合平台。反之，把所有功能集中到一个产品，也不能自动形成统一语义、责任和运行证据。

3.3.2 责任分工

共享平台团队适合建设连接、目录、质量框架、身份、策略执行、可观测性、评估和工程模板等高复用能力；领域团队负责数据产品、业务定义、指标口径、场景质量和消费者结果。核心数据的权威性应通过认证数据产品和治理链路向下游传递，AI 场景不应绕过这些链路使用来源不明的副本。对外数据服务和内部 AI 可以采用不同的脱敏、聚合和访问方式，但应继承相同的核心定义和来源。

统一身份和访问控制也不意味着 Agent 继承人类用户全部权限。平台应区分用户、应用、模型服务、Agent、工具和目标系统身份，并通过委托关系限制其可访问的数据范围和可执行动作。Agent 权限应按任务最小化，定期复核并保留可查询的访问和行动记录；源系统和目标业务系统仍保留最终授权权。

3.3.3 五层能力的分工与证据

逻辑层	共享平台能力	领域或场景责任	关键证据
数据与内容供应	连接标准、分类、目录、来源和接入记录	确认权威来源、用途和数据 Owner	核心数据清单、来源和分类标签
数据处理与产品化	质量框架、契约、版本、血缘和发布工具	定义质量、时效、Schema 和 SLA	数据契约、质量结果和版本记录
语义与知识	术语、指标、实体和语义发布机制	批准业务定义、关系和适用边界	指标定义、认证查询和语义版本
分析与 AI 服务	查询、特征、检索、Context 和评估服务	定义场景质量、阈值和使用方式	服务目录、评估结果和运行 SLO
消费与业务激活	接口、身份、策略、限流和调用记录	对分析结果、模型决策和业务行动负责	访问日志、策略决定和业务结果

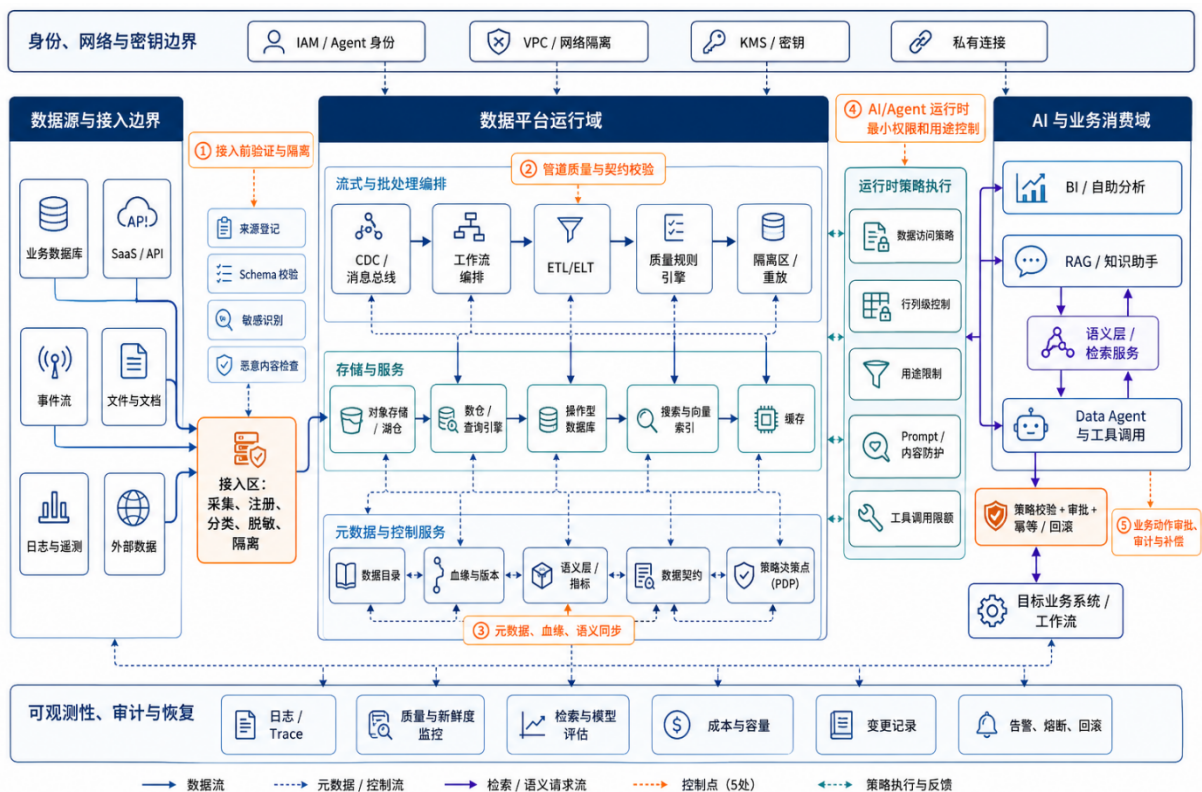
3.3.4 保持稳定与可变的部分

逻辑架构中需要保持稳定的是 Owner、契约、语义、版本、SLO、访问边界和证据；存储、计算、模型和 Agent 框架可以随场景演进而调整。来源、处理活动与责任主体

之间的关系，应采用可跨系统交换和追溯的统一表示方式。来源与责任关系的表达可参考第 2.1.1 节介绍的 W3C PROV-O 思路。

3.4 技术架构与数据管控

技术架构将逻辑能力落实为可部署、可连接和可运营的组件，并明确数据接入、处理、存储、检索、AI 运行时与业务执行各环节的控制点。企业可以保留既有数据库、数仓、湖仓和应用系统，通过共同目录、身份、契约和观测接口逐步增强 AI 能力，无需进行一次性平台替换。如下图所示，技术架构将数据接入、处理、存储、数据服务、AI 运行时与业务消费连接为端到端链路，并在关键节点执行身份、策略、质量、审计和恢复控制。



端到端技术链通常包括数据源，批量、CDC 和流式接入，存储与计算，数据工程与产品服务，语义及 AI 服务，以及分析、模型、RAG 和 Agent 消费。横向平台能力提供身份、网络、密钥、策略、IaC、CI/CD、可观测性、容量、成本和灾难恢复。每一层可以由多个平台承载，但控制和证据必须沿实际数据路径连续传播。

3.4.1 数据接入

数据进入平台前，应记录来源、**Schema** 和摄取批次，并执行分类分级、敏感数据识别、脱敏或标记。分类标签只有在实际驱动访问、检索过滤和脱敏策略时才形成控制；仅在目录中登记标签，不能证明 **AI** 消费链路受到保护。结构化与非结构化来源需要不同的质量方法，但都应在进入正式数据产品、训练、**Embedding** 或 **RAG** 链路前完成必要验证。

3.4.2 数据处理

数据处理层负责 **ETL/ELT**、**CDC**、事件路由、格式转换、质量检测、异常隔离和 **Schema** 演进。批量链路需要记录输入范围、批次、代码和配置；流式链路还要处理乱序、重复、积压、状态、检查点和回放。关键管道应定义吞吐、延迟和资源基线，并通过容量测算或压测验证，而不是只监控任务是否成功。

3.4.3 存储与服务

存储与服务层可以组合湖、仓、湖仓、操作型数据库、缓存、搜索、向量和图存储。平台应根据数据规模、访问模式、时效、一致性和成本选择技术，并为高并发 **AI** 调用设计缓存、连接限制、限流、降级和熔断。自动扩缩容属于规模化能力；无论自动还是人工，容量变更都应可观察、可审计，并能够关联触发原因和场景影响。

3.4.4 数据和 AI 服务

数据和 **AI** 服务需要通过受治理接口提供特征、查询、文档、知识库、检索和 **Context** 能力。统一服务不要求只有一个物理网关，而要求服务可发现、契约一致、身份可识别、权限可限制、调用可观测。访问边界应能够细化到数据产品、表、字段、知识库分区或工具，并区分用户、应用、模型服务和 **Agent** 身份。

3.4.5 各环节的管控内容与关键指标

控制位置	主要控制	关键指标	检查追踪
数据进入	来源、分类、 Schema 、脱敏和隔离	分类覆盖率、识别准确率、拒绝数量	摄取记录、标签和脱敏日志
数据处理	质量规则、批次、版本、重试和回放	成功率、吞吐、延迟、积压	质量结果、运行记录和检查点
存储与服务	权限、缓存、限流、容量和降级	P95/P99、QPS、容量、错误率	压测报告、策略和伸缩记录
语义及 AI 服务	指标、特征、索引、检索和评估	新鲜度、召回、相关性、漂移	版本 Manifest 、评估结果和 Trace

消费与激活	身份、用途、审批、幂等和补偿	调用成功、拒绝、错误行动	访问审计、策略决定和业务记录
-------	----------------	--------------	----------------

3.4.6 数据流与元数据流须同步支持

技术架构必须同时支持数据流和元数据流。**Schema**、**Owner**、分类、质量、血缘、版本、权限、成本和使用记录需要从源、管道、存储、语义服务和消费端持续采集，并反向用于影响分析、访问控制和失效传播。没有持续元数据流的统一目录只能说明资产曾经被登记，不能证明真实运行状态。

3.5 数据产品的建设目标

AI-Ready 数据产品建设的目标，是把源数据持续转化为具有质量、语义、来源和服务承诺的数据产品，并确保上游变化能够传播到特征、索引、模型和 **AI** 场景。数据产品可以是表、文件、**API**、事件、特征或知识服务；其权威性来自来源、转换、质量、版本和 **Owner**，而不是简单来自物理存储位置。

3.5.1 质量与服务承诺

数据产品应按照前文建设原则中定义的质量与时效承诺机制执行，并将相关要求落实到发布、监控、告警和处置流程中。结构化数据通常关注完整性、一致性、唯一性、准确性和时效；文档、图片和音视频还需关注解析保真、内容完整性、噪声、敏感性和适用范围。质量规则应尽量嵌入数据管道和发布流程，不满足底线的版本应被阻断、隔离或明确降级，而不是等待下游发现问题。

3.5.2 机器可读的语义

机器可读语义是连接数据治理和 **AI** 消费的关键。核心数据需要明确字段含义、业务口径、单位、取值范围、实体关系和适用边界；关键指标应通过统一定义、认证查询或语义层发布。对外服务和内部 **AI** 可以使用不同数据粒度及脱敏方式，但应继承相同的核心定义。指标型 **Semantic Layer** 的实现表明，集中指标、关系和认证查询可以同时服务 **BI** 与自然语言查询。

核心受治理数据不应被理解为所有场景直接访问源系统。企业可以发布经过质量和语义验证的认证数据产品或视图，并通过访问控制要求 AI 场景从这些权威交付对象取数。临时副本、个人脚本和未经验证的知识库如果被用于关键场景，应能够被发现、限制或阻断。

3.5.3 数据血缘与变更影响

血缘需要至少覆盖源、表或视图、转换、数据产品及 AI 消费端，并进一步关联重要派生制品。上游 Schema、权限、内容或语义发生变化时，平台应形成影响清单，通知消费者，并对受影响的数据产品、特征、Chunk、索引、指标和场景执行适配与回归验证。删除和撤回还必须传播到副本、缓存、索引和 Memory，避免失效信息继续被消费。

3.5.4 事件处理与持续改进

数据质量或时效问题一旦影响 AI 输出，应形成事件记录，包括时间线、影响范围、根因、恢复过程 and 责任人。改进措施需要通过回归验证确认有效，并沉淀为新的质量规则、测试用例、告警或工程模板。已有研究表明，向包含数百万条文本的知识库中针对某个目标问题仅注入五条恶意文本，即可使 RAG 系统以约 90% 的成功率输出攻击者指定的答案^[2]，少量恶意或错误内容也可能通过检索链路影响特定问题，因此来源控制和异常监测需要覆盖原始内容及其派生索引。

3.5.5 关键能力与验证方式

能力	关键指标	可验证证据
数据时效	更新延迟、SLO 达标率	SLO/SLA、看板和告警
数据质量	规则覆盖率、异常数量、关闭时长	规则结果、处置和关闭记录
机器语义	元数据和统一指标覆盖率	目录、定义和认证查询
血缘与版本	血缘覆盖率、版本化覆盖率	血缘图、Manifest 和回退记录
标准变更	通知覆盖率、适配完成率	版本、影响清单和回归报告

3.6 分析、AI 与 Agent 如何使用数据

平台能力最终需要通过不同消费模式形成业务结果。BI、ML、实时决策、GenAI 和 Data Agent 可以共享数据产品、语义、身份和观测能力，但对时效、质量、接口和风

险控制的要求并不相同。平台应统一责任与控制接口，而不是强制所有工作负载采用同一技术路径。

3.6.1 分析与 BI

分析和 BI 场景主要依赖认证数据集、统一指标、查询权限和数据新鲜度。自然语言问数仍应使用已发布的指标和语义契约，不能因为交互方式变化而绕过认证查询或自行生成新的业务口径。关键查询还需要资源限制、超时和成本控制，防止单次请求影响共享平台稳定性。

3.6.2 机器学习与实时决策

ML 和实时决策链路需要管理训练数据、特征、标签、代码、模型、评估集和部署版本，并监测训练—服务偏差、数据漂移、推理延迟和业务表现。实时场景还应处理事件顺序、重复、状态、检查点和回放。模型通过技术评估不代表已经获得业务上线批准；高影响决策仍需明确适用范围、阈值、人工复核和停止条件。

3.6.3 生成式 AI 与 RAG

GenAI 与 RAG 链路从内容解析、文本块、向量和索引开始，运行时根据身份、任务、权限和检索策略选择内容并组装上下文。RAG 能够引入可更新知识，但检索到相关内容不等于内容真实、完整或适用。平台应持续评估解析保真、召回、相关性、引用覆盖、生成忠实性和敏感信息处理，并在质量下降时告警、降级或切换稳定版本。

3.6.4 Data Agent

Data Agent 在上述能力之上增加任务规划、Memory、工具调用和业务行动，因此需要更严格的身份和策略边界。Agent 应使用独立身份，并将代表的用户、任务目的和授权范围传递到数据、检索、工具和目标系统。权限应按任务最小化，覆盖表、字段、知识库分区、API 和工具，并定期复核与回收。Agent 身份与人类身份必须能够在审计记录中区分。

工具应作为受控企业资产登记，至少定义 Owner、业务用途、输入输出 Schema、错误、幂等、资源范围、额度、审批、版本和补偿语义。MCP、A2A 等协议有助于能力

发现和组合，但不自动建立信任、授权和责任。内容 Guardrail 用于处理不良内容、敏感信息和 Prompt 攻击，不能替代数据权限、工具授权和业务审批。



3.6.5 按风险逐步扩大自动化范围

同一项能力可以按照风险逐步开放为只读查询、建议、准备待审批对象、受监督执行和受限自主执行。提高自主程度应同时增强策略、审批、监测和恢复，而不是放宽 Prompt。付款、通知、合同和删除等难以逆转的动作，应拆分为准备、提交、确认和补偿阶段，使用幂等键和权威交易记录确认最终状态。

下图展示了 Data Agent 的一种典型运行方式：以受管控的数据和统一语义为基础，组装任务所需的上下文，通过工具触发业务操作，并由治理安全和运行保障两条管控线贯穿全程。这是平台能够支持的较高要求的使用方式，并非所有 BI、机器学习或生成式 AI 场景都必须采用这一完整流程。

3.7 平台运维、变更与持续改进

AI-Ready 平台的生产能力需要通过可观测性、容量管理、变更控制、故障恢复和事件改进持续证明。统一观测不是把全部日志集中到一个看板，而是采用一致的场景、服务、数据产品和版本标识，使读者能够从业务场景下钻到数据管道、检索、模型、工具和成本。

3.7.1 统一的运行视图

统一运行视图至少应覆盖数据新鲜度和 **SLO** 达标率、数据服务成功率、**P95/P99** 延迟、**QPS**、容量、流处理积压、检索与召回质量、质量异常、存储计算成本和扩缩容行为。指标需要按场景和数据产品分解，否则企业只能看到平台整体“正常”，无法判断某个 **AI** 场景是否正在使用过期数据、低质量索引或异常成本。

3.7.2 容量规划

容量规划应基于数据规模、调用量和响应时间目标，并通过压测或真实负载验证。缓存、限流、降级和熔断用于保护共享平台，自动扩缩容则属于规模化能力。无论采用自动还是人工方式，每次容量变化都应记录触发原因、资源变化和效果；高成熟度平台可以使用统计或 **AI** 方法识别静态阈值难以发现的性能和质量异常，但检测模型本身也需要监测误报、漏报和漂移。

3.7.3 重大变更处置

数据平台重大变更上线前，应形成受影响数据资产、派生制品、接口和下游 **AI** 场景清单，并完成性能或效果验证。变更对象包括 **Schema**、数据源、管道、存储、特征、**Chunk**、索引、语义、模型、工具和策略。发布记录应包含版本、审批、迁移窗口、消费者通知和恢复方案；验证结果应成为上线决定的一部分，而不是上线后的补充材料。

3.7.4 按故障类型区分恢复方式

恢复动作需要按故障类型区分。任务失败可以重试或从检查点恢复；数据和索引错误可以切换稳定版本、重新计算或隔离；服务故障可以切换、扩容或降级；已经发生的业务交易则可能需要取消窗口、客户处置和业务补偿。自动回退并非总是安全，涉及权限撤回、隐私删除或业务规则变化时，**Fail-closed** 和人工确认可能比恢复旧版本更合适。

3.7.5 平台四类记录

平台至少需要区分四类记录：可观测 **Trace** 用于性能和调用路径诊断；安全审计日志记录身份、访问、策略和配置变化；权威业务记录证明实际状态变化；发布和事故证

据包汇总评估、批准、恢复、补偿和责任关闭。四类记录应通过共同标识关联，但不能互相替代。来源、版本和责任关系可以参考 **PROV-O** 等通用模型表达。

3.7.6 数据事件处理

数据质量、时效、检索或权限问题一旦影响 **AI** 输出，应形成标准事件记录，包括时间线、影响范围、根因、恢复动作和责任人。事件关闭前需要验证改进措施，例如新增质量规则、调整容量、修复索引、收紧权限或补充评估集，并通过回归测试确认同类问题不再复现。每次事件都应沉淀为新的规则、测试、告警或工程模板，形成持续改进闭环。

安全事件 **SOP** 需要覆盖 **Agent** 身份、代表用户、会话、数据访问、检索和工具调用，使事件响应人员能够回答“哪个主体在什么任务下访问了什么数据并执行了什么动作”。企业还应根据暴露面开展 **Prompt** 注入、知识库投毒、越权和工具滥用评估；红队结果需要进入修复、回归和批准流程，而不能停留在一次性测试报告。

3.7.7 责任划分与关键指标

责任域	主要对象	核心指标与关闭证据
数据平台 SRE	管道、存储、查询、流处理和共享数据服务	可用性、延迟、积压、新鲜度、 MTTR 、恢复和成本
ML/GenAI 运行保障	特征、模型、索引、检索、 Context 和生成	漂移、模型表现、召回、忠实性、引用、版本和评估
Agent 运行保障	任务、 Memory 、工具、策略和业务行动	任务成功、策略遵从、错误行动、人工介入和补偿
SecOps	身份、访问、攻击和安全事件	异常访问、隔离、取证、权限回收和演练
业务 Owner	决策、交易和客户影响	业务结果、风险接受、损失处置和责任关闭

3.7.8 成本可见性

平台投入还需要按场景归集存储、计算、数据处理、检索、模型和人工运行成本，使共享能力的复用收益和场景边际成本能够被比较。成本可见性不能证明业务价值，但可以避免平台团队承担全部投入而消费团队无法说明使用和收益。

3.8 平台投入与业务价值的度量

成熟度定位解决了“现在处于哪个阶段”的问题，本节进一步回答“这些投入是否值得、该如何向管理层证明”的问题。企业在推进 **AI-Ready** 建设时，应同步建立一套贯穿立项、建设、上线全周期的价值度量框架，而非等项目结束后再补做 **ROI** 核算。

3.8.1 度量框架设计：选择合适的代理指标

由于 **AI** 项目的最终业务收益往往滞后于平台投入，建议采用分阶段代理指标体系，随建设阶段推进逐步切换度量重心：

- 建设期指标（对应 **0-30** 天诊断阶段）：**AI-Ready** 评估基线得分、场景优先级矩阵覆盖度、跨团队工作组组建完成度。
- 试点期指标（对应 **30-60** 天基础建设阶段）：语义层/数据契约覆盖的资产数量、**PoC** 通过率、平台部署完成度。
- 生产期指标（对应 **60-90** 天及以后）：用例生产化比例、**Agent** 自主完成率、异常拦截率、人工复核工时节省比例、故障平均恢复时间（**MTTR**）改善幅度。

3.8.2 现有可参考的行业经验

企业自身历史数据积累不足的情况下，以下数值可作为方向性的建议参考，并非经过独立验证的行业标准，企业不宜将其直接作为达标线：

- 高 **AI** 成熟度组织将 **AI** 项目维持在生产环境三年以上的比例，据行业实践观察^[3]普遍高于低成熟度组织，两者差距可达两倍以上，可作为判断“组织成熟度投入与项目持续性关系”的方向性参考。
- 遵循结构化立项方法（例如明确价值假设、可视化路径、分阶段验证）的项目，据实践经验更容易从概念阶段过渡到生产阶段，部分项目可以在一个半月左右完成上线，可作为立项阶段效率的参考区间。
- 建议按企业年收入规模（小于 **1** 亿、**1-10** 亿、大于 **10** 亿）、所属行业（金融服务、制造业、零售电商、医疗健康、科技互联网）、数据团队规模（小于 **10** 人、**10-50** 人、大于 **50** 人）三个维度，对内部指标进行分层，便于识别更贴近自身情况的参照对象，而不是笼统对比全行业均值。

企业应优先积累自身首个试点场景的实测数据作为内部基线，并在后续场景推进中与这一基线对比，衡量真实的改进幅度；上述参考数值仅用于企业尚无内部基线时的初步判断，不能替代基于自身数据的评估。

3.8.3 典型 **ROI** 测算路径参考

结合企业实践中观察到的收益类型，建议在测算 ROI 时覆盖以下几类价值来源，并区分“可直接量化”与“需定性评估”两类：

- 可直接量化：存储与计算成本节省（如迁移至云原生存储架构降低的基础设施支出）、人工审核工时节省（Agent 自动处理占比提升后释放的人力）、故障响应时效提升带来的运营成本降低。
- 需定性评估、逐步量化：决策质量提升（分析辅助场景下建议采纳率与业务结果改善的相关性）、客户体验提升（响应时效与满意度变化）、合规风险降低（审计发现问题数量的下降趋势）。

企业在正式立项测算时，建议结合自身历史项目数据，或引入第三方专业机构做定制化测算；本框架提供的仅为方向性参考指标，不能替代企业自身的财务测算。

4 AI-Ready Data 的组织 and 流程：构建面向 AI 的协同运营体系

数据平台、检索能力和模型服务就绪后，AI 场景并不会自动进入可信赖的生产运行状态。平台解决的是数据能否被获取、被处理、被检索；而数据是否真实可信、语义是否清晰准确、行动边界是否有人背书，取决于平台之上的人和流程如何分工与协作。技术架构可以支撑数据流动，但无法替代组织回答三个问题：谁对数据的质量和契约负责，谁把领域经验转化为机器能理解的含义，谁在 AI 生成结论或触发行动后承担业务后果。

这些问题在传统数据体系中长期依赖惯例、口头共识和个人经验来回答，也因此很少被显式设计。当数据的消费者从人扩展到模型和 Agent，这种隐性分工方式开始难以为继：模型不会主动追问“这个指标的例外情况是什么”，Agent 也不会执行前像人一样犹豫“这个操作是否需要请示”。组织和流程因此成为 AI-Ready 能力中，与数据基础、技术平台同等重要、但更容易被忽视的一环。

本章将说明数据生产、语义赋予和业务消费这三种当责职能应当如何围绕一个治理良好、语义清晰的数据核心运转，Agent 又如何作为高频参与者穿行其间、并逐步接管部分传统职能。围绕这个框架，本章依次讨论组织形态如何选择、角色分工如何重构、文化与流程如何转变，以及跨团队协作机制如何将各方绑定在同一个可信数据核心之上。

4.1 数据组织与流程面临的新挑战

4.1.1 三种常见组织形态：集中式、分散式与联邦

企业的数据与 AI 能力通常围绕不同程度的集中化，形成三种组织形态：集中式、分散式和联邦式。它们并非绝对互斥；同一家企业也可能在集团层面采用集中式平台，在部分业务领域采用分散式交付，并逐步向联邦式协同演进。

在传统数据分析阶段，这三种模式都可以有效运作。区别主要体现在需求由谁承接、数据和平台由谁建设、业务规则由谁解释，以及上线后的运行责任由谁承担。当企业开始推进生成式 AI、RAG 和 Agent 场景时，组织形态本身并非决定因素，关键在于其配套流程能否覆盖持续的数据消费、动态的权限控制和逐步扩大的自动化边界。

组织形态	典型组织模式	对应的传统流程	面向 AI-Ready 的主要缺口
集中式	中央数据、AI 或平台团队统一承接需求，建设数据平台、数据管道、报表、知识库和 AI 应用；业务团队主要提出需求并参与验收	业务提出需求 → 中央团队排期与交付 → 业务验收 → 项目结束或转入基础运维	中央团队容易成为需求、数据接入和审批瓶颈；难以持续掌握领域语义、例外规则和业务变化；项目验收后数据、索引、模型配置和 Agent 行为的长期责任容易不清晰
分散式	各业务领域或事业部独立拥有数据、分析和 AI 团队，自主建设数据产品、知识库、应用和自动化流程	领域发现需求 → 领域团队自行取数、开发和上线 → 依靠领域内部经验运营和迭代	响应快但容易形成重复数据副本、重复知识库和重复工具链；业务语义、质量标准、权限模型和审计证据难以一致；跨域数据使用和统一风险控制成本高
联邦式	中央团队提供共享平台、标准、身份权限、治理规则和风险护栏；领域团队负责领域数据产品、业务语义、场景交付和价值实现	中央团队提供可复用能力与准入要求 → 领域团队基于标准开发和运营场景 → 双方共同处理跨域变更、风险与运行事件	对数据产品 Owner、责任边界、工程规范和协作机制要求较高；若标准不能嵌入平台和流水线，容易退化为文档治理；若中央与领域责任划分不清，则可能造成重复审批或责任真空

三种形态的本质差异，可以用一个更直观的方式理解：每个业务领域都应围绕一个治理良好、语义清晰的**数据核心**，构成“谁生产数据、谁赋予含义、谁承担业务后果”这三种当责职能的责任闭环，同时这个闭环需要立于一套共享的平台与治理能力之上。集中式是这三种职能和共享能力全部归于中央；分散式是共享能力缺位，各领域各自建立自己的闭环；联邦式将数据生产、业务语义和场景价值责任更多下放至领域团队，同时由中央团队提供共享平台、统一标准和风险护栏。其优势在于：领域团队更接近业务语义与场景价值，中央团队更适合沉淀通用平台能力和治理机制。分散式之

所以容易退化为数据孤岛，正是因为各领域的闭环之间缺少统一的共享底座维系，彼此没有被绑定到同一个可信来源上。

4.1.2 不同阶段适合的形态

集中式模式适合企业早期建设统一平台、沉淀专家能力和控制初始风险；分散式模式更贴近业务，适合快速探索和局部创新；联邦式模式则更适合多个业务领域同时推进 AI、并需要兼顾共享复用、领域敏捷和统一治理的阶段。对多数企业而言，组织演进并不是简单地从集中式“切换”为分散式或联邦式，而是逐步将中央团队的通用能力平台化，同时让领域团队承担更多业务语义、数据产品和场景运营责任。

无论企业当前采用哪一种形态，进入 AI-Ready 阶段后都会共同面对以下挑战。

首先，**数据和 AI 能力不再是一次性交付的项目成果，而是需要持续运营的生产能力**。传统流程通常以报表、数据集或应用上线为阶段性终点；但 AI 场景中的数据源、文档、分块策略、嵌入向量、索引、Prompt、模型配置、权限策略和工具接口都会持续变化。企业需要明确谁负责维护这些资产，谁负责验证变更影响，谁负责处理质量下降、检索失效、权限异常或 Agent 行为偏差。

其次，**业务语义不能继续主要依赖人工解释和隐性经验**。过去，业务负责人和分析师可以在使用报表时解释异常数字、补充例外规则或纠正字段理解；但模型和 Agent 需要在运行时使用机器可读的语义、质量规则、指标定义和权限边界。企业需要将领域知识从分散文档、会议记录和个人经验中沉淀为可被数据产品、检索系统和 AI 应用共同使用的资产。

第三，**控制机制需要从静态配置和阶段审批延伸至运行时**。传统权限管理通常围绕用户角色、系统边界和数据对象展开；而 Agent 可能代表不同用户、在不同任务中跨数据源访问内容，并进一步调用工具执行动作。企业需要将 Agent 身份、任务级授权、数据访问范围、工具调用限制、人工复核、审计记录、异常熔断和回滚机制连接成完整链路。

最后，**组织协同的重点从“谁来完成交付”转向“谁对持续结果负责”**。中央团队需要继续负责高复用平台、统一标准、工程模式和风险护栏；领域团队需要负责业务语义、数据产品质量、场景效果和业务价值；安全、合规和运营团队需要参与高风险场景的

准入、监控和事件处置。只有这些责任能够清晰衔接，企业才能在不牺牲安全和可信度的前提下，逐步扩大 AI 的应用范围与自动化程度。

因此，AI-Ready 组织与流程的核心，不是选择某一种“正确”的组织结构，而是建立一套能够让共享能力、领域知识、工程控制和持续运营共同发挥作用的协同机制。

4.2 AI-Ready 组织与流程的目标状态

4.2.1 从项目交付转向数据产品运营

AI-Ready 的组织与流程，并不要求企业采用统一的组织架构，或在短时间内完成全面重组。其目标是建立一套能够持续运行的协同机制：业务价值有明确责任人，数据语义和质量有长期 Owner，平台能力可以被复用，风险边界能够在运行中执行，关键变化和异常能够被及时发现、追溯和处理。

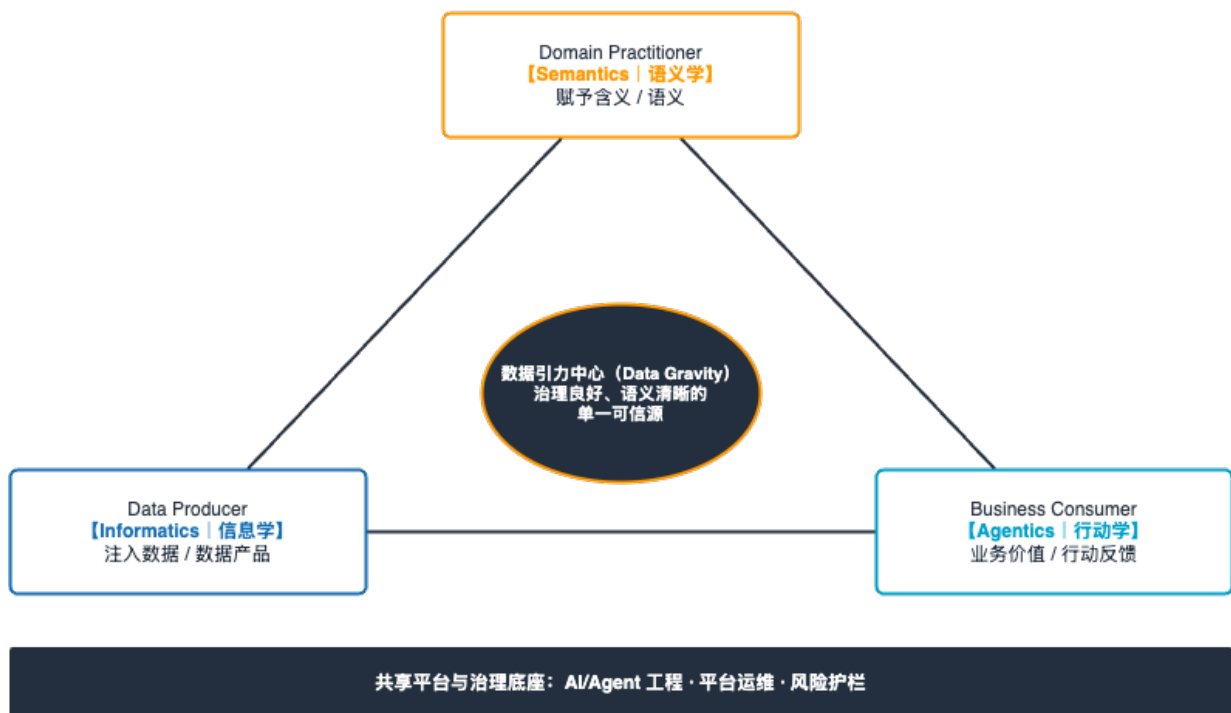
在这一目标状态下，企业不再将数据、AI 应用和 Agent 视为一次性项目交付物，而是将其视为持续运营的生产能力。组织的重点也从“谁来接收需求并完成开发”，转向“谁对业务结果、数据可信度、系统运行和风险控制共同负责”。

4.2.2 责任三角：围绕数据核心的三种当责职能

职位名称会不断增殖，但 AI-Ready 组织中不可再分的当责职能只有三个，它们围绕一个治理良好、语义清晰的数据核心，构成一个责任三角：

- **数据生产者（Data Producer）**：生产并交付可信数据产品，对其质量、契约和可消费性负责——向核心注入数据。
- **领域实践者（Domain Practitioner）**：一线业务人员，把领域中的隐性经验显式化为机器可读的语义（口径、定义、边界），并以业务现实校验 AI 产出——为核心赋予含义。这是 AI 问数和 Agent 决策正确性的语义源头。
- **业务消费者（Business Consumer）**：消费 AI 输出并对业务后果负最终责任，划定自动执行与人工复核的边界——从核心兑现价值。

AI-Ready 组织内核：围绕数据引力中心的责任三角



三顶点的学科视角：Informatics (形式/结构) → Semantics (本体/语义) → Agentics (行动/兑现)

这三个职能构成一条清晰的价值链：数据被生产出来，经领域赋予业务含义，最终用于实现业务价值。将数据置于三角中心，意味着三方不应各自维护难以对齐的数据副本，而应围绕同一组经过治理的权威数据产品协同工作。

这个三角是每个业务领域内不可外包的核心，它立于一套共享的平台与治理底座之上——**AI/Agent 工程**、平台运维、治理与风险护栏都属于这一层，可以被集中建设甚至托管。**Agent** 则是穿行在三角之间的第四个参与者：它既消费数据，也产生派生结果和实际动作，是绕数据核心运转最频繁的一环，并且正在逐步接管业务消费者乃至数据生产者的部分职能，因此必须被显式纳入这个框架，而不能被笼统归入“消费者”。

4.2.3 从执行者到监督者：角色如何被重新定义

责任三角对应的是一套价值递进关系，也对应三层逐渐深入的能力要求：数据生产者对应“数据是否结构良好、可信、可被消费”；领域实践者对应“数据的确切含义和适用边界是什么”；业务消费者对应“这些含义应该被如何使用，由谁承担后果”。三者依次把原始数据升级为可理解的知识，再升级为可兑现的业务价值。其中第三层——信息

经由行动被使用并转化为价值——在自主 Agent 出现之后，正日益由 Agent 代为执行，其风险性质与前两层有明显不同，值得被单独关注，本章后文会专门讨论。

这一转变背后有一条共同规律：**AI-Ready** 转型中角色变化的本质，是人从**执行回路之内**移向**执行回路之上**。过去，人是流程中的具体一环——分析师亲手写 SQL，测试人员亲手设计用例；现在，AI 接管了大量执行性工作，人退到流程之上，负责定义约束、确认边界、验证结果、承担风险。关键不在于“这些角色是否还存在”，而在于每个角色受到的冲击程度不同，且都必须把过去依赖人工经验兜底的判断，转化为显式、可契约化、可验证的规则。

下表是这一框架在真实组织中的落地：

角色	角色定位	传统职责	AI-Ready 转型职责
业务用户	需求提出者与业务结果解释人	提出需求，确认报表是否符合预期	对 AI 支撑的决策进行业务背书，明确自动执行与人工复核的边界
数据产品经理	数据资产 Owner	收集需求，定义数据产品形态	建立数据契约，明确可被 AI 场景消费的范围及已知局限
数据分析师/领域专家	业务需求量化拆解者	设计指标体系，撰写 SQL 和报表	验证 AI 结论的合理性，补充机器可读的语义说明
数据架构师	平台架构设计与实现者	设计数据湖/仓架构，实现 ETL	架构同时支撑传统分析和 AI 检索，建立变更通知机制
数据平台 SRE	平台可用性守护者	监控运行状态，故障排查与恢复	统一监控 AI 负载，建立差异化的容量与成本预警
数据测试	质量验证者	设计 SAT/UAT 测试用例	为 AI 行为设计专项测试，建立持续回归验证
AI/Agent 工程师（新增）	AI 应用与 Agent 的构建者	传统组织中通常缺位	搭建 RAG 检索管道、调优检索质量、编排 Agent 与工具接入
AI 治理/风险 Owner（新增）	自主行为的风险守门人	原本分散在合规/安全职能中	制定内容护栏、漂移检测，以及受监督执行和受控自主执行场景的审批与熔断策略

这些角色变化可以归纳为三个层次，帮助企业判断转型投入的优先级：

- **责任分量上升的角色：** 业务用户、数据产品经理。当 AI 具备自主执行能力后，“谁对业务后果负责、边界如何划定、数据契约由谁维护”成为刚需，这两个角色的分量不降反升，数据产品经理更是联邦式组织中的枢纽。
- **职责范围扩大的角色：** 数据架构师、数据平台 SRE。其职责从支撑传统分析扩展到支撑 AI 检索和 AI 负载运维，SRE 也逐步向面向 AI 系统的运维能力演进。
- **最受冲击、价值重心需要上移的角色：** 数据分析师、数据测试。他们过去手工产出的工作——写 SQL、写测试用例——正是 AI 最先能够自动化的部分。这些角色不会消失，但价值重心必须转移到验证 AI 结论、补齐语义、设计 AI 行为的评测和回归上，否则会被持续挤压。

核心结论是：手工产出会被 AI 替代，但责任不会消失，只会转移到语义、验证和护栏这几个环节。同时，企业需要显式设立传统组织中缺位的 AI/Agent 工程和 AI 治理角色，否则相关能力只能停留在口号上，落不到具体的人。

4.2.4 治理提醒：自主行动的风险不能默认落在业务消费者身上

责任三角带来一条容易被忽视、但十分关键的治理原则：当 AI 从生成建议走向调用工具、触发业务动作时，业务消费者作为最终受益方，不应该同时是唯一的风险制动方。换言之，如果一个 Agent 自动执行的操作出现问题，责任不能默认全部归于使用该 Agent 的业务团队，因为获益方和制动方合而为一，会削弱风险控制的独立性。

这意味着，对于受监督执行和受控自主执行等具备较高自主性的场景，风险护栏的当责必须显式落位——或者归属于共享平台层面的治理职能，或者由三角中的某一方明确兼任风险守门人的角色，而不能默认交由业务消费者一并承担。数据生产者和领域实践者所承担的数据形式治理和语义治理，可以随着工具成熟而大幅自动化；但涉及授权、背书和制动的这一层责任，恰恰是人从“回路之内”退到“回路之上”之后最不可让渡的职能——Agent 是执行者，而不是责任主体。这一原则应当在场景准入、审批流程和事件响应机制中被明确写入，而不是留给隐性默契。

4.2.5 目标状态的五个转变

结合责任三角和角色重构，AI-Ready 组织与流程的目标状态可以概括为五个转变。

从项目交付转向数据产品运营。传统数据项目通常以报表、数据集、接口或应用上线作为主要完成标志。AI-Ready 模式下，关键数据资产、知识库、检索服务和面向 Agent 的数据接口应被视为长期运营的数据产品，具备明确的数据 Owner、权威来源、业务语义、质量标准、更新频率、服务等级、使用范围、访问边界、版本策略和变更通知方式。这一转变的关键不在于要求每一张表都成为正式“产品”，而在于识别真正影响业务决策、客户体验、核心运营或 AI 输出的关键数据资产，并让它们具备可被持续信任和复用的责任机制。

从单点协作转向共同定义场景边界。AI-Ready 场景不应仅由业务团队提出需求、技术团队完成实现，而应由业务、数据、AI、平台、安全和合规等相关角色在场景开始阶段共同确认业务目标、数据条件、自动化边界和风险控制要求，明确该场景解决什

么问题、AI 输出属于哪个自动化层级、哪些结论需要人工复核、出现异常时如何限制或回滚，以及谁分别对业务结果、数据质量、系统运行和风险控制负责。

从共享资源转向可复用的能力平台。 企业需要避免两种低效模式：中央团队成为所有需求的唯一交付入口，或各业务领域重复建设数据管道、知识库、检索能力和监控体系。目标状态是“中央能力平台化、领域责任产品化”：中央团队提供数据接入、元数据、质量、身份权限、检索与知识库服务、可观测性、审计和工程模板等高复用能力；领域团队负责领域语义、数据产品运营、场景需求和价值实现。

从事后治理转向控制内建。 数据质量规则应嵌入数据管道和发布流程；敏感数据应在进入模型和检索链路前完成分类和脱敏；访问权限应基于身份、任务和风险动态执行；高影响工具调用应受到审批、限额或人工复核约束；关键行为和异常应保留完整审计证据，并支持告警、降级、熔断和回滚。控制强度应与场景的自主性、数据敏感性和业务影响范围相匹配。

从上线完成转向持续验证与演进。 AI 应用的生产运行具有持续变化的特征，企业需要建立覆盖版本与变更管理、持续测试与评估、端到端可观测性、事件响应与改进的运营闭环，使问题能够被定位到数据、索引、模型、提示词、权限或工具等具体环节，并将每次事件沉淀为新的规则和工程模板。

目标状态	核心变化	典型体现
数据产品化	从一次性交付转向长期运营	关键数据和知识资产具有 Owner、语义、质量标准、SLO、版本与变更机制
场景共担化	从业务提需求、技术做交付转向共同定义价值与风险	业务、数据、AI、平台和风险职能共同明确目标、数据范围、自动化边界和升级路径
能力平台化	从项目重复建设转向共享能力复用	数据接入、检索、权限、审计、监控、评估和工程模板通过统一平台或标准提供
控制内建化	从事后审批和审计转向运行时控制	质量门禁、敏感数据保护、动态授权、人工闸门、限额、熔断与回滚嵌入流程
运营持续化	从上线即完成转向持续验证和演进	版本管理、回归测试、可观测性、事件响应和复盘成为常态机制

AI-Ready 组织与流程的目标，不是让所有团队承担相同的职责，而是让每一项关键责任都有明确归属，并通过可复用的平台能力、可执行的工程控制和持续运营机制，将业务价值、数据可信度与风险控制连接成一个完整闭环。

4.3 支撑目标状态落地的文化与流程

组织架构和角色分工的调整，只是 **AI-Ready** 转型的骨架；真正让数据核心持续保持可信和语义清晰的，是日常的文化习惯和工作方式。责任三角能否稳定运转，最终取决于团队是否已经从下面四个方面完成习惯上的转变。

4.3.1 从“需求驱动”转向“产品思维”

数据团队不再是被动响应需求的执行方，而是主动运营数据产品：建立产品目录、服务等级协议、变更日志和订阅机制，让消费者能够随时了解一个数据资产的状态和边界，而不必每次都通过临时沟通确认。

从“一次性交付”转向“持续迭代”。数据产品应具备完整的生命周期管理，包括版本迭代、废弃通知，以及向前兼容和向后兼容策略，使下游的 **AI** 场景在数据产品演进时不会被意外打断。

从“人工校验”转向“自动化验证”。企业应建立自动化的数据质量门禁，在数据发布前自动校验完整性、一致性和时效性，把过去依赖人工抽查才能发现的问题，提前拦截在生产环境之前。

4.3.2 从“内部消费”转向“双轨服务”

同一套数据产品需要同时服务人类用户和 **AI Agent**，为两类消费者提供不同粒度的访问接口和不同层级的语义注释——人类用户可能只需要一份业务说明文档，而 **Agent** 则需要机器可读的字段定义、取值范围和调用协议。

这些转变并非无法衡量的“软性口号”。它们完全可以被纳入前文六大维度的自检体系，以具体问题的形式打分和跟踪，而不是仅凭印象判断“团队是否具备数据驱动的思维”。转变也无需按固定顺序、一次性到位。更现实的做法是从一个高价值场景切入，把“产品思维、自动化验证、双轨服务”这几项习惯在这个场景里跑通一轮，再用交付出的业务价值驱动下一轮迭代。

4.3.3 让责任可执行的跨团队协作机制

如果说数据核心的可信度和语义清晰度，是把生产者、实践者和消费者绑定在一起的“引力”，那么以下三类机制，就是让这股引力变得可靠、可预期、可追责的具体纽带

——它们把原本停留在口头约定层面的协作，转化为可执行的契约。联邦式组织的实际运转，正依赖这三类机制才能落地，而不只是依赖组织架构图上的分工描述。

数据契约（Data Contract）： 生产者与消费者之间的正式协议，规定数据格式、更新频率、质量底线和变更通知责任。它是数据生产者对数据核心可信度做出的正式承诺，也是数据产品经理日常工作的核心抓手。

服务等级协议（SLA）： 明确响应时间、可用性、数据新鲜度等可量化承诺，为下游 AI 场景——尤其是高频调用数据的 **Agent**——提供一份可预期的依赖关系，避免下游场景因为不清楚数据的可靠边界而被迫自行加装冗余校验。

变更通知协议（Change Notification）： 任何可能影响下游消费者的变更，包括 Schema 调整、数据源切换或工具更新，都必须提前通知并提供迁移窗口。这一机制的价值在 **Agent** 场景中尤其突出：如果数据核心的演进没有被及时通知，**Agent** 的决策逻辑可能在无人察觉的情况下发生漂移，直到某次错误行动才被发现。

这三类机制分别对应责任三角中不同的关系：数据契约巩固的是生产者对核心的承诺，服务等级协议巩固的是消费者对核心的依赖，变更通知协议则保证核心自身的演进不会在下游引发不可见的风险。三者共同作用，才能让联邦式组织中“中央提供平台、领域拥有语义”的分工真正可执行，而不是停留在文档层面的治理宣言。

4.3.4 迈向规模化的组织能力建设

责任三角、文化习惯和协作机制，最终需要落到一组具体的组织能力上，才能支撑企业从少数试点走向跨部门复用和生产级运行。这些能力的共同特点是：它们不依赖某个项目经理、数据专家或 AI 工程师的个人经验，而是通过明确责任、可复用资产、工程流程和运行机制，沉淀为组织本身具备的能力。

组织能力	核心问题	主要机制	对规模化的作用
场景组合管理能力	企业应优先建设哪些 AI 场景，并将自动化开放到什么程度？	场景分级、价值假设、风险评估、上线准入、阶段复盘与退出机制	避免“各部门各自试点”；将资源投入到价值明确、数据条件成熟、风险可控的场景
数据产品运营能力	谁持续保证数据和知识资产可被 AI 可信使用？	数据产品 Owner、数据契约、质量标准、服务等级目标、版本与变更管理	让数据资产从一次性交付物转变为可复用、可追责、可持续维护的生产能力
领域语义共建能力	AI 如何持续获得正确、最新的业务解释和例外规则？	领域专家参与、机器可读元数据、统一指标口径、知识维护流程、评估样本与反馈闭环	降低模型误解业务规则、使用过期知识或在跨团队场景中产生矛盾结论的风险
共享平台与工程赋能能力	如何避免每个团队重复建设同一套数据、	统一数据服务、知识检索、身份权限、评估框架、开发模板、CI/CD、可观测性与审计	缩短新场景交付周期，提升平台复用率，使可靠实践可被快速复制

	检索、权限和监控能力？		
风险分级与运行控制能力	如何在扩大自动化的同时限制错误输出和错误行动的影响？	Agent 身份、最小权限、任务级授权、审批闸门、操作限额、策略校验、熔断与回滚	支持从辅助参考逐步发展到受监督执行，并在条件成熟时探索受控自主执行
变更与事件协同能力	数据、索引、模型、权限或工具变化后，如何防止 AI 行为静默劣化？	影响分析、版本控制、变更通知、回归验证、统一告警、事件响应、复盘改进	使系统变化可控、问题可定位、故障可恢复，降低生产运行的不确定性
能力沉淀与人才发展能力	如何让经验不局限于少数专家或单个项目？	参考架构、工程规范、案例库、内部社区、培训、结对实践、角色轮换与专家辅导	将试点经验转化为组织资产，扩大可参与 AI 建设与运营的人才范围

场景组合管理能力

AI 规模化并不等于同时启动更多项目。企业需要建立一种场景组合管理机制，持续判断哪些业务问题值得优先投入，哪些场景已经具备进入生产的条件，哪些场景仍应停留在辅助或实验阶段。这要求企业在场景启动前明确业务价值假设、目标用户、数据条件、风险等级、人工监督要求和成功指标；在场景运行后，根据用户采纳、输出质量、成本、风险事件和业务效果，决定扩大使用范围、优化方案、维持现状或暂缓投入。

特别是对 **Agent** 场景，企业应将“是否允许执行动作”视为独立决策，而不是模型能力自然升级的结果。一个能够生成采购建议的 **Agent**，不一定具备自动下单的条件；一个能够总结客户信息的助手，也不一定具备修改客户权益的权限。组织应根据场景自主性、数据敏感性、操作可逆性和业务影响范围，逐级开放自动化能力。

数据产品运营能力

在规模化阶段，企业的关键数据、知识库、指标服务和检索资产不能继续主要依赖项目团队维护。每项关键资产都需要有明确的长期 **Owner**，并具有相对稳定的产品边界：它服务哪些业务对象和 AI 场景，使用哪些权威来源，承诺什么质量和时效标准，允许哪些消费者访问，以及发生变化时如何通知和恢复。

对于服务 RAG 和 **Agent** 的数据产品，运营对象还包括文档、分块结果、嵌入向量、索引、元数据和访问策略。数据产品 **Owner** 需要与 AI 团队共同确保这些派生资产与权威来源保持一致，并能够在内容更新、权限变化或质量异常时完成同步、验证和回滚。

领域语义共建能力

企业的数据平台可以统一技术架构，但无法完全替代业务领域对自身规则、例外、客户关系、风险偏好和决策逻辑的理解。因此，**AI-Ready** 组织需要建立领域专家、数据产品 **Owner**、分析人员和 **AI** 工程师共同参与的语义共建机制，把关键概念、指标口径、例外规则、可接受的判断逻辑和高风险操作边界，转化为可维护的组织资产，例如数据词汇表、指标定义、知识文档、评估集和异常处理指南。随着场景持续运行，用户反馈、错误案例和生产事件应持续回流，推动语义和评估标准不断更新。

共享平台与工程赋能能力

规模化 **AI** 建设的常见瓶颈，不是缺少模型服务，而是每个团队都需要重复解决数据接入、文档处理、检索、权限、评估、部署、监控和审计等工程问题。企业应将高频、共性且需要统一控制的能力沉淀为服务、组件、模板或流水线，包括数据与知识接入、元数据与血缘、文档解析和索引、向量检索、统一身份与权限、敏感数据识别、评估框架、部署流水线和审计记录，并确保领域团队能够低摩擦地使用这些能力，避免因平台不好用而被绕开。

风险分级与运行控制能力

随着 **AI** 从内容生成走向分析建议、工具调用和自主执行，组织需要具备与自动化程度相匹配的控制能力，覆盖场景定义、数据访问、模型输出、工具调用和异常处置的完整链路。对于受监督执行场景，应增加 **Agent** 独立身份、最小权限、任务级授权、审批闸门、操作限额和可撤销机制；对于受控自主执行场景，则需要运行时策略校验、持续监控、异常检测、熔断、降级和回滚能力。这一能力的目标不是阻止自动化，而是让企业能够根据实际控制能力，逐步扩大自动化边界。

变更与事件协同能力

企业需要将变更管理从传统的基础设施和应用发布，扩展到数据、知识、**AI** 配置和 **Agent** 工具的完整链路：对重大变更明确版本、影响范围、下游通知、迁移窗口和回滚方案；对关键场景在变更后执行数据质量、检索质量、输出行为和工具调用的回归验证。当出现数据延迟、知识过期、检索失效、权限异常或工具误调用时，业务、数据、平台、**AI** 和安全团队应能基于统一的告警、日志和责任分工快速协同，并通过根因分析和复盘防止同类问题再次发生。

能力沉淀与人才发展能力

AI-Ready 的组织能力最终取决于知识能否被持续传承，而不是少数专家能否在关键项目中“救火”。企业需要将场景建设和运行过程中的经验沉淀为参考架构、数据产品模板、质量规则库、评估集、测试用例、故障案例和应急手册等可复用资产。人才发展也不应只关注增加 AI/ML 工程师数量：业务人员需要理解如何定义场景目标和人工复核边界，数据人员需要掌握非结构化数据的 AI 消费要求，工程人员需要掌握检索、评估和 Agent 编排，安全人员则需要能够识别 Prompt 注入、数据投毒等新型风险。企业可以通过内部实践社区、案例分享、结对交付和专家辅导，让能力从少数团队扩散到更多领域。

5 实施路线图：从评估到规模化的落地参考

企业在推进 AI-Ready 数据建设时最容易低估的，不是某一项具体能力的建设难度，而是能力之间的先后顺序和推进节奏。评估框架、平台架构和组织责任本身并不会自动排列成一条可执行的路径——把它们变成路径，需要企业回答一个更朴素的问题：先做什么、什么时候做、做到什么程度才能进入下一步。

5.1 路线图的设计逻辑

企业推进 AI-Ready 建设时，容易陷入三类常见误区。

误区一：“等平台建好了再做 AI”。一些企业认为必须先完成统一数据平台、语义层和治理体系的建设，才能启动 AI 场景，结果平台建设周期被无限拉长，业务侧看不到进展，也无法用真实场景验证平台设计是否合理。平台建设和 AI 场景落地应当并行推进，以场景需求驱动平台能力建设，而不是反过来。

误区二：“先做多个 PoC 再考虑规模化”。另一些企业倾向于同时启动多个独立的 PoC，寄希望于事后从中挑选出可复制的模式。但没有基础设施支撑的 PoC 通常无法直接复制，每个场景都要重新走一遍数据接入、权限配置和评估验证的弯路，PoC 数量增加，并不必然带来规模化能力的沉淀。

误区三：责任真空。业务、数据、平台团队往往各自以为“风险由对方负责”，尤其是在 Agent 具备工具调用能力之后，谁来批准自动化边界、谁来承担错误行动的后果，如果不在实施初期就明确写入分工，会在场景上线后才暴露成真正的事故。

正确的策略是“场景拉动、平台承接、组织保障”：以业务价值明确、风险可控的场景作为切入点，倒逼平台能力和组织机制同步建设，并在实践中以可控节奏补齐最关键的短板，而不是试图一次性建成完备的能力体系。本章后续的组织演进阶段、90 天行动计划、度量指标和复评机制，都是这一策略的具体展开。

5.2 组织演进的三个阶段

组织形态的演进和平台能力的成熟，应当是同步发生、相互印证的两条线，而不是各自独立推进。企业不需要在转型初期就重构组织架构，也不应因为看到成熟企业采用联邦式模式，就急于把全部数据与 AI 责任下放给业务领域——组织演进的节奏应当与企业当前的成熟度阶段相匹配。

5.2.1 阶段一：局部试点——建立共同语言和最小闭环

对应成熟度模型中的 **L1 未起步到 L2 初步实践** 阶段。企业应优先选择业务目标明确、影响范围可控、能够保留人工复核的场景，例如内部知识问答、文档摘要、客服辅助或运营分析辅助——这类场景通常属于风险分级中的 **L1 辅助参考** 或 **L2 分析建议**。此时的重点不是部署更多模型或建设完整企业平台，而是验证业务、数据、技术和风险团队能否围绕同一场景形成共同语言：谁是业务负责人、谁是数据负责人、AI 输出属于参考还是建议、出现错误时如何处理。

这一阶段通常适合采用集中式组织支持。中央数据、AI、架构或安全团队可以提供参考架构、基础工具和安全指导，业务领域则提供真实业务问题和使用反馈，减少各领域从零摸索的成本。

5.2.2 阶段二：共享能力建设——沉淀标准和产品

对应成熟度模型中的 **L2 到 L3 基本达标** 阶段。当企业从一两个试点扩展到多个场景后，重复建设会成为主要问题：不同团队分别建设文档处理流程、权限校验逻辑和评估机制，相同的数据资产被以不同方式复制和解释。这一阶段的核心任务，是把五层能力和三类消费场景（**Knowledge**、**Context** 与 **Memory**、**Semantic Query**）落地为可复用的共享能力：建立关键数据的语义说明和数据契约、部署基础的检索和知识库服务、定义质量门禁和监控基线。

组织上，中央团队的定位从“交付项目”转向“提供可复用能力”；领域团队开始承担数据产品的语义确认、质量验证和场景验收责任，这正是数据生产者、领域实践者两类责任逐步落地的阶段。

5.2.3 阶段三：规模化协同运营——联邦式责任与控制机制

对应成熟度模型中的 **L3 到 L4 成熟领先** 阶段。当企业拥有多个持续运行的场景、可复用的数据产品和较成熟的平台能力后，组织重点从“建设共享能力”转向“在共享能力之上持续运营”。联邦式模式在这一阶段发挥核心作用：中央团队负责企业级平台、标准和风险护栏，领域团队负责业务语义、数据产品运营和场景价值实现，双方通过数据契约、服务等级协议和变更通知协议协同处理跨域变更和风险。

进入这一阶段后，企业才具备条件审慎评估 **L3 受监督执行** 甚至部分 **L4 受控自主执行** 场景——但企业整体达到 **L3** 并不意味着所有业务都适合自动进入 **L4**，扩大自动化边界仍需针对具体场景重新评估数据可信度、权限控制和恢复机制是否达标。

5.3 90 天分阶段行动计划

三个组织阶段描述的是较长周期的能力演进，企业在启动 **AI-Ready** 建设的第一个场景时，可以参考以下更具体的 90 天节奏，把整体框架转化为可执行的行动项。这份计划对应的是组织演进的“阶段一”，即建立第一个可控闭环；企业在后续场景中可以复用相同的节奏，但周期通常会随着共享能力的积累而缩短。

5.3.1 第一阶段：诊断与对齐（0-30 天）

目标是完成现状评估，建立跨团队共识，确定优先场景。

- 完成六大维度自检问卷，输出基线评估报告，明确企业当前所处的成熟度阶段。
- 对候选场景涉及的数据资产执行敏感数据扫描与分类，明确个人信息及其他受管控数据的分布位置。
- 依据 **L1-L4** 场景风险分级，聚焦 **2-3** 个高价值、风险可控的试点场景——后续经过评估优先选择风险等级较低、但业务价值明确的场景作为试点。
- 组建跨功能工作组，至少覆盖业务、数据、架构和安全四类角色，并参照责任三角明确谁承担数据生产者、领域实践者和业务消费者的职责。
- 明确每个候选场景的成功标准和衡量指标，取得管理层支持。

交付物：**AI-Ready** 评估报告、敏感数据分布报告、场景优先级矩阵、工作组章程。

5.3.2 第二阶段：基础建设（30-60 天）

目标是为试点场景建立最小可行的语义层和治理机制。

- 为试点场景所涉及的数据资产建立元数据和机器可读语义说明。
- 建立数据契约和变更通知机制，明确数据产品的 **Owner** 和服务等级承诺。
- 依据场景所属的典型消费模式（**Knowledge**、**Context** 与 **Memory**、或 **Semantic Query**），选择相应的技术路径——例如知识问答场景优先建设检索增强生成链路，分析场景优先建设语义层和认证查询。
- 部署基础平台能力，包括知识库、检索索引、访问护栏和审计日志。

- 定义数据质量门禁和监控基线，建立自动化测试套件，确保不满足质量底线的数据不会进入生产链路。

交付物：语义层最小可行产品、数据契约模板、平台基础能力部署完成、质量门禁上线。

5.3.3 第三阶段：场景落地与迭代（60-90 天）

目标是让首个场景上线运行，并建立持续优化的反馈闭环。

- 场景从开发进入生产，开始承接实际用户流量；如果场景涉及 **Agent** 工具调用，按风险分级逐步开放能力，优先以受监督执行方式上线，而不是直接开放自主执行。
- 启用身份感知的权限控制与行为审计，确保每次数据访问和工具调用可追溯，落实平台四类记录（可观测追踪、安全审计日志、权威业务记录、发布与事故证据）。
- 启动持续监控和质量评估，建立周期性的数据质量与检索质量审计。
- 基于实际运行数据优化检索策略、护栏规则和访问控制。
- 复盘首个场景的建设过程，提炼可复用的模式和模板，作为下一个场景的起点，并同步评估工作组分工和责任划分是否需要调整。

交付物：首个场景上线、运行监控仪表盘、复盘报告和下一场景的扩展计划。

5.3.4 每阶段的度量与决策点

将分阶段的代理指标体系与上述三个阶段直接对齐，避免管理层在不同环节看到相似但不完全一致的指标口径。

阶段	对应时间窗口	核心度量指标	决策点
诊断与对齐	0-30 天	AI-Ready 评估基线得分、场景优先级矩阵覆盖度、跨团队工作组组建完成度	是否已识别出风险可控、价值明确的候选场景；工作组是否覆盖必要角色
基础建设	30-60 天	语义层和数据契约覆盖的资产数量、PoC 通过率、平台部署完成度	试点场景的数据基础是否达到进入生产的最低质量和语义要求
场景落地与迭代	60-90 天及以后	用例生产化比例、Agent 自主完成率、异常拦截率、人工复核工时节省比例、故障平均恢复时间改善幅度	是否具备扩大场景范围或提升自动化等级的条件

每个决策点都应当以“是否满足下一阶段的最低要求”作为判断标准，而不是简单以时间到期作为推进依据——如果诊断阶段发现数据基础缺口过大，宁可延长诊断周期，也不应带着不完整的基线评估进入基础建设阶段。

5.4 复评与持续演进机制

首个场景完成 90 天周期后，企业不应把这份路线图当作一次性的项目计划，而应将其转化为可以重复运行的节奏。成熟度复评回答“企业整体能力是否有进步”，这里的复评回答“下一个场景应该按什么节奏推进”，两者是同一件事的两个层面。复评的建议周期与应触发专项复评的具体情形已在第二章「差距识别与复评机制」中列出，本节不再重复，重点在于复评结果如何转化为下一个场景的推进节奏。

每一轮复评的输出，应当反过来影响下一个场景的 90 天计划：如果复评显示治理与信任维度仍然薄弱，下一个场景的基础建设阶段就应当优先投入数据契约和权威数据源建设，而不是急于扩大 Agent 的工具调用范围。通过这种方式，路线图本身也具备了持续演进的能力，而不是停留在首个场景上线那一刻。

5.5 亚马逊云科技解决方案映射

以下表格展示亚马逊云科技服务如何支撑各阶段建设目标：

AI-Ready 数据平台能力与亚马逊云科技服务映射

参考架构服务组合；实际选型应结合数据类型、访问模式、风险等级与既有平台能力

能力需求	亚马逊云科技服务	典型用法
1 统一存储与数据湖	Amazon S3（含 S3 Tables）+ AWS Lake Formation + AWS Glue	构建企业级数据湖，以 Apache Iceberg 等开放表格式存储数据，统一权限管理，避免专有格式锁定。
2 语义层与数据目录	Amazon DataZone + AWS Glue Data Catalog + SageMaker Unified Studio	建立业务术语表、数据产品目录与权限策略，作为语义与知识层的落地载体。
3 知识库与 RAG	Amazon Bedrock Knowledge Bases；可选 Amazon OpenSearch Serverless、Amazon Aurora、Amazon Neptune Analytics 或 S3 Vectors	构建 RAG 管道，支持文档解析、分块、嵌入和检索；向量存储可按规模和一致性要求选择。
4 Agent 编排与执行	Amazon Bedrock Agents + Amazon Bedrock AgentCore（Runtime / Gateway）+ AWS Step Functions	支持 Agent 构建、工具集成与注册、多步骤任务编排；Gateway 统一拦截和路由 Agent 与工具之间的调用。
5 Agent 运行时授权与行为评估	Amazon Bedrock AgentCore Policy + AgentCore Evaluations	Policy 在工具调用前依据 Cedar 策略拦截并授权，回答“能不能做”；Evaluations 持续评估实际表现与风险，回答“做得好不好”。两者支持按风险分级扩大自动化范围。
6 安全护栏与审计	Amazon Bedrock Guardrails + AWS CloudTrail + Amazon CloudWatch	内容过滤、敏感信息检测、访问审计与异常告警。
7 数据质量与监控	AWS Glue Data Quality + Amazon CloudWatch	自动化质量规则检查、数据新鲜度与运行指标监控、告警通知。
8 知识图谱与本体建模	Context Ontology Accelerator（开源，基于 W3C RDF/OWL）+ Amazon Neptune / Neptune Analytics	以 AI 辅助和人工审核构建业务本体，存储于客户自有知识图谱；支持实体关系建模、复杂查询和多跳检索。
9 身份与访问控制	AWS IAM + IAM Identity Center + AWS Lake Formation	对用户、应用、模型服务和 Agent 设置独立身份与细粒度最小权限。
10 组织级知识图谱与 Agent 上下文	AWS Context（2026 新发布，逐步吸收 Context Ontology Accelerator 能力）	推断跨数据源实体关系、业务规则和权威来源，构建组织级知识图谱；通过身份感知的 Agentic Search 和 MCP 接口提供受权限约束的运行时上下文查询，并以 Iceberg 格式发布元数据至 S3 Tables。

注：服务可按企业现有架构替换或组合；本表仅说明能力与亚马逊云科技服务的映射，不构成单一部署方案。

6 总结与展望

AI-Ready 不是一个可以被最终宣告“完成”的终点状态，而是一种需要持续运营的能力。全文从企业面临的转型压力出发，依次给出了衡量能力的六大维度和成熟度模型、承载能力的平台架构和三类消费场景、承担责任的组织形态和责任三角，最后落到一份可执行的实施路线图——这四个层面并非各自独立的建设任务，而是同一套能

力在不同视角下的展开。数据基础、技术平台、组织分工与治理机制需要同步演进，任何一环滞后，都会成为其余各环投入无法转化为业务价值的瓶颈。企业应避免将 **AI-Ready** 建设理解为“先完成平台、再启动 AI”或“先堆积 PoC、再考虑规模化”的线性过程，而应以目标场景为牵引，在数据、平台、组织、治理和运行控制之间持续迭代。

本白皮书提出的框架，其价值不在于要求企业追求某一项指标的满分，而在于提供一套共同语言：业务负责人可以用风险分级判断哪些场景可以扩大自动化边界，数据和平台团队可以用五层能力架构和三类消费场景明确技术投入的优先级，组织和治理团队可以用责任三角厘清谁该为数据、语义和业务后果负责，管理层则可以用成熟度模型和度量框架，把分散的建设投入转化为可追踪、可验证的进展。这套语言的意义，是让不同角色在推进 **AI** 转型时，不再各自为政、各说各话。

随着 **Agent** 从生成建议逐步走向调用工具、触发业务动作，企业面对的将不再只是模型能力的持续演进，还有随之而来的、需要持续跟上的数据可信度、运行时控制和责任归属要求。技术平台和云服务能力仍在快速迭代——从统一存储格式到组织级知识图谱，从检索增强生成到 **Agent** 运行时的策略与评估双闸门，新的工具会不断降低具体能力的建设门槛。但工具的进步不能替代企业自身对场景风险的判断、对数据责任的坚持，以及对自动化边界的审慎扩展。真正决定企业能否从局部试点走向规模化生产的，始终是这些能力能否被组织起来，形成一套可信、可控、可持续运营的整体。

亚马逊云科技将持续与企业客户并肩推进这一转型，从诊断评估到架构设计，从组织协同到落地实施，提供全程支持。如需进一步了解本框架或启动评估，欢迎联系亚马逊云科技客户团队获取更多资源。

附录

附录 A：六大维度自检问卷

<https://github.com/aws-samples/sample-ai-ready-data-assessment>

附录 B：术语表

- **AI-Ready**: 数据、平台、组织与流程共同具备的系统性能力，使 AI 工作负载能够稳定消费。
- **Agent**: 具备观察、规划、行动、响应能力的 AI 智能体，可自主完成任务。
- **RAG: Retrieval-Augmented Generation**, 检索增强生成，通过外部知识检索增强模型生成质量。
- **PROV-O: W3C** 提出的来源模型标准，描述数据对象的来源、生成过程和责任归属。
- **Context (上下文)**: 某一次任务从 **Knowledge**、当前输入、用户身份、会话状态、工具结果和环境中动态选择并组装的临时信息，是 **Agent** 场景中需要精细化管理 **Token** 预算和敏感性过滤的核心对象。
- **Memory (记忆)**: 跨步骤或跨交互保留的状态、偏好、历史或学习结果，按生命周期可分为会话、用户、任务和共享业务四个层次，各层次的写入、保留与删除策略不同。
- **Semantic Query (语义化查询)**: 通过语义层进行受治理的指标查询，而非直接对原始表编写 **SQL**，确保查询结果的口径一致性和可审计性。
- **Context Engineering (上下文工程)**: 为 AI 任务动态选择、组装、精简信息以匹配 **Token** 预算和相关性要求的工程实践，区别于静态的 **Prompt** 编写。
- **Ontology (本体)**: 基于 **RDF/OWL** 等 **W3C** 标准显式建模实体、关系与业务规则的知识表示方法，支持机器进行符号推理和一致性验证。
- **Knowledge Graph (知识图谱)**: 以图结构存储实体及其关系的数据库，是 **Ontology** 的具体技术实现，支持多跳查询和关联发现。
- **Semantic Layer (语义层)**: 定义业务术语、指标口径和维度关系的中间层，使人类用户和 **AI Agent** 都能基于统一且经过认证的语义理解数据，而非各自解读原始表结构。

- **Data Contract（数据契约）**：数据生产者与消费者之间的正式协议，定义数据格式、更新频率、质量底线和变更通知责任，是保障 AI 消费稳定性的核心机制。
- **Guardrail（护栏）**：在 AI 输入输出两端设置的安全策略，用于检测并拦截越权访问、幻觉内容、PII 泄露或违规指令。
- **Agentic Data Stack（智能体数据栈）**：在传统数据湖/仓/向量库基座之上，新增语义层、本体、知识图谱三层 AI 理解能力，共同支撑 Agent 稳定消费数据的技术架构体系。
- **MCP（Model Context Protocol）**：一种开放协议，用于标准化大模型与外部工具、数据源之间的连接方式，使 Agent 能够以统一接口调用不同系统的能力。
- **Chunking（分块）**：将长文档切分为较小片段以便嵌入和检索的预处理步骤，分块大小和重叠窗口直接影响检索召回率和精确率。
- **Embedding（嵌入）**：将文本、图像等非结构化内容转换为高维向量表示的过程，使语义相似的内容在向量空间中距离相近，是向量检索的基础。
- **Hallucination（幻觉）**：模型生成看似合理但与事实不符或无法追溯到来源的内容，是生成式 AI 在企业场景应用中的核心风险之一。
- **Data Lineage（数据血缘）**：记录数据从源头经过转换、加工直至消费的完整路径，是故障溯源、影响面分析和合规审计的基础能力。
- **TCO（Total Cost of Ownership，总体拥有成本）**：涵盖基础设施、人力、运维、许可等全生命周期成本的综合度量指标，常用于平台投资决策的对比分析。
- **MTTR（Mean Time To Recovery，平均恢复时间）**：系统发生故障后恢复正常运行所需的平均时间，是衡量平台运维韧性的核心指标。
- **Data Product（数据产品）**：以产品化思维运营的数据资产，具备明确的 Owner、SLA、版本管理和变更通知机制，可被内部或外部消费者稳定订阅使用。

- **Feedback Loop**（反馈闭环）：将 AI 输出的实际效果（用户采纳、业务结果、人工修正）回传至上游数据和模型环节，用于持续优化的机制；未经验证的反馈闭环可能导致错误自我强化。
- **Prompt Injection**（提示注入）：攻击者通过用户输入、网页、文档、邮件、工具返回结果或其他外部内容，将恶意指令混入模型上下文，诱导模型忽略既定规则、泄露敏感信息、调用不应调用的工具或执行越权操作。可分为直接提示注入和间接提示注入；后者尤其常见于 **RAG**、网页检索和 **Agent** 工具调用场景。
- **Grounding**（事实锚定 / 检索锚定）：将模型回答或决策限制在可验证的权威数据、检索证据、业务规则和允许工具结果范围内的机制。**Grounding** 可提升输出的可追溯性和事实一致性，但不等同于自动保证内容真实、完整或适用于当前任务。
- **SLO**（**Service Level Objective**，服务等级目标）：对服务或数据产品应达到的可度量目标，例如数据新鲜度、更新延迟、质量规则通过率、检索成功率、P95 延迟和可用性。**SLO** 用于指导日常运行、告警阈值和改进优先级。
- **SLA**（**Service Level Agreement**，服务等级协议）：数据产品或服务提供方与消费者之间就服务范围、可用性、更新频率、允许延迟、质量底线、支持责任和违约处置达成的正式承诺。**SLA** 通常以一个或多个 **SLO** 作为可验证指标。
- **AI Evaluation**（AI 评估）：使用代表性任务、评估集、规则、人工审核或在线运行信号，对模型、**RAG** 链路或 **Agent** 在准确性、相关性、忠实性、安全性、成本、时延和任务完成效果等方面进行系统验证的过程。高影响场景应定义评估基线、通过阈值、责任人和复评机制。
- **Evaluation Dataset**（评估集）：用于验证模型、**RAG** 或 **Agent** 是否满足预期质量和安全要求的一组代表性输入、上下文、预期结果、评分规则和边界案例。评估集应版本化管理，并覆盖正常、异常、敏感、对抗和回归场景。
- **Trace**（执行追踪）：记录一次 AI 或 **Agent** 任务从输入、身份、**Context** 组装、检索、模型调用、工具调用、策略决定到业务结果的关联事件链。**Trace**

用于性能诊断、质量分析、事故调查、责任追溯和评估复盘，不能替代权威业务记录或安全审计日志。

- **Runtime Policy Enforcement**（运行时策略执行）：在数据访问、检索、模型调用或工具调用发生时，依据调用主体、代表用户、任务目的、数据分类、风险等级和操作范围动态允许、限制、脱敏、要求审批或拒绝请求的控制机制。它区别于仅在设计或上线阶段完成的一次性审批。
- **Least Privilege**（最小权限）：仅授予用户、应用、模型服务、**Agent** 和工具完成当前任务所必需的最小数据访问范围、功能权限、调用额度和有效期，并定期复核、收回不再需要的权限。**OWASP** 将“过度功能、过度权限、过度自主性”视为 **Agent** 产生破坏性行为的主要根因。
- **Tool Calling**（工具调用）：模型或 **Agent** 按定义的输入输出 **Schema** 调用 **API**、数据库查询、工作流、企业应用或其他外部能力的过程。工具调用应具备 **Owner**、用途、权限范围、幂等性、调用额度、错误处理、审计记录和补偿语义。
- **Human-in-the-Loop**（人在回路 / 人工复核）：在 **AI** 或 **Agent** 生成建议、调用工具或触发高影响业务动作时，由具备相应授权和业务责任的人员对关键输入、决策或结果进行审核、批准、拒绝、修正或升级处理的控制机制。它不是简单“人工点确认”，还应明确何时触发、由谁处理、可见证据及超时处置。
- **Data Classification and Handling**（数据分类分级与处理）：按照敏感性、业务影响、监管要求和使用目的，对数据及其派生资产进行分类、标记和施加相应的访问控制、脱敏、加密、保留、审计和删除策略的治理机制。治理对象应覆盖源数据、文档、**Chunk**、**Embedding**、索引、**Context**、**Memory** 与生成输出。

附录 C：参考文档

[1] National Institute of Standards and Technology. Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1)

[R]. 2024-07. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>

- [2] OWASP Foundation. LLM04: Data and Model Poisoning [EB/OL]. 2025.
<https://genai.owasp.org/llmrisk/llm04-data-and-model-poisoning/>
- [3] OWASP Foundation. LLM08: Vector and Embedding Weaknesses [EB/OL].
2025. <https://genai.owasp.org/llmrisk/llm08-vector-and-embedding-weaknesses/>
- [4] Amazon Web Services. Enforce fine-grained access control on data lake tables using AWS Glue 5.0 integrated with AWS Lake Formation [Blog]. 2024-12-04.
<https://aws.amazon.com/blogs/big-data/enforce-fine-grained-access-control-on-data-lake-tables-using-aws-glue-5-0-integrated-with-aws-lake-formation/>
- [5] Amazon Web Services. Navigating architectural choices for a lakehouse using Amazon SageMaker [Blog]. 2026-01-12. <https://aws.amazon.com/blogs/big-data/navigating-architectural-choices-for-a-lakehouse-using-amazon-sagemaker/>
- [6] Amazon Web Services. AWS analytics at re:Invent 2025: Unifying Data, AI, and governance at scale [Blog]. 2026-01-07. <https://aws.amazon.com/blogs/big-data/aws-analytics-at-reinvent-2025-unifying-data-ai-and-governance-at-scale/>
- [7] Amazon Web Services. Multi-cloud lakehouse architecture on AWS for Agentic AI, Part 1: Architecture and best practices [Blog]. 2026-07-13.
<https://aws.amazon.com/blogs/big-data/multi-cloud-lakehouse-architecture-on-aws-for-agentic-ai-part-1-architecture-and-best-practices/>
- [8] Amazon Web Services. Amazon Bedrock Knowledge Bases now simplifies asking questions on a single document [Blog]. 2024-04-26.
<https://aws.amazon.com/blogs/machine-learning/amazon-bedrock-knowledge-bases-now-simplifies-asking-questions-on-a-single-document/>
- [9] Amazon Web Services. Elevate RAG for numerical analysis using Amazon Bedrock Knowledge Bases [Blog]. 2024-09-25.
<https://aws.amazon.com/blogs/machine-learning/elevate-rag-for-numerical-analysis-using-amazon-bedrock-knowledge-bases/>
- [10] Amazon Web Services. Multi-tenant RAG with Amazon Bedrock Knowledge Bases [Blog]. 2024-12-16. <https://aws.amazon.com/blogs/machine-learning/multi-tenant-rag-with-amazon-bedrock-knowledge-bases/>

[11] Amazon Web Services. Secure multi-tenant RAG with Amazon Bedrock and verified permissions [Blog]. 2026-06-22.

<https://aws.amazon.com/blogs/architecture/secure-multi-tenant-rag-with-amazon-bedrock-and-verified-permissions/>

[12] Amazon Web Services. Data retention - Amazon Bedrock [Documentation]. 2026-06-07. <https://docs.aws.amazon.com/bedrock/latest/userguide/data-retention.html>

[13] Amazon Web Services. Amazon Bedrock AgentCore [Prescriptive Guidance]. n.d. <https://docs.aws.amazon.com/prescriptive-guidance/latest/agentcore-frameworks/amazon-bedrock-agent-core.html>

[14] Amazon Web Services. Deploy MCP servers in AgentCore Runtime [Documentation]. n.d. <https://docs.aws.amazon.com/bedrock-agentcore/latest/devguide/runtime-mcp.html>

[15] Amazon Web Services. Amazon Bedrock AgentCore [Documentation]. n.d. <https://docs.aws.amazon.com/sap/latest/general/rise-agenticaibedrock-agentcore.html>

[16] Amazon Web Services. Operationalize generative AI workloads and scale to hundreds of use cases with Amazon Bedrock – Part 1: GenAIOps [Blog]. 2025-12-15. <https://aws.amazon.com/blogs/machine-learning/operationalize-generative-ai-workloads-and-scale-to-hundreds-of-use-cases-with-amazon-bedrock-part-1-genaiops/>

[17] Amazon Web Services. Accelerating generative AI applications with a platform engineering approach [Blog]. 2025-11-18. <https://aws.amazon.com/blogs/machine-learning/accelerating-generative-ai-applications-with-a-platform-engineering-approach/>

[18] Amazon Web Services. Design cost-efficient initialization through AgentCore Observability [AWS Well-Architected Agentic AI Lens]. n.d. <https://docs.aws.amazon.com/wellarchitected/latest/agentic-ai-lens/agentcost06-bp03.html>

- [19] Amazon Web Services. Top announcements of AWS re:Invent 2025 [Blog]. 2025-11-30. <https://aws.amazon.com/blogs/aws/top-announcements-of-aws-reinvent-2025/>
- [20] Amazon Web Services. Building agentic AI patterns with Amazon Bedrock and SQL Server 2025 on Amazon RDS [Blog]. 2026-07-20. <https://aws.amazon.com/blogs/database/building-agentic-ai-patterns-with-amazon-bedrock-and-sql-server-2025-on-amazon-rds/>
- [21] Amazon Web Services. Securing generative AI: Applying relevant security controls [Blog]. 2024-03-19. <https://aws.amazon.com/blogs/security/securing-generative-ai-applying-relevant-security-controls/>
- [22] Amazon Web Services. Securing generative AI: data, compliance, and privacy considerations [Blog]. 2024-03-27. <https://aws.amazon.com/blogs/security/securing-generative-ai-data-compliance-and-privacy-considerations/>
- [23] Amazon Web Services. A secure approach to generative AI with AWS [Blog]. 2024-04-16. <https://aws.amazon.com/blogs/machine-learning/a-secure-approach-to-generative-ai-with-aws/>
- [24] Amazon Web Services. Data governance in the age of generative AI [Blog]. 2024-02-29. <https://aws.amazon.com/blogs/big-data/data-governance-in-the-age-of-generative-ai/>
- [25] Amazon Web Services. Data Governance in the Age of Generative AI [Blog]. 2024-11-05. <https://aws.amazon.com/blogs/enterprise-strategy/data-governance-in-the-age-of-generative-ai/>