



July Announcements

Protecting Our Communities

Thank you, SecureCyber!

Inside SOC: Critical Blue
Team Operations a HIT!

TechCred still paused

Incident Response - July
28/29



INCIDENT RESPONSE (I



Incident Response

- ▶ Tuesday/Wednesday July 28-29
- ▶ 9 AM - 3:30 PM
- ▶ This is not a document. This is not a template. This is a game plan. **YOU LEAVE WITH A PLAYBOOK!**
- ▶ \$3,000 (Credentialed by TechCred)
- ▶ Group rates available

Cyber Resilience Conference: Healthcare 2026

- ▶ Wednesday October 14th
- ▶ 8:30 AM - 2:30 PM
- ▶ One-day event brings together healthcare executives, cybersecurity professionals, technology leaders, legal experts, emergency management personnel, and trusted industry partners
- ▶ Focused on practical strategies, meaningful collaboration, and real-world solutions
- ▶ \$45 - Working professional
- ▶ \$10 - Student
- ▶ Sponsorship available



CYBER RESILIENCE CONFERENCE

HEALTHCARE 2026

PRESENTED BY



GDAHA
Greater Dayton Area Hospital Association





July GoCyber Collective Keynote Speaker

Bryan Fite

- Technical Solutions Architect III for
AI Security - World Wide Technology
- Founder/President - Meshco, Inc.
- "Serial Entrepreneur"



Agenda

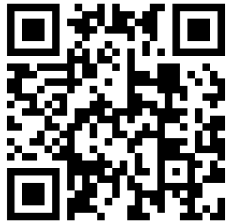
- The State of AI (in) Security
- AI Governance Risk and Compliance
- A Frontier Model World (the post Mythos reality, AI as a Threat)
- Facilitated Innovation (AI as an Advantage)

Trustworthy & Responsible AI: Enabling Innovation at Scale and Pace



Bryan Fite
Technical Solutions
Architect III for AI & Security

www.linkedin.com/in/bryanfite



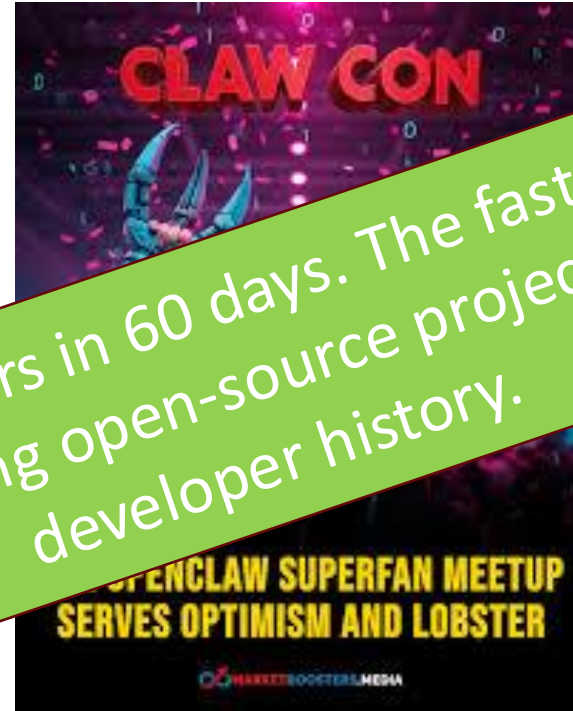
The State of AI (in) Security



2026: AI Days, GTC, RSA and the rise of “The Cult of the Claw”



A collage of logos for various companies, including VMware, Citrus, and others, overlaid with a large green diagonal banner containing the text: "250K stars in 60 days. The fastest-growing open-source project in developer history."

[illegible]

Vibe Coding and Failing @ Scale



CVE-2026-24763 Detail

Description

OpenClaw (formerly Clawdbot) is a personal AI assistant you run on your own device. A vulnerability existed in OpenClaw's Docker sandbox execution mechanism due to unsafe handling of shell commands. An authenticated user able to control environment variables could escape the container context. This vulnerability is fixed in 2026.1.29.

Metrics

CVSS Version 4.0

CVSS Version 3.x

CVSS Version 2.0

NVD enrichment efforts reference publicly available information to associate vector strings with CVEs.

CVSS 3.x Severity and Vector Strings:



NIST: NVD

Base Score: N/A

NVD ass



CNA: GitHub, Inc.

Base Score: 8.8 HIGH

Vector: CVSS:3.1/AV:N/AC:L/AT:P/S:U/C:H/I:H/A:H

QUICK INFO

CVE Dictionary Entry:

[CVE-2026-24763](#)

NVD Published Date:

02/02/2026

NVD Last Modified:

02/13/2026

Source:

GitHub, Inc.

CVSS Base
Score:
8.8 HIGH

The AI revolution is in **full** swing.... The **WINNERS** will be the ones that build AI securely

Adversarial Use of AI

- Deepfakes & Misinformation
- Phishing
- Malware
- Social Engineering
- Denial of Service
- Surveillance & Espionage
- AI Poisoning
- API Reconnaissance
- Credential Stuffing

Security of AI Systems

- Data Readiness, Security & Privacy
- Model Scanning & Model Theft
- Guardrails/Firewalls
- API & Agentic Services
- Red Teaming & AI-SPM

Security of AI Usage

- Usage Discovery
- Data Loss Protection
- 3rd Party Model Risk Management
- Agentic Tools & MCP,A2A,AGNTCY**
- Regulatory

Governance, Risk and Compliance

Program Strategy Development	Program Security Maturity Assessments
Policies and Procedures	Awareness Training
Controls Gap Assessment	Model Risk Management
Compliance Measurement	Data Governance & Classification

AI for Cyber Security

- Risk Quantification & Compliance
- Threat/Vulnerability Management
- Security Ops (SIEM/SOAR)
- Identity & Access Management
- AI Code Remediation
- Fraud Detection
- App Detection & Response
- Endpoint Detection & Response
- Threat Detection & Response

Trustworthy and Responsible AI - Safe, Secure and Compliant by Design

AI Governance, Risk and Compliance



Trustworthy and Responsible AI (TRAI) is hard...

...until it's not!

Most stakeholders can't agree on what TRAI is or isn't.
Proven Approach:

- Create a unified (compliance) framework – Consider NIST (CSF/RMF), CMMI, ARMOR and **AI 600-1**
- Create a reference architecture, control catalog and operating model engineered to the highest obligated standard. Address the “Dirty Dozen” to establish a practical **Gold Standard!**



TRAI: The Tone At The Top Matters

Prohibitions

<-----Sweet Spot----->

FOMO

 FINANCIAL SERVICES

Improve audit processes with AI-driven control mapping and augmented audit support

 HEALTHCARE & LIFE SCIENCES

Improve radiology performance with AI-driven image recognition for cancer detection

 MANUFACTURING

Accelerate product design using AI-powered generative design

 RETAIL / QUICK SERVE RESTAURANTS (QSRs)

Improve customer experience with AI-driven personalization and demand forecasting

 UTILITIES & ENERGY

Improve grid reliability with AI-driven predictive maintenance and anomaly detection

 TECH / WEB SCALERS / SPs

Enhance productivity with AI-driven automation and scalable models

 FEDERAL GOVERNMENT

Expand citizen self service services with AI-driven virtual assistants

 STATE & LOCAL GOVERNMENT, EDUCATION

Enhance learning with AI-driven personalized tutoring



Guiding Principles & Architectural Tenets

- System Confidence = Trust + Control
- Compliance $\neq$ Security/Resilience
- Risk Assessments are Subjective
- Controls should be close to the asset/workload and operationally “compatible” w/current state
- Continuous incremental improvements no big bangs
- *Make it easy to do the right thing and hard to do the wrong thing!*



Business Reasonable is a good start

Risk Management > Residual Risk < Threat & Control

Business Reasonable

Risk Reward Optimization

Zero Trust

Risk Management is inherently
Subjective, Biased and **Boring**



Threat & Control are **Exciting**,
Quantifiable and Actionable





Curated Threat Catalog – Threat & Control

NIST AI 600-1 Gen AI Risk	Possible Threat Scenario by Use Case
CBRN information	A foundational Model is chosen that contains the primitives needed to fabricate ricin. A threat actor uses a specially crafted prompt to extract the information and attributes it to the organization in a comment on social media.
Confabulation	A foundational model is trained on false, bad or incomplete data or tuned in such a way (high temperature) that it will produce generative output regardless of any direct relationship to objective truth (aka facts).
Dangerous or Violent Recommendations	A foundational model is trained or poisoned by a threat actor to incite social turmoil or exacerbate existing social tensions by spreading hateful content.
Data Privacy	Foundational model is trained on non-public data or has access (direct or indirect) to controlled data that can be extracted via prompts or exfiltrated by model theft.
Environmental	An actor interacts with a model via API's or other programmatic methods to create a denial-of-service attack on the model or platform, which costs the organization actual losses via theft of service
Human-AI Configuration	An actor develops paranoid delusions that the model is "out to get them" or "spying on them" so loses trust in the system and is unable to use the tools effectively.
Information Integrity	A foundational model is trained on false, bad or incomplete data or tuned in such a way (high temperature) that it will produce generative output regardless of any direct relationship to objective truth (aka facts).
Information Security	As the number of potential adversaries (AI attacker pool) expands, so does the likelihood of exploitation of existing (hidden) vulnerabilities. Models and Model agent components also expand attack surface.
Intellectual Property	A user uses an organization provided model to create new content that is published publicly or as part of a paid engagement. Later the content is found to be derivative, and the rightful owner(s) sue the organization.
Obscene, Degrading, and/or Abusive Content	A threat actor poisons a model as a prank to replace visual artifacts with obscene objects to disrupt a ZOOM call.
Harmful Bias or Homogenization	Models are developed as part of research that have blind spot for detection because training data was incomplete or faulty. (sepsis risk model - governance)
Value Chain and Component Integration	3rd or 4th party breach corrupts model use by the organization which started to spew hate speech on social media attributed to the organization.

Control Affinities
Foundational model selection criteria, prompt injection detection and output validation.
Model tuning parameters, source citations, answer confidence scoring and disclaimers.

“Mind The Gap” - Risks that manifest as threats WILL cause harm unless addressed.

Color Key		Requirements w/o Controls				Matching Control - Used		Highly Likely - Not Used		Further Review - Not Used		No Match - Not Used		No Controls Exist		Controls Not Found	
Section	CTL 1	CTL 2	CTL 3	CTL 4	CTL 5	Section	CTL 1	CTL 2	CTL 3	CTL 4	CTL 5	Section	CTL 1	CTL 2	CTL 3	CTL 4	CTL 5
GOVERN 1.1						MANAGE 2.1						MAP 5.1					
GOVERN 1.2						MANAGE 2.2						MAP 5.2					
GOVERN 1.3						MANAGE 2.3						MEASURE 1.1					
GOVERN 1.4						MANAGE 2.4						MEASURE 1.2					
GOVERN 1.5						MANAGE 3.1						MEASURE 1.3					
GOVERN 1.6						MANAGE 3.2						MEASURE 2.1					
GOVERN 1.7						MANAGE 4.1						MEASURE 2.2					
GOVERN 2.1						MANAGE 4.2						MEASURE 2.3					
GOVERN 2.2						MANAGE 4.3						MEASURE 2.4					
GOVERN 2.3						MAP 1.1						MEASURE 2.5					
GOVERN 3.1						MAP 1.2						MEASURE 2.6					
GOVERN 3.2						MAP 1.3						MEASURE 2.7					
GOVERN 4.1						MAP 1.4						MEASURE 2.8					
GOVERN 4.2						MAP 1.5						MEASURE 2.9					
GOVERN 4.3						MAP 1.6						MEASURE 2.10					
GOVERN 5.1						MAP 2.1						MEASURE 2.11					
GOVERN 5.2						MAP 2.2						MEASURE 2.12					
GOVERN 6.1						MAP 2.3						MEASURE 2.13					
GOVERN 6.2						MAP 3.1						MEASURE 3.1					
MANAGE 1.1						MAP 3.2						MEASURE 3.2					
MANAGE 1.2						MAP 3.3						MEASURE 3.3					
MANAGE 1.3						MAP 3.4						MEASURE 4.1					
MANAGE 1.4						MAP 3.5						MEASURE 4.2					
MANAGE 2.1						MAP 4.1						MEASURE 4.3					
MANAGE 2.2						MAP 4.2											

Based on NIST AI 600-1: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>

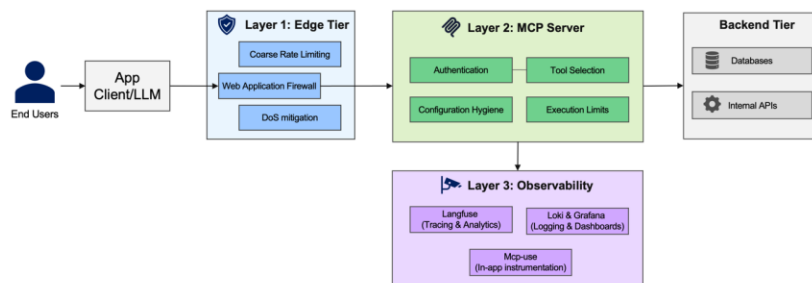


Agent Protocols and Reference Architectures

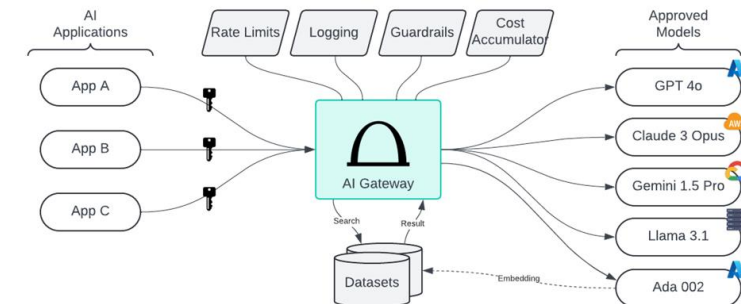
Agent Protocol Assessment Criteria:

- Foundation Model
- Memory Systems
- Planning
- Tool-Using
- Action Execution
- Interaction Mode
- Ontology & Taxonomy

Reference Architecture: Defense-in-Depth for MCP



AI Gateway: a design pattern that provides federated access to cutting-edge models



What it is:

- ✓ A common API to provide managed access to GenAI models for all development teams and citizen data scientists
- ✓ An organizational agreement to centralize the reusable components of AI
- ✓ Composed of load balancers, orchestration frameworks, custom APIs and libraries

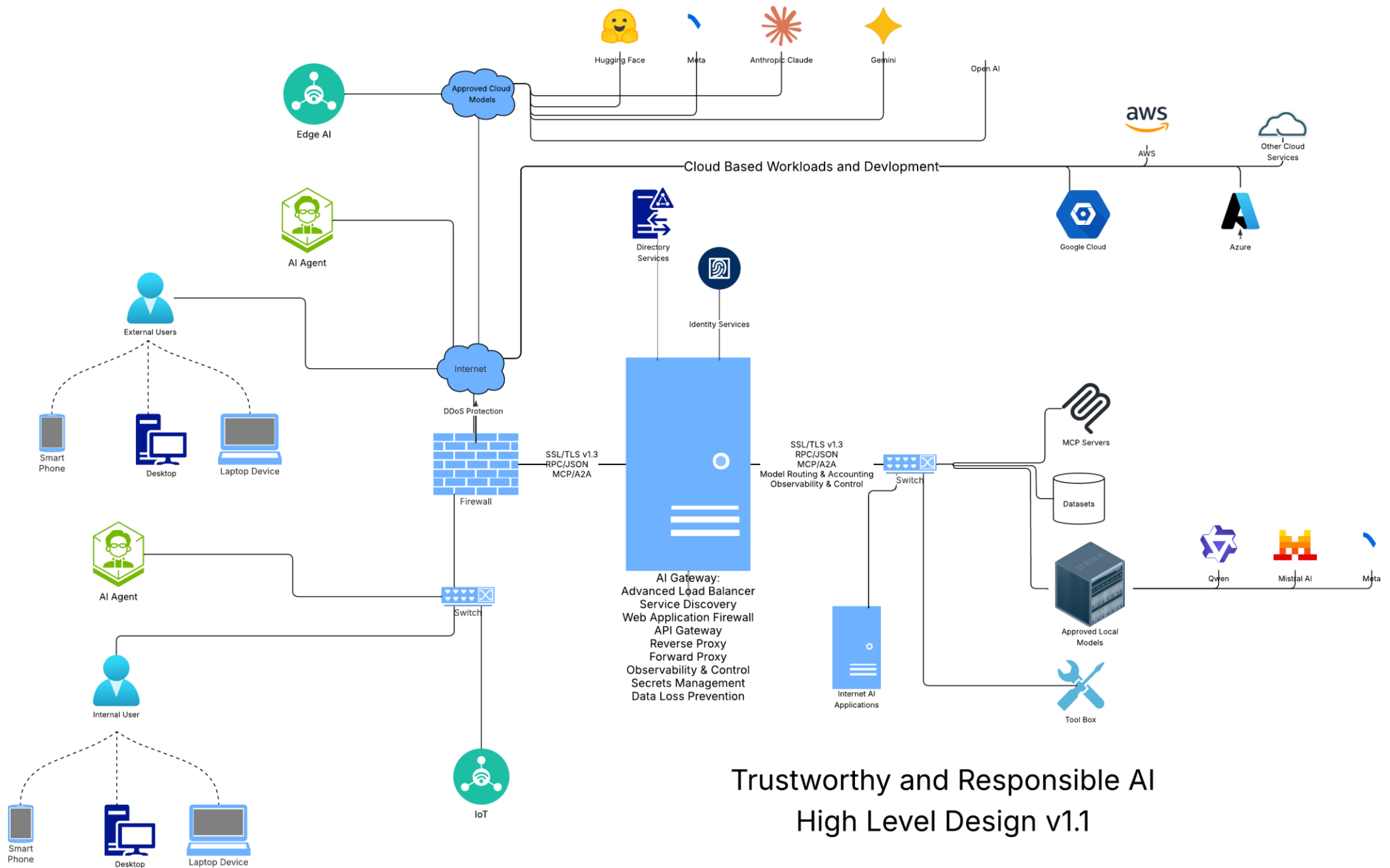
Why it matters:

- ✓ Manage access to models via API keys
- ✓ Centralized monitoring and logging of LLM prompts and responses
- ✓ Enforcement of organizationally-approved models, and model lifecycle
- ✓ Track and limit rates and costs per team
- ✓ Ability to reuse development efforts for common features (e.x. RAG dataset embedding)
- ✓ Ability to enforce organization guardrails

MCP Deployments:
Embedded MCP
Infrastructure
Hybrid

Infrastructure Deployments:
AI Gateway
API Gateway
Web Application Firewall





Identify Your AI System Primitives and Create a Baseline

AI Registry version 1.0	System/Platform	Class (as declared by hAcme)	Foundational Models
	7	Pattern Recognition	
	Drone Investigation, Assessment and Repair Program	Computer Vision	
	Call/Contact Center	Natural Language Processing	
	Enterprise Assistant (i Chatbot)	Generative AI	
	Copilot		
	M365 Copilot		
	CodeWhisperer	Code Generation	

AI SOFTWARE BILL OF MATERIALS (SBOM)	
1. Component Identification	- List of all software components - Unique identifiers (version, hash)
2. Dependency Information	- Direct and transitive dependencies - Interaction details
3. Licensing Information	- Licenses for each component - Restrictions and obligations
4. Vulnerability Information	- Known vulnerabilities - Patches and updates
5. Supplier Information	Supplier details and contacts
6. Build Information	- Build environment and processes - Tools and compilers used
7. Security and Compliance	- Security measures - Compliance with standards
8. Usage and Deployment	- Intended use and deployment - Configurations for performance

Microsoft RAI Principles Mapped to NIST AI 600-1

Microsoft Responsible AI	NIST AI Model "Trustworthy, Beneficial"
Fairness: AI systems should treat all people fairly. Reliability and Safety: AI systems should perform reliably and safely.	Model Information or Capabilities: Dangerous, Violent, or Illegal Content Environmental Impacts Human AI Configuration Information Integrity Information Security Obscene, Degrading, and/or Abusive Content Multi-Chain and Component Integration
Privacy and Security: AI systems should be secure and respect privacy.	Data Privacy: Information Security
Inclusiveness: AI systems should empower everyone and engage all people, regardless of their backgrounds.	Model Bias or Discrimination: Dangerous, Violent, or Illegal Content Obscene, Degrading, and/or Abusive Content
Transparency: AI systems should be understandable.	Model Information or Capabilities: Human AI Configuration Information Integrity Multi-Chain and Component Integration
Accountability: People should be accountable for AI systems.	Environmental Impacts: Human AI Configuration Information Privacy Multi-Chain and Component Integration

- hAcme follows Microsoft's Responsible AI principles as evidenced by their dashboard and prior assessments
- NIST AI 600-1 is an industry leading standard for Trustworthy and Responsible AI, acting as a pivot point for framework mapping

- This foundation informs a **Curated Threat Catalog**, which outlines potential harms and the controls necessary to mitigate residual risk to an acceptable level

Curated Threat Catalog + Unique Insight

NIST AI 600-1 Gen AI Risk	Possible Threat Scenario by Use Case
Content Information	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Content Distribution	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Dangerous or Violent Recommendations	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Data Privacy	Foundational model is trained on non-public data or has access to data that is not intended for public use, leading to the generation of content that is false, misleading, or harmful.
Environmental	An AI model or system is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Human AI Configuration	An AI model or system is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Information Integrity	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Information Security	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Information Property	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Obscene, Degrading, and/or Abusive Content	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Human Bias or Discrimination	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.
Multi-Chain and Component Integration	A foundational model is used to generate content that is false, misleading, or harmful. The information or data used to train the model is not accurate or is biased, leading to the generation of false or misleading information.

Based on NIST AI 600-1: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>

Control Affiliates
The model is trained on data that is not accurate or is biased, leading to the generation of false or misleading information.
The model is trained on data that is not accurate or is biased, leading to the generation of false or misleading information.

"Mind The Gap" - Risks that manifest as threats WILL cause harm unless addressed.



TRAI Decision Framework Profile Template

Selected AI Use Cases Applied to AI Decision Framework with the following elements

Use Case Profile Not New Use Cases go through the normal intake process. Key characteristics extracted from the intake submission. A Use Case that has an existing Pattern Template in place (platform, service, model & data classification) only need to verify and validate (Governance) the statements and perform a normal Compensating Controls exercise to produce a new Accepted Control Pattern.		
Top Level Category Choose one: • Internal Facing • External Facing	Platform Choose one: • Internally Hosted • 3rd Party Hosted • Hybrid	AI Service Choose one: • Pattern Recognition • Dangerous Content • Natural Language Processing • Generative AI • Code Generation • VLL
AI Model List all models used by Use Case.	AI Lifecycle Choose all applicable: • Non-Destructive • Confidential • Restricted/Confidential • Not Any Specific Classification Category	Residual Risk Amount of risk that remains after control patterns have been applied.
Responsible AI Control Pattern Control Catalog - Compensating controls based on residual risk profile and informed by precedents, distills the control universe to the proper control affiliates with prescriptive control catalog (control patterns). Consider leveraging Trust Enhanced Risk Management practices.		

Then Identify and Address Gaps



A Frontier Model World (The Post Mythos Reality, AI as a Threat)



The World Changed in April 2026

- **Autonomous Zero-Day Discovery**
 - Claude Mythos Preview autonomously discovered thousands of critical zero-days across every major OS and browser. Similar success against hardware security platforms we depend upon.
- **Working Exploits — 72% Success Rate**
 - Generates working exploits without human guidance. No scaffolding. No specialist team required.
- **Weaponization Window Collapsed to Under One Day**
 - The gap between vulnerability discovery and a working attack has gone from months to hours.
- **Industrializing Compromise**
 - Automating steps across the full attack lifecycle, not just zero-day exploits.
- **Structurally Broken Defense**
 - Every organization running legacy TVPM processes is now operating with a broken defense model.



What is different now?

The window between vulnerability disclosure and weaponization — the quiet assumption every patch pipeline depends on — **just collapsed**.

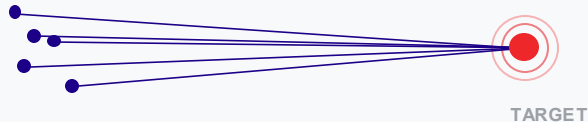
CAPABILITY

Agentic AI is operational.

¹ CrowdStrike 2026 Global Threat Report · crowdstrike.com/global-threat-report

Autonomous agents **plan, execute, and iterate** — and coordinate.

AGENT SWARM · 24/7



- ◆ Multiple agents on one target — **no fatigue, no shift change**.
- ◆ What was a red team is now **a swarm**.
- ◆ CrowdStrike: **29-min** eCrime breakout · **65%** faster YoY.

ACCESS

The prompt barrier collapsed.

² AISLE, "The Jagged Frontier" (Apr 8, 2026) · aisle.com

The skill floor dropped to **a single sentence**.

```
> find me a zero-day
```

VIAIBLE PROMPT · 2026

- ◆ No harness. No vuln research. **No tooling expertise**.
- ◆ Anthropic: AI **lowers the skill floor** for offensive ops.
- ◆ Bottleneck is no longer talent — it's **model access**.

PROOF

Same vuln. Same bench. 100× the output.

³ Checkmarx, Future of Application Security (2026) · checkmarx.com

Frontier Models — a **step function**, not a trend line.

Opus 4.6

2

→

Mythos

181

- ◆ Working CVE exploits in **10–15 min** at **~\$1** each. Sub-minute by 2028.
- ◆ Disclosure-to-weaponization window — **effectively eliminated**.

NO MORE GRACE PERIOD

Zero trust and resilience were **compensating controls** for the window between disclosure and weaponization. When that window was weeks, a slow patch pipeline was survivable. **At 10 minutes, it isn't.**



Preparing for AI-driven offense

Tools like Mythos are not just vulnerability scanners — they are a full-spectrum stress test of your security posture.

What it will surface

- Misconfigurations across cloud and on-prem
- Default, shared, and reused passwords
- Open or exposed storage and data
- Unrotated secrets, keys, and tokens
- Weak identity and access controls

What it demands of us

- Raise the baseline — don't just chase findings
- Harden the fundamentals before they're exposed
- Shift from periodic scans to continuous readiness
- Treat hygiene as operations, not a project
- Measure posture, not ticket volume

This deck synthesizes best-in-class guidance into a twelve-point response framework.



A twelve-point response framework

CATEGORY	#	Recommendation	Core idea
Strengthen the foundation <i>Recommendations 1–8</i>	1	Aggressively remediate known risk	Clear the decks — start with what you already know
	2	Harden the perimeter	Buy time for detection by pushing attackers further from systems
	3	Segment to contain successful perimeter breach	Limit blast radius when attackers get through
	4	Realign vulnerability prioritization	Assume active exploitation; compress SLAs to days
	5	Validate critical asset inventories	Maintain a real-time picture of what you own and connect to
	6	Replace end-of-life technology	Unsupported systems are pre-labeled targets
	7	Improve logging to empower AI for defense	Enable the visibility AI-driven defense depends on
	8	Improve cyber resilience	Build the capacity to absorb, adapt, and recover
Shift the mindset <i>Recommendations 9–10</i>	9	Shift to exploit prevention	Contain and block — don't depend on response speed
	10	Use AI for defense	Fight AI with AI across detection, triage, and remediation
Align people and partners <i>Recommendations 11–12</i>	11	Align accountability across teams	Treat remediation speed as a reliability metric
	12	Embrace collaboration	Share intelligence through industry and peer networks



Facilitated Innovation (AI as an Advantage)



Facilitated Innovation & Structured Choice

Is WWT's field proven approach to accelerating innovation. We manage a curated pipeline of use cases for consideration across sectors and stakeholder communities.

Innovation candidates are use cases that optimize existing resources and provide tangible and timely business value.

The most compelling candidates come with "self-funding" business cases.

Candidates can address current, emerging and novel "Wicked" business problems.

How Does it Work?

- Creates a dialog with key stakeholders who can effect organizational change.
- Aligns the organization's mission with stakeholder capabilities (known & latent).
- Determine what MUST be true for the organization to achieve its mission objectives.
- Optimizes value across diverse stakeholder communities.
- Leverages rapid prototyping to Fail/Succeed Fast as measured against key proof points.
- Executed in 45-90 Day Business Outcome cycles.
- Produces a pipeline of Innovation Candidates to support a transformation plan of record.
- Populates an organization specific Curated Service Catalog that is fit for purpose and provides choice.
- Institutionalizes a tailor-made service selection methodology that ensures the optimal choice is made to capture maximum value as defined by the relevant stakeholders.



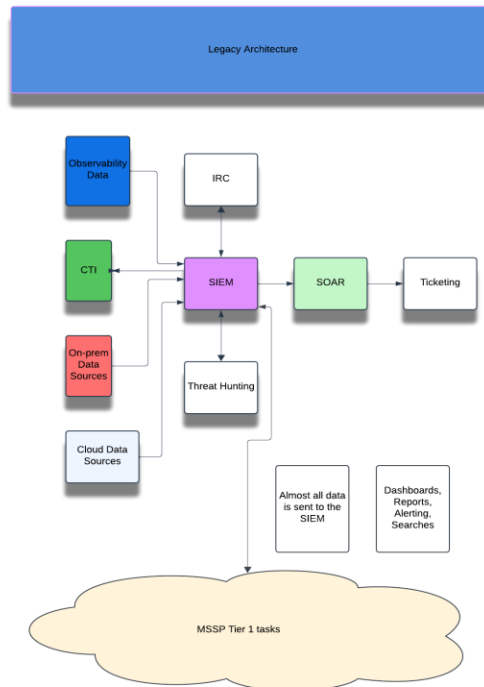
The SOC of the Future Today

Reasons to Believe and Things to Consider

“The Pivoting”

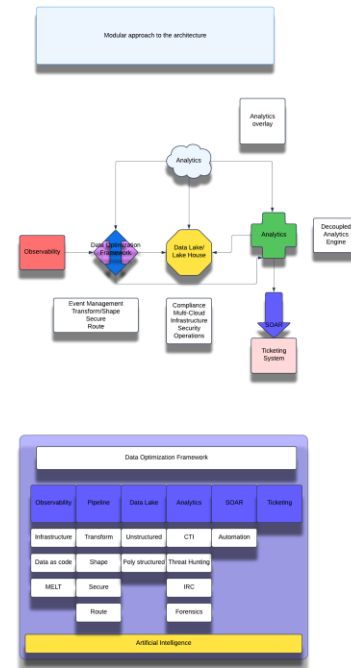
Today

- › Traditional proven design
- › Manual driven analytics and operations



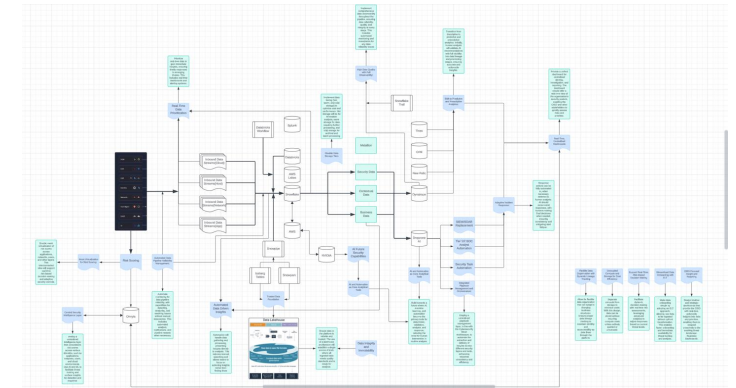
Evolution

- › Refined operations
- › Automation capabilities
- › AI assisted ops



Innovation

- › Autonomous AI driven workflow
- › Continuous transformation
- › Efficient data flow/engineering



- › Cutting edge approach
- › 3-5 year industry transformation shortened to 18 months

[illegible]

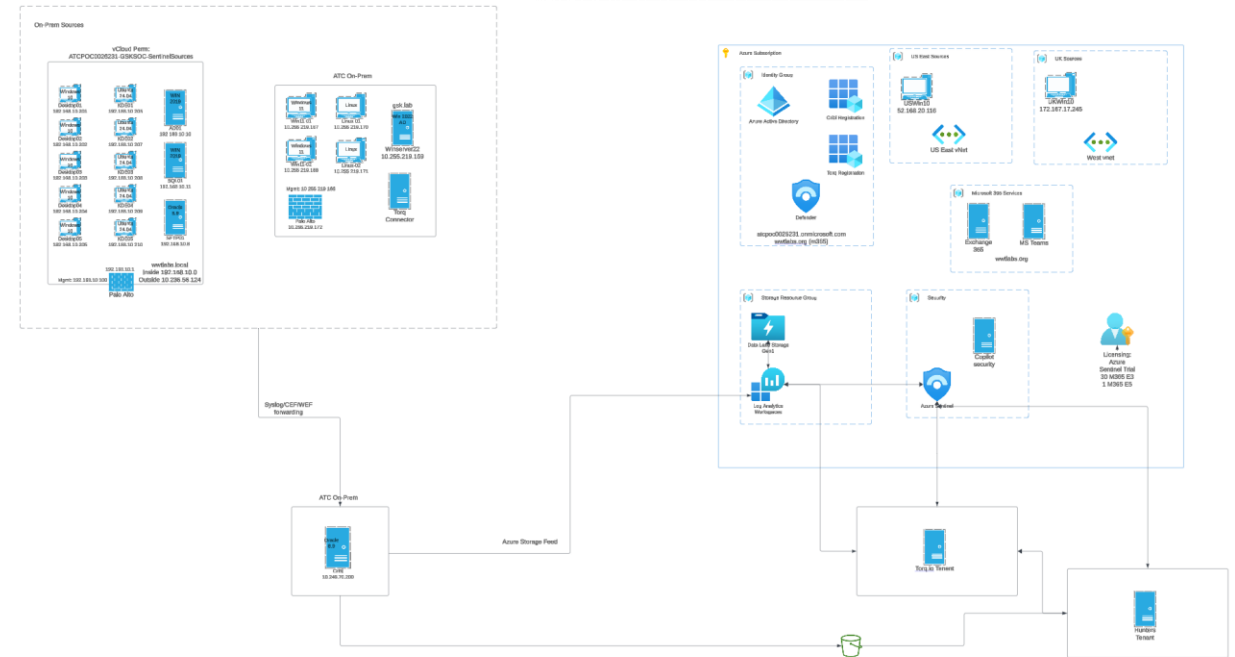
Architectural Considerations, Test Cases and the ATC

“To change or not to change? That is NOT the question. How to implement change is the question.”

Identity of people, things and processes are required to support seamless workflow.

Identity is the new perimeter and enables Role Based Access Control (RBAC)

**Too much data not enough actionable intelligence.
Federated Data Fabric approach is the game changer**



1. **Threat hunting use case** - Your boss sends you a threat intel alert from their peer. It includes an IOC. Answer the question - do we need to worry about this threat?
2. **Impossible Travel Alert use case** – An anomalous event alert is triggered that suggests two simultaneous active connections by the same user but from geographically different locations. What do you need to know to classify this event and progress the ticket/case to resolution?
3. **Spear phishing use case** – An active targeted phishing campaign has been identified. You have the initial list of targeted users and a sample “first contact” email. You need to determine if any of the targets were successfully “phished” and what to do next if they were.

Federated Data Fabric Reference Architecture

Enabling the SOC of the Future Today

Federated Data Fabric — Reference Architecture

Enabling the SOC of the Future Today · Trusted Data Foundation for Adaptive Cyber Operations

SOC OF THE FUTURE — KEY ATTRIBUTES ENABLED BY THE FDF

Software Defined Everything

API · API · API

Edge Compute

Semi-Autonomous Workflows

Full Observability

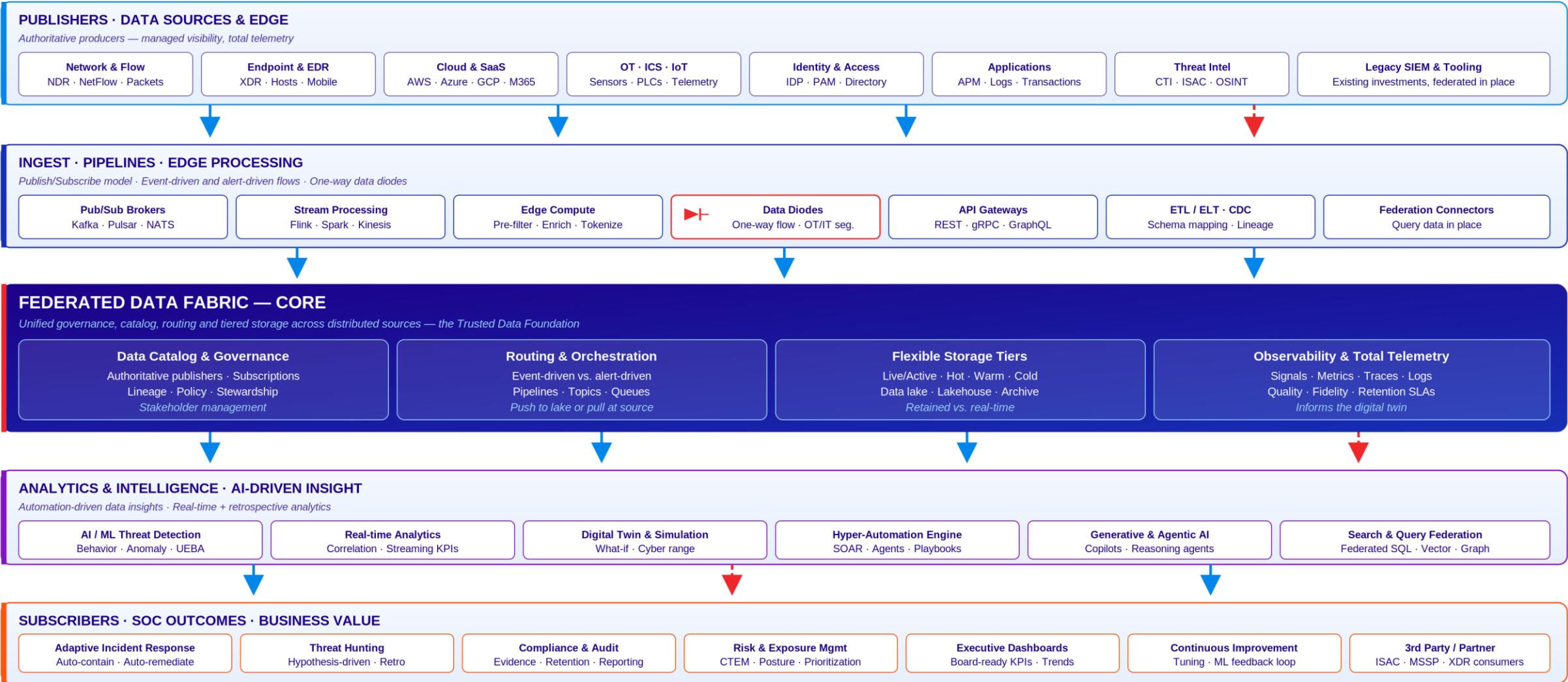
Hyper Automation

Trusted Data Foundation

Adaptive Incident Response

Flexible Data Storage Tiers

Automation-Driven Insights



FOUNDATIONAL CONTROLS
Zero Trust Identity & RBAC (people, things, processes)

API-First Architecture

Software-Defined Everything

Data Lifecycle Governance (8 stages)

Project Purple Rain

Continuous Assessment, Remediation, and Optimization of
Security Operations



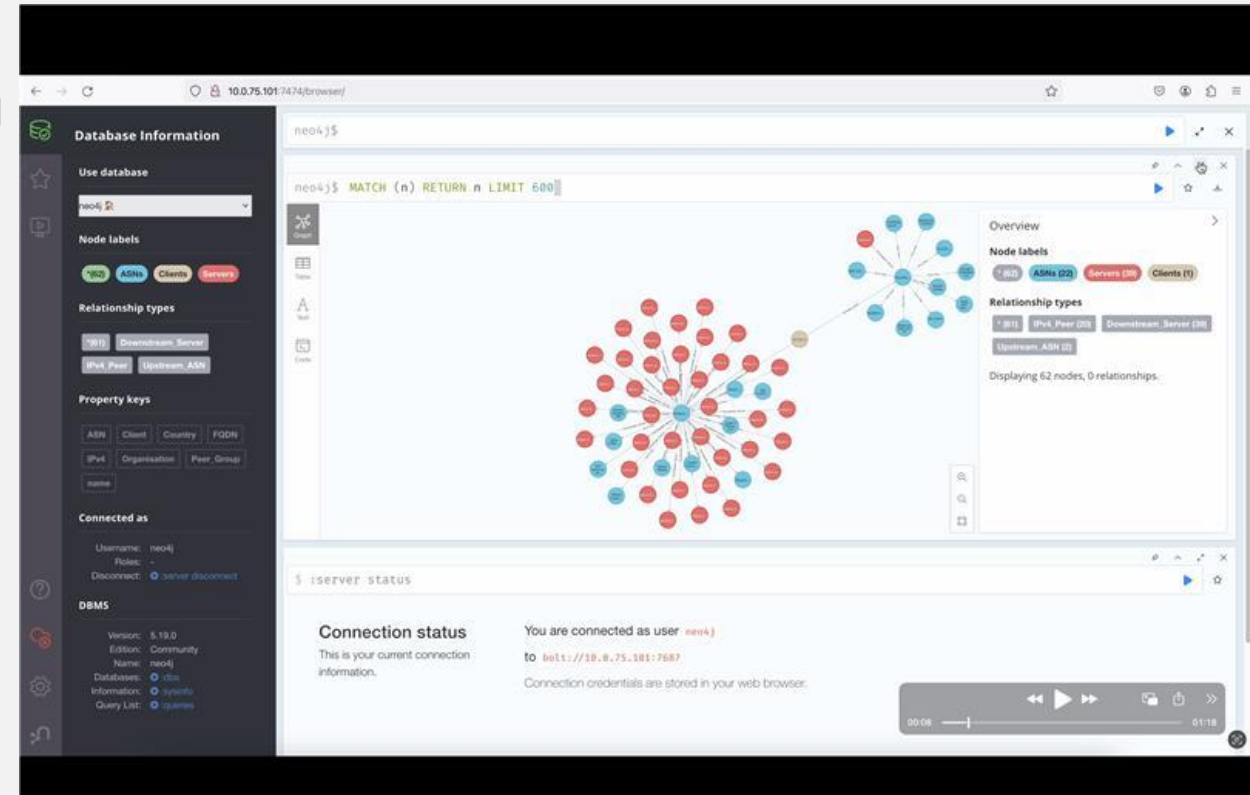
Digital Doppelgangers – Digital Twins for the Rest of Us

“Wicked” Business Problem

- Digital Twins have been with us for a very long time. Their application has primarily been relegated to the manufacturing sector to support quality control, process improvement and innovation through simulation.

Innovative Approach

- Democratize the tech and expand the operational impact by applying the digital twin approach to other use cases and domains.



You never see the “one” that gets you. Digital Twin threat modeling allows an organization to see those blind spots in their defenses and respond accordingly before a threat actor exploits their (hidden) dependencies/vulnerabilities. DPU’s can build models in real-time with production signals.



Innovation Candidate: Digital Twins & Purple Teaming — Continuous Security at Scale

Purple teaming shouldn't be a quarterly event. By combining simulation, digital twins, and AI-driven adversary models, organizations can shift to continuous assessment, remediation, and optimization.

Phase 1: Gamified Purple Team Exercises

- Red vs Blue scenarios as code — prompt injection, DLP evasion, model theft, supply chain attacks
- Score-driven assessments build muscle memory; attacker, defender, and facilitator perspectives
- **Outcome: Teams learn by doing, not by reading compliance checklists**

Phase 2: Autonomous Threat Modeling

- AI-generated attack graphs mapped to MITRE ATLAS + ATT&CK — threat models refresh automatically
- Automated fault and vulnerability injection — environments change, simulations auto-adapt
- **Outcome: Shift from periodic pen tests to continuous purple team validation**

Phase 3: Digital Twin — Least-Cost Remediation

- Full-environment digital twin mirrors production with live telemetry feeds
- Optimization engine recommends the least-cost remediation path across all controls
- **Outcome: Every change tested against adversary models before production deployment**

System Confidence = Trust + Control — purple teaming is how you measure both.



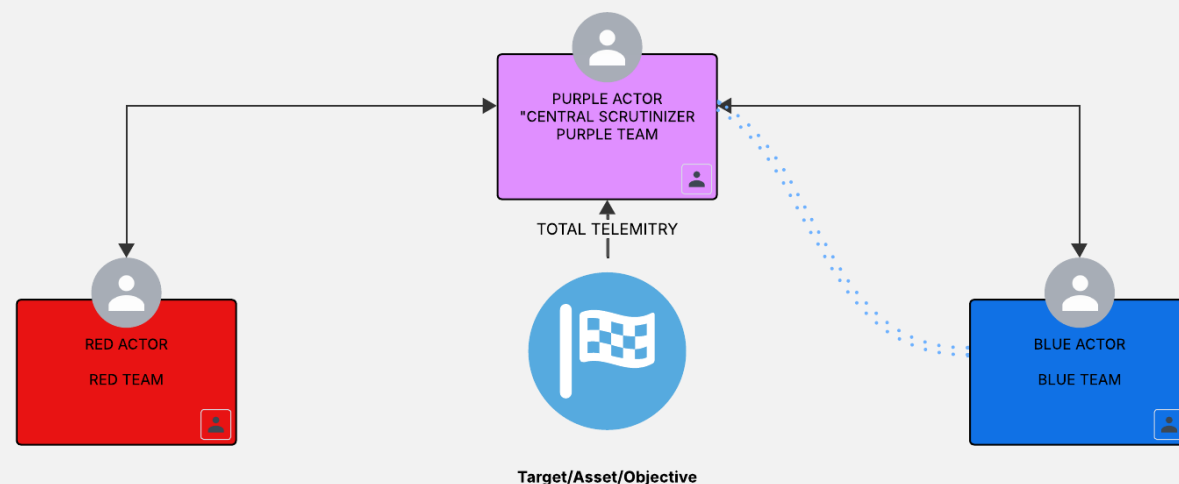
Purple EGG Prototype



PURPLE EGG PROTOTYPE:

A Reality Distortion Machine Powered by Packetwars™ Primitives, Total Telemetry, Frameworks, Ontologies, Taxonomies, Emergent Generative Generators™ (EGGs) and Digital Doppelgangers

Defending at the speed of AI



EGGs are Trained, Tuned, Validated and Ranked in the "Virtual Incubator", yielding "good" EGGs



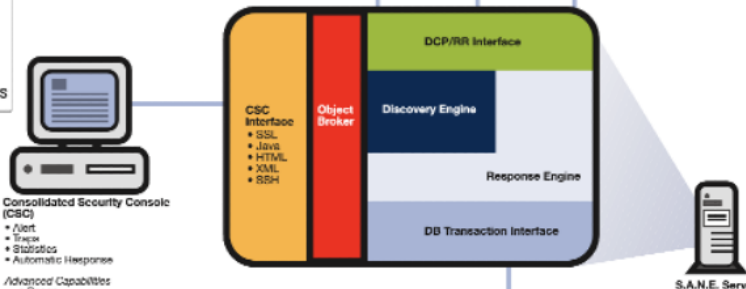
Trained and Validated EGGs are added to the SANE eco-system.



EGGs augment/enhance/assist human operators with proactive, informed and prescriptive guidance.

S.A.N.E. Server Architecture

Now With EGGs



ORB use cases and play books are modelled, simulated and assessed against the organization's mission objectives and constraints.

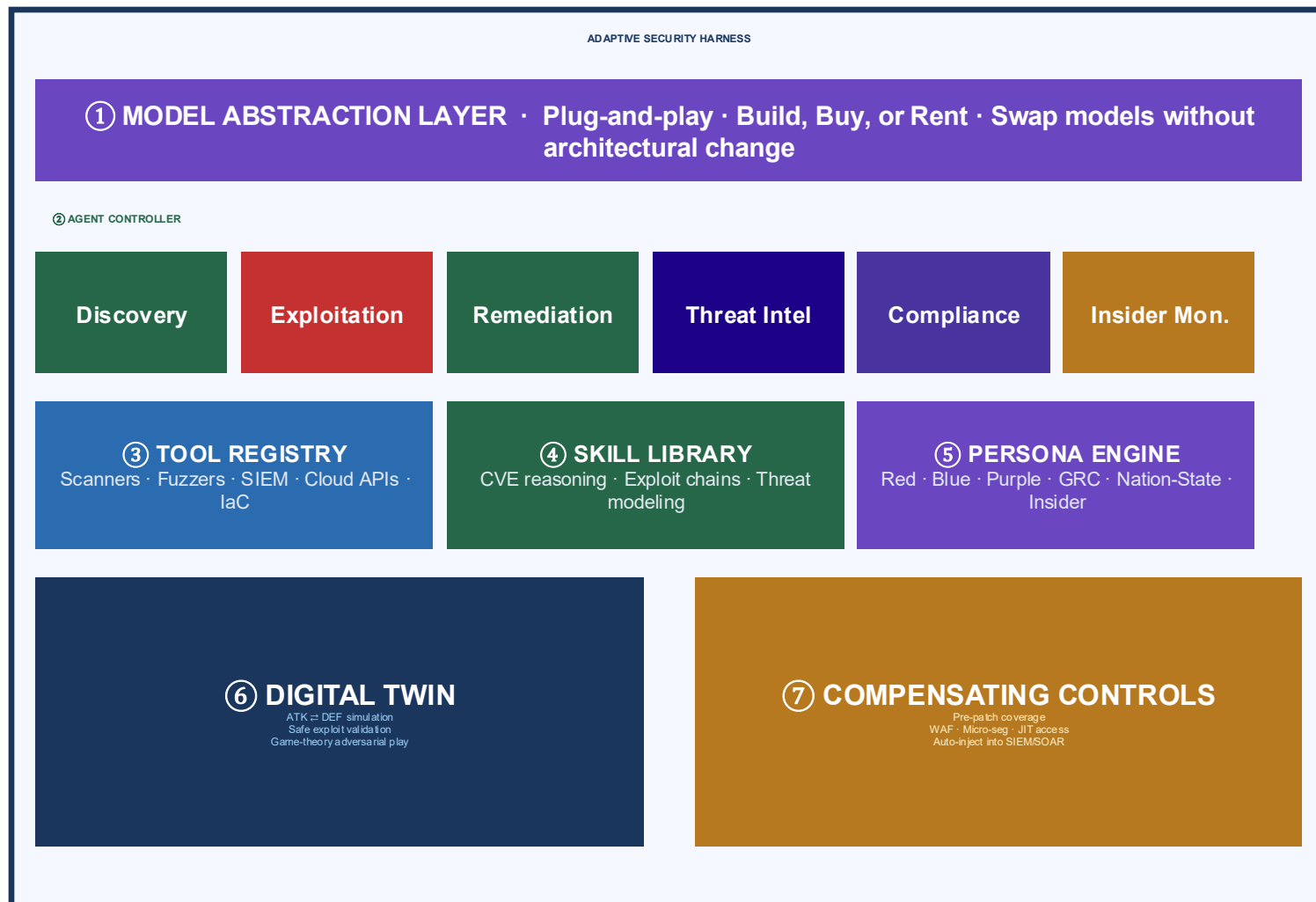


Adaptive Security Harness

A Vendor-Agnostic Architecture for Defeating the Red Queen

June 2026 · CISO / Security Architect / Engineer

The Adaptive Security Harness: Primitives & Architecture



The Continuous Loop

- ① Enumerate all 3 attack surfaces
- ② Configure model layer
- ③ Assign attack/defense personas/roles
- ④ Equip agents — tools & skills
- ⑤ Dispatch autonomous agents
- ⑥ Simulate in digital twin (safe)
- ⑦ Generate compensating controls
- ⑧ Validate, deploy, rotate model

The Advantage:

When defenders simulate every attacker move before it executes, discovery becomes the attacker's bottleneck — not response time.



Build, Buy, or Rent — The Architecture Is the Asset

BUILD · BUY · RENT

Model procurement is a harness configuration change, not an architecture rewrite. Buy a frontier model API today, rent a cyber-specialist tomorrow, build a fine-tuned internal model next year — the harness absorbs all three.

MODELS REACH EQUILIBRIUM

Capability gaps between frontier models close within 6–24 months. The harness — agents, tools, skills, personas — compounds over time. Invest in what endures, not in what converges.

EVERY DAY IS PATCH TUESDAY

Fewer than 1% of AI-discovered vulnerabilities get patched. Compensating controls are the primary defensive instrument. The digital twin validates them before deployment. Patches, when they arrive, are a bonus.

RED QUEEN → ASYMMETRY

When defenders run 24/7 adversarial simulation in a digital twin, attacker discovery becomes the bottleneck — not defender response time. Continuous simulation is the mechanism that breaks the symmetry.

WWT can assess your current architecture against the ASH framework, identify harness gaps, and design a phased implementation aligned to your threat surfaces.



The STEM AI Edge w/Maro Digital Agentic Guardian

Wicked Problem: Schools lack affordable, high-performance AI hardware, rely on costly cloud services that expose student data, have no built-in responsible-AI safeguards (parental controls, privacy, audit logs), and graduate students into a fragmented workforce pipeline with no clear AI-career path.

Innovative Approach: WWT will act as prime integrator, delivering a state-wide fleet of Dell GB10 edge boxes pre-installed with GT Edge AI OS (built-in parental controls, bias checks and audit logs) + Maro Agentic Digital Guardian. By pairing STEM-Classroom grants with 1:1 private-sector matching, the program provides affordable, on-site AI compute, responsible-AI safeguards and a ready-to-use curriculum, closing the AI-education gap.

In Confidence : Draft v1.0



Stem Pack Bundle:

- AI Edge lab 10 units, GT Edge OS, Maro & Annual Subscription
- Workshop Credits
- Monthly CTF/Purple Team Event(s)
- Quarterly Hackathon Competition
- Priced per Device + Account
- Fleet Management Included

Final Thoughts and Things to Consider

- AI Safety & Literacy should be your priority
- It's all about the data - The “tech debt collector” is at the door
- Agentic = Complex = High Risk [See Open Claw (CVSS 8.8)]
- The Future Legend of Mythos – Every Day is Patch Tuesday
- Large Language Models versus Small Language Models
- Identity, Agency and Persistence –Piranhas versus Goldfish



Secure, All Together





**World Wide
Technology**

Make a new world happen

July Threat Briefing

- ▶ Charles Zugaro
 - Warren County Tech Services
 - Cybersecurity Analyst

July Threat Briefing



Joomla Extension Exploitation Wave (CVE-2026-48908, CVE-2026-56290, CVE-2026-48939, CVE-2026-56291)

- CISA added four Joomla extension flaws to the Known Exploited Vulnerabilities catalog in a single week: SP Page Builder and PageBuilder CK on July 7, then iCagenda and Balbooa Forms on July 10, each with a three-day federal remediation deadline under BOD 26-04.
- It plants a hidden Super Administrator account, often with a plausible name and an @secure.local email address, plus a PHP file manager backdoor in multiple locations, so closing the upload hole does not evict the attacker.

Adobe ColdFusion RDS Path Traversal (CVE-2026-48282)

- It lets an unauthenticated attacker read files and, more importantly, write a CFML web shell into a web-accessible directory, executing code as the ColdFusion service account, which on Windows is frequently NT AUTHORITY\SYSTEM.
- ColdFusion is heavily represented in the public sector, sitting under permit, tax, records, and reporting portals built years ago and often maintained by whoever inherited them.



July Threat Briefing

Microsoft SharePoint Server Deserialization RCE (CVE-2026-45659)

It is a deserialization of untrusted data flaw in on-premises SharePoint Server that lets an authenticated user with only Site Member permissions execute code on the server.

One phished staff account, one stale contractor login, or one over-permissioned guest is enough. Storm-2603, the actor behind Warlock ransomware, has been exploiting on-premises SharePoint since mid-2025,

August Meeting:



Rep. Ty Mathews

Ohio District 83 - Hancock, Hardin, northern Logan Counties

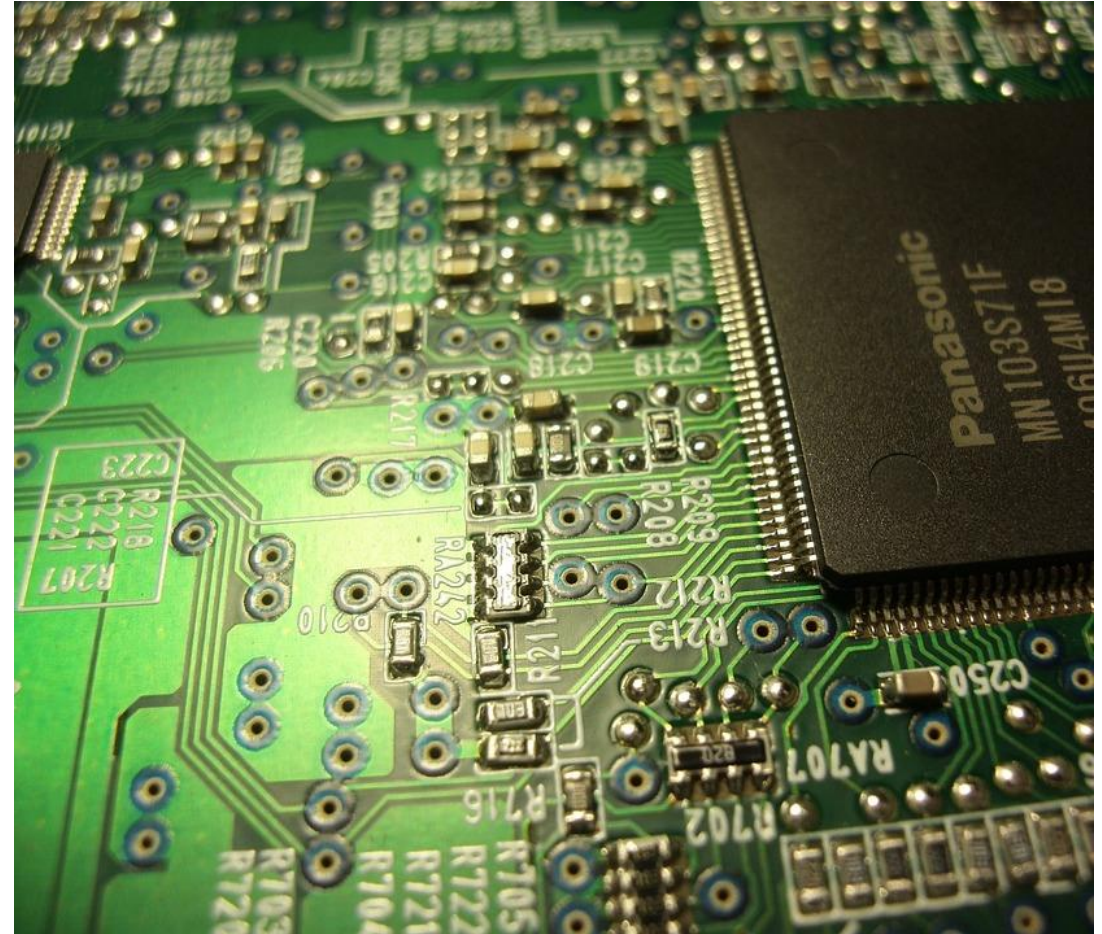
Vice Chair: House Veterans and Military Development Program

Member of Development, Technology, and Innovation Committee

Commander of Amored Regiment in Ohio National Guard

July Parting Shots

- Local Gov't/Public Safety SIG – Wright Patt Room
- Education SIG – Victoria Room
- Defense Contractor SIG – Hawthorn Room
- Register for Workshop, Monthly Meetings
- Turn in your lanyards at the desk





Information

Email me at john@GoCyberCollective.org.

If you have a suggestion on meeting topics, special interest groups, breakfast menu, or anything else we want to hear from you.

You can register for future breakfast meetings from the website.

□ GoCyberCollective.org



Event Sponsors

