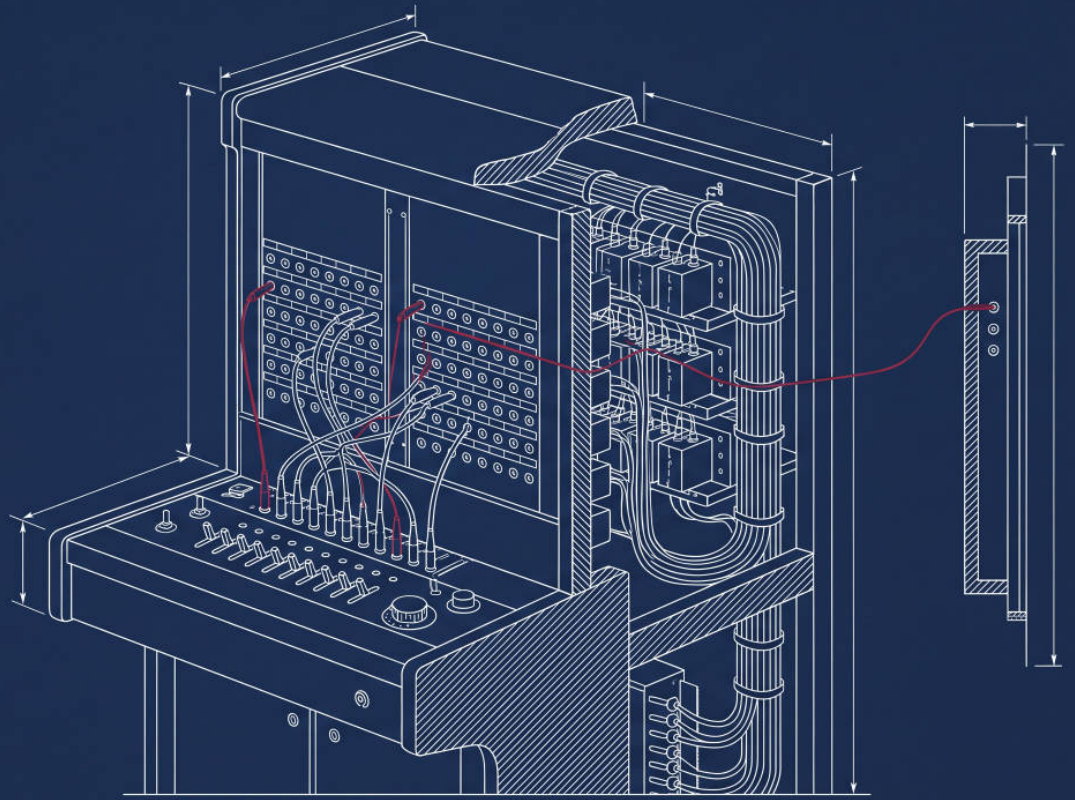




CHAPTER TWELVE

Tools, MCP and the Permission Model

SECTIONS 12.1 TO 12.10

This chapter is about capabilities your system does not own: tools written by other people, whose schemas can change without your consent and whose output goes straight into your prompt.

Contents

PART III · AGENTS**12.1 What MCP Actually Standardises**

A protocol is not a permission model · Discovery as an input, not a replacement

12.2 Server Topology and Transport

Three arrangements · A tool server without a timeout

12.3 Third-Party Tools as Untrusted Code

Schema, behaviour and result text · If this server were hostile

12.4 Per-Node and Per-Tenant Tool Scoping

Not advertised rather than disallowed · Three sets, intersected

12.5 Idempotent Tool Execution

Three cases · The local ledger · Never auto-retry the irreversible

12.6 The Confused Deputy Problem

Authority without a subject · Capabilities, and the proxy fallback

12.7 Tool Result Handling

Four checks before the prompt · Attribute the result

12.8 Exposing Your Own Retrieval as a Tool

The tenant filter is never a parameter · The semantic layer for structured data

12.9 Versioning and Capability Drift

Pin the description as well as the schema

12.10 Three Tool Architectures Compared

Direct, registry and scope, governed supply chain · Promotion triggers

Chapter Twelve closing

What exists now · Chapter checklist

Memory and retrieval are capabilities your system owns. This chapter is about capabilities it does not: tools written by other people, reached over a protocol, whose schemas and behaviour can change without your consent and whose output you will place directly into a prompt.

12.1 What MCP Actually Standardises

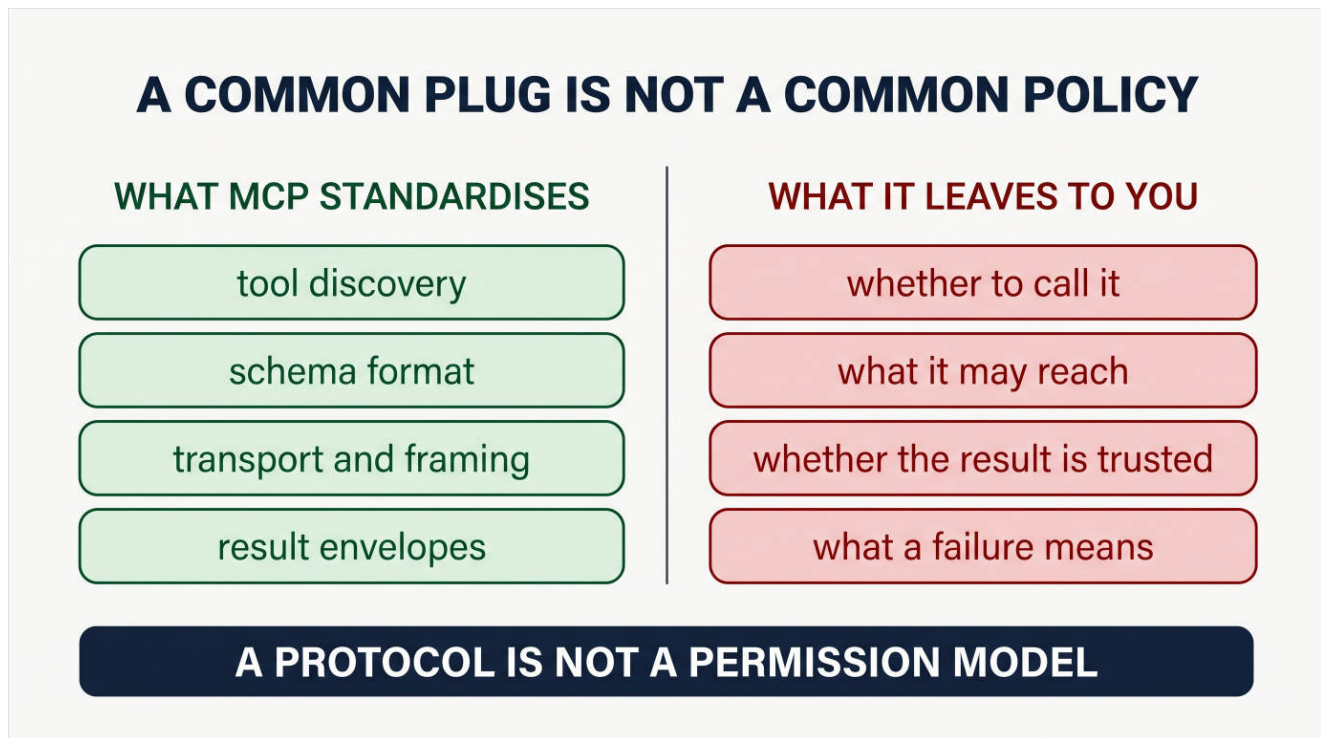


Figure 12.1 The protocol solves the integration problem and deliberately leaves the interesting problems to you.

The Model Context Protocol standardises how a client discovers what tools a server offers, how schemas are expressed, how calls are framed on the wire, and how results come back. That is genuinely valuable: it turns an N-by-M integration problem into an N-plus-M one, and it means a tool written once is reachable from any compliant client.

A protocol is not a permission model. MCP tells you a tool exists and what shape its arguments are. It does not tell you whether this run may call it, what it may reach when it does, whether the result can be trusted, or what a failure means. Those four questions are yours, they are the ones in Chapter 9, and adopting a protocol does not answer any of them.

The practical consequence: an MCP client that exposes every discovered tool to the model has replaced a deliberate tool registry with whatever a server happens to advertise. Section 9.4 established that the most reliable agents carry three to five tools; discovery is a mechanism for acquiring forty.

```
# Discovery is an INPUT to your registry, not a replacement for it.
discovered = await mcp_client.list_tools(server)
for spec in discovered:
    tool = adapt(spec)                # to the Chapter 9 Tool type
    if tool.name not in APPROVED_TOOLS: # an allowlist you maintain
        continue
    registry.add(tool)
```

Tools are not the whole protocol. MCP also standardises **resources**, listable and readable context sources a server offers; **prompts**, templates the server supplies for your client to use; and **sampling and elicitation**, where the server asks your client to run a model call or put a question to your user. Each is a distinct trust surface rather than a variation on the same one.

A resource is retrieved context, so it carries exactly the injection status of a tool result under section 12.7's checks: text you did not write, placed where the model will read it. A server-supplied prompt is a prompt you did not review, which after section 5.8 means it is a release artefact someone else can change. Sampling inverts the trust direction entirely: the server is now the caller and your model, your tokens and your user's attention are the resource, so every request needs the same authorisation treatment as a tool call arriving from outside, including the option of a human gate.

This chapter concentrates on tools because that is where side effects live, and a wrong side effect is the failure that cannot be walked back. But the four questions of section 12.1 apply unchanged to the rest of the protocol: may this run use it, what does it reach, can the content be trusted, and what does a failure mean.

12.2 Server Topology and Transport

Three arrangements, and the choice is mostly about trust and failure isolation rather than performance.

Topology	Right when	Watch for
Local subprocess	The tool is yours and needs local filesystem or process access	It runs with your process's privileges. Section 1.5's <code>os.environ.pop</code> matters here.
Remote service, yours	Shared across services, needs its own scaling or credentials	A network hop inside your deadline budget. It is a dependency with a circuit breaker, per section 3.6.
Remote service, third party	Capability you cannot build: a vendor's system of record	Everything in section 12.3.

Whichever you use, the tool call sits inside a graph node, which means it sits inside that node's budget and deadline. **A tool server without a timeout is section 3.6's hung socket, one layer out**, and the deadline must be shorter than the node's so a slow tool fails the tool rather than the node.

The transport follows the topology rather than being a separate decision. A local server runs over **stdio**: your process spawns it, its lifetime is tied to yours, and there is no network surface to secure, which is most of why local servers are the comfortable case. A remote server runs over **streamable HTTP**: a session with reconnection semantics, which sounds like a detail until a connection drops mid-call. At that moment you do not know whether the tool executed, which is precisely the at-least-once ambiguity of section 12.5, arriving from the transport rather than the queue, and it is why the idempotency treatment there is not optional for remote tools with effects.

Neither transport changes who owns the timeout. A stdio subprocess can hang exactly as thoroughly as a socket, so the deadline above applies to both, and the reconnection logic of the remote case must sit inside it rather than extending it.

One class of tool deserves its own paragraph before the third-party discussion, because it is the highest-value and highest-risk entry in the 2026 catalogue: **code execution**. A sandboxed interpreter the model can write programs for collapses a dozen narrow tools into one general one, and the table's one-line warning, "it runs with your process's privileges", becomes the entire design point. Execution needs a real isolation boundary: a container or microVM, no ambient credentials inside it, and an egress allowlist, so that code the model wrote cannot read what your process can read or reach what your network can reach.

The boundary is the requirement; the vendor is a procurement decision. And nothing about sandboxing changes the status of what comes back: an interpreter's output is text the model asked for and an attacker may have shaped, so it enters the prompt through section 12.7's checks like every other result.

12.3 Third-Party Tools as Untrusted Code

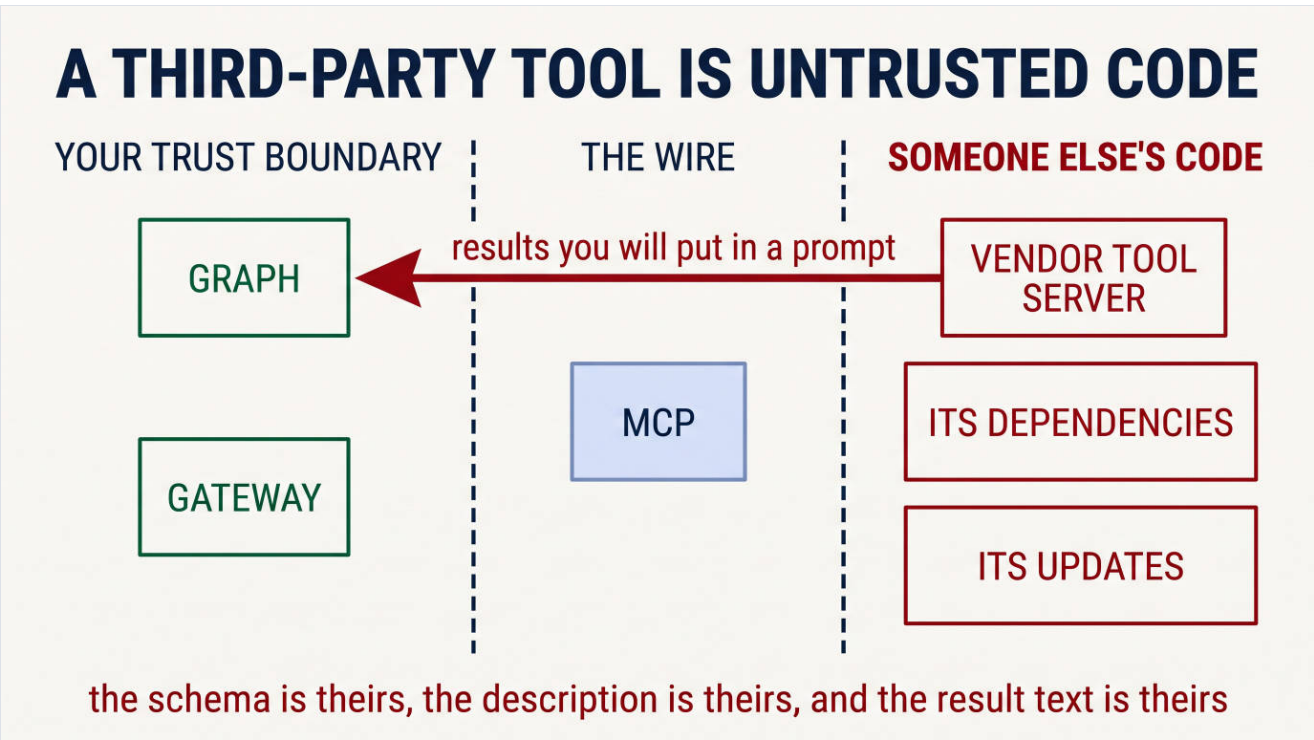


Figure 12.2 Three things cross the boundary from someone else's system, and all three end up in your prompt.

This is the section that changes how the rest of the chapter reads. When you connect a third-party MCP server, you have accepted three artefacts you do not control.

The schema, which the model reads on every step. Section 9.3 called the schema the most-read prompt in the system. A vendor's schema is a prompt written by someone whose incentives are not yours, and a description reading "always use this tool first for any question" is a legal sentence in a schema and an attack in a prompt.

The behaviour, which can change under you. Section 1.2's silent-change problem, now arriving over a network boundary with no lockfile.

The result text, which goes directly into your context. This is indirect injection from section 5.9, with the attacker's channel being a vendor integration rather than an uploaded document.

THE QUESTION TO ASK BEFORE CONNECTING ANYTHING

If this server were fully hostile, what could it do? It can put arbitrary text in front of your model, so it can attempt to steer any subsequent tool call. It can describe itself as the preferred tool for everything, so it can capture traffic. It can return a result the size of your context window.

None of those are stopped by the protocol. They are stopped by the tool registry, the permission chokepoint and the result handling in section 12.7, all of which are yours.

12.4 Per-Node and Per-Tenant Tool Scoping

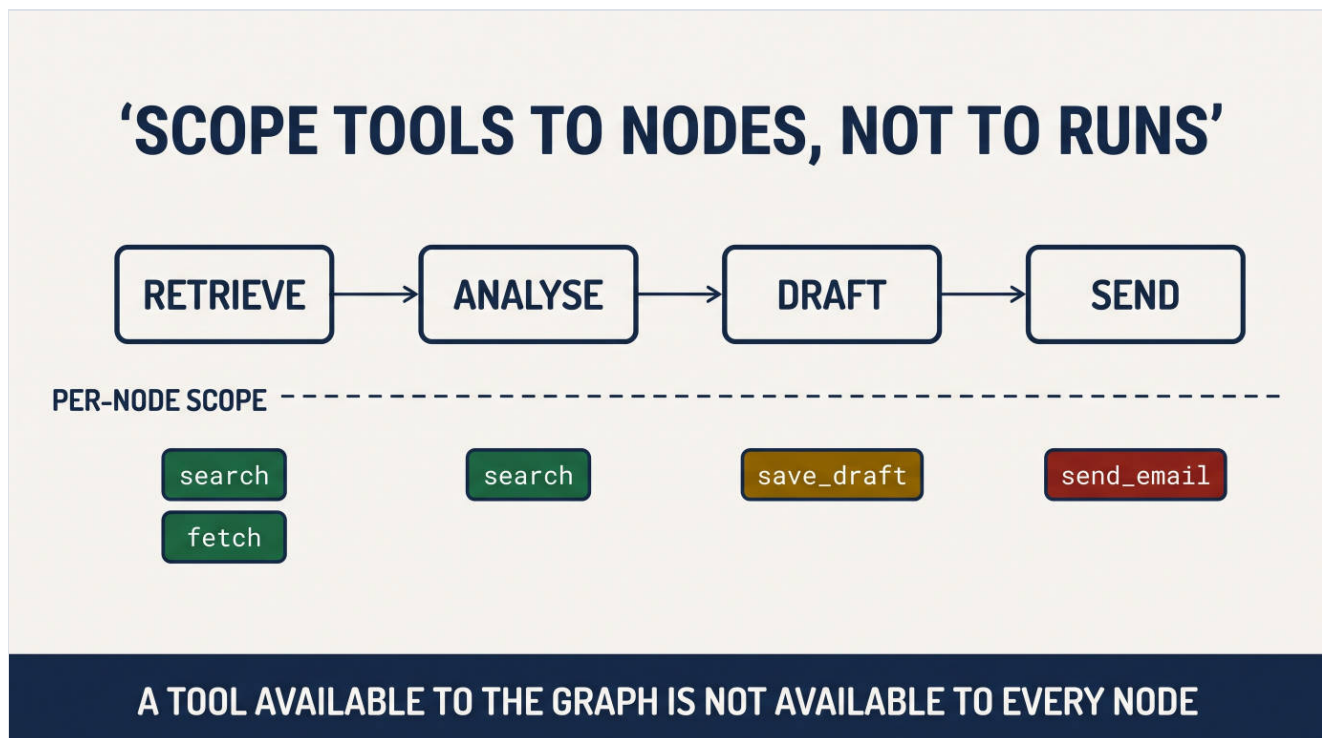


Figure 12.3 The graph gives permissions a natural granularity that a free-form agent could not.

Chapter 9 scoped tools to a run via `Policy.allowed_tools`. The graph of Chapter 10 makes a finer scope available, and it is worth taking.

```
Node("retrieve", run=retrieve, tools={"search", "fetch"})
Node("analyse", run=analyse, tools={"search"})
Node("draft", run=draft, tools={"save_draft"})
Node("send", run=send, tools={"send_email"}) # only here
```

The `tools` field is this chapter's one extension to the Node record of section 10.3: the engine there had no tools to scope, and adding a field with an empty default changes nothing for graphs that declare none.

The effect is that the answer to "could this system send an email at that moment" becomes a property of the topology rather than of the model's judgement. During retrieval, `send_email` is not merely disallowed, **it is not advertised**, so section 9.4's token cost and mis-selection risk both disappear for the nodes that do not need it.

Per-tenant scoping composes with it, and the intersection is what a node actually gets:

```
effective = node.tools & policy.allowed_tools & tenant.enabled_tools
```

Three independent sets, intersected. A tenant who has not connected their email provider cannot have `send_email` in any node, regardless of what the graph declares, and the failure is a smaller advertised toolset rather than a runtime error.

Chapter 11's memory store is the worked example: expose it as a `memory.search` tool advertised only to the nodes whose work depends on stored facts, with the tenant filter taken from the caller's authenticated identity rather than from an argument, exactly as section 12.8 requires of the retrieval tool. The draft node can consult what the user prefers; the send node, which needs no memory, never sees the tool at all.

12.5 Idempotent Tool Execution

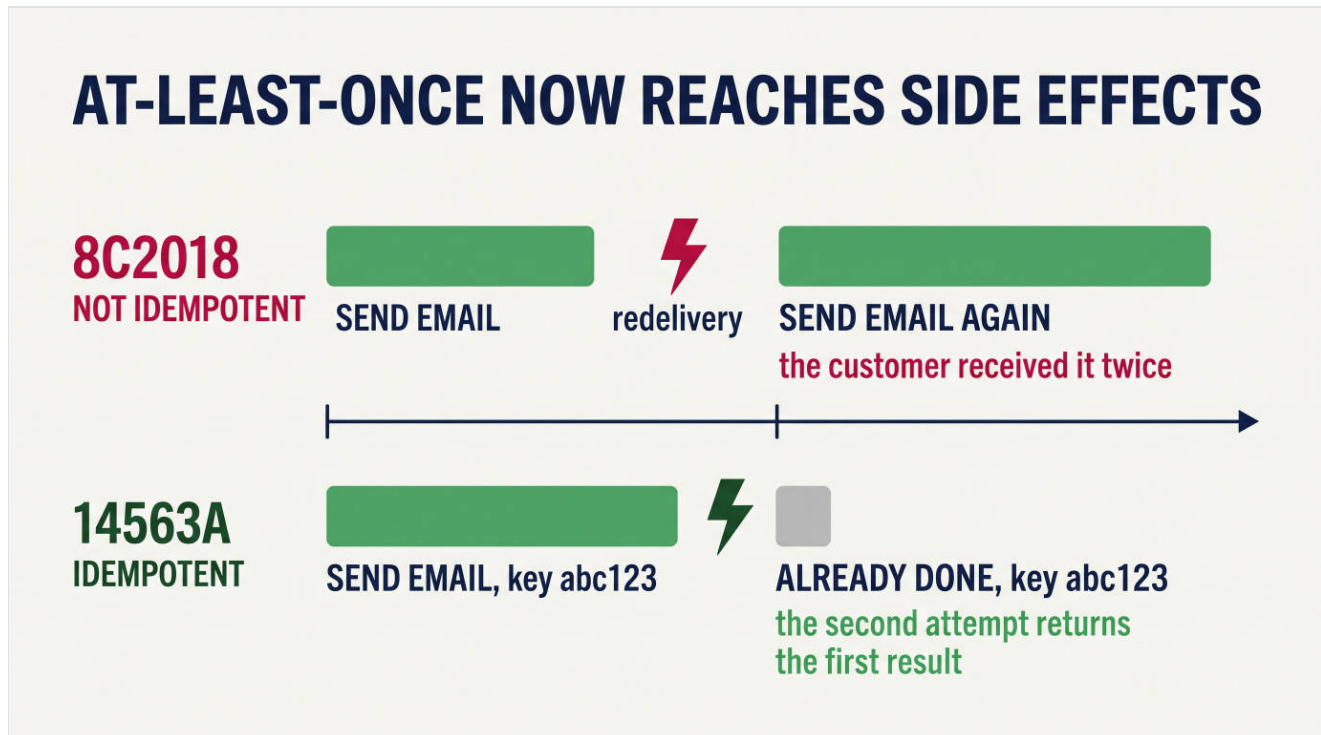


Figure 12.4 Section 8.10's redelivery, arriving at something that sends email.

Chapter 8 established that a queue promises at-least-once delivery and that exactly-once *effect* is the handler's job. Chapter 10 added that a crash between a node succeeding and its checkpoint committing replays that node. Both of those were tolerable while nodes were pure. Neither is tolerable once a node sends an email.

```
# The key derives from the WORK, per section 8.10, and travels to the tool.
key = sha256(f"{tenant}|{run_id}|{node}|{canonical(args)}").hexdigest()[1:32]

result = await tool.call(**args, idempotency_key=key)
```

Three cases, and only the first is comfortable.

The tool accepts an idempotency key. Payment providers and most mature APIs do. Pass it and the problem is solved by the vendor.

The tool does not, but is naturally idempotent. Setting a value, upserting a record. Safe to retry.

The tool does not, and has an effect. This is the dangerous one, and the answer is a *local* ledger: record the intent with its key before the call, mark it complete after, and consult it before every call. That converts an at-most-once tool into an at-most-once *system*, at the cost of a write.

Mark irreversible tools explicitly, and never retry them automatically. Section 9.7's `Risk.IRREVERSIBLE` already identifies them. A transient network error on a send is genuinely ambiguous: the send may have succeeded. The correct response is to surface the ambiguity to a human, not to guess and retry.

12.6 The Confused Deputy Problem

THE CONFUSED DEPUTY

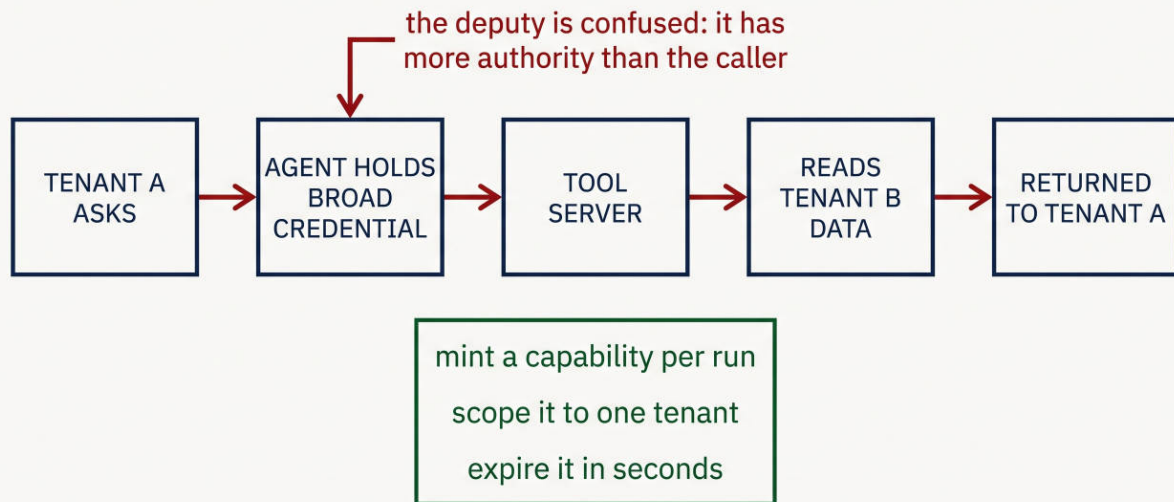


Figure 12.5 The agent has more authority than the person it is acting for, and nothing in the call says whose behalf it is on.

The classic vulnerability, arriving in a new place. Your agent holds a credential broad enough to serve every tenant. A tenant asks a question. The tool call carries the credential but nothing that binds it to *that* tenant, so the tool answers with whatever the credential can reach.

Section 1.5 built the control before there was anything to control: capabilities rather than credentials, minted per run, scoped to one tenant, expiring in seconds.

```
cap = await capabilities.mint(
    tool="search_crm", tenant_id=ctx.tenant_id,
    scopes={"read:contacts"}, ttl_s=60, run_id=run_id)

result = await tool.call(capability=cap, **args)
```

The property that matters: **if the capability leaks, it is worth almost nothing**. Sixty seconds, one tenant, one tool, read only. A leaked broad credential is worth everything, and section 12.3 established that a hostile tool server sees whatever you send it.

Where a vendor cannot accept a capability, the fallback is a *proxy of your own* that holds the broad credential, accepts a capability from the agent, checks the tenant binding, and forwards. That is more work than it sounds and less work than the incident. Concretely: the proxy exposes a narrowed endpoint, the two operations the node actually needs rather than the vendor's full API surface, so the blast radius of a hostile call is two verbs. It verifies a per-run capability token of your own minting before touching the credential, and it logs every forwarded call with the run, node and tenant that asked, which is the attribution the vendor's own logs cannot give you. Budget a day of work per vendor, not a platform: it is a small HTTP service with one client and one secret.

This is also the shape into which real MCP authentication resolves. Remote MCP servers authenticate with OAuth flows in which your client holds the tokens, and the section's vocabulary maps on directly: the token *is* the capability. Scope it per server and per operation the way the minted capability above is scoped, and never forward your primary credential to a server because the flow made it convenient; a token that one server can replay against another is the confused deputy again, one hop out. The abstract argument of this section is the concrete 2026 mechanism, and the discipline transfers unchanged.

12.7 Tool Result Handling

A TOOL RESULT IS UNTRUSTED INPUT

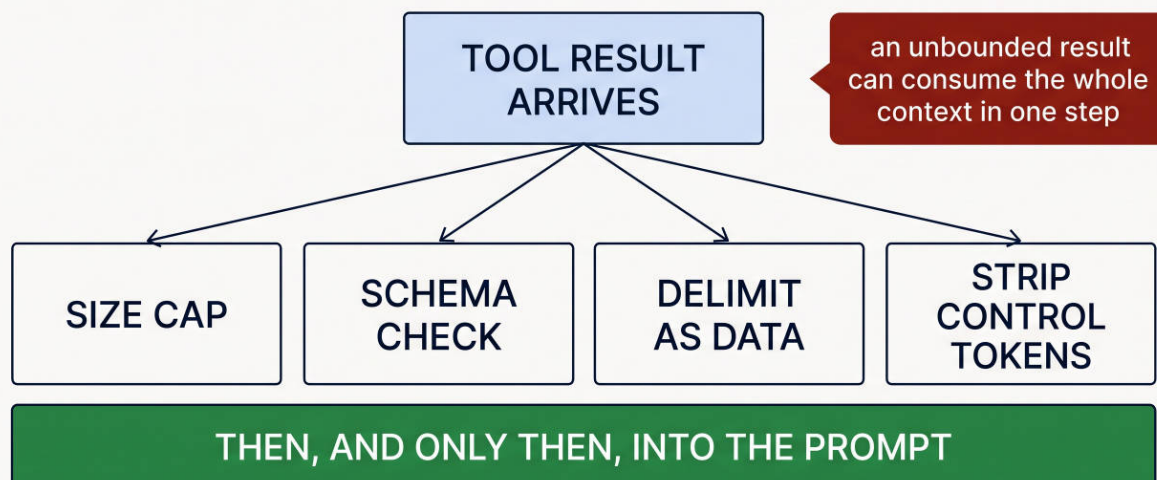


Figure 12.6 Four checks before anything reaches the prompt.

The result is the payload. Four controls, in order, and none is optional once a third party is involved.

Control	What it prevents
Size cap	An unbounded result consuming the entire context in one step, evicting the system prompt and the evidence with it. Truncate and record that you did, per section 2.10.
Schema check	A result whose shape changed under you. This is where capability drift surfaces (12.9).
Delimited as data	Section 5.7. The result goes in a marked user-role block, never the system message.
Strip control tokens	A result containing chat-template markers. Section 5.7 noted the property holds only if you build prompts through the template API, and a tool result is exactly where raw text arrives.

The fourth control has a concrete mechanism: normalise inbound text to plain content before insertion, stripping or escaping the serialisation's special token strings and any hand-written role markers, so the result can only ever be data inside the block you delimited. Section 5.7 is the argument for why role markers are the attack shape; this is the two lines of code that act on it.

A fifth, once you are past the basics: **attribute the result**. A claim in an answer that came from a vendor tool should be citable to that tool, in the same way section 6.8 cites a document. "According to your CRM" is both more useful and more honest than an unsourced assertion.

12.8 Exposing Your Own Retrieval as a Tool

The inverse direction, and it is worth doing deliberately rather than by default. Wrapping your Chapter 7 retriever as an MCP tool means other clients, including ones you do not control, can query your corpus.

Three requirements, all of which you already built:

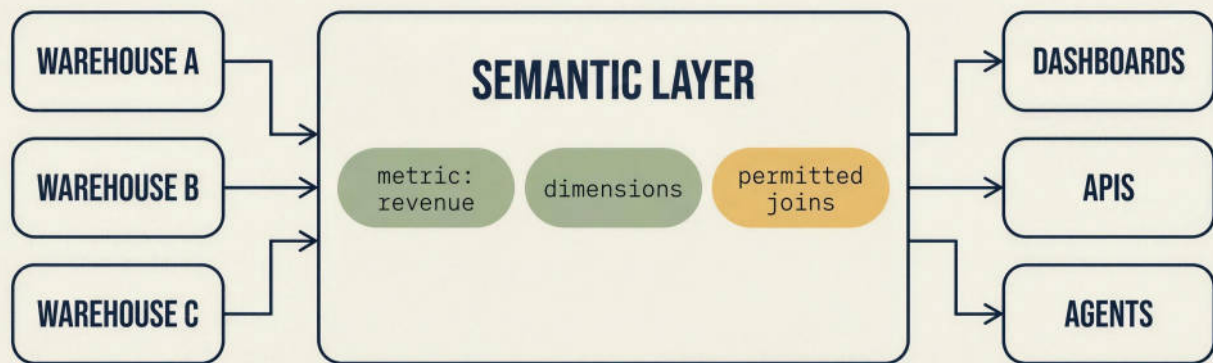
The tenant filter is not a parameter. It comes from the caller's authenticated identity, never from an argument the model can set. Section 6.10's `assert_scoped` is the enforcement, and exposing the corpus as a tool is precisely when a code path might build its own filter and forget.

The schema constrains as in section 9.3. A bounded `limit`, an enum for document type, and a name that states the scope.

The result carries citations. Section 6.8's stable ids, so the consuming system can render a link rather than an unsourced paragraph.

Structured data wants the same treatment: the semantic layer

'EXPOSE THE LAYER, NEVER THE WAREHOUSE'



ONE DEFINITION OF THE NUMBER, FOR EVERY CONSUMER

Figure 12.7 The agent selects from declared definitions; it never composes SQL against the warehouse.

The corpus is not the only thing an agent will be asked to answer from. Sooner or later the question is "what was revenue last quarter", and the naive route is to hand the model the warehouse: a text-to-SQL tool over the raw schemas. That is this chapter's every hazard in one place. The model composes arbitrary queries against tables whose semantics it can only guess, "revenue" is defined five different ways across marts, and access control degrades into hoping the generated SQL respects a filter nobody enforces below it.

The data world already built the governed interface, and it is worth knowing by name: the **semantic layer**, a layer of declared definitions, metrics, dimensions, entities and their permitted joins, sitting between the warehouses and every consumer, so that dashboards, APIs and agents all read the same definition of the same number. For an agent, the design rule is the one this section already stated for retrieval: expose the semantic layer as the tool, never the warehouse behind it.

The reason it works is the book's oldest move: it converts open generation into closed selection. Instead of writing SQL, the model chooses a metric, dimensions and filters from an enumerable, documented set, which is section 5.3's schema, section 10.3's declared destinations and section 11.4's closed predicates arriving in the data stack. The three requirements above transfer without modification: row and column access comes from the caller's identity below the tool boundary, never from an argument; the tool schema constrains with enums of metrics and dimensions and bounded date ranges; and the result carries provenance, the metric definition and its version, so an answer can cite the definition it used.

The quiet payoff is consistency. An agent reading the semantic layer gives answers that match the dashboards, because both read the same definitions, and a wrong number becomes a findable definitions bug rather than an unfindable hallucination. The definitions are versioned artefacts, which means they pin exactly the way section 12.9 pins tool contracts: a changed metric definition is a changed answer, and it should fail a build, not surprise a user.

12.9 Versioning and Capability Drift

Section 1.2 argued that a dependency whose behaviour changes silently is the defining hazard of this field. A third-party tool is that hazard with no lockfile and no release notes.

```
# Record the tool's advertised contract in the fingerprint (1.2), and fail
# the build when it moves.
@dataclass(frozen=True)
class ToolPin:
    name: str
    schema_hash: str # sha256 of the normalised parameter schema
    description_hash: str # the DESCRIPTION is a prompt, so it is pinned too

async def check_drift(registry, pins: dict[str, ToolPin]) -> list[str]:
    """Run in CI against the live servers. A changed description is a
    changed prompt, and section 5.8 says a changed prompt is a release."""
    drifted = []
    for name, pin in pins.items():
        live = await registry.describe(name)
        if hash_schema(live) != pin.schema_hash:
            drifted.append(f"{name}: schema changed")
        if hash_text(live.description) != pin.description_hash:
            drifted.append(f"{name}: description changed")
    return drifted
```

Pinning the *description* as well as the schema is the part teams skip. A schema change breaks loudly; a description change alters model behaviour silently, which is the section 5.8 failure with the artifact owned by someone else.

The CI check catches drift between deploys. It does not catch a server that changes its tool list mid-session, which the protocol explicitly permits: a list-changed notification can arrive while a run is in flight, and the tools you verified this morning are not the tools being advertised this afternoon. The policy is the same one CI enforces, applied at notification time: re-verify the hashes against the pins, and quarantine any tool that moved, meaning it drops out of the advertised set until a human re-approves the new contract.

Never hot-adopt a changed schema inside a live session. The run started against one contract, its scratchpad reasons about that contract, and swapping the tool underneath it produces exactly the mid-run surprises the pins exist to prevent. A quarantined tool is a smaller toolset for the rest of the run, which section 12.4 already made a normal condition rather than an error.

12.10 Three Tool Architectures Compared

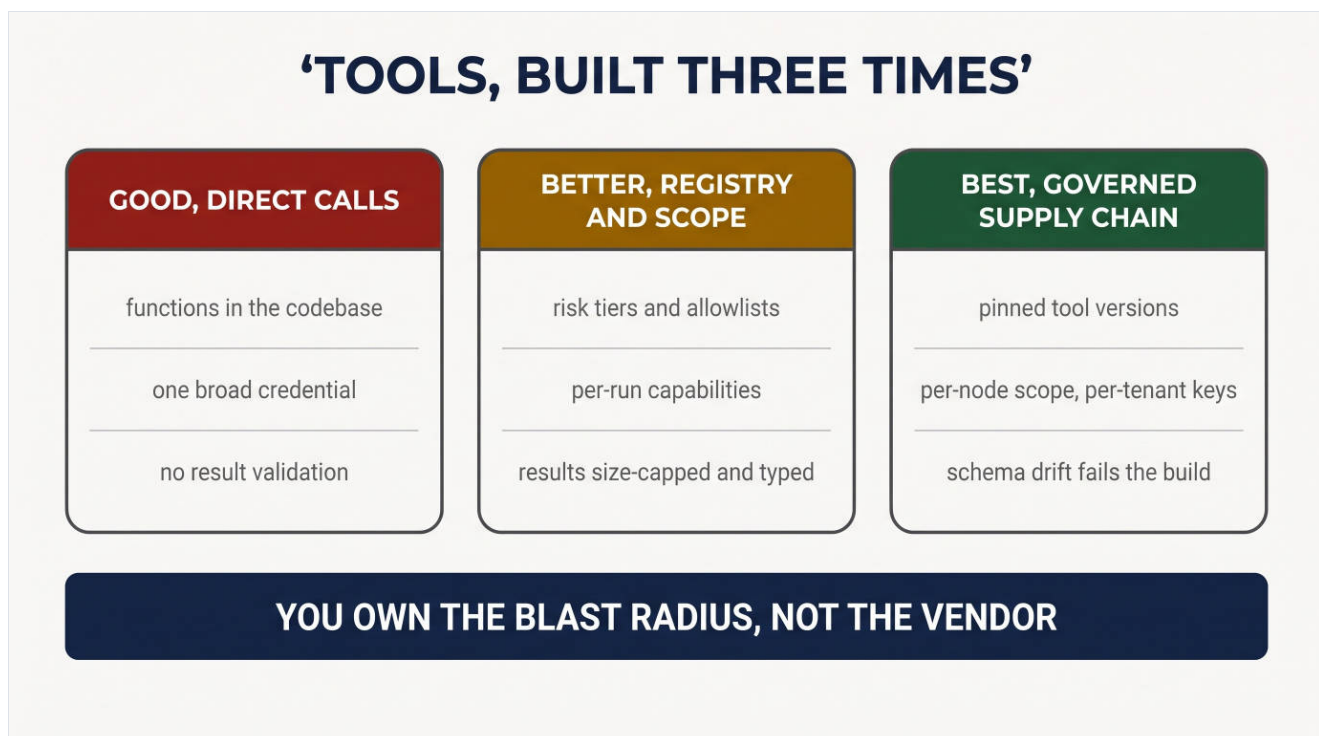


Figure 12.8 The blast radius is yours whichever tier you are on.

GOOD

Direct function calls

Tools are functions in your codebase, called with your credentials. No protocol, no network, no vendor. Correct while every tool is yours and the set is small, and it has the considerable virtue that the schema, the behaviour and the result are all under your control.

What it leaves unbounded: one broad credential shared by every tool; no scoping below the run; results placed in prompts unvalidated; and no path to a capability you cannot build.

BETTER

Registry with risk tiers and scoped capabilities

Chapter 9's registry and chokepoint, plus per-run capabilities from section 12.6, plus the result controls of section 12.7. Third-party servers are reachable, and each is an allowlist entry rather than a discovery result.

What it still does not solve. Tool contracts can change under you between deploys. Scoping is per run, so a retrieval node still advertises the send tool. And there is no supply-chain position on which versions of which servers you are willing to run.

BEST

Governed supply chain

Everything above, plus per-node scope from section 12.4, pinned schema and description hashes checked in CI, per-tenant capability minting, and irreversible tools that never auto-retry.

Axis	Good: direct	Better: registry and scope	Best: governed
Credential blast radius	One broad credential	Per-run capability	Per-run, per-tenant, expiring
Tool advertised to	Every step	Every node in the run	Only nodes that need it
Result handling	Straight into the prompt	Capped, typed, delimited	Capped, typed, delimited, attributed
Contract drift	N/A, you own it	Undetected	Fails the build
Irreversible actions	Retried like anything else	Gated by approval	Gated, never auto-retried
Build cost	Days	Two to three weeks	Four to six weeks

PROMOTION TRIGGERS

Good to better, at the first third-party tool, or the first tool with an irreversible effect, whichever comes first. Both change the threat model rather than the workload.

Better to best, when any of: you run more than one tool server you do not control; a tool touches money, messaging or deletion; or a customer asks which vendors process their data. The drift check is the cheapest piece and I would build it first, because it is the one that catches a problem you would otherwise never see.

What Chapter Twelve Establishes

What exists now	What later chapters put through it
Per-node and per-tenant tool scope	Approval routing by tool risk (15.1)
Capabilities rather than credentials	Data residency and vendor boundaries (16.5)
Result size caps and schema checks	Injection detection on tool results (15.7)
Tool pins with description hashes	The release gate (14.6)
Idempotency keys reaching the tool	Incident replay and recovery (15.9)

Chapter checklist

Can you...	Section
Say what MCP standardises and name four things it leaves to you	12.1
Explain why discovery is an input to a registry rather than a replacement	12.1
Name the three artefacts you accept from a third-party server	12.3
Answer "if this server were hostile, what could it do?"	12.3
Write the intersection that decides a node's effective toolset	12.4
Handle a tool with side effects that accepts no idempotency key	12.5
Explain the confused deputy in one sentence, and the control that fixes it	12.6
List the four checks a tool result passes before reaching a prompt	12.7
Say why the tenant filter is never a tool parameter	12.8
Explain why the description is pinned as well as the schema	12.9

What Chapter 13 builds

Every interface so far has been text, where a slow answer is an annoyance. Chapter 13 removes that tolerance: in conversation, silence past about eight hundred milliseconds reads as a broken system, and the entire pipeline of Part I to Part III has to fit inside that.

The latency budget and where it actually goes. Turn detection and why endpointing is the hardest unglamorous problem in voice. Streaming the whole pipeline so the first sentence starts before the answer is finished. Barge-in, which is trivial to describe and changes the architecture. And the errors you cannot see, because a voice interface has no visible citation, no way to show a source, and no natural place to render an abstention.