



EXECUTIVE SECURITY SUMMARY

AI Risk Posture — Executive Briefing

Prepared for **Northwind Trading Co.**

REPORTING PERIOD

Q1 2026 · as of 12 Aug 2026

SCOPE

320 endpoints · fleet-wide

CLASSIFICATION

Confidential

PREPARED BY

Shadow AI Discovery

REPORT ID

SAD-EXEC-2026-0142

AUDIENCE

Board & senior management

CONFIDENTIAL — prepared for the named recipient. Do not distribute without authorization.

SECTION 1

At a glance

64

AI RISK POSTURE / 100

39 open risks · higher is safer

320

Endpoints monitored

1,284

AI artifacts discovered

17

Critical + high risks

92%

Findings AI-validated

Across 320 monitored endpoints, the platform discovered **1,284 AI artifacts** — unsanctioned LLM tools, coding agents, MCP servers and exposed API keys — and scored **39 open risks** (4 critical, 13 high). The dominant exposures are **secrets pasted into public LLMs** and **unregistered MCP servers on vulnerable versions**. Posture improved 6 points this quarter as high-severity items were remediated, but critical secret-exposure remains above target.

RISK SEVERITY MIX



■ Critical 4 ■ High 13 ■ Medium 12 ■ Low 7 ■ Info 3

KEY EXPOSURES

Critical

API keys in ChatGPT prompts

14 endpoints · secrets pasted into public LLM

Critical

Unregistered MCP server (vulnerable)

SHAI-2026-001 · 6 endpoints

High

Unapproved coding agents

31 endpoints · code + credential access

High

Typosquatted LLM package

server-filesyssem · 3 endpoints

High

Offboarded-user AI access

9 identities still active

Medium

Shadow browser AI extensions

48 endpoints · data egress risk

PRIORITY RECOMMENDATIONS

- 1 Contain critical secret exposure** 0–2 weeks
Rotate the 14 keys detected in LLM prompts and enable inline AI-DLP to block secret egress at the endpoint.
- 2 Quarantine the vulnerable MCP servers** 0–2 weeks
Isolate the 6 endpoints running the unregistered MCP server and upgrade past SHAI-2026-001.
- 3 Bring coding agents under policy** 2–6 weeks
Approve a sanctioned agent, block the rest by policy, and route exceptions through review.
- 4 Close offboarded AI access** 4–8 weeks
Wire directory sync so departed users lose AI-tool access automatically.

COMPLIANCE SNAPSHOT

 EU AI Act Partial · 12 gaps	 NIST AI RMF Partial · 8 gaps
 OWASP LLM Top 10 6 controls at risk	 ISO 42001 Partial · 9 gaps

WHERE YOU STAND

0 · at risk  100 · secure

What "64 / 100" means

Managed but exposed. Visibility across the fleet is strong (92% of findings AI-validated), but active data-leak paths to public LLMs keep you below the 75 threshold for a low-risk posture.

What's driving the score down

Four **critical** issues — API keys in public LLMs and vulnerable unregistered MCP servers — cap the score. Remove those and the posture rises sharply.

Trajectory

Improving. Up 6 points this quarter (58 → 64) as 11 high-severity items were closed. The direction is right; the critical tail is the blocker.

If you act on 1–2

Projected > 75 within 30 days, clearing every critical finding and moving Northwind into the low-risk band.

METHODOLOGY

Endpoint sensors continuously discover AI usage across the fleet. Inventory counts and detection sources are authoritative; an AI reasoning layer adds confidence-scored vulnerability analysis and remediation priority, and every finding is validated to reduce false positives. Posture is a 0–100 composite weighted by open-risk severity. Full evidence is in the accompanying Detailed Findings Report.