

AGENTIC AI READINESS

Agentic AI Architecture & Readiness Checklist

Assess the building blocks, controls, and capabilities needed for effective AI agents

1. What Is Agentic AI?

Agentic AI describes AI systems that go beyond answering a single question or generating a single piece of content. Instead, an agent is given a goal, and it plans a sequence of steps, uses tools to gather information or take action, observes the results, and adjusts its approach until the goal is met with as much or as little human involvement as you design in.

How Agentic AI Differs from Traditional Generative AI

Traditional generative AI is typically reactive and single-turn: you provide a prompt, and the model returns a response. Each interaction is largely self-contained.

Agentic AI is goal-driven and multi-step. Rather than producing one output, the system works through an extended task, deciding what to do next based on what has already happened, and it often takes real actions in external systems rather than just producing text.

The Core Agent Loop

Most agentic systems operate on a repeating cycle. Understanding this loop is the foundation for everything else in this checklist, because each stage introduces its own set of architectural and governance requirements:

- ☐ Planning — the agent breaks a goal down into a sequence of steps or sub-tasks.
- ☐ Tool Use — the agent selects and calls the tools, APIs, or data sources it needs.

- ☐ Action — the agent executes a step, which may include writing data, sending a message, or triggering a workflow.
- ☐ Observation — the agent reviews the outcome of the action against what it expected.
- ☐ Adjustment — the agent revises its plan based on the observation and continues, escalates, or stops.

Levels of Autonomy

Not every agentic system needs — or should have — the same degree of independence. Autonomy exists on a spectrum, and choosing the right level for a given use case is one of the most important design decisions an organization will make:

- ☐ Recommendation only — the agent proposes an action but a human decides.
- ☐ Human approval before action — the agent prepares the action, but nothing executes without sign-off.
- ☐ Automated execution for predefined tasks — the agent acts independently within a narrow, well-tested scope.
- ☐ Higher autonomy within boundaries — the agent makes a broader range of decisions, but is still constrained by policies, permissions, and monitoring.

Example — Two systems, two levels of autonomy

A customer-support agent that drafts refund responses for a human to approve is operating at the 'human approval before action' level. A logistics agent that automatically reroutes a shipment when a carrier delay is detected — within pre-approved carriers and cost thresholds — is operating at 'automated execution for predefined tasks.' Both are agentic; they simply sit at different points on the autonomy spectrum.

2. Agentic AI Architecture: Key Building Blocks

Every agentic AI system is built from the same core set of components, regardless of the specific use case. Before evaluating oversight, security, or governance, it's worth confirming that each of these building blocks is deliberately designed — rather than assumed or bolted on after the fact.

Component	Key Questions	Ready?
Foundation / Reasoning Model	Do we have an appropriate model for the task?	☐

Tools & APIs	Can the agent securely interact with required systems?	☐
Memory	Does the system need short- or long-term memory?	☐
Planning & Orchestration	How will tasks be broken down and coordinated?	☐
Data Sources	Does the agent have access to reliable, relevant data?	☐
Retrieval	Can the agent retrieve the right context when needed?	☐
Security & Permissions	Are access and action permissions clearly defined?	☐
Monitoring & Evaluation	Can agent behavior and outcomes be monitored?	☐
Human Approval	Which actions require human review or approval?	☐

What Each Component Actually Covers

Foundation / Reasoning Model — the underlying language model that powers the agent's reasoning and decision-making. A model well-suited to conversational support may not be well-suited to multi-step financial analysis, so match model capability to task complexity.

Tools & APIs — the interfaces the agent uses to act in the world, from internal databases to third-party services. Every tool the agent can call is also a capability it could misuse, so each one deserves deliberate scoping.

Memory — short-term memory lets an agent track context within a single task; long-term memory lets it recall information across sessions. Not every agent needs both — over-provisioning memory can introduce unnecessary privacy and consistency risk.

Planning & Orchestration — the logic that decomposes a goal into steps and sequences tool calls. This is where multi-agent systems also coordinate handoffs between specialized sub-agents.

Data Sources — the systems of record the agent draws on. Stale, incomplete, or ungoverned data sources are one of the most common causes of agent failure in production.

Retrieval — the mechanism (such as a vector search or knowledge base lookup) that surfaces relevant context at the right moment, rather than relying solely on what's in the prompt.

Security & Permissions — the guardrails that determine what the agent is allowed to access and do, independent of what it's technically capable of.

Monitoring & Evaluation — the visibility layer that lets your team see what the agent did, why, and whether the outcome was correct.

Human Approval — the checkpoints where a person reviews or authorizes an action before, during, or after execution.

Example — Architecture in practice

An internal IT-helpdesk agent might use a mid-sized reasoning model, short-term memory only, three narrowly-scoped tools (ticketing system, password-reset API, knowledge base), and require human approval solely for account-lockout overrides. Each component was sized to the task rather than maximized by default.

3. Autonomy & Human Oversight Checklist

Agentic AI does not automatically mean full autonomy. One of the most important governance decisions an organization makes is defining exactly where a human must stay in the loop — and ensuring that boundary is enforced technically, not just documented in a policy.

Assess

- ☐ Tasks appropriate for autonomous execution are clearly identified
- ☐ Human approval requirements are defined
- ☐ Sensitive actions require explicit authorization
- ☐ Escalation rules are documented
- ☐ The agent operates within defined boundaries
- ☐ Users can intervene when necessary

- ☐ High-impact decisions remain subject to appropriate human oversight

Autonomy Level

Which level best describes your planned system?

- ☐ Recommendation only
- ☐ Human approval before action
- ☐ Automated execution for predefined tasks
- ☐ Higher autonomy within predefined boundaries

Example — Defining a boundary, not just a policy

A finance agent might be authorized to auto-approve expense reimbursements under \$200, but any request above that threshold — or any request flagged as unusual — routes to a human approver automatically, with no ability for the agent to override the threshold itself.

4. Tools, APIs & External System

Readiness

An agent is only as safe and effective as the tools it's allowed to use. This section assesses whether the systems your agent will interact with are actually ready to be touched by an autonomous process.

- ☐ Required tools have been identified
- ☐ APIs are available and documented
- ☐ Tool permissions are clearly defined
- ☐ Agents can access only necessary systems
- ☐ Tool failures have defined fallback procedures
- ☐ External actions can be logged
- ☐ High-risk tools require additional authorization

Example — Fallback procedures matter

If a shipping-carrier API times out mid-task, a well-designed agent should pause and flag the order for manual handling rather than silently retrying indefinitely or guessing at a status — a fallback procedure defined in advance, not improvised at runtime.

5. Data & Context Readiness

An agent's plans and actions are only as good as the information behind them. This section assesses whether the agent has access to data that is accurate, current, appropriately governed, and scoped to what the task actually requires.

- ☐ Required data sources have been identified

- ☐ Data is accessible to the agent
- ☐ Data quality has been assessed
- ☐ Sensitive information is appropriately protected
- ☐ Retrieval mechanisms are defined where needed
- ☐ Context provided to the agent is relevant and current
- ☐ Data access follows organizational permissions

Example — Why data quality gets checked first

A sales agent that recommends next steps based on a CRM field that's six months stale will confidently produce a plausible-sounding — and wrong — recommendation. Confirming data freshness and quality before deployment is far cheaper than diagnosing it after a bad outcome.

6. Security & Permission Controls

Because agentic systems can take real actions, security controls need to constrain what the agent is allowed to do — not just what it's technically capable of doing. This section connects directly to the reality that capability and permission are two different things.

- ☐ Role-based access controls are implemented
- ☐ Agent permissions are explicitly defined

- ☐ Least-privilege access is applied
- ☐ Sensitive actions require additional controls
- ☐ Tool/API access is authenticated
- ☐ Agent activity is logged
- ☐ Security risks are regularly assessed

Example — Least privilege in practice

An agent that only needs to read order status should be granted read-only access to the orders table — not the broader database credential a human engineer might use — so that even a flawed plan can't cause unintended write actions.

7. Monitoring, Evaluation & Governance

An agent shouldn't simply be deployed and forgotten. Sustained oversight is what separates a well-run agentic system from one that quietly drifts out of alignment with its intended purpose.

Monitoring

- ☐ Agent actions are logged
- ☐ Failures can be identified

- ☐ Unexpected behavior can be detected

Evaluation

- ☐ Agent performance metrics are defined
- ☐ Outputs are evaluated for accuracy
- ☐ Tool usage is monitored
- ☐ Escalations are tracked

Governance

- ☐ Ownership is clearly assigned
- ☐ Operating boundaries are documented
- ☐ Approval processes are defined
- ☐ Review processes are established

Example — Monitoring catches drift early

A content-moderation agent that starts flagging an unusual spike in a single category might be reacting to a change in user behavior — or to a bug in its own reasoning. Logging and regular evaluation are what make it possible to tell the difference before it becomes a larger problem.

8. MCP & Agent Connectivity Readiness

The Model Context Protocol (MCP) is an increasingly common way for agents to connect to external tools and data sources through a standardized interface. It's a useful piece of the connectivity picture, but — as covered in the broader discussion of agentic AI myths — it is one implementation detail among many, not a requirement or a guarantee of readiness on its own. Adopting MCP doesn't substitute for the security, permission, and monitoring practices covered elsewhere in this checklist.

- ☐ Required external tools/resources are identified
- ☐ MCP is relevant to the intended use case
- ☐ MCP servers/resources are appropriately secured
- ☐ Authorization requirements are understood
- ☐ Tool and resource access is controlled
- ☐ MCP-related activity can be monitored

Example — MCP is a connector, not a control

Connecting an agent to a file-storage system via MCP still requires the same least-privilege thinking as any other integration: scope the connection to the specific folders the agent needs, not the entire drive, and log what it reads and writes.

9. Agentic AI Readiness Scorecard

Use this scorecard to translate the checklists above into a single readiness snapshot.

Score each area using the scale below, then total your results.

☐ 0 = Not Started

☐ 1 = Partially Ready

☐ 2 = Ready

Area	Score
Architecture	___ / 2
Tools & APIs	___ / 2
Data & Context	___ / 2
Autonomy & Human Oversight	___ / 2
Security & Permissions	___ / 2
Monitoring & Evaluation	___ / 2
Governance	___ / 2

MCP / Connectivity	___ / 2
Total	___ / 16

Interpreting Your Score

0–5: Early Stage — several foundational capabilities still need to be defined. This is a normal starting point; use the sections above to prioritize what to build first.

6–11: Developing — core capabilities are emerging, but some areas require further preparation before broader deployment.

12–16: Foundation in Place — most foundational capabilities have been identified and established.

This scorecard is intended as a self-assessment to guide planning and prioritization — not as a claim that a given score means an organization is objectively “ready” for production deployment. Readiness ultimately depends on the specific risk profile of each use case.

10. Final Pre-Deployment Checklist

Before moving an agentic AI system toward deployment, confirm each of the following:

- ☐ Business objective clearly defined
- ☐ Agent responsibilities clearly defined

- ☐ Tools and permissions documented
- ☐ Data sources validated
- ☐ Human oversight established
- ☐ Security controls implemented
- ☐ Escalation rules defined
- ☐ Monitoring established
- ☐ Evaluation criteria established
- ☐ Governance ownership assigned
- ☐ Failure and fallback procedures documented
- ☐ Deployment boundaries clearly defined

AGENTIC AI PROFESSIONAL CERTIFICATION

THE AGENTIC AI PROFESSIONAL CERTIFICATION PROGRAM IS DESIGNED TO EQUIP PROFESSIONALS WITH THE KNOWLEDGE TO BUILD, MANAGE, AND OPTIMIZE AUTONOMOUS AI AGENTS THAT CAN REASON, PLAN, AND EXECUTE COMPLEX TASKS.



ABOUT GSDC CERTIFICATION



EBOOK

Extensive and exclusive Ebook created by world's experts to help you with understanding core concepts.



EBOOK

Extensive and exclusive Ebook created by world's experts to help you with understanding core concepts.



CREATED BY EXPERTS

GSDC certifications are created and authored by world's leading experts in the field.



LEARNING MATERIALS

Get access to learning materials such as videos, ebooks, templates, and practice exams, which will help you clear the certification exam.

LEARNING OBJECTIVE

- Identify real-world applications of Agentic AI
- Learn from practical use case studies
- Explore ethical and compliance aspects of Agentic AI

Enroll now with the code **LEARN20** To avail **20%** discount

Enroll Now



www.gsdccouncil.org