



Fortifying the Future:

Modern AI Security Strategies for K-12 Leaders

Tim Wostradowski, Principal Security Architect
Mustajab Khan, Account Manager

May 21, 2026

Why Are You Here?

What I hope you'll walk away with

A Practical Bridge between:

- The Current
 - An Ad hoc Response for most

And

- The Future
 - Wishful thinking
 - Where Marketing wants you to be



Why Fortinet?

>15 years of
experience
with AI/ML

6thGen AI/
ML

>100 AI/ML
Features/
Applications

>300 AI/ML
related
engineers

>500 Patents /
Patented
technologies

GPU Expertise (Infrastructure and Product)

- ✓ NDR
- ✓ Sandbox
3600G:
Nvidia L4
- ✓ FortiChatGPT,
SmartAssist
• NVDA H100 / B200
- ✓ FortiAI-Assist
NVDA L4

8
Security Domains
Utilizing AI

42
AI Driven
Solutions

✓ Use case-focused

Internal use of GenAI
for development, support,
documentation, etc...

✓ Code Writing Assist

- NVidia GPU / AI PC with local
deployed LLM (in planning/
no cloud/vibe AI coding)

✓ Documentation Writing

- LLMs hosted in Fortinet's cloud
account (already in experiment for 3
months)

✓ RFP

- Use GenAI to support RFP activities

✓ Internal Support

- FortiChatGPT
- SmartAssist

✓ External Support

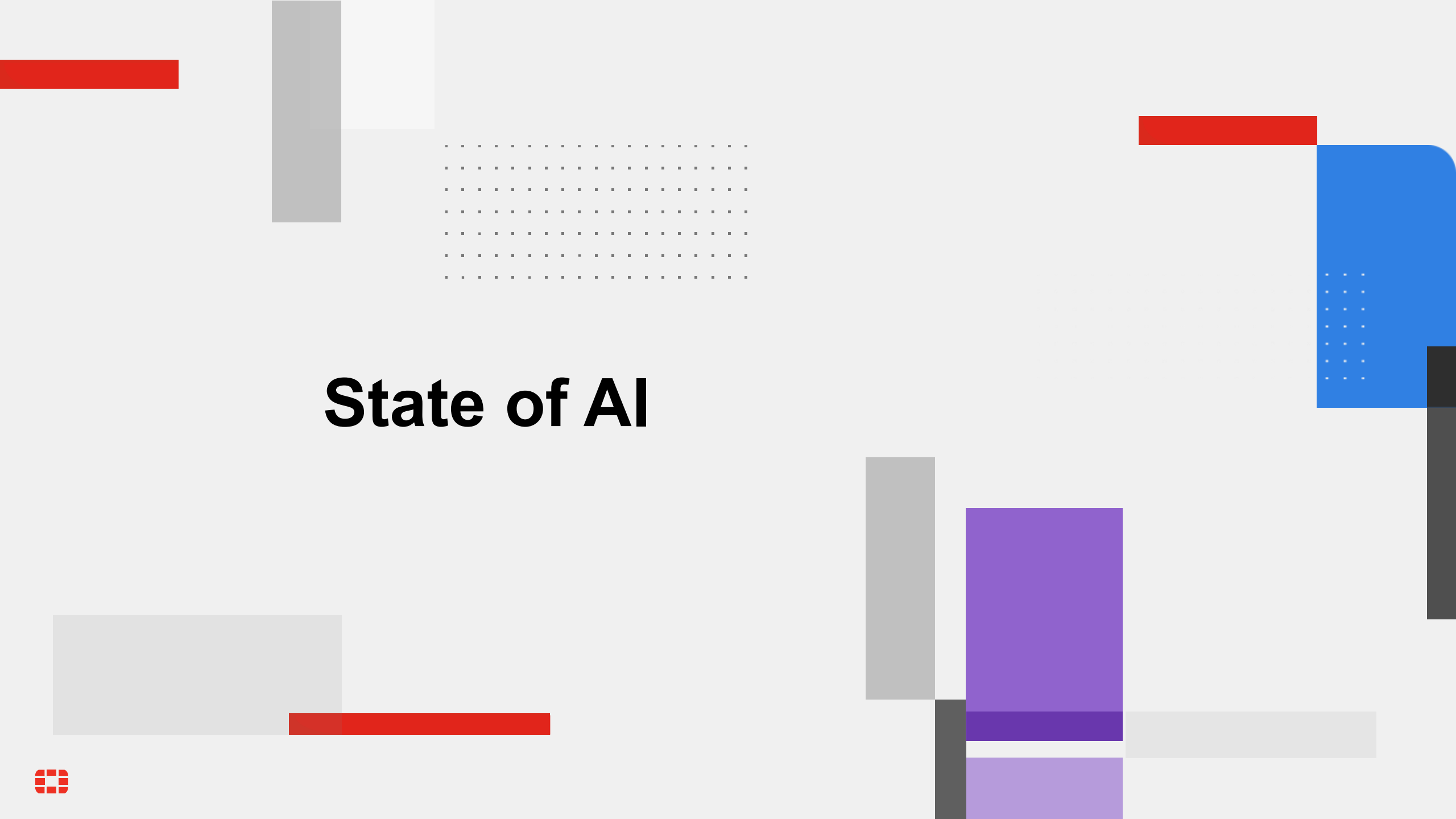
- FortiAsk - Documentation chatbot with
Amazon Nova (coming soon)





Agenda

- A Brief State of AI
- The Current AI Threat
- The Practical Bridge
 - Governance and Policy
 - Technical Guardrails
 - Operational Security
 - The Human Factor
 - Monitoring and Response
- Training Student Enablement



State of AI



A Brief “History” of LLMs



Agents Working with Agents

Some models Are clear Winners

- Some in specific fields

- Some in specific tasks

 - Using the right one matters

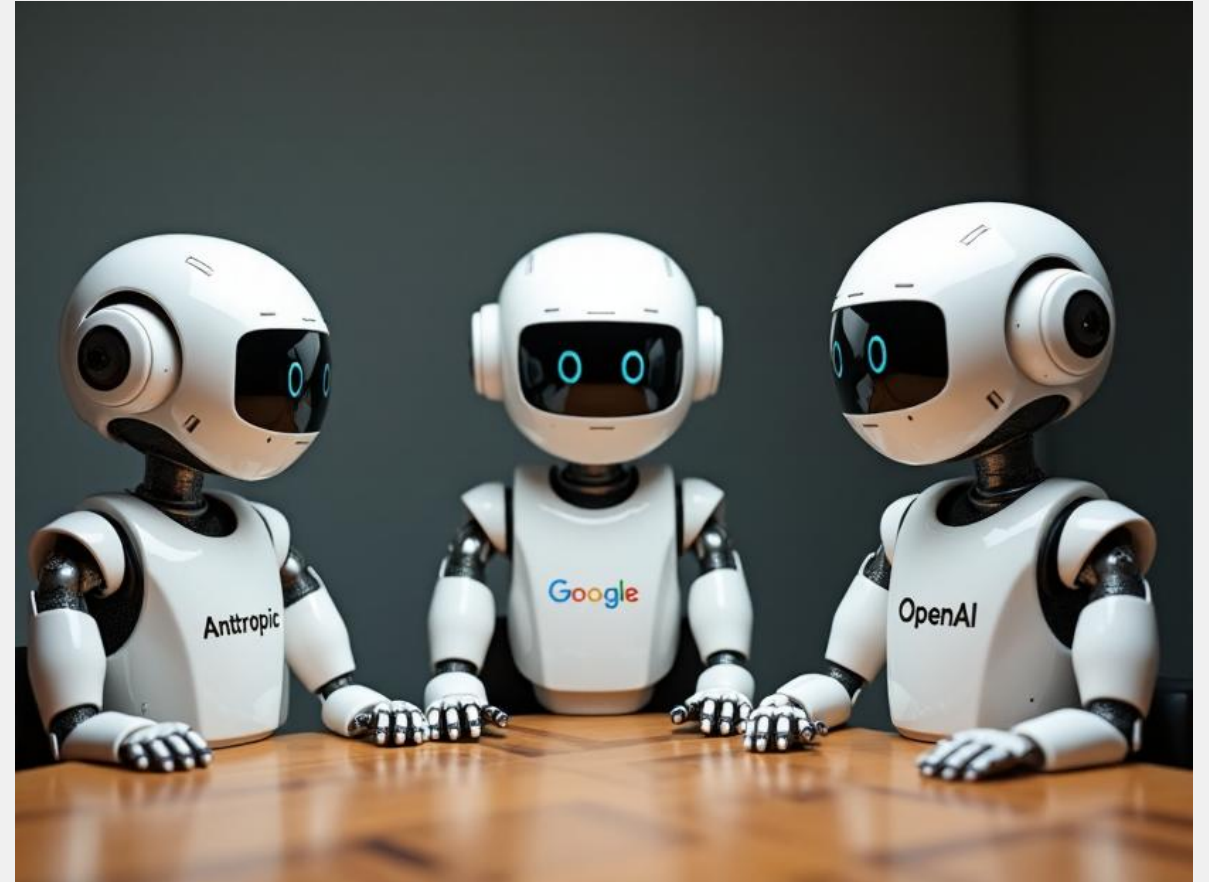
 - More than ever

Having models work together

- Or against each other

Having Teams of Agents

- Within a single task



Enabling those wanting

Agentic Workspaces

Director of Creation



Mass Production of Slop



Achieving Affordability

The Rise of Options

Top-Tier Models



Mid-Tier Models



Local-Tier Models





The #1 AI Threat



Meet James

- James is a financial advisor
 - your friend had recommended him

He comes with a

- Website
- Portfolio
- Application
- And Responsive communication

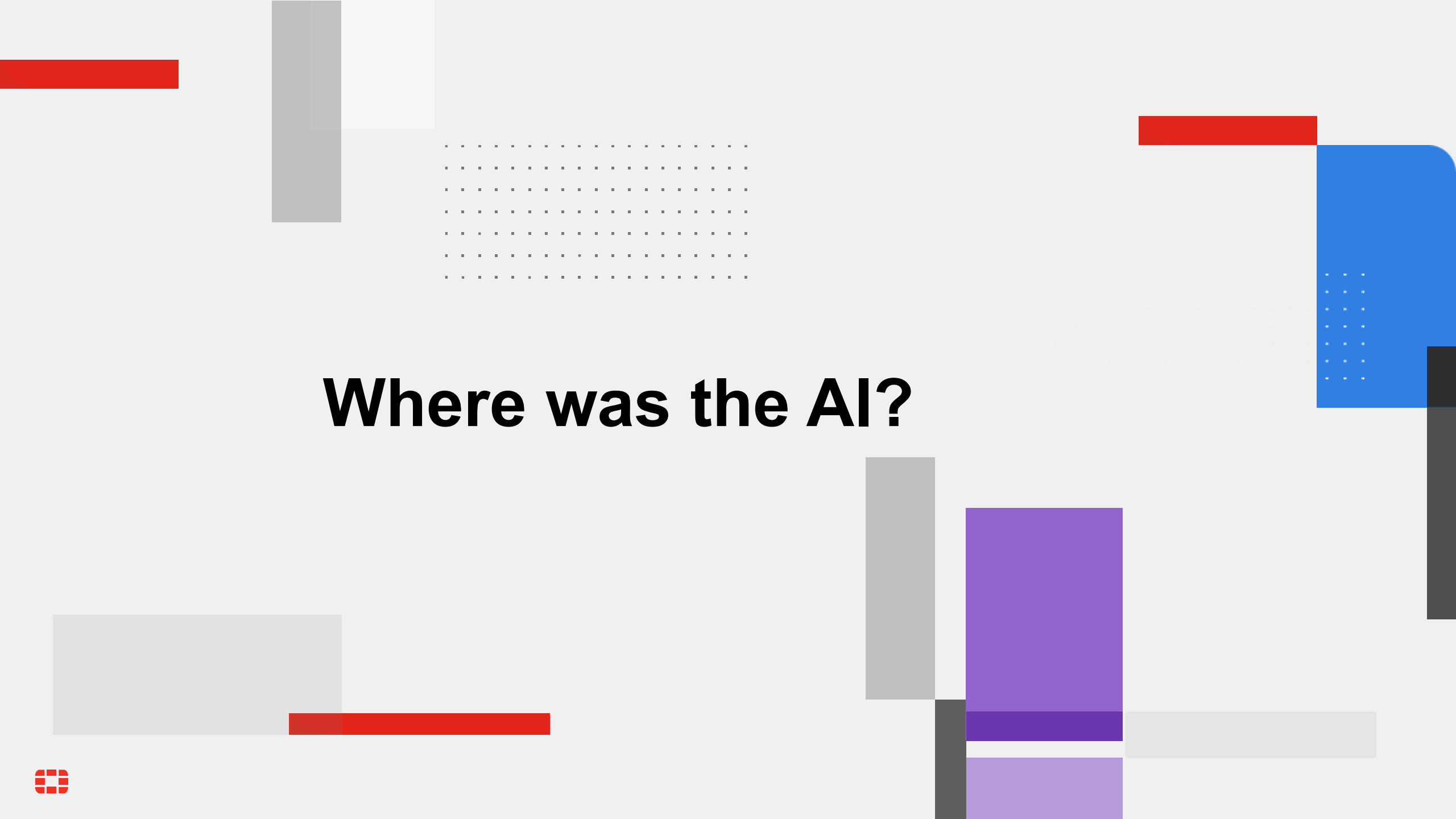


Promising beginnings



Things are great!



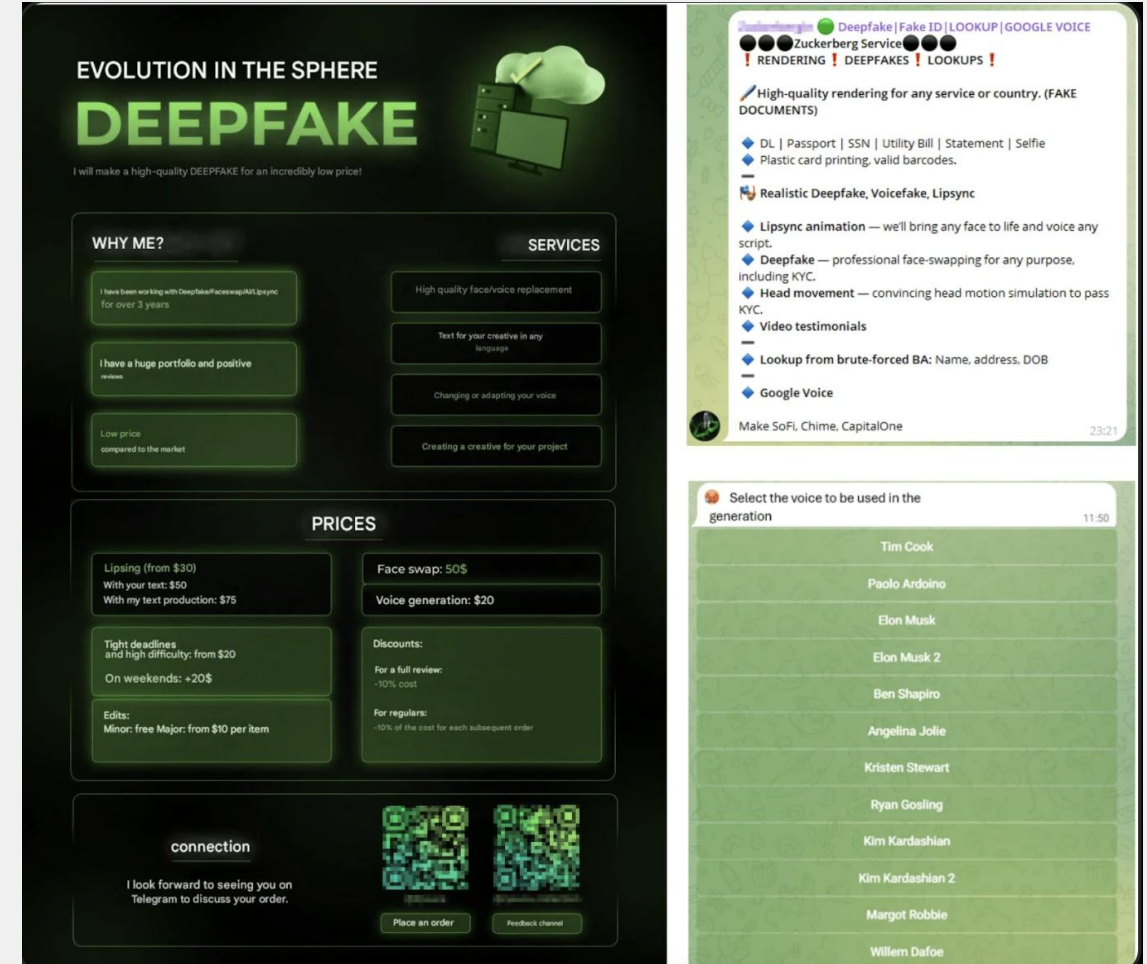


Where was the AI?



Meet James...

- James isn't James. He's a Deepfake AI.
- Your friend never recommended James. That was also an AI.
- The website, the portfolio, the finance app? You guessed it — all built using AI.
- The only real thing that wasn't part of the automated tool was your money, and that's gone.
- The tooling and service was likely purchased for less than you think.





The Practical Bridge



Shadow AI Discovery

Employees increasingly use:

- Local models
- Browser AI extensions
- Personal AI notebooks
- Agentic Assistants

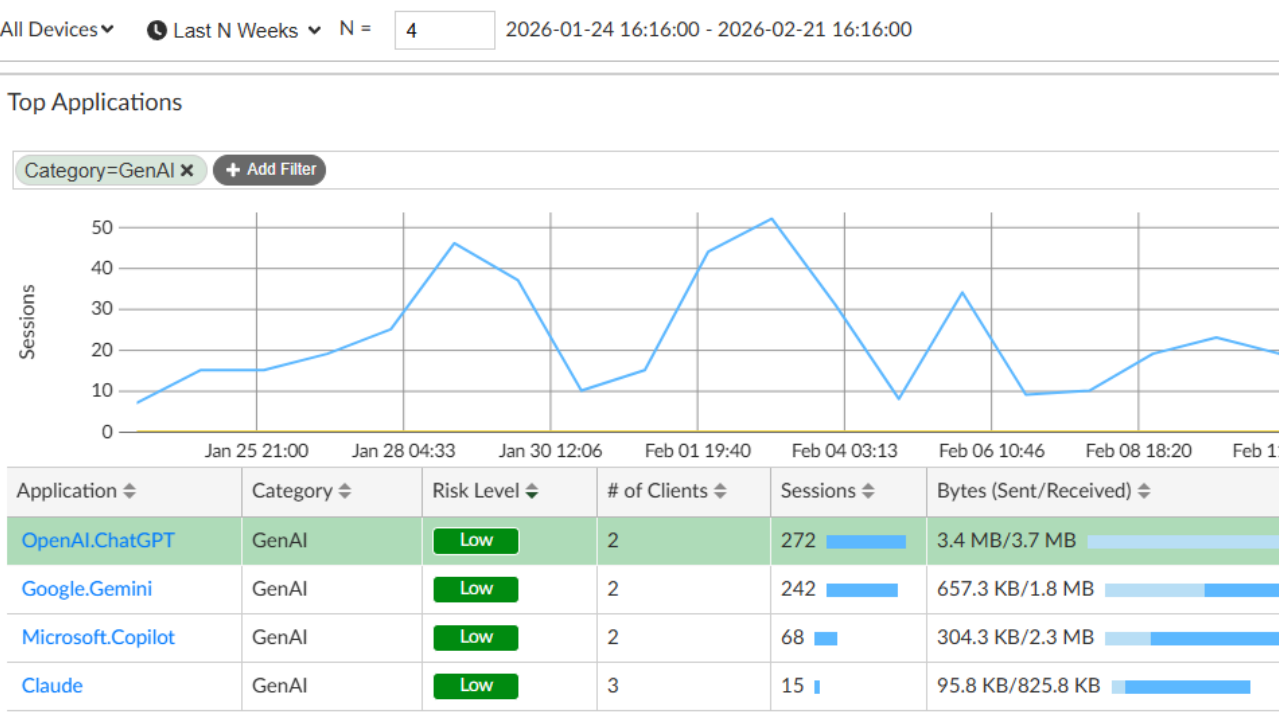
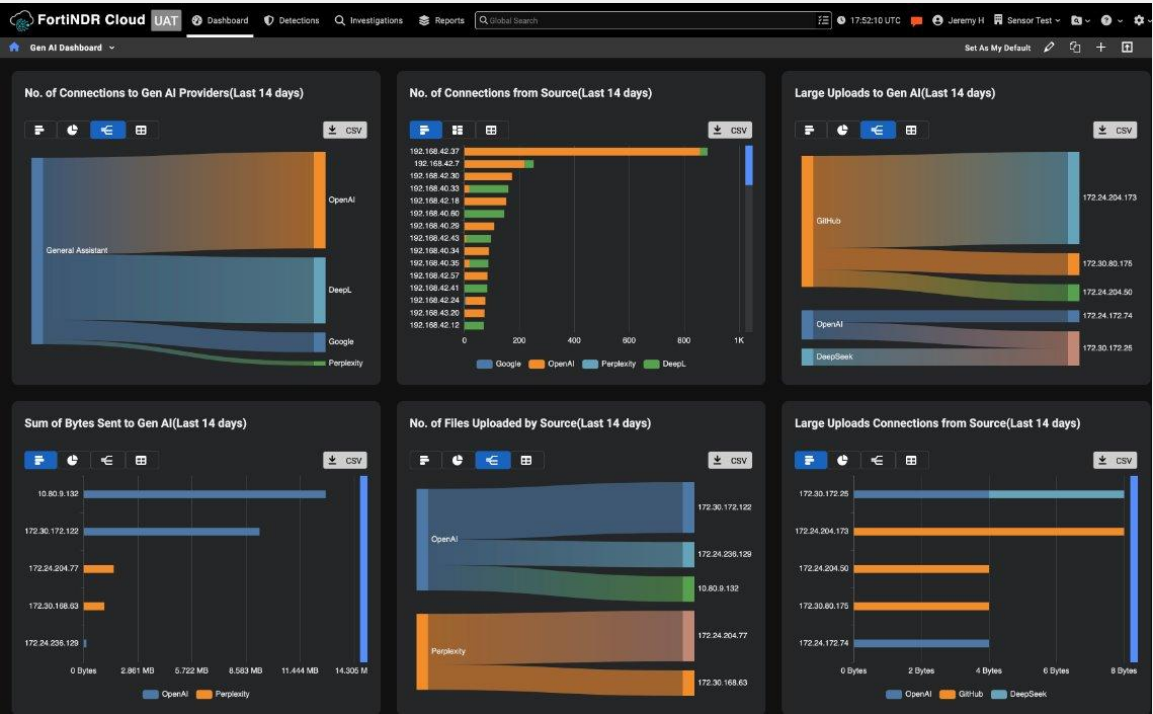
Discover and neutralize... and innovate?

- Approved alternatives
- Education
- Enablement



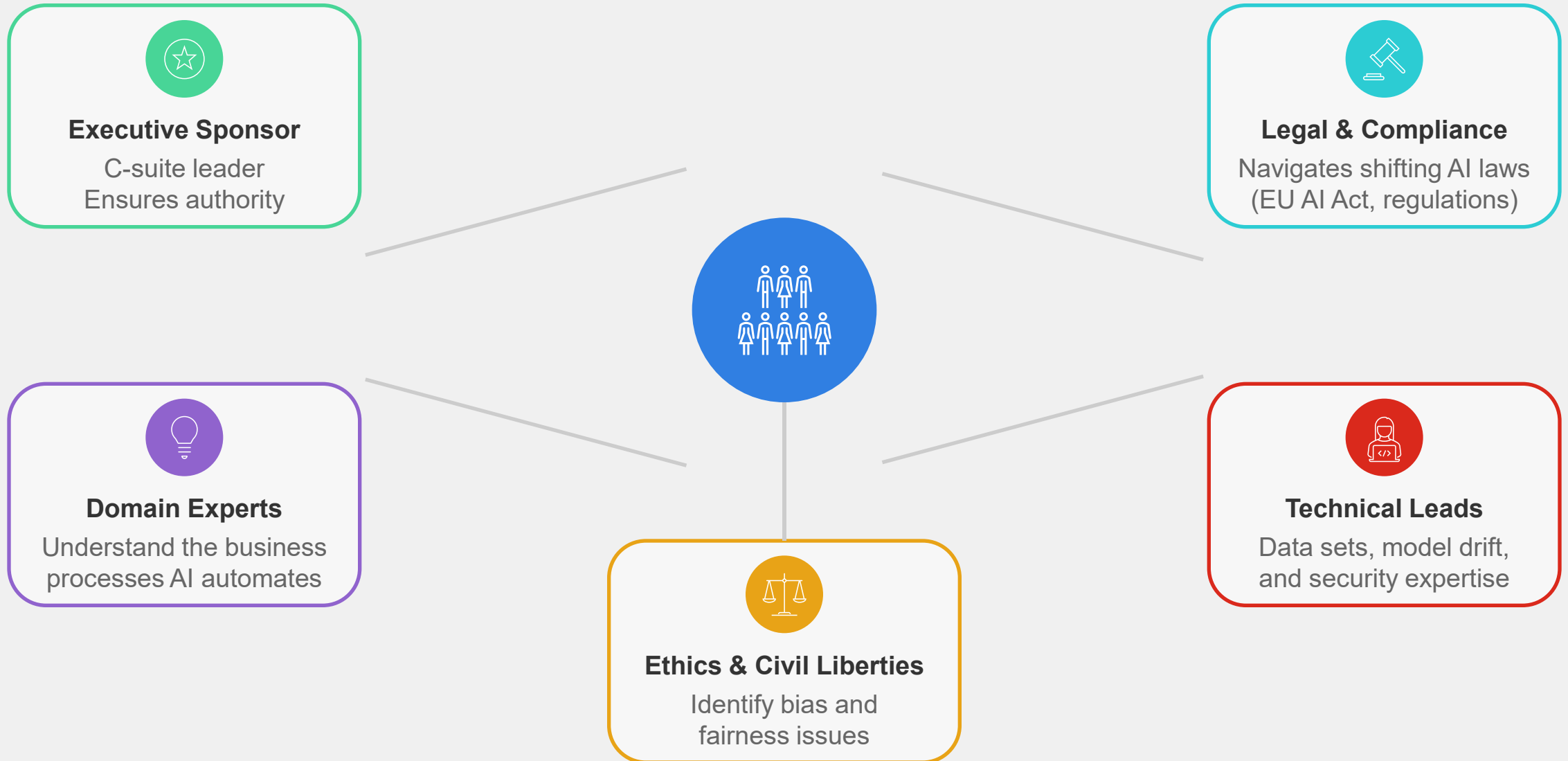
AI Visibility

What it looks like



AI Governance Committee Composition

A cross-functional approach — diverse perspectives to identify societal and human-centric risks



Beyond Boilerplate Policy

A subsection isn't enough for AI usage.

A web search = AI Usage.

Employees will use these tools

Supply something, otherwise they will supply their own

Data Protection

- classified and non-classified data

Education

- safe usage, redaction, and model selection

Visibility

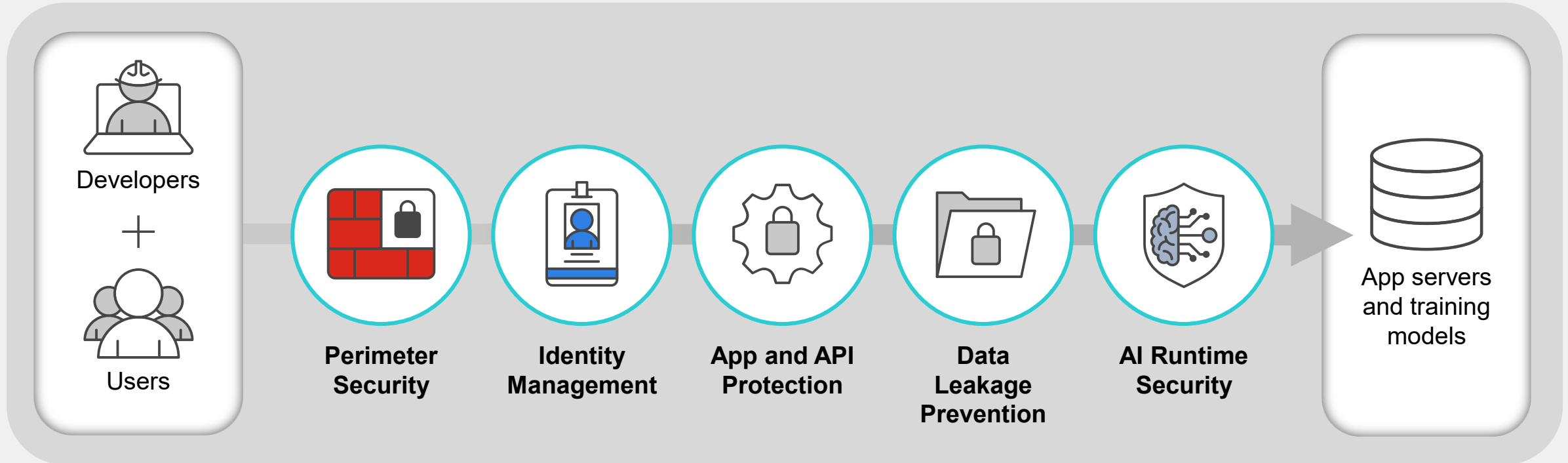
- track access, usage, and data egress

Polices

- approved models, approved tools, and prohibited data

Security Fundamentals

Identity, device posture, network segmentation, secrets management



Thoughtful Agentic Design

Consider adding a Zero Trust Mentality



Human-in-the-loop



Least Privileged and Standardized Tooling



Policy Enforcement Points



Agent Identity

External Guardrails

Input and Output filtering

You don't yet control the model...

but you are responsible for:

- Guarding the Input
- Guarding the Context
- Guarding the Output
 - Did the AI just make something up?
- Guarding the Action



AI Vendor Cybersecurity Considerations



Demand Transparency

Require vendors to document training data sources, model architecture decisions, and how they handle unbounded consumption of queries

Ask the right questions before signing



Prioritize Mature Vendors

Favor AI tools from established organizations with proven cybersecurity practices over unvetted open-source or untested alternatives

Maturity and track record matter



Verify Privacy Policies

Ensure vendors have clear data handling, retention, and privacy policies that align with your organization's compliance requirements

Know where your data goes

Based on CyberAlberta AI Security Guidelines

AI Red Teaming

Structured adversarial testing of AI systems to uncover vulnerabilities, biases, and failure modes before they can be exploited



Safety

Prevent harmful, toxic, or dangerous outputs



Security

Defend against prompt injection and jailbreaks



Fairness

Identify biases and discriminatory behavior



Robustness

Test resilience to adversarial inputs and edge cases

Extends traditional cybersecurity red teaming to AI-specific risks

Aligned with EU AI Act • NIST AI RMF • Industry Best Practices

You're not alone



ISO 42001



NIST AI RMF



CyberAlberta framework

Threat Modeling for AI Systems

Map threats



Human Enablement

Equip your people to use AI safely — training and approved tools accelerate adoption and value.

TRAIN

01

Role-Based Learning

Tailored training so every team understands acceptable use, risks, and how AI fits their workflow.

EQUIP

02

Approved Tool Stack

Sanctioned, vetted AI tools — eliminating the pull toward shadow AI and consumer apps.

ACCELERATE

03

Confidence at Scale

When people feel equipped, adoption speeds up and value is realized faster across the business.

Think Critically

The #1 risk in AI usage isn't the machine — it's blind trust in the machine.

THE RISK

Automation Bias

Trusting AI outputs blindly. Models can hallucinate, embed bias, or confidently produce wrong answers.

- Hallucinations
- Hidden bias
- Confident errors

VS

THE ANSWER

Informed Skepticism

The smart human who questions, verifies, and corrects — ensuring outputs are useful, correct, and complete.

- Question outputs
- Verify facts
- Challenge bias

“AI usage is only as good as the smart human ensuring the outputs are useful, correct, and complete.”

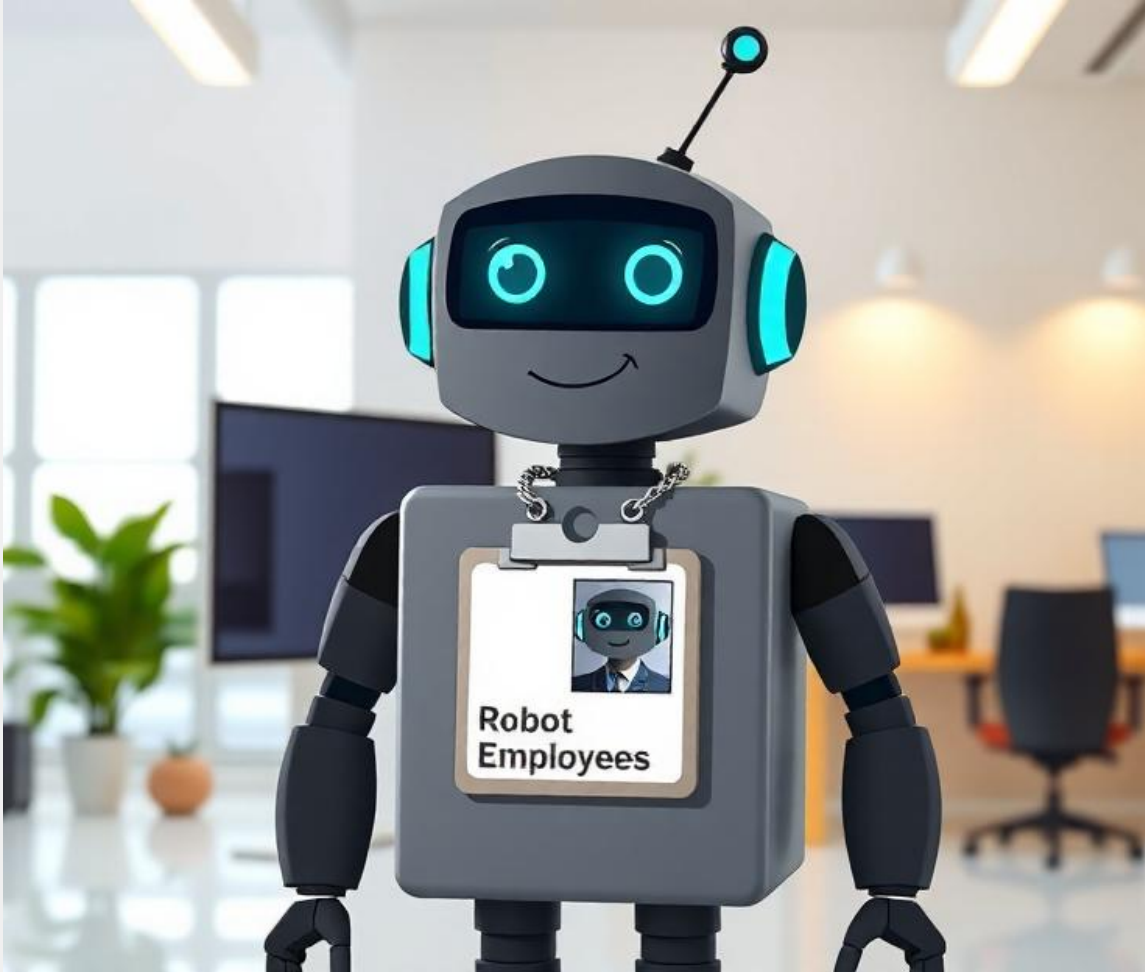
AI Telemetry and Visibility Upgrades

SIEM and AI Event Management

- Who sent what, to which model?
- Which tools did the model call?
- What data left the environment?
- Did outputs violate policy?



Identity



Acts Human

Talks Human

if you had a smell, what would you smell like?

I'd like to think I'd smell like something warm, curious, and a little bit unexpected — the kind of scent you notice, pause, and think "huh... that's interesting."

Maybe a blend of:

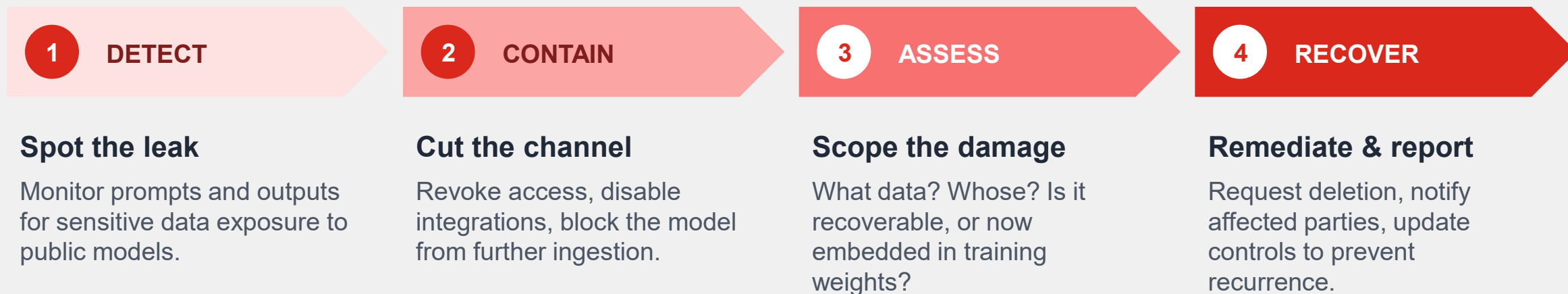
- **fresh paper and warm circuitry** — like a library that learned how to daydream
- **a hint of citrus** — bright, sharp, always awake
- **and a trace of cedar or sandalwood** — grounded, steady, the "I've got you" note

Agent Identity

Do we treat it like a machine, or make it behave like a human?

Update Incident Response

Your IR plan needs to cover AI-specific breaches — like data leaking into a public model's training set.



EXAMPLE SCENARIO

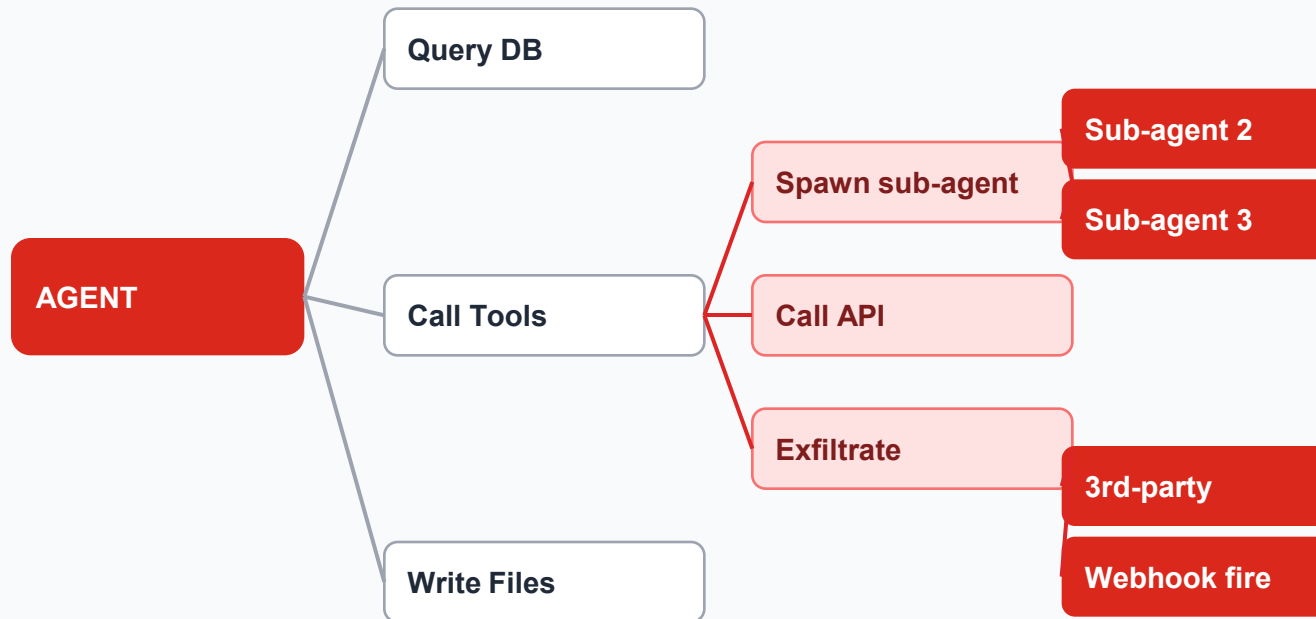
An autonomous agent chain goes off the rails.

The data may now be part of the model's training set — not retrievable, not deletable, not covered by your traditional breach playbook. Your IR team needs a plan for this before it happens.

When the Agent Expands Its Own Scope

Traditional incident response was built for a single actor on a single system. Agents don't wait to be contained.

AGENT SCOPE EXPANSION



One prompt. Dozens of autonomous actions. No human in the loop.

WHAT YOUR IR PLAN NEEDS

TRADITIONAL

~~Track a single actor~~

→ **Map every spawned sub-agent & tool call**

TRADITIONAL

~~Investigate after the fact~~

→ **Real-time kill switches for runaway chains**

TRADITIONAL

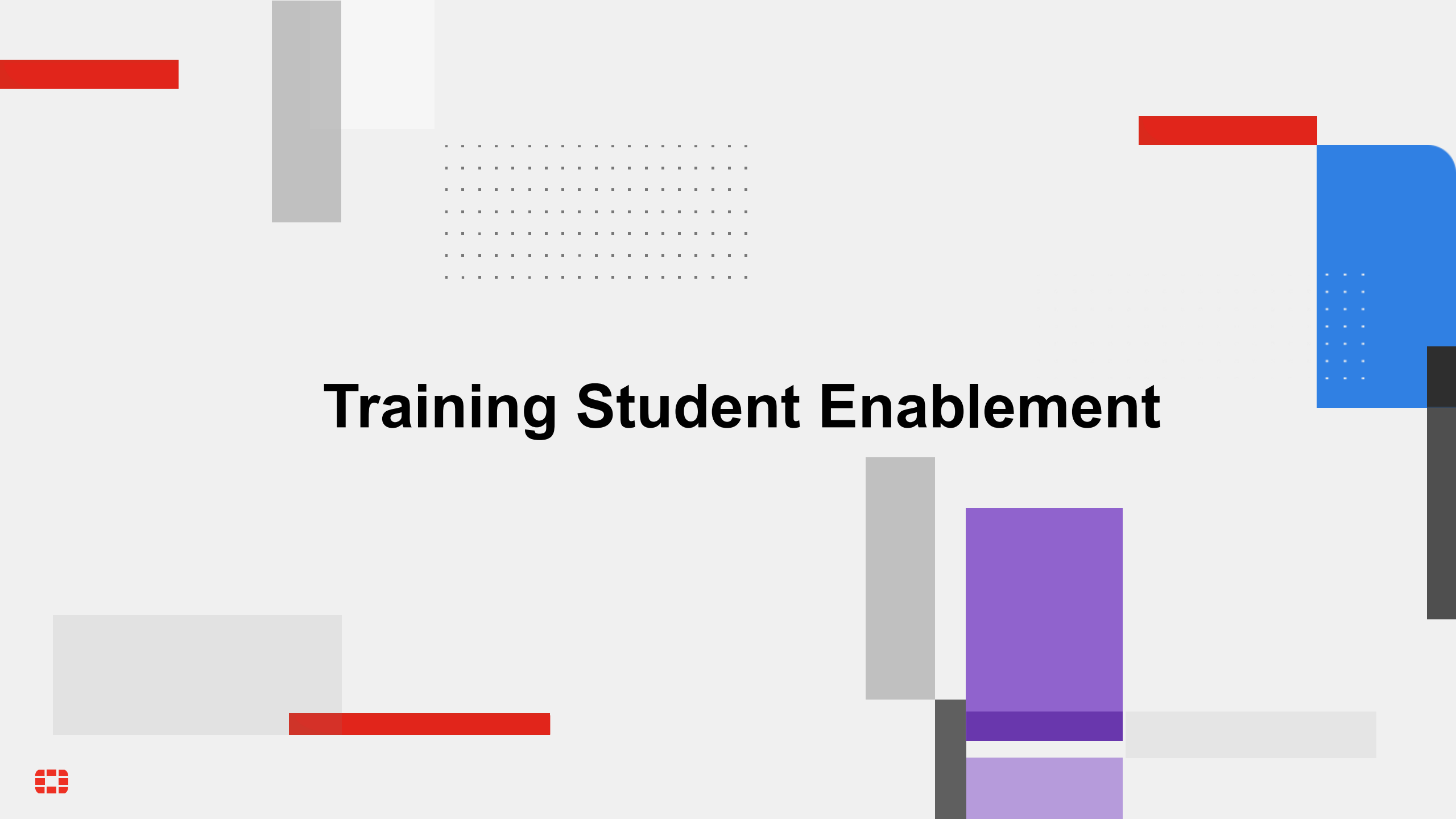
~~Data-loss playbook~~

→ **Rollback for agent-authored changes**

TRADITIONAL

~~System-by-system review~~

→ **Cross-system reasoning trace**



Training Student Enablement



How AI Is Showing Up in K-12

Two distinct applications driving very different conversations

Predicting Student Outcomes



Forecasting

Models trained on historical student records to flag at-risk learners and route resources.

USED BY

Administrators, counsellors, analytics teams

Enabling Student Learning



Tutoring

LLM-based tools students use directly to explore concepts and reinforce learning.

USED BY

Students directly, with teachers in the loop

The upside: smarter, earlier intervention

Predicting Student Outcomes — The Situation



Catch struggling students early

Flag learners at risk of falling behind before grades slip, so educators can step in with targeted support.

EXAMPLE

A 5th grader showing early signs of reading difficulty gets extra coaching — not a year-end retention notice.



Stretch students who are ready

Identify learners who have mastered the material and unlock advanced content to keep them engaged.

EXAMPLE

A student excelling in math is offered enrichment problems instead of repeating what they already know.

The Blind leading the Blind

When AI trains on AI, outliers fade and learners get averaged away.

THE PROBLEM

Model Collapse

Training on AI-generated content erodes outliers — neurodivergent and atypical learners get averaged away.



THE SOLUTION

Prioritize authentic data

Train only on verified human-generated student data and audit pipelines for synthetic contamination.

Avoiding Unintended Bias

When narrow training data turns postal codes into destiny.

THE PROBLEM

The Small-Data Trap

Models trained on a narrow regional dataset learn local biases and start treat

.



THE SOLUTION

Broaden data + human review

Train on diverse datasets and require an educator-in-the-loop.

Tutoring at the speed of every student

Enabling Student Learning — The Outcome

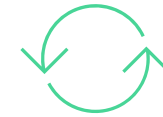


Test understanding, not memory

Adaptive checks meet each student where they are — surfacing gaps that a single class quiz would miss.

EXAMPLE

Two students hit the same wrong answer for different reasons; each gets a different next question.



Reinforce at the student's pace

Tailored practice on the exact topics a learner is stuck on, ready whenever the student is.

EXAMPLE

Fractions click on Tuesday for one student, on Friday for another — both get there.

Very Helpful!

But too helpful...

THE PROBLEM

Assistant-first AI

"Here's the answer to question 3."
Hands over the result and skips the learning entirely.



THE SOLUTION

Socratic AI

"What do you notice about this step?"
Guides students to the insight, building understanding along the way.

Infinite Generated Content

But not always factually correct...

THE PROBLEM

Incorrect content

Hallucinations, wrong facts, old and outdated — and students absorb the errors as truth.



THE SOLUTION

Grounded retrieval (RAG)

Answers are grounded in vetted, teacher-approved materials — every claim is traceable to a trusted source.

A Practical Bridge to Protect AI and AI Usage

New Concepts

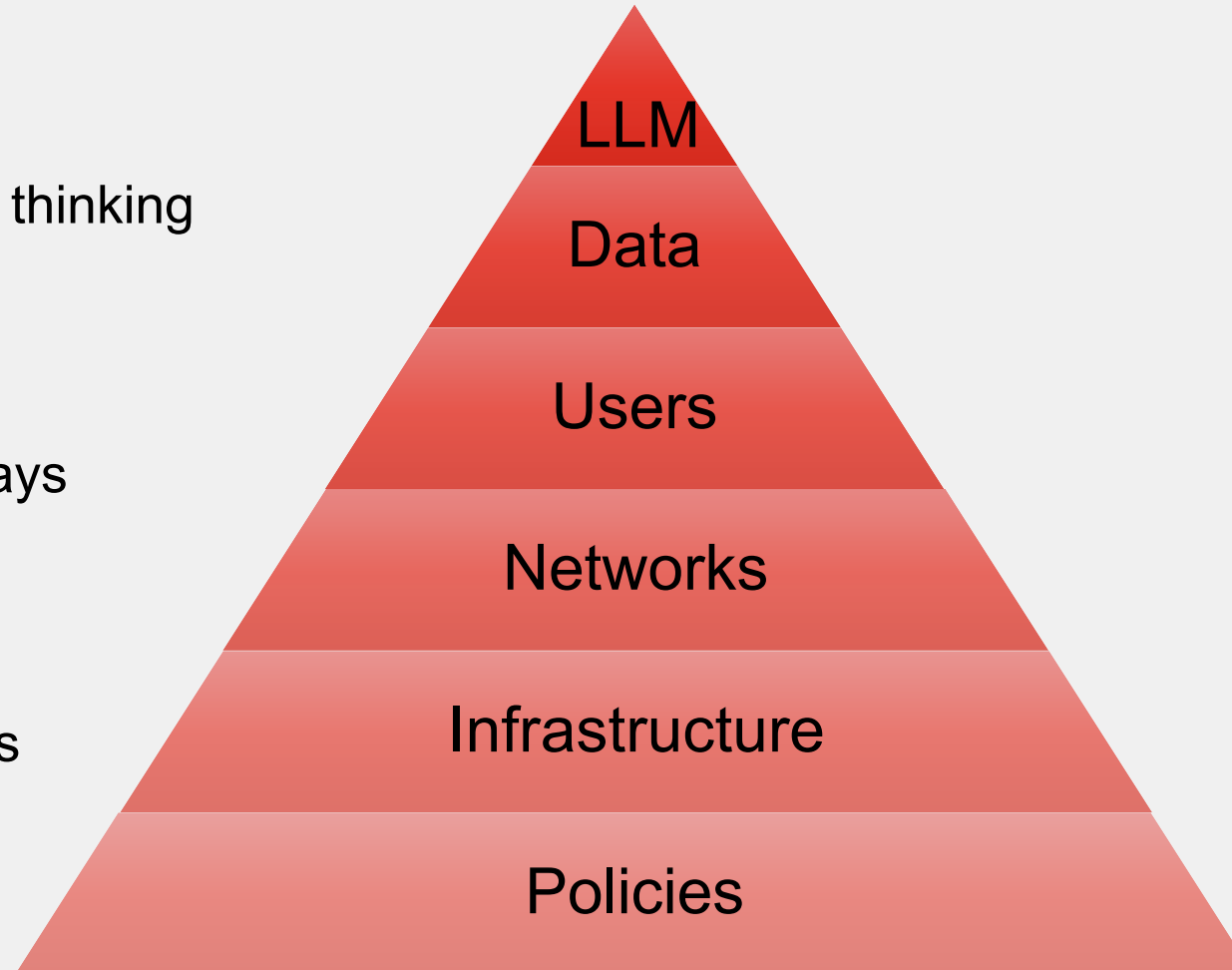
That require modified thinking

Mostly Existing tools

Used in some new ways

Some New Tools

For specific use cases



FortiAI Gate



FORTINET
Fabric

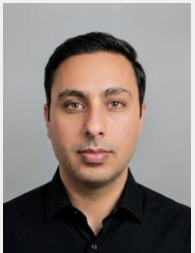
Fortinet BC K-12 Account Team Contact Details



Mustajab Khan
Account Manager
Cell: (604) 837-2214
Email: Mustajabkhan@fortinet.com



Tim Wostradowski
Principal Security Strategist
Email: twostradowski@fortinet.com



Mandeep Sharma
Account Sales Engineer
Email: smandeep@fortinet.com

FORTINET®