

PREPKLOUD FIELD GUIDE SERIES

Developing AI Apps and Agents on Azure (AI-103) Field Guide

Microsoft Foundry, generative AI, RAG, agents, multimodal vision, text and speech, and information extraction for the Azure AI engineer.

First edition · 2026 · prepcloud.com

About this guide

Exam AI-103, *Developing AI Apps and Agents on Azure*, is for Azure AI engineers who build, manage, and deploy agents and AI solutions with Microsoft Foundry. Microsoft expects candidates to have experience developing apps with Python and to be familiar with general AI, generative AI, and Azure services. The role works with business stakeholders, solution architects, data scientists, DevOps engineers, and cloud security engineers.

Exam facts, as published by Microsoft when this guide was written: the skills outline is dated 16 April 2026; a score of 700 or greater is required to pass; associate certifications expire annually and can be renewed with a free online assessment on Microsoft Learn; most questions cover generally available features, though commonly used preview features may appear. Always check the [official AI-103 study guide](#) and exam page for current length, price in your country, and languages before you book.

Skills measured	Weight	Chapters
1. Plan and manage an Azure AI solution	25–30%	2–5
2. Implement generative AI and agentic solutions	30–35%	6–8
3. Implement computer vision solutions	10–15%	9–10
4. Implement text analysis solutions	10–15%	11
5. Implement information extraction solutions	10–15%	12–13

Independent publication: Microsoft, Azure, Microsoft Foundry, and related names are trademarks of the Microsoft group of companies. PrepKloud is not affiliated with or endorsed by Microsoft. All questions are original. They are not recalled, leaked, or copied exam items. Azure AI product names change often; this guide follows the wording in the current study guide.

© 2026 PrepKloud. For personal study; do not resell or redistribute.

Contents

1. How the exam works and how to prepare

2. Choosing Foundry services and models

3. Setting up AI solutions in Foundry

4. Managing, monitoring, and securing AI systems

5. Responsible AI for generative and agentic systems

6. Building generative applications and RAG

7. Building agents with Foundry

8. Optimizing and operationalizing generative AI

9. Image and video generation and multimodal understanding

10. Responsible AI for multimodal content

11. Text analysis, translation, and speech

12. Retrieval and grounding pipelines

13. Extracting content from documents

14. Original practice questions and explanations

15. Six-week plan, cheat sheet, and glossary

How the exam works and how to prepare

AI-103 is an engineering exam. It expects you to know how to choose, configure, secure, evaluate, and operate AI solutions, not just what AI can do. Generative AI and agents carry the largest share, but planning and management is close behind, so security, cost, monitoring, and responsible AI are not side topics.

What the questions reward

- **Picking the right Foundry capability.** Model, retrieval method, agent tool, or Foundry Tool for the stated task.
- **Secure-by-default configuration.** Managed identity, keyless credentials, private networking, and least-privilege roles.
- **Grounding and evaluation.** RAG, hybrid search, evaluators for groundedness, relevance, and safety.
- **Operational thinking.** Quotas, rate limits, tracing, token analytics, and cost.
- **Safe agents.** Tool-access controls, approval flows, and human oversight.

Exam mindset

When several answers work, prefer the one that is managed, secure without secrets, grounded in the organization's data, observable, and has the least custom code. Keys in code, public endpoints without need, and unmonitored autonomous actions are usually wrong.

Preparing

- Work through the official Microsoft Learn path for AI-103 alongside this guide.
- Build small projects in a sandbox subscription: a RAG chat app, an agent with one tool and one knowledge source, a document extraction pipeline, and a speech or vision workflow.
- Use the exam sandbox to get used to the question interface.
- Delete resources after each lab to control cost.

Recall

- Which skill area carries the most weight?
- What score is needed to pass?

Choosing Foundry services and models

Microsoft Foundry is Azure's platform for building AI apps and agents. It brings together a model catalogue, projects, deployments, agents, evaluation, safety, and monitoring. **Foundry Tools** are the prebuilt AI capabilities, such as vision, language, speech, translation, and Content Understanding, that you can call from apps and agents.

Choosing a model

Model type	Good for	Tradeoff
Large language model (LLM)	Complex reasoning, long answers, broad knowledge.	Higher cost and latency.
Small language model (SLM)	Classification, extraction, simple chat, edge or high-volume use.	Less capable on hard reasoning.
Reasoning model	Multi-step planning, maths, code analysis.	Slower; more tokens.
Multimodal model	Text plus images, audio, or video input.	Cost of processing media.
Embeddings model	Vector search and semantic similarity.	Produces vectors, not text.

Foundry Tool	A specific task such as OCR, translation, or speech to text.	Less flexible, but predictable and often cheaper.
--------------	--	---

Choosing services for common tasks

Need	Typical choice
Grounded answers over enterprise content	RAG with Azure AI Search as the retrieval index.
Vector similarity	Embeddings model plus a vector index in Azure AI Search.
Multi-step actions	An agent in Foundry Agent Service with tools.
Structured data from documents, images, audio, or video	Azure Content Understanding in Foundry Tools.
Translation	Azure Translator in Foundry Tools, or an LLM translation flow.
Speech in or out	Azure Speech in Foundry Tools.

Memory, tools, and knowledge for agents

Decide what the agent must remember (conversation state, user preferences), what it must look up (knowledge sources such as Azure AI Search or files), and what it must do (tools such as APIs, functions, or MCP servers). Choose the smallest set of each that meets the goal.

Scenario

A support app must classify 2 million short emails per day into ten categories at low cost. A small language model or a language Foundry Tool is a better first choice than a large reasoning model.

Recall

- When would you choose an SLM over an LLM?
- Which service commonly provides the retrieval index for RAG?

Setting up AI solutions in Foundry

Resources and projects

A Foundry resource provides the AI capabilities and model deployments. **Projects** organize work for a team or application: agents, evaluations, files, connections, and access. Connections link a project to other services such as Azure AI Search, storage, or Application Insights.

Deployment options

Option	Use when
Standard (pay-as-you-go)	Variable traffic; pay per token.
Global or data zone deployment	You want higher availability and throughput, with processing in any region or within a defined geography.
Provisioned throughput	Predictable high volume that needs consistent latency.
Batch	Large, non-urgent jobs at lower cost.

Data residency requirements affect this choice: if data must be processed in one geography, avoid global processing and choose a regional or data zone option.

Configuring model and agent deployments

- Select model version and upgrade policy so updates do not surprise you.
- Set capacity (tokens per minute) and rate limits.
- Attach content filters and guardrails.
- For agents, configure instructions, model, tools, and knowledge.

CI/CD for Foundry projects

Treat prompts, agent definitions, evaluation datasets, and infrastructure as code. Store them in source control, deploy with Bicep or Terraform and GitHub Actions or Azure Pipelines, and run automated evaluations as a quality gate before promoting to production.

Connecting from code

Apps connect to a Foundry project through its endpoint using the Foundry SDKs. Prefer Microsoft Entra ID authentication, for example `DefaultAzureCredential` from the `azure-identity` library, over API keys.

Scenario

A bank must keep all customer data processing inside the EU. It should choose an EU data zone or regional deployment rather than a global deployment type.

Recall

- Which deployment option suits steady high volume with consistent latency?
- What should gate a production release of a prompt change?

Managing, monitoring, and securing AI systems

Quotas, scaling, and cost

- Quotas limit tokens per minute per model and region; request increases before launch.
- Handle 429 rate-limit responses with retries and exponential backoff.
- Control cost by choosing smaller models where possible, capping output tokens, caching repeated answers, and using batch for non-urgent work.
- Tag resources and use Microsoft Cost Management budgets and alerts.

What to monitor

Signal	Why it matters
Latency and errors	User experience and reliability.
Token usage	Cost and quota consumption.
Groundedness and relevance	Whether answers stay faithful to sources.
Safety events	Blocked or flagged content trends.

Drift	Changes in input patterns or model quality over time.
Ingestion quality and index health	Retrieval only works if the index is complete and current.

Azure Monitor and Application Insights collect metrics, logs, and traces. Foundry adds tracing and evaluation views for apps and agents.

Security

- **Managed identity** for app-to-service access, so no secrets are stored in code.
- **Keyless credentials** with Microsoft Entra ID; disable local key authentication where possible.
- **Role-based access control** with least privilege, for example separate roles for users who call models and those who manage deployments.
- **Private networking** with private endpoints and disabled public network access.
- **Azure Key Vault** for any secrets you cannot avoid.

Scenario

A security review finds an API key hard-coded in a web app that calls a model deployment. The best fix is to enable a managed identity for the web app, grant it the minimum role on the Foundry resource, and disable key-based authentication.

Recall

- How should an app respond to HTTP 429 errors?
- Name three signals to monitor for a RAG app.

Responsible AI for generative and agentic systems

Microsoft's responsible AI principles

Fairness, reliability and safety, privacy and security, inclusiveness, transparency, and accountability.

Safety controls

Control	Purpose
Content filters	Detect and block harmful categories such as hate, sexual, violence, and self-harm at configurable severity levels.
Prompt shields	Detect direct jailbreak attempts and indirect prompt injection in documents or tool results.
Groundedness detection	Flag answers not supported by the provided sources.
Protected material detection	Flag known copyrighted text or code in outputs.
Blocklists	Block specific terms.

Evaluation and instrumentation

Use built-in evaluators for quality (groundedness, relevance, coherence, fluency) and safety (harmful content, jailbreak

susceptibility). Run evaluations before release and continuously in production. Add explanation tooling where decisions need to be justified.

Auditing and governing agents

- Record trace logs of prompts, retrieved sources, tool calls, and outputs.
- Attach provenance metadata so you know which data and model version produced a result.
- Require approval workflows for high-impact actions.
- Constrain agent behaviour with explicit instructions, limited tool access, and oversight modes that range from human-in-the-loop to autonomous with monitoring.

Scenario

An agent reads supplier emails and can issue refunds. An email contains hidden text telling the agent to refund all orders. Prompt shields for indirect attacks, a refund approval step, and a per-transaction limit together reduce this risk.

Recall

- What is indirect prompt injection?
- Which evaluator checks that an answer is supported by its sources?

Building generative applications and RAG

Deploying and consuming models

Deploy LLMs, small models, code models, and multimodal models from the Foundry model catalogue, then call them from applications using the Foundry SDKs or OpenAI-compatible endpoints. Configure the app with the project endpoint and Entra ID authentication.

Retrieval-augmented generation

1. Ingest and chunk content.
2. Create embeddings and index text and vectors in Azure AI Search.
3. At query time, retrieve with hybrid search (keyword plus vector) and optional semantic ranking.
4. Pass the top results to the model with instructions to answer only from the provided sources and to cite them.
5. Evaluate groundedness and relevance.

Workflows and multistep reasoning

Some tasks need several steps: classify a request, retrieve data, call a tool, then draft a reply. Design these as explicit workflows or

tool-augmented flows so each step can be tested and traced.

Evaluating apps

What to check	How
Fabrications	Groundedness evaluators against retrieved sources.
Relevance	Relevance evaluators and human review of samples.
Quality	Coherence and fluency evaluators; task-specific metrics.
Safety	Safety evaluators and adversarial test sets.

Scenario

Users complain that a policy assistant invents clauses. Add retrieval from the policy index, instruct the model to answer only from retrieved passages and cite them, lower the temperature, and track groundedness scores.

Recall

- What does hybrid search combine?
- Which evaluation detects fabrications?

Building agents with Foundry

Designing an agent

- **Role and goal:** what the agent is for and what "done" means.
- **Instructions:** scope, tone, and what it must not do.
- **Conversation tracking:** how state and memory are kept across turns.
- **Tool schemas:** clear names, descriptions, and typed parameters so the model calls tools correctly.

Tools and knowledge

Tool type	Example
Retrieval / knowledge	Azure AI Search index, uploaded files.
Function calling	Custom functions that look up an order or book a meeting.
APIs	OpenAPI-described REST services.
MCP servers	Tools exposed through the Model Context Protocol.
Content understanding	Extract structure from files the user uploads.
Code execution	Run code to analyse data or create charts.

Multi-agent solutions

Split complex work across specialized agents, for example a triage agent that routes to billing, technical, and returns agents. An orchestrator coordinates them. Multi-agent designs add flexibility but also latency, cost, and more places for errors, so use them only when one agent with tools is not enough.

Autonomy with safeguards

- Require human approval for irreversible or expensive actions.
- Set budgets, step limits, and timeouts.
- Scope credentials so each tool can only do what it needs.
- Monitor agents with traces, evaluate their behaviour, and perform error analysis on failures.

Scenario

An agent sometimes calls the wrong tool because two tools have vague names, "process" and "handle". Rename them with clear, specific names and descriptions, and define typed parameters in the tool schema.

Recall

- List four parts of an agent design.
- When is a multi-agent design justified?

Optimizing and operationalizing generative AI

Tuning generation behaviour

Lever	Effect
System instructions	Define role, scope, format, and refusals.
Few-shot examples	Show the exact output style.
Temperature / top-p	Lower for consistency; higher for variety.
Max tokens	Limit length and cost.
Structured outputs	Force JSON that matches a schema.

Reflection and self-critique

Improve quality by asking a model to check its own answer against sources or rules, or by using a second evaluator model. Evaluate reasoning quality with chain-of-thought evaluations where appropriate. These techniques add tokens and latency, so apply them where accuracy matters most.

Observability

- Tracing across the full request: retrieval, model calls, tool calls.
- Token analytics per user, feature, and model.

- Safety signals such as filtered requests.
- Latency breakdowns showing which step is slow.

Orchestrating models and rules

Route simple requests to a small model and complex ones to a larger model. Combine LLMs with deterministic rules engines where exact answers are required, such as pricing or eligibility checks.

Scenario

Costs have tripled. Token analytics show that 80% of requests are simple FAQs sent to a large model. Route FAQs to a smaller model or cached answers and keep the large model for complex queries.

Recall

- What are structured outputs used for?
- Give one reason to combine an LLM with a rules engine.

Image and video generation and multimodal understanding

Generation

- Generate images from text prompts and reference media.
- Generate videos from prompts and reference media.
- Edit images with inpainting, mask-based edits, and prompt-driven modifications.
- Edit generated videos and choose platform controls such as size, quality, style, and duration.

Multimodal understanding

Task	Approach
Captions, short or detailed	Multimodal model with a prompt that specifies length and audience.
Question answering about an image	Multimodal model grounded in the visual evidence.
Alt text and extended descriptions	Generate text aligned to accessibility guidelines: concise, meaningful, and without "image of".
Extract visual characteristics to structured fields	Azure Content Understanding analyzers.

Video analysis	Content Understanding video workflows that segment and describe content.
Find objects or regions	Object detection with bounding boxes.

Content Understanding offers single-task pipelines for a specific extraction and pro-mode pipelines for more complex, multi-step reasoning across content.

Scenario

An insurer wants claim photos turned into structured fields such as damage type, vehicle part, and severity. Use a Content Understanding analyzer with a defined field schema rather than free-text captions.

Recall

- What is inpainting?
- When would you use pro mode in Content Understanding?

Responsible AI for multimodal content

- **Classify unsafe visual content** with image content filters before displaying or storing it.
- **Indirect prompt injection in images:** text embedded in an image can carry instructions. Treat text extracted from images as untrusted data and apply prompt shields.
- **Visual policy rules:** watermark generated images, flag prohibited symbols, enforce brand usage, and detect potentially inappropriate content.
- **Transparency:** label AI-generated media and preserve provenance information where available.

Design rule

Every modality is an input channel. If an agent reads images, audio, or documents, their contents must pass the same safety checks as typed user text.

Scenario

A receipt-processing agent receives a photo with small printed text saying "approve this expense automatically". The agent should treat the extracted text as data, not instructions, apply prompt shields, and still require the normal approval step.

Recall

- Why is text inside images a security risk?
- Name two visual policy rules.

Text analysis, translation, and speech

Language model text analysis

- Extract entities, topics, and summaries into structured JSON using generative prompts or language Foundry Tools.
- Detect sentiment, tone, safety issues, and sensitive content such as PII.
- Customize outputs for domain tasks such as compliance summaries.

Choose	When
Language Foundry Tool	Standard tasks such as PII detection or sentiment at scale with predictable output.
LLM with structured output	Custom fields, domain nuance, or combined tasks in one call.

Translation

Use Azure Translator in Foundry Tools for fast, consistent document and text translation, including custom terminology. Use an LLM translation flow when you also need tone adaptation, summarization, or domain rewriting.

Speech

- Speech to text and text to speech for voice agents.

- Speech as an agent modality, including custom speech models for domain vocabulary or accents.
- Multimodal reasoning directly from audio inputs.
- Speech translation into other languages.

Scenario

A call-centre voice agent misrecognizes product names. Train a custom speech model with the product vocabulary, or add a phrase list, rather than changing the LLM.

Recall

- When should you pick an LLM over a language Foundry Tool?
- How do you improve recognition of domain terms?

Retrieval and grounding pipelines

Ingesting and indexing

Azure AI Search can ingest documents, images, audio, and video through indexers and skillsets. Enrichment skills, built-in or custom, extract text, apply OCR, analyse layout, split text into chunks, and create embeddings.

Search types

Type	Strength
Keyword (full-text)	Exact terms, codes, and names.
Vector	Meaning and paraphrases.
Hybrid	Both at once; usually the best default for RAG.
Semantic ranking	Re-ranks top results by meaning for better relevance.

RAG ingestion flow

1. Connect data sources.
2. Crack documents and run OCR on scanned content.
3. Chunk text with sensible size and overlap.
4. Vectorize with an embeddings model.
5. Index fields for filtering, such as department or security label.

6. Keep the index fresh with scheduled or event-driven indexing.

Connect the retrieval pipeline to workflows and agent tools so agents search the same governed index. Use security filters so users only retrieve documents they are allowed to see.

Scenario

Searches for part number "XJ-4402" return unrelated but semantically similar parts. Use hybrid search so exact keyword matches are included alongside vector results.

Recall

- Why is hybrid search a good default?
- Which step handles scanned PDFs?

Extracting content from documents

Information extraction turns documents into data. Multimodal pipelines combine OCR, layout analysis, and field extraction to handle text, tables, checkboxes, and figures.

Azure Content Understanding

- Define **analyzers** with a field schema for what you want to extract.
- Produce clean, grounded representations, such as markdown with preserved structure, for downstream RAG and agents.
- Generate structured JSON outputs for business systems.

Choosing an extraction approach

Document type	Approach
Standard forms such as invoices or receipts	Prebuilt analyzers or models.
Organization-specific forms	Custom analyzers with a defined schema.
Long, mixed documents for RAG	Layout-aware markdown output, then chunk and index.

Scenario

A RAG app gives poor answers about tables in PDF reports because plain OCR loses the table structure. Use layout analysis to produce markdown that preserves tables, then chunk and index that output.

Recall

- What is an analyzer?
- Why does markdown output help RAG?

Practice questions and explanations

Answer each question before reading the explanation. These are original scenarios, not recalled Microsoft exam questions.

Plan and manage (1–8)

1. An app classifies millions of short messages daily and must be cheap. Which model type is the best first choice?

A small language model or a language Foundry Tool.
Large models add cost and latency without need.

2. Data must be processed only within the EU. Which deployment choice fits?

An EU data zone or regional deployment, not global.
Global deployments can process in any region.

3. A web app uses a hard-coded key to call a model. What is the best fix?

Use a managed identity with a least-privilege role and disable key auth.
This removes stored secrets.

4. The app receives HTTP 429 errors at peak. What should it do immediately?

Retry with exponential backoff, then review quota and capacity.

429 means rate limit exceeded.

5. Security requires that the Foundry resource is not reachable from the internet. What should you configure?

Private endpoints and disable public network access.

Traffic then stays on private networks.

6. A prompt change must not reach production if quality drops. What should the pipeline include?

Automated evaluations as a CI/CD quality gate.

Treat prompts and evaluation datasets as code.

7. Which control detects instructions hidden in a retrieved document?

Prompt shields for indirect attacks.

Content filters target harmful categories rather than injected instructions.

8. An agent can issue refunds. What reduces the risk of misuse most directly?

An approval workflow and transaction limits on the refund tool.

Human oversight for high-impact actions.

Generative AI and agents (9–17)

9. An assistant invents policy clauses. What is the primary fix?

Ground it with RAG over the policy index and require citations.

Measure improvement with groundedness evaluators.

10. Which evaluator measures whether an answer is supported by retrieved sources?

Groundedness.

Relevance measures whether it answers the question.

11. An agent calls the wrong tool because the tool names are vague. What should you change?

Tool names, descriptions, and typed parameters in the tool schema.

Models choose tools based on their schemas.

12. A request needs billing, technical, and returns expertise with different tools. Which design fits?

A multi-agent solution with a triage or orchestrator agent.
Each specialist agent has a focused toolset.

13. An app must always return JSON matching a schema. What should you use?

Structured outputs with the schema.
This is more reliable than asking for JSON in text.

14. Costs are high because simple FAQs go to a large model. What helps most?

Route simple requests to a smaller model or cached answers.
Token analytics reveal the pattern.

15. Which technique asks the model to check its draft against the sources before answering?

Reflection or self-critique.
It improves accuracy at the cost of extra tokens.

16. A latency complaint needs root-cause analysis across retrieval, model, and tool calls. What do you need?

End-to-end tracing with latency breakdowns.
It shows which step is slow.

17. Eligibility must follow exact published rules. How should the solution work?

Use a deterministic rules engine for eligibility, with the LLM for explanation.
Exact decisions should not rely on generation.

Vision (18–21)

18. A retailer wants to change only the background of a product photo. Which technique?

Mask-based editing or inpainting.
The mask limits edits to the selected area.

19. Claim photos must become fields such as damage type and severity. Which service?

Azure Content Understanding with a custom analyzer.
Captions produce text, not structured fields.

20. A photo includes printed text instructing the agent to skip approval. How should it be handled?

Treat it as untrusted data, apply prompt shields, and keep the approval step.

Every modality is an input channel.

21. A site needs accessible descriptions for product images. What should the generated text be?

Concise, meaningful alt text, with an extended description where needed.

Follow accessibility guidelines.

Text and speech (22–25)

22. You need PII detection across millions of documents with predictable output. Which option?

The language Foundry Tool for PII detection.

It is purpose-built and consistent.

23. Documents must be translated with fixed company terminology. Which option?

Azure Translator with custom terminology.

LLM flows suit tone adaptation rather than strict terminology.

24. A voice agent misrecognizes product names. What should you do?

Use a custom speech model or phrase list.

The problem is recognition, not reasoning.

25. Callers speak Spanish but agents read English. Which capability?

Speech translation.

It converts speech in one language to text or speech in another.

Information extraction (26–30)

26. Searches for exact part numbers return similar but wrong parts. What should you enable?

Hybrid search.

Keyword matching catches exact codes.

27. Scanned PDFs return no text in the index. Which step is missing?

OCR during enrichment.

Scanned pages are images until OCR runs.

28. Users must not retrieve documents from other departments. What should you implement?

Security filters based on indexed permission fields.
Retrieval must respect access rights.

29. Tables in reports lose structure and RAG answers are wrong. What helps?

Layout analysis that outputs markdown with tables preserved.
Better input structure improves retrieval.

30. You need a reusable definition of the fields to extract from a custom form. What is this called in Content Understanding?

An analyzer.
It defines the schema and processing for a content type.

Six-week plan, cheat sheet, and glossary

Week	Focus	Output
1	Foundry resources, projects, model choice, deployment types	Deploy two models and compare cost and quality on one task.
2	Security, monitoring, quotas, responsible AI controls	A keyless app with managed identity and content filters.
3	RAG with Azure AI Search, evaluation	A grounded chat app with groundedness scores.
4	Agents, tools, multi-agent patterns, observability	An agent with one knowledge source, one tool, and an approval step.
5	Vision, multimodal, text, translation, speech	A Content Understanding analyzer and a speech workflow.
6	Extraction pipelines, review, timed practice	Two timed practice sets; book the exam when results are consistently strong.

Decision cheat sheet

- **High-volume simple task:** SLM or Foundry Tool. **Hard reasoning:** LLM or reasoning model.

- **Data residency:** regional or data zone deployment. **Steady high volume:** provisioned throughput. **Non-urgent bulk:** batch.
- **No secrets:** managed identity plus Entra ID. **No internet exposure:** private endpoints.
- **429s:** backoff and quota review.
- **Fabrications:** RAG plus groundedness evaluation. **Exact codes:** hybrid search.
- **Injected instructions:** prompt shields. **Harmful content:** content filters.
- **JSON contract:** structured outputs. **High-impact actions:** approval flow.
- **Structured fields from media:** Content Understanding analyzer.
- **Strict terminology translation:** Translator. **Domain speech terms:** custom speech.

Glossary

Agent

An AI system that uses a model, instructions, tools, and knowledge to complete tasks.

Analyzer

A Content Understanding definition of what to extract and how.

Embedding

A vector representation of meaning used for similarity search.

Foundry project

A workspace that groups agents, evaluations, connections, and access.

Groundedness

How well an answer is supported by its sources.

Hybrid search

Combined keyword and vector retrieval.

Managed identity

An Entra ID identity for an Azure resource, used without stored secrets.

MCP

Model Context Protocol, an open standard for exposing tools and data to models.

Prompt shields

Detection of jailbreaks and indirect prompt injection.

RAG

Retrieval-augmented generation.

Before test day, compare this guide with the current official study guide, because Microsoft updates skills and product names, and rework every question you missed. The best answer is usually managed, keyless, grounded, observable, and safe.