

The AI Cyber Defense Vendor Scorecard

Fifteen questions to ask every vendor – including us. Tuskira

- ☐ 01 What context does your AI reason over when it makes a decision, and how current is it?
- ☐ 02 Do you resolve entities across tools, or operate on each tool's data separately?
- ☐ 03 Does your context include our control state and detection coverage, or only exposures?
- ☐ 04 Show me a low-severity finding you escalated and a critical one you deprioritized, with the reasoning.
- ☐ 05 Do you validate exploitability in our environment, or score based on threat intel alone?
- ☐ 06 What actions can you execute, and through which of our existing tools?
- ☐ 07 Can you close exposure with a compensating control before a patch ships?
- ☐ 08 What runs autonomously, what requires approval, and who defines that boundary?
- ☐ 09 Show me the full evidence trail and audit log for one real verdict.
- ☐ 10 What happens when the AI is wrong, and which actions are reversible?
- ☐ 11 What are your reference customers' findings-to-actionable ratio, time-to-verdict, and time-to-close?
- ☐ 12 What does deployment require from my team, and when do we see the first validated result?

AND THREE QUESTIONS ABOUT THE AI ITSELF

- ☐ 13 Which model providers do you depend on, can we change them, and what happens if a model degrades?
- ☐ 14 Exactly what data leaves our environment, to which providers, and under what controls?
- ☐ 15 How is the AI system itself secured – prompts, credentials, agent identities – and where does it refuse to decide?

Scoring: a crisp answer with evidence = 1 point. A redirect to model quality or a chat-window demo = 0.
13+: you're looking at an operating model. Below 9: you're looking at a chatbot.