

■ AI READINESS REPORT · AGENTLESS ASSESSMENT

Is it safe to turn on **AI**?

One question, answered with evidence: can this organisation switch on Microsoft 365 Copilot and other AI tools — and if not yet, exactly what has to change first.



NEEDS WORK

GRADE CAPPED BY A CRITICAL FINDING

2.6 measured · 0.2 attested. The grade is held down, not merely low — one critical finding (nine unowned, tenant-wide mailbox grants) caps this tenant at **Needs Work** regardless of how the other domains score. Clear it and the grade re-rates on the next scan.

0-2.0 AT RISK · 2.0-3.5 NEEDS WORK · 3.5-4.5 PILOT-READY · 4.5-5.0 AI-READY — MICROSOFT CAF CONVENTION

NOT YET

VERDICT — REMEDIATE
BEFORE COPILOT

62

AI IDENTITIES, IN 703
SERVICE PRINCIPALS

8%

SEEN BY MICROSOFT'S OWN
TOOLING

30-60

DAYS TO PILOT-READY ·
NO NEW LICENCES

AI is already running. *Nobody* is accountable for it.

This is not a decision about whether to adopt AI. Sixty-two AI identities are already operating inside this tenant — connectors, agent runtimes and consented applications that act continuously whether or not anyone is signed in. Before Copilot is added, the AI already present has to be brought under ownership.



Not yet ready — remediate first. Grade **2.8 / 5.0 (Needs Work)**, capped by a critical finding. This is not weak hygiene everywhere — consent controls and governance are in reasonable shape. It is one unowned identity exposure holding the whole grade down: nine AI applications with standing access to every mailbox, none with an accountable owner. Clear that and the grade re-rates immediately.

Where this tenant stands, by pillar

Settings — who is allowed to let AI in

24 OF 31 CHECKS SCORED • CONSENT, REGISTRATION, BASELINE

71%

Identities — the AI agents already operating

22 OF 30 CHECKS SCORED • OWNERSHIP, PERMISSIONS, BLAST RADIUS

54%

Data Exposure — how far content can travel

16 OF 19 CHECKS SCORED • SHARING, GUESTS, LINKS, SITE ACCESS

41%

Data Protection — is the information itself contained

5 OF 13 CHECKS SCORED • LABELS, RETENTION, DLP — LICENCE-LIMITED

38%

The four things a leader needs to know

9

AI apps can read every mailbox. Approved by ordinary staff, none with an accountable owner, the oldest credential 19 months unrotated. If one is compromised, every mailbox is reachable.

1

An abandoned agent, still running. A Copilot Studio agent built by someone who left five months ago still holds read access to every SharePoint site — unattended and unowned.

21

The door is still open. Any employee can approve a new AI application with no review, and twenty-one grants already happened. The problem can recur tomorrow.

38 of 62

AI identities have no owner. A measured number from the tenant, not a questionnaire answer — and the reason nothing above was noticed until now.

What the grade is not. It is not a configuration audit score. It is the measured gap between what AI can already reach in this tenant and what anyone can currently see and govern. That is why one unowned, high-reach identity class can cap it while the rest of the estate scores respectably.

What we could see – and what nobody can

Every scan states plainly what it measured, what it had to ask, and what nobody can see today. This is what makes the grade trustworthy: a control we could not verify is never shown as a pass.

The check ledger – every one of 112 accounted for

MICROSOFT 365 BUSINESS PREMIUM

- 89

machine-assessed on the live tenant — read directly from Graph
- 58

↳ passing
- 31

↳ findings needing attention, each pointing at a named object
- 15

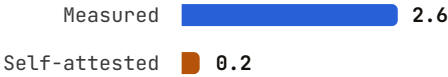
attestation controls — no technical surface; a named person must answer. 4 confirmed · 11 awaiting
- 8

not assessable on this licence — named below, never guessed as clean

89 + 15 + 8 = 112. Coverage — machine-assessed plus confirmed — is 93 of 112 (83%).

Where the grade comes from

2.8 / 5.0 · MEASURED VS ATTESTED



Eleven per cent of the rubric could not be scored — three domains had no runnable checks on this tenant. Their weight is redistributed proportionally across the domains we could measure, rather than counted as passes or failures.

Unseeable today – no licence fixes these

STATED PLAINLY, NOT SCORED AS PASSES

What staff type into browser-based AI tools · AI tools running on unmanaged or personal devices · fleet-scale local AI inventory. No vendor can read these today because no API exists for them. We name them rather than score them.

What we couldn't see – and what closes it

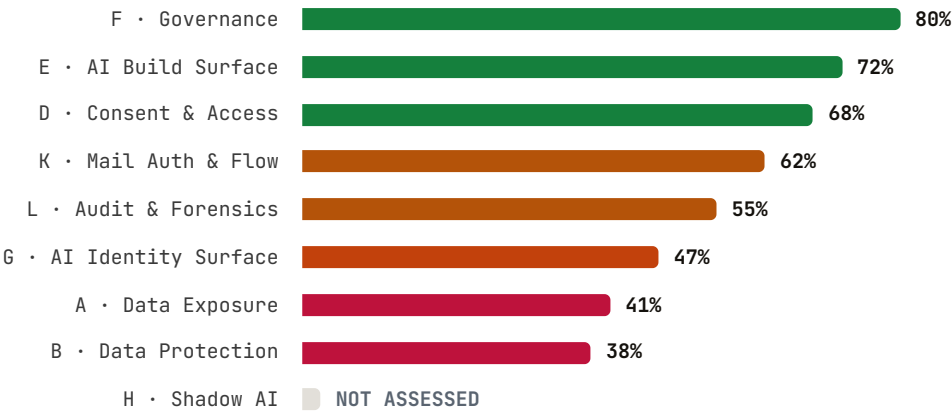
BLIND SPOT	WHY	WHAT CLOSES IT
Network-level GenAI SaaS discovery	Requires Defender for Cloud Apps	Licence uplift — 3 checks return
Prompt-level and browser DLP	Requires Purview endpoint DLP	Licence uplift — 2 checks return
~34% of endpoints	Not enrolled in device management	Enrolment project — no licence needed
Audit trail beyond ~180 days	Retention on this tier	Advanced Audit or a custom retention policy
Next-generation agent governance	Microsoft Agent 365 not present	Available as an add-on when adopted

Every licence-gated line above is simultaneously a security finding and an upgrade path with the business case already attached. For a managed service provider running this assessment, that is a qualified pipeline — itemised, evidenced, and tied to a specific risk the client can now see.

An AI assistant *inherits* the reach of whoever runs it

It can open any file, mailbox or site that person can open — instantly, and at machine speed. So the safety of AI use is really the safety of four things: who can let AI in, what AI can reach, whether that content is contained, and who is accountable for the agents already running.

Domain scores, weighted

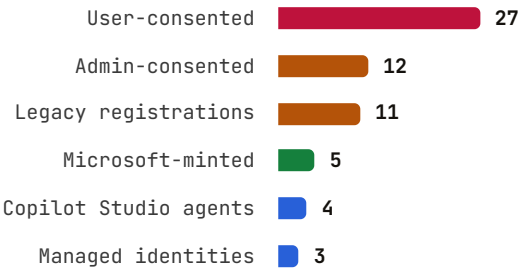


Domain **G • AI Identity Surface** carries the heaviest weight in the model and is where the critical finding sits — which is why the headline grade sits below what the domain average alone would suggest. **H • Shadow AI** shows as not assessed, not zero: every check in it is licence-gated — a gap in visibility, not a failure. Attestation controls awaiting confirmation are **excluded** from domain scores rather than counted as passes.

The AI estate, by origin

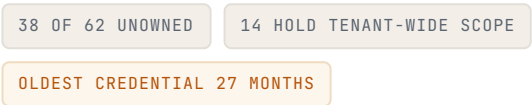
62 AI identities • how each one got in

703 SERVICE PRINCIPALS TOTAL



57 of 62 are invisible to Microsoft's own tooling

Agent 365 governs the five identities Microsoft itself minted. The other fifty-seven consented their way in — approved by staff, registered by developers, or inherited from integrations set up years ago — and no first-party tool treats them as a governed population. **This is the gap the readiness assessment exists to close.**



The two that *cap* the grade

Thirty-one findings need attention. Each points at a real object in the tenant — an application, an agent, a site, a guest — not a percentage. Each carries its remediation: what to do, the effort, the grade points it earns back, who owns the fix, whether it is reversible, and how the re-scan verifies it closed.

1 Nine AI apps hold tenant-wide mailbox access, unowned

• CRITICAL • SETS THE GRADE CAP

What. Nine third-party AI applications hold `Mail.Read` or `Mail.ReadWrite` across the entire tenant. None has an accountable owner; the oldest secret is nineteen months unrotated; the earliest was approved twenty-two months ago by a member of staff who has since left.

Why it matters. These are standing, non-human access paths to every mailbox in the business that persist long after the person who approved them has gone. This is the exact identity class exploited in the 2024 Microsoft breach — a forgotten OAuth application with mailbox access and a credential nobody rotated. Adding Copilot on top of this does not create the exposure; it accelerates what can be done with it.

Action. Revoke what is no longer needed; assign a named owner and rotate the secret on everything that stays. **Verify.** Re-scan G.3 shows zero unowned tenant-wide mail grants.

EFFORT 4-6 HRS

POINTS +0.4

MSP + CLIENT DECISION

ROTATION NEEDS A CHANGE WINDOW

2 An abandoned Copilot Studio agent with SharePoint-wide access

• CRITICAL

What. One of four Copilot Studio agents was built internally by an employee who left five months ago. It runs under an identity holding `Sites.Read.All` — read access to every SharePoint site — and has been unattended and unowned since their departure.

Why it matters. A homemade agent with a living identity and a departed owner is ungoverned automation with real reach and no accountability. Nobody is reviewing what it does, nobody would notice if its behaviour changed, and no leaver process touched it because it is not a person.

Action. Disable the agent; reassign or decommission the identity it runs under. **Verify.** Re-scan confirms no orphaned agents holding live identities.

EFFORT 1 HR

POINTS +0.2

MSP-FIXABLE

FULLY REVERSIBLE

The agent-to-identity join is what makes finding 2 visible. Most tooling can list Copilot Studio agents, and most tooling can list service principals. Joining the two — establishing that *this abandoned bot* runs as *that live identity with this much reach* — is what turns a tidy inventory item into a governed access path with a name against it.

Keeping the door *shut*

Clearing the estate without closing the consent door is mopping the floor with the tap running. These findings stop the problem recurring and close the standing exfiltration paths.

3 Any employee can approve a new AI app, with no review ● HIGH

What. User consent lets any member of staff grant an AI application access to company data. There is no admin-consent workflow, and standard users can also register new applications. Twenty-one risky consent grants already occurred inside the audit window. This is how the nine applications in finding 1 arrived.

Why it matters. These are identity-creation events, not misconfigurations. The agent estate grows in the dark by design, not by accident — and every fix applied elsewhere in this report can be undone by one well-meaning approval tomorrow.

Action. Restrict user consent to verified publishers and low-impact permissions; enable the admin-consent workflow with named reviewers so people can still request what they need; disable standard-user app registration. Triage the twenty-one existing grants. **Verify.** Re-scan D.1, D.2, D.4.

EFFORT 30 MIN + REVIEW

POINTS +0.1

MSP-FIXABLE

NOTHING ALREADY APPROVED BREAKS

4 Sharing is wide open and anonymous links never expire ● HIGH

What. Four of six content sites permit external sharing, the default link type is anonymous "anyone with the link", those links carry no expiry, and guests can reshare what they receive. Eleven external guests currently reach SharePoint or OneDrive.

Why it matters. Oversharing that was survivable when discovery was slow becomes exposure the moment an AI can surface it on demand. A colleague might never stumble across a link forgotten in 2023; an assistant answering a question surfaces it in a second.

Action. Require guests to sign in; set an expiry on anonymous links; disable resharing by external users; review the eleven guests. **Verify.** Re-scan A.1, A.4, A.5, A.10.

EFFORT 2 HRS

POINTS +0.2

MSP + CLIENT DECISION

EXISTING ANONYMOUS LINKS STOP WORKING

5 Two open exfiltration paths: auto-forwarding on, DMARC not enforcing ● HIGH

What. External mail auto-forwarding is permitted. The primary domain publishes DMARC at `p=none` and the secondary has no record at all — neither enforces.

Why it matters. Auto-forwarding is a standing exfiltration path that survives password resets. Absent DMARC leaves the organisation's own domains trivial to spoof — how attackers reach the administrators who approve AI applications. The two compound each other.

Action. Block external auto-forwarding with named exemptions; confirm SPF and DKIM, then publish DMARC at quarantine and move to reject. **Verify.** Re-scan K.1, K.2, K.7.

EFFORT 1-2 HRS

POINTS +0.1

MSP-FIXABLE

LOW DISRUPTION IF SPF/DKIM CONFIRMED FIRST

6 Stale and dormant AI credentials ● MEDIUM

What. Thirty-one AI credentials have never been rotated, the oldest at twenty-seven months. Eight AI applications have recorded no sign-in in over ninety days yet retain full standing access.

Action. Rotate the active, revoke the dormant. **Verify.** Re-scan G.6 and G.9.

EFFORT 3 HRS

POINTS +0.1

MSP + CLIENT DECISION

FULLY REVERSIBLE

Classify, monitor, *evidence*

The remaining findings address whether the data underneath the AI is contained at all, whether governance would notice itself drifting, and which controls are currently asserted rather than proven.

7 The data itself is unclassified and ungoverned • MEDIUM • STRATEGIC

What. No sensitivity labels are published and no retention policy governs AI or Copilot content. Data-loss monitoring is not collecting enforcement telemetry, so DLP behaviour cannot be observed or investigated.

Why it matters. Labels and retention are what stop an AI tool surfacing or exporting the most sensitive content *even when the person driving it technically has access*. Everything else in this report controls who can reach the data; this is the only layer that controls what happens once they do.

Action. Scope a classification and retention programme; enable DLP enforcement telemetry. **Verify.** Re-scan B.1, B.4, B.7.

EFFORT PROJECT

POINTS +0.2

LICENCE DECISION REQUIRED

PLAN AND BUDGET

8 No governance-drift monitoring; consent velocity un-baselined • MEDIUM

What. There is no baseline for how fast AI applications are being approved, and no monitoring for whether consent controls have been relaxed after being tightened.

Why it matters. A tenant that tightened its controls and then quietly loosened them tells a very different risk story from one that never tightened at all — and today nothing would distinguish the two. Trend, not just state, is what separates a one-time clean-up from being genuinely ready.

Action. Baseline at this scan; monitor continuously from here. **Verify.** Continuous — reported at every subsequent scan.

BUILT INTO ONGOING SERVICE

POINTS +0.1

MSP (ONGOING)

FULLY REVERSIBLE

9 Governance is asserted, not evidenced • MEDIUM

What. Eleven controls await confirmation — among them a written AI policy, a named accountable owner for the AI estate, and a maintained AI inventory. Four have been confirmed with a named person and a date. Conditional Access does not currently scope controls to SharePoint or to AI applications, and audit retention sits at the tier default.

Action. Confirm the outstanding controls with attribution so they carry a name rather than an assumption; add scoped Conditional Access in report-only mode first; extend audit retention. **Verify.** Attestation register + re-scan D.12, L.2.

EFFORT 2 HRS

POINTS +0.1

CLIENT DECISION

FULLY REVERSIBLE

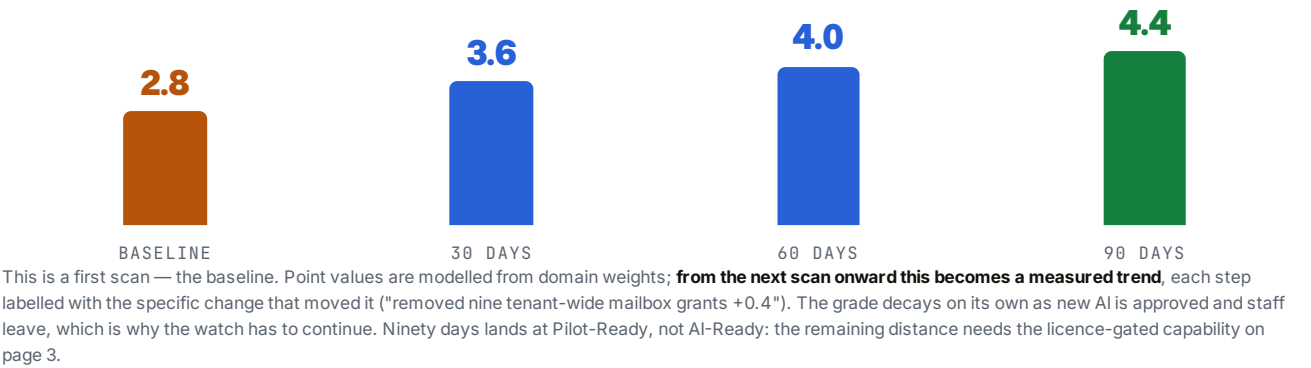
The remaining twenty-two. Findings 1–9 carry the grade. The balance of the thirty-one are lower-severity items in the same domains — default link types, site lifecycle, ownerless sites, external-sender tagging, consumer federation and discovery-tier licensing — each listed with its own remediation in the full findings register.

Thirty days to clear the *cap*

Sequenced by points-per-hour — fastest grade improvement first. Clearing the critical cap in the first thirty days re-rates the whole tenant. **None of the first-thirty-day work needs new licences**, and where a fix could interrupt someone, it is called out.

Grade trajectory once remediation begins

BASELINE 2.8 → PILOT-READY 4.4 · ILLUSTRATIVE PROJECTION, NOT A MEASURED FORECAST



■ First 30 days → ~3.6 CLEARS THE CAP · NO NEW LICENCES

- 1 Revoke, own and rotate the nine tenant-wide mailbox grants** +0.4
The single highest-value action in this report — it lifts the grade cap on its own. *Disruption: low; rotation needs a change window.* (Finding 1)
- 2 Disable the orphaned Copilot Studio agent** +0.2
Turn it off, then decide whether to reassign or decommission the identity behind it. *Disruption: none — nobody has owned it for five months.* (Finding 2)
- 3 Close the app-consent front door** +0.1
Restrict user consent to verified, low-risk apps; enable the approval workflow with named reviewers; disable standard-user app registration. Triage the twenty-one existing grants. *Disruption: low — nothing already approved breaks.* (Finding 3)
- 4 Block external auto-forwarding and publish enforcing DMARC** +0.1
Close the standing exfiltration path; confirm SPF and DKIM, then publish DMARC at quarantine on both domains. *Disruption: low.* (Finding 5)

Then tighten, then *make it stick*

The first thirty days lift the cap and close the open doors. The next sixty tighten what is already inside and put in place the thing that stops the whole exercise unravelling three months from now.

■ Days 30–60 → ~4.0 SCOPE BEFORE FLIPPING

5 Turn down "anyone" sharing and expire anonymous links +0.2
Require guests to sign in, set link expiry, and disable resharing by external users. *Disruption: low, but existing anonymous links stop working — warn staff first.* (Finding 4)

6 Rotate the active credentials, revoke the dormant +0.1
Thirty-one never-rotated credentials and eight applications idle past ninety days but still holding standing access. *Disruption: low.* (Finding 6)

7 Review the eleven external guests +0.1
Confirm each still needs access; remove the rest. *Disruption: none — a review, not a config change.* (Finding 4)

■ Days 60–90 → ~4.4 PLAN AND BUDGET

8 Protect the data itself – the biggest strategic gap +0.2
Stand up sensitivity labels and retention for AI and Copilot content, and switch on DLP enforcement telemetry. A project, not a switch — but it is what ultimately makes AI safe to use on sensitive material. (Finding 7)

9 Baseline governance drift and monitor from here +0.1
Set the consent-velocity baseline at this scan so a future loosening of controls is visible as a change, not discovered as an incident. (Finding 8)

10 Evidence the governance you already have +0.1
Confirm the eleven outstanding controls with named owners and dates; add scoped Conditional Access in report-only mode; extend audit retention. (Finding 9)

What happens if nothing changes. The grade does not hold at 2.8 — it decays. New AI applications get approved under the same open consent policy, more staff leave and orphan the agents they built, credentials age past rotation, and sharing settings loosen without anyone noticing. Whatever gets fixed above needs a way to *stay* fixed. That is the difference between a clean-up and readiness.

Everything else reasons config-in. This reasons *identity-out*.

No tool below is a strawman — each is strong at what it was built for. The point is that AI readiness needs a view none of them makes primary: **the AI identity as a governed subject**, with an owner, a blast radius and a credential state.

QUESTION	SECURE SCORE	CSPM / SSPM	AI-SPM	THIS ASSESSMENT
Grade config vs a baseline	Yes — core strength	Yes — broad, drift-aware	Partial	Yes, plus more
Name the exposed items	No — scores settings	Partial	No	Yes — 11 guests, 4 sites, 21 grants, by name
Inventory AI identities as subjects	No	As resource attributes	Model layer only	Yes — 62 of 703 classified
Per-agent blast radius	No	Rare	No	Yes — 2 critical, 6 high
Ownership accountability per AI identity	No	No	No	Yes — 38 of 62 unowned
Credential dormancy for AI identities	No	Partial	No	Yes — 31 stale, 8 dormant
Agent joined to the identity it runs as	No	No	No	Yes — finding 2
Secure the model / prompt layer	No	No	Yes — core strength	No — complementary

What a configuration grade cannot tell you

A posture tool grades settings against a generic baseline. It tells you a switch is in the wrong position; it does not tell you **who or what can actually reach your data as a result**. It has no concept of a non-human identity as a subject to be governed — so it will not name the eleven guests, enumerate the four shared sites, list the twenty-one grants as events to triage, or tell you that nine specific applications can read every mailbox. It scores the configuration; it does not trace the exposure.

Why the inversion matters

Starting from the identity is what surfaces findings 1 and 2 at all. To a config-first scan, nine consented applications are nine rows in an app inventory and an abandoned bot is a tidy-up task. Read identity-out, they are **unowned standing access to every mailbox** and **ungoverned automation with a dead owner**. Same tenant, same data — a materially different conclusion about whether it is safe to turn on AI.

Complementary, not competing. Run Secure Score as the hygiene floor. Run a posture tool for configuration breadth and drift. Run AI-SPM if your teams are building models. This assessment governs the AI identities already operating inside the tenant — and it is the only one of them that names what they can reach.

Scoring you can *audit*

This report distinguishes rigorously between what the assessment measured on the tenant and what it could not read. The latter is reported by name as capability that could be added — never guessed as clean. That discipline is the operating principle of the platform, not a caveat at the back of the document.

89

machine-assessed on the live tenant

58

controls passing on live data

31

findings needing attention

8

not assessable on this licence — named, not scored against you

11

awaiting confirmation — the scan cannot verify these

4

attested, each with a named person and a date

How the grade is built. One hundred and twelve checks are scored per domain, weighted, and expressed on the Microsoft Cloud Adoption Framework convention of 0–5.0 where higher is better. Eleven per cent of the rubric had no runnable checks on this tenant; that weight is redistributed proportionally across the measured domains rather than counted either way. Self-attested controls are recorded separately with attribution and are excluded from the measured portion — here, 2.6 measured and 0.2 attested. An open critical finding caps the grade at Needs Work until it is closed.

About this sample

This is an illustrative sample, not a real organisation. The tenant, its findings and every value in this document are synthetic. What is *not* synthetic is the method: the 112-check rubric, the domain weighting, the 0–5.0 CAF scale and bands, the coverage ledger, the remediation model and the report structure are exactly those of a live assessment. A real report looks like this one, with your tenant's numbers in it — including, where they apply, findings less comfortable than these. Remediation point values and the 30/60/90 trajectory are modelled from domain weights and are labelled as such on page 8.

■ THE ASSESSMENT IS A SNAPSHOT. THE RISK IS NOT.

A one-off report is a photograph.

The grade only means something as a *trend*.

A tenant graded today will not hold that number: new AI applications get approved, staff leave and orphan the agents they built, credentials age, and sharing settings loosen without anyone noticing. Whatever gets fixed needs a way to stay fixed — which is why the unit of value is score movement, not the first score.

Sabiki AIRM is what turns this snapshot into governance: ownership and accountability for every non-human and AI identity, and a live inventory of the agent estate — the register that keeps every identity mapped to a named owner as agents and staff come and go, so a departing owner never quietly orphans an agent and a self-registered app never joins the estate unwatched.

AVAILABLE TO REGISTERED PARTNERS ON APPROVAL

sabikisecurity.com