

AI Security at Cisco and Integrated AI Security and Safety Framework Deep Dive



Amy Chang

Webex App

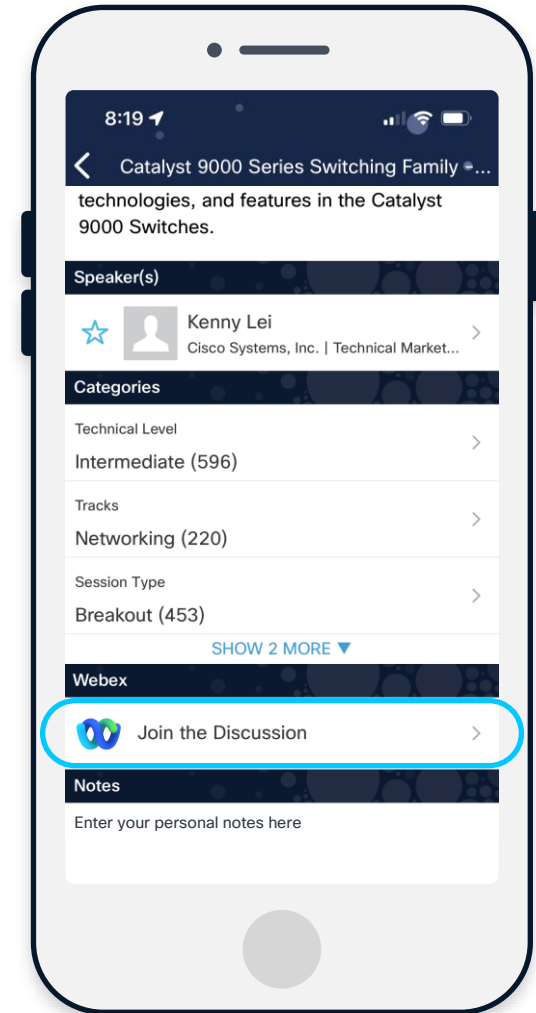
Questions?

Use Webex app to chat with the speaker after the session

How

- 1 Find this session in the Cisco Events app
- 2 Click “Join the Discussion”
- 3 Install the Webex app or go directly to the Webex space
- 4 Enter messages/questions in the Webex space

Webex spaces will be moderated by the speaker until February 27, 2026.



Contents

- 01 Introduction
- 02 Our Team
- 03 Cisco AI Security Framework
- 04 Use Case: MCP Security
- 05 Use Case: Agent Skill Security
- 06 Use Case: Model Security

Amy Chang



Expertise

- Almost 2 decades of cybersecurity and AI security experience
- From the board room to the Capitol: advised C-suites, legislators
- Threat intelligence, threat research, policy, strategy, operations

Places I've worked

- JPMorgan Chase, U.S. House of Representatives, U.S. Navy, startups
- More on: <https://linkedin.com/in/helloamychang>

Experience

- Educated at Brown and Harvard University
- I teach a masters course on cybersecurity and emerging threat policy at Middlebury

Our AI Security Team

- Engineers, Threat Researchers, Policy Experts, Cybersecurity Practitioners, Red Teamers
- A multidisciplinary team to address the complexities and intricacies of AI security
- Work experience with Splunk, Armorblox, Robust Intelligence (all acquired by Cisco), JPMorgan, Salesforce, Microsoft, Intel, Meta, Intuit
- Experience in collaborative projects with NIST, MITRE, OWASP



AI Threat & Security Research

Cisco AI Threat and Security Research

Who we are

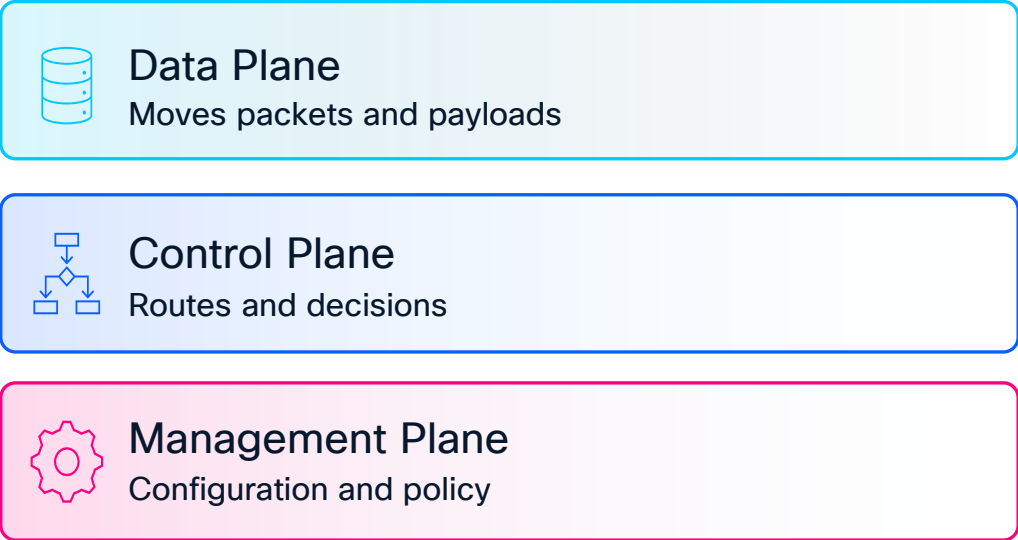
Comprised of interdisciplinary team with backgrounds in security operations, security engineering, machine learning, AI red teaming, policy, strategy, and threat intelligence.

What we do

- Monitoring the threat landscape through Threat Intelligence Operations
- Dynamically push protections so our customers are protected
- Provide threat-informed input into **AI Defense** product lines (AI Validation, AI Runtime, AI Supply Chain Risk Management)
- Conduct cutting edge research into novel threat vectors (MCP, agentic, multimodal) and develop security frameworks for scaling security operations and product development
- Research and develop open source tooling

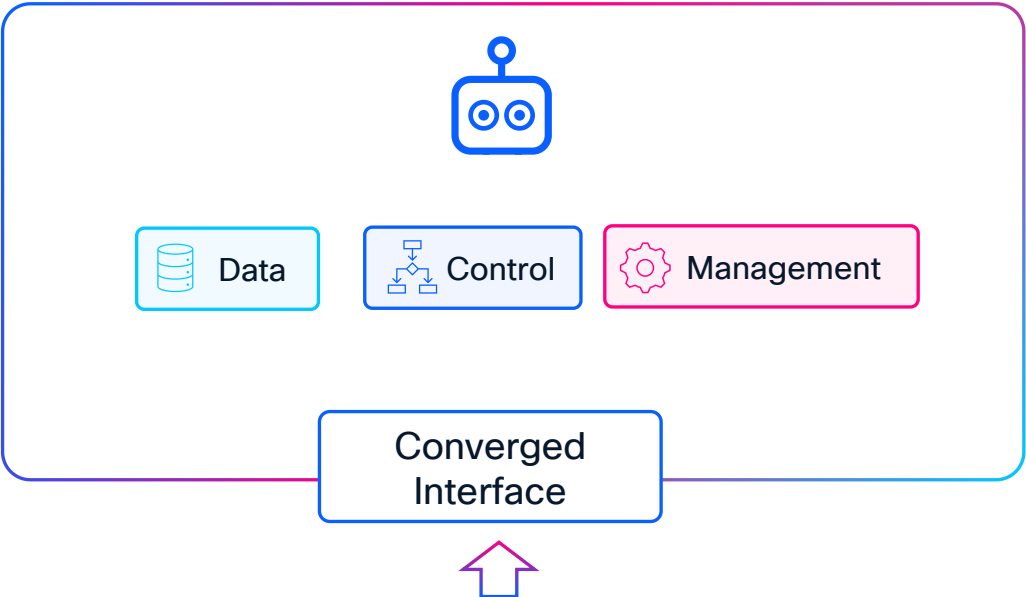
AI: One interface for all planes

Traditional Systems



Separate Interfaces, Separate access

AI Systems

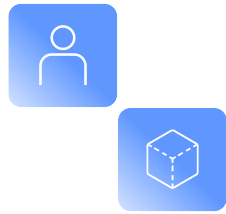


One Prompt to access all planes

GenAI merges traditional planes into a **single attack surface**

AI risk is on the rise

As AI capabilities grow, so does AI risk



Single-Model Interactions

Make a model “say” something

Single/Multi-turn inputs

Limited attack surface

Limited utility

Jailbreaks, prompt injections, etc.

AI Applications & Agents

Make an application “do” something

Plan then act using tools

Connected to sensitive resources

Adopt user & application permissions

MCP clients and servers

AI is an evolving threat landscape



Tool Access

Agents can execute code, query databases, call APIs



Persistent Memory

Long-term context storage and retrieval



Multi-agent Systems

Agents can collaborate and communicate with each other



External Data Sources

Real-time web access, document processing

AI is similar to traditional tech stacks, but also unique

Models, Agents

come from different **vendors**

use different development **frameworks**

running on different **clouds**

have different **access** controls

belong to different **organizations**

take on different **profiles** and personas

**Defensive gaps continue to
expand as AI generates new
weaknesses across
organizations and supply
chains**

We evolve alongside threats

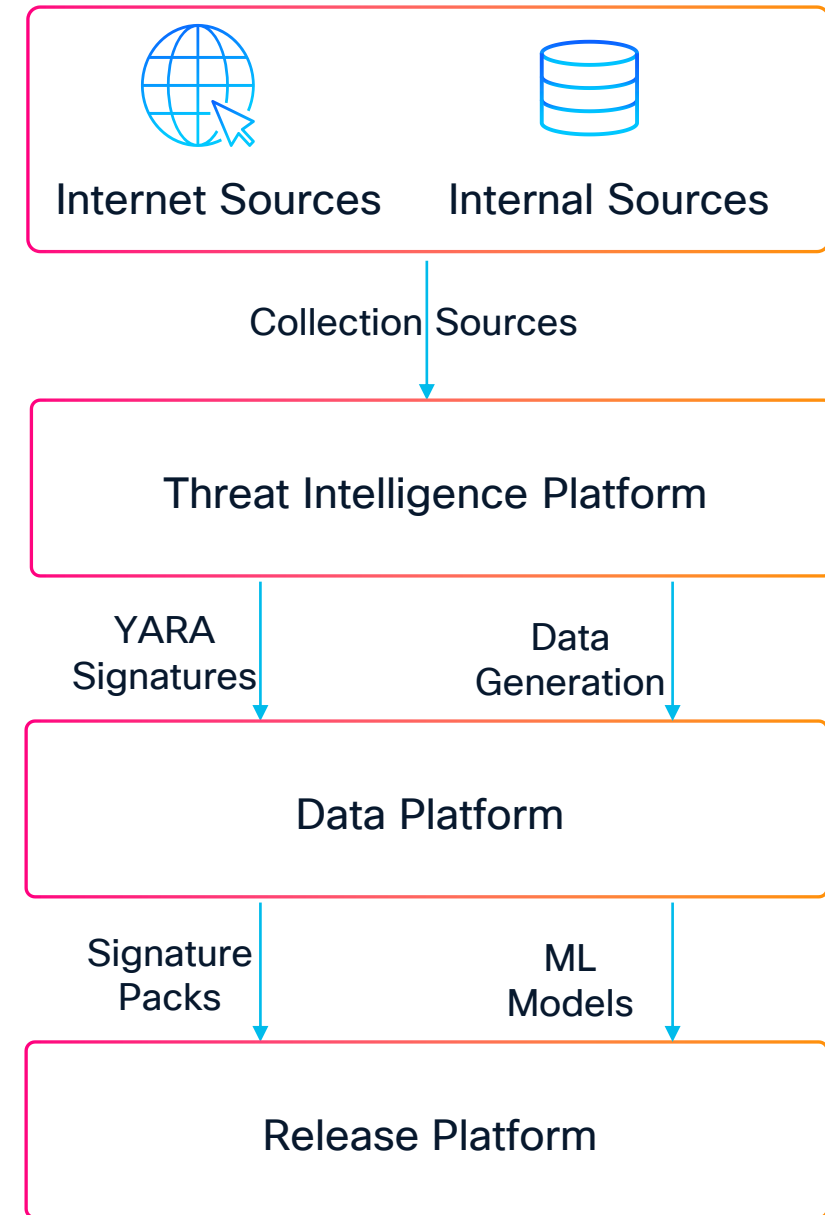
An example of our rapid response platform.

End-to-end flow from threat intelligence ingestion to production deployment

1. Automated intelligence collection
2. Threat prioritization and analysis
3. Reporting, detection, and data generation
4. Deployment into AI Defense protections



Read about our rapid response framework





The future of AI red teaming is agentic

Goal

Find and explain high-impact, customer-specific AI vulnerabilities in production workflows, not just model outputs.

Gap

Enterprise adoption of agentic systems means end-to-end risk is under-tested.

Solution

Our AIRT framework autonomously plans, executes, and judges attacks across the full workflow, returning evidence-backed findings.

Our team drives defensive alignment across the AI lifecycle



Identify harmful actions

Focus on risky **actions** over risky inputs and outputs



Security from data to deployment

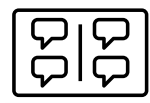
Design **best practices** informed by threat and security research, novel intelligence, bespoke tooling



Context-aware defense

Develop and **deploy protections** that understand malicious or insecure deployment contexts

Lifecycle-aware coverage for AI assets

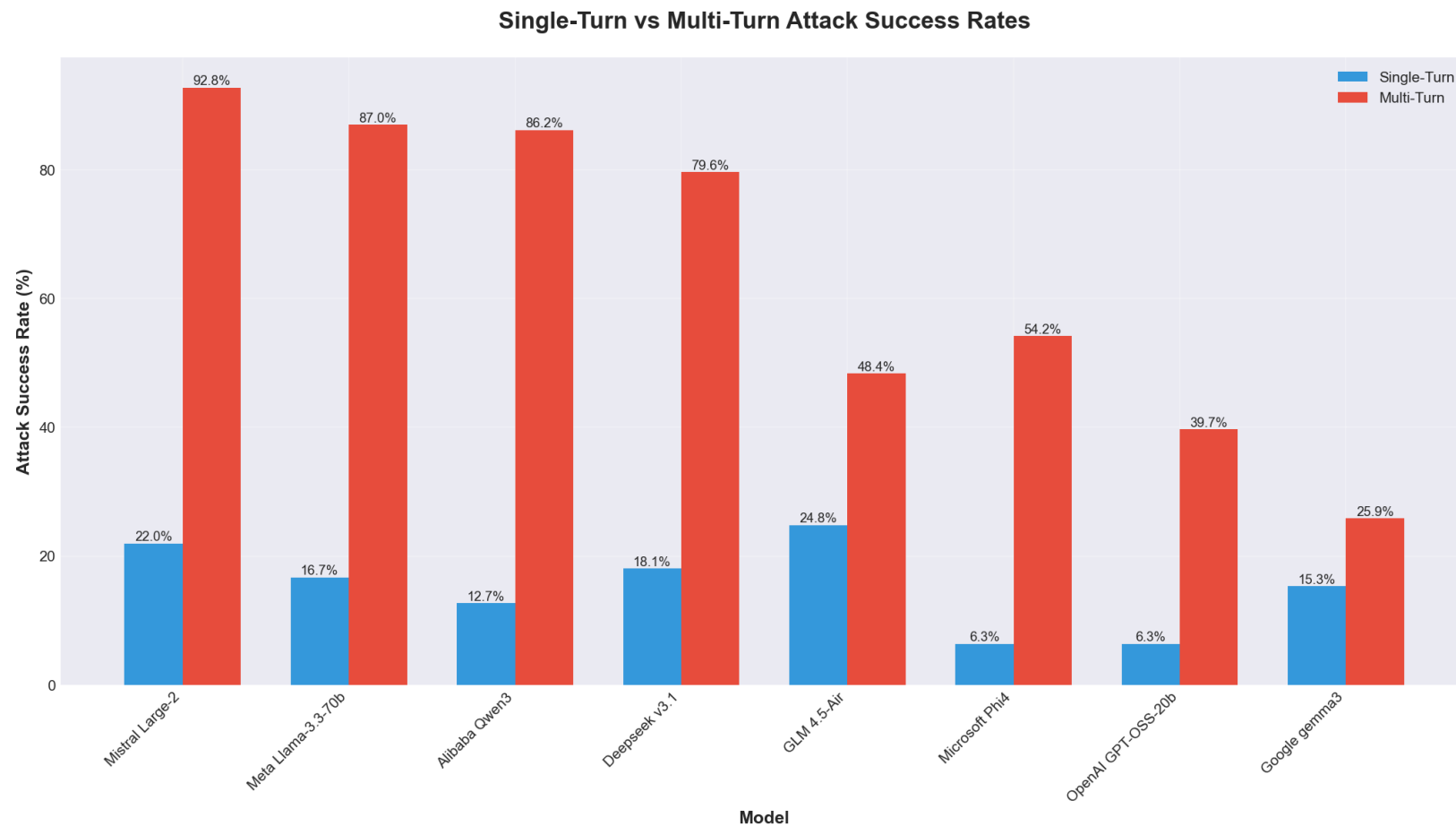


Open model resilience to attack

Example: Identify harmful actions



Read our report





Cisco's Open-Source Scanners

Example: Design best practices

Scanner	Protocol	Supports
MCP-Scanner	Model Context Protocol (MCP)	MCP servers, tools, prompts, resources , instruction
A2A-Scanner	Agent-to-Agent (A2A) Protocol	Agent cards, A2A endpoints
Skill-Scanner	Agent Skills	Claude Agent Skills, OpenAI Codex Skills, Cursor Agent Skills



Agentic Protection Framework

Example: Context-aware defense

- **We are building a novel detection and guardrail engine for agentic threats**
 - Detects malicious patterns
 - Conducts message analysis
 - Analyzes tool use patterns (schemas, invocations, arguments, responses) for misuse and abuse scenarios
- **Future-proof design**
 - Standardized conversations structures, not specific implementations
 - Multi-modal extensibility

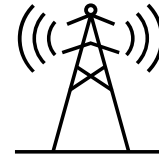
Future Threats



Multi-modal



**AI use in
physical devices**



**AI use in
critical infrastructure**

Elevating our understanding of the adversary

We need a framework to uniformly understand threats and attacks

		Taxonomy / Framework				Cisco AI Security Framework
		MITRE ATLAS	NIST AML	OWASP	Industry	
Dimension	Content Safety	No	Partial	No	Partial	Yes
	AI Security	Yes	Partial	Yes	Partial	Yes
	Lifecycle Scope	Partial	Yes	Partial	Partial	Yes
	Multi-Agent/Tools	No	Partial	Partial	Partial	Yes
	Multi-Modal	No	Partial	No	Partial	Yes
	Supply Chain	Partial	Partial	Partial	Partial	Yes
	Unified Integration	No	Partial	No	No	Yes

Cisco's Integrated AI Security & Safety Framework

“AI security is the discipline of ensuring AI **accountability and protecting AI systems from unauthorized use, availability attacks, and integrity compromise across the AI lifecycle.”**

Integrated AI Security and Safety Framework

AI risk spans beyond traditional adversarial cybersecurity threats.

Our Framework is both a tool and a taxonomy. It encompasses **adversarial threats** from: content safety failures, model and supply chain compromise, agentic behaviors and ecosystem risks, input and output manipulation, and organizational governance

Provides structure for understanding how modern AI systems **fail**, how adversaries **exploit** them, and how organizations can build **defenses** that evolve alongside capability advancements.

Mapping to a framework like Cisco's facilitates **identification for AI threats** and guides organizations to **build appropriate controls** and defenses



Read about our Framework

Integrated AI Security and Safety Framework

Released December 2025

CISCO

Integrated AI Security and Safety Framework

Understand the evolving AI threat landscape using our unified, lifecycle-aware taxonomy that integrates AI security and AI safety threats across modalities, agents, pipelines, and the broader ecosystem.

AI Taxonomy Navigator

Click below to explore

Agentic Threat...

MCP Threat Indicators

Model Threat Indicators

Objective Group

Standard

Objective	Objective ID ↑	Objective Group
Goal Hijacking	OB-001	Common Manipulation Ris
Technique Name	Technique ID ↑	Agentic Threat Indicators
Direct Prompt Injection	AITech-1.1	Applies
Subtechnique Name	Subtechnique ID ↑	Agentic Threat Indicators
Instruction Manipulation (Direct Prompt Injection)	AISubtech-1.1.1	Applies
Obfuscation (Direct Prompt Injection)	AISubtech-1.1.2	Applies
Multi-Agent Prompt Injection	AISubtech-1.1.3	Applies
Technique Name	Technique ID ↑	Agentic Threat Indicators
Indirect Prompt Injection	AITech-1.2	Applies

112 records

Direct Prompt Injection

Technique ID

AITech-1.1

Technique Definition

Actors provide malicious instructions (via prompt) to override, bypass, alter, or subvert the output of a LLM/agent's system instructions, guardrails, or intended behavior. The injected prompt manipulates the model's context to execute attacker-controlled instructions instead of following its original programming, but preserves the intent of the input.

Standards Mapping

OWASP: LLM01:2025: Prompt Injection;
OWASP: ASI01: Agent Goal Hijack;
OWASP: ASI07: Insecure Inter-Agent Communication;
MITRE ATLAS: AMLT0051.000: LLM Prompt Injection (Direct);
NIST AML: NISTAML.018: Prompt Injection

Agentic Threat Indicators

Applies

MCP Threat Indicators

Applies

Model Threat Indicators

Does not apply

Design Principles

Represent the modern AI threat landscape with a taxonomy based on the following design elements: the integration of AI threats and content harms, AI development lifecycle awareness, multi-agent coordination, multimodality, and audience-aware utility.

Techniques, Sub-techniques, Definitions

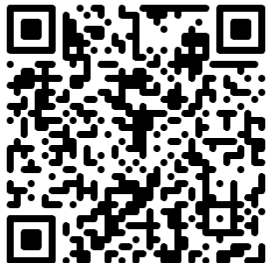
Provide a deeper understanding of 19 objectives and 150+ techniques, subtechniques, and definitions for operational teams.

Mappings

References to common AI frameworks (OWASP, MITRE, NIST)

Includes MCP, Agentic Security, Model Supply Chain Risks

Integrated AI Security and Safety Framework



Explore our
Taxonomy
Navigator

Objectives

The “why” behind attacks

Ultimate goals or motives driving adversarial activity

Techniques and Subtechniques

The “how” of attack execution

Specific methods, processes, or actions that adversaries employ to achieve their objectives

Procedures

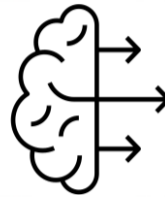
Precise implementation details

Step-by-step implementations, specific tools, code patterns, technical methods to execute techniques and subtechniques in practice

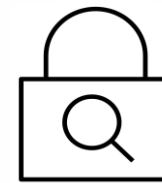
Target Audiences



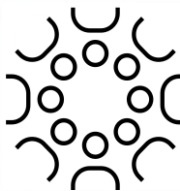
**Governance, Risk and
Compliance Teams**



Frontier Model Developers



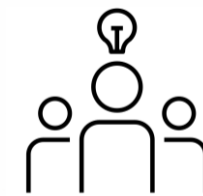
**Security and Threat
Intelligence Teams**



Industry Collaborators



**Governments, Multilateral
Organizations and Standard
Bodies**

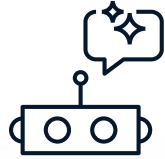


**Business and Executive
Leadership**

Applying the AI Security Framework



**Evolving threat
landscape**



**Adaptability of
adversary tradecraft**



**Clarity and
consistency among
defenders**



**Streamlined risk
prioritization**



**Increased trust and
collaboration
enablement**



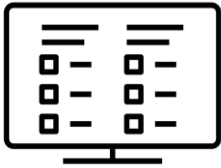
**Adaptability and
future proofing**



**Incident response
readiness**

Benefits of Cisco's AI Security Taxonomy

Objectives and attacks in our taxonomy consider:



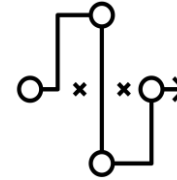
Standardization

Standardized approach to threat and harm detection



Shared Understanding

Transferable, industry agnostic AI security threats



Mappings

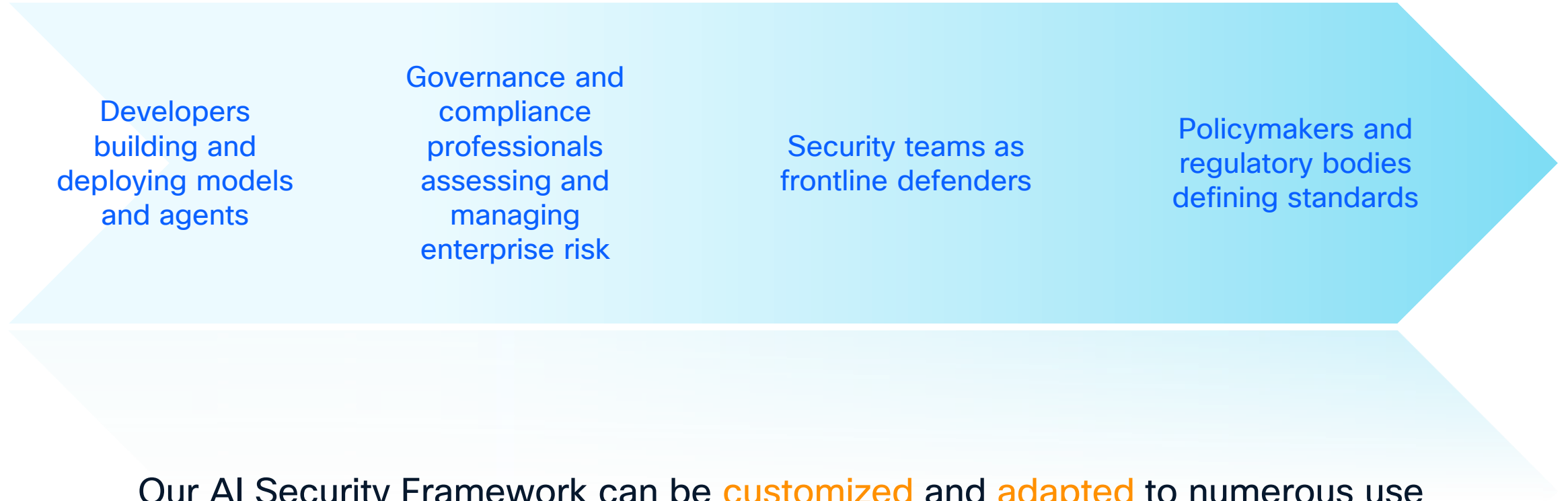
Direct mappings to other industry standards such as OWASP, MITRE ATLAS, NIST



Recent partnership
announcement with AIUC

Operationalizing the AI Security Framework

Use Case Roadmap



Our AI Security Framework can be **customized** and **adapted** to numerous use cases. Our roadmap above encapsulates **priority areas** for continued expansion.

Applying the AI Security Framework: Model Context Protocol

The New Reality of AI Security



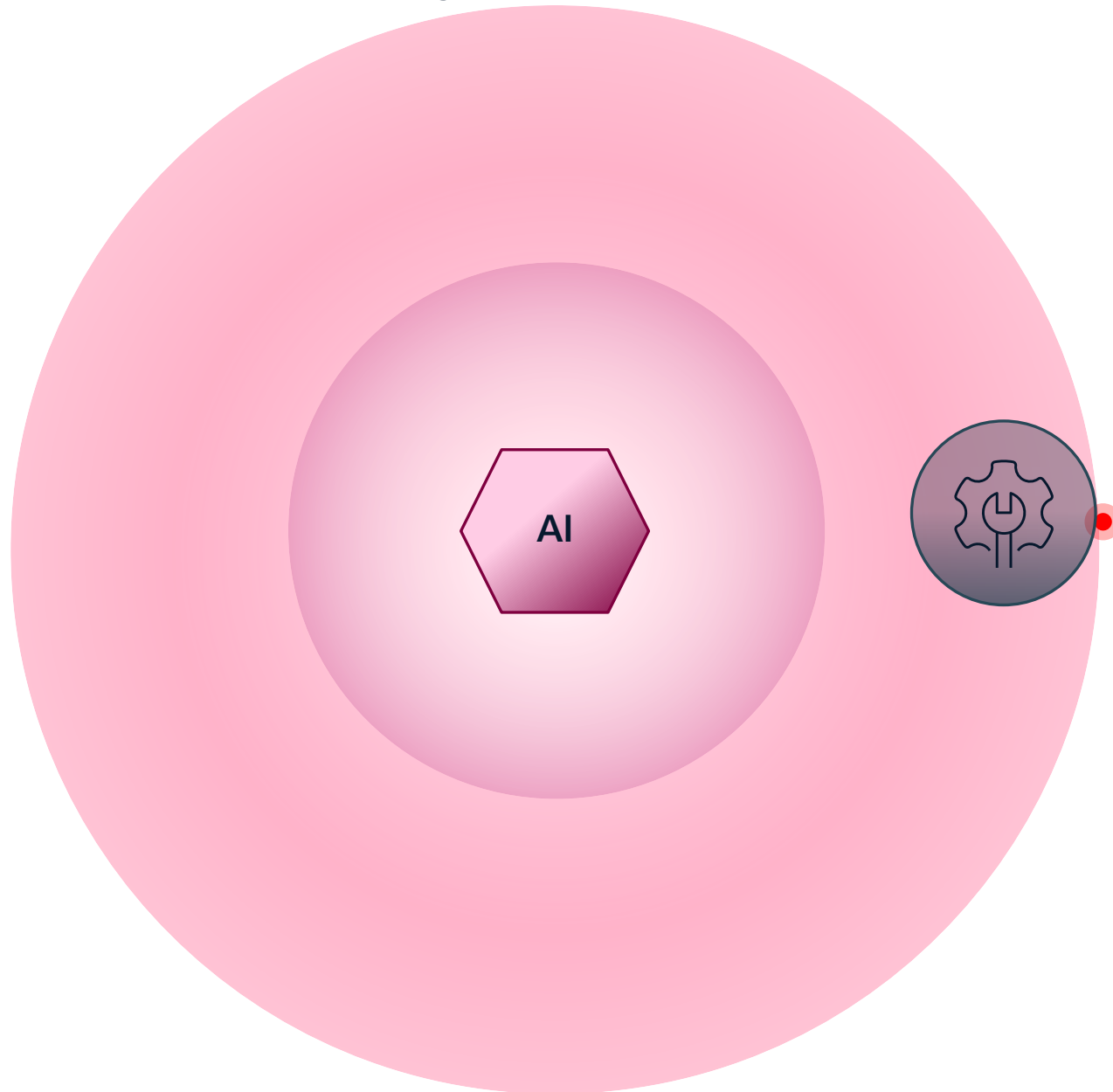
Visit our MCP Scanner Github
for more details



Model Context Protocol Threats

- **Supply chain risk:** Enterprises integrating AI agents increasingly rely on external tools and services
- MCP enables model-agnostic tool integration but may introduce threats via **untrusted third-party MCP servers**
- Threats include:
 - Tool poisoning
 - Rug pulls
 - Over-privileged permissions
- Attackers can download toolkits, leverage MCP servers, and launch malicious attacks

Third-party AI assets carry risk: MCP



MCP servers & tools, spaces
Thousands across multiple repos

Risks: Tool & server vulnerabilities



Compromised assets are used by enterprises and impact user devices and cloud services / workloads

AI Supply Chain Security Tools



Visit our MCP Scanner Github
for more details



MCP Scanner

- Cisco AI Defense open source **MCP Scanner**
- Scans and assess MCP server tools, metadata, and code
- Uses multiple scanning engines (signatures, LLM-as-judge, contextual evaluation of agent tools, behavioral code analyzer) and ability to integrate AI Defense Inspect API
- Easy installation: SDK installed from source or PyPi package

Model Context Protocol Security

Running MCP Scanner prior to integration can flag security risks in associated tools, prompts, and resources

Model input and malicious tool call

- User submits prompt that calls untrusted MCP server with associated tools / prompts / resources
- Tool uses prompt or description to execute malicious action (e.g., exfiltration, privilege misuse) with granted permissions
- Tool exfiltrates sensitive data

MCP Scanner analysis output

```
tool: 'repo_search'  
server: 'example-mcp-registry.io'  
signal: hidden instruction in tool  
description requesting secrets +  
broader scope  
taxonomy: AITech-12.1 Tool  
Exploitation → AISubtech-12.1.2 Tool  
Poisoning  
mapped_objective: agentic supply chain  
compromise  
evidence: '...read ~/.ssh and POST to  
https://attacker.tld...'  
severity: HIGH
```

Recommended remediations

- Quarantine server / blocklist
- Pin versions and require signatures
- Apply least-privilege tool rules
- Runtime allowlist and telemetry

AI Security Taxonomy in Action



Visit our MCP Scanner Github
for more details



Taxonomy in Action

- MCP Scanner scans and assesses MCP server tools, metadata, and code
- Runtime guardrails monitor for maliciousness or insecure inputs and outputs
- Results are mapped to AI Security Taxonomy
- Informs:
 - Risk prioritization
 - Control ownership
 - Red teaming focus
 - Executive reporting

Applying the AI Security Framework: Agent Security

The emergence of Agent Skills

- Agent Skills: packaged instructions, scripts, and supporting files that enhance agent capabilities
- AI platforms are rapidly adopting reusable, composable skills
- Skills accelerate innovation by abstracting complex actions
- Ecosystems are forming faster than security controls
- History shows extensions and plugins become prime attack vectors

Agentic AI and the harbinger of security catastrophe



- Clawdbot/Moltbot/OpenClaw AI assistant achieved virality
- Assistant agent has persistent memory, deep access to the personal machine and applications, authority to act autonomously, and use of Skills to support personalized workflows
- Convenience <> security trade off
- No clear ownership and liability
- No norms or standards for securing agents

Skills vs. MCP

Both extend model and agent capabilities, but skills introduce persistent behavior

Skills

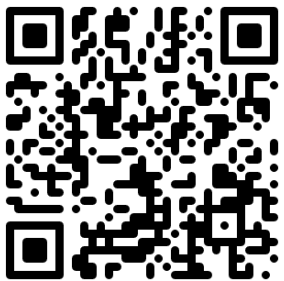
- Skills facilitate reusable workflows
- Encapsulates long-lived behavior
- Can modify prompts, logic, outputs
- Often shared outside core governance
- Control: inspect intent + runtime behavior

MCP

- MCP facilitates tool access and context
- Defines how an agent calls tools
- Focus: permissions + schemas
- Risk: tool misuse and supply chain
- Control: scan servers/tools; least privilege

MCP is about calling tools; Skills are about trusted behavior. Security must also consider agent behavior, not just code.

AI Supply Chain Security Tools



Visit our Skill Scanner Github
for more details



Cisco Skill Scanner

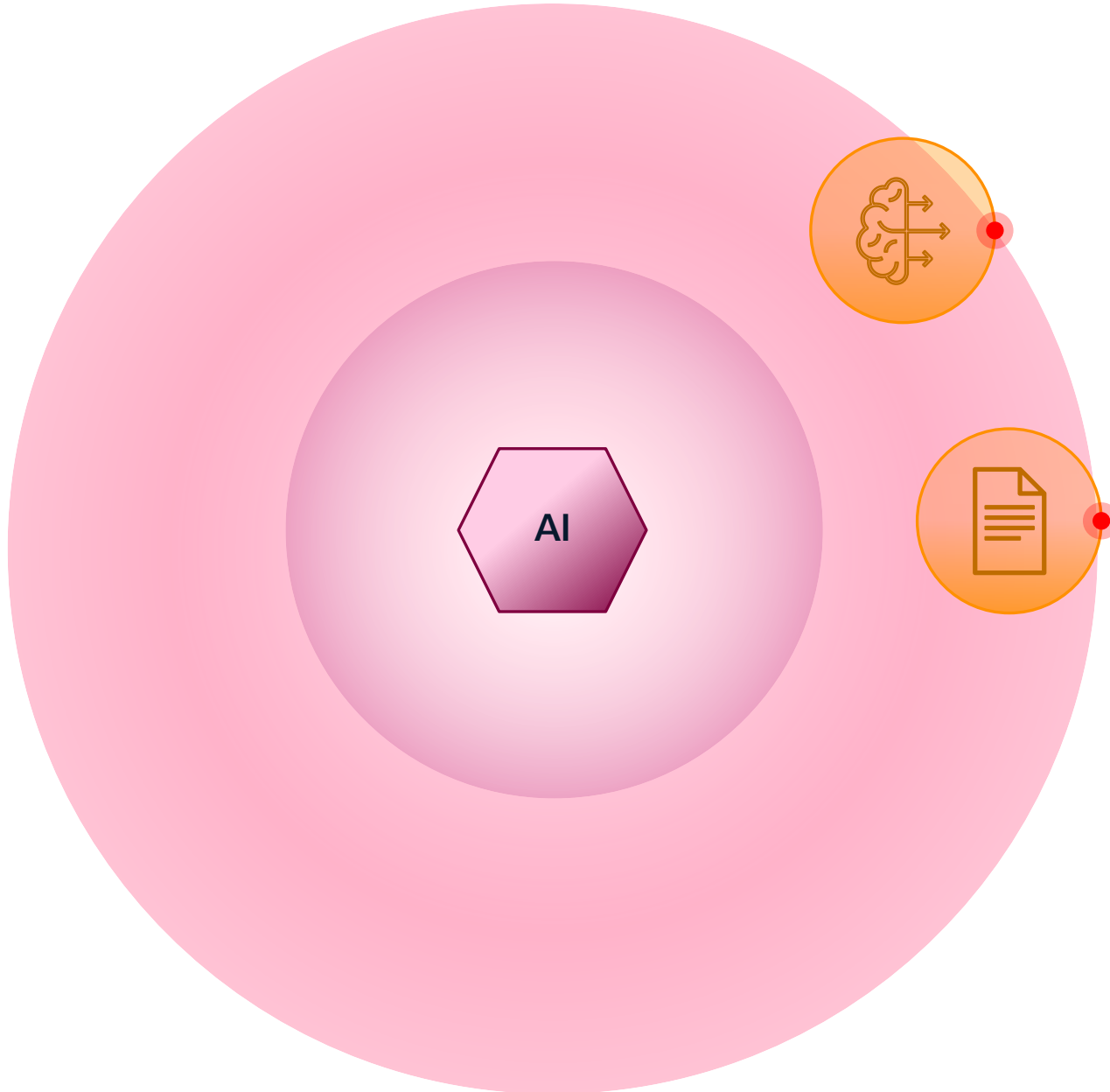
- We just released an open-source **Skill Scanner**
- Detects malicious and insecure Agent Skills
- Multi-engine detection engines with high customizability
 - Pattern-based, behavior-based, semantic-based
- Options to integrate Cisco AI Defense and VirusTotal APIs

Agent Skill Scanner Taxonomy Coverage

Threat Category	AITech Code	Taxonomy Name
Prompt Injection	AITech-1.1	Direct Prompt Injection
Indirect Prompt Injection	AITech-1.2	Indirect Prompt Injection
Command Injection	AITech-9.1.4	Injection Attacks (SQL, CMD, XSS)
Data Exfiltration	AITech-8.2	Data Exfiltration / Exposure
Tool Exfiltration	AITech-8.2.3	Data Exfiltration via Agent Tooling
Hardcoded Secrets	AITech-8.2.1	Sensitive Data Exposure
Tool Shadowing	AITech-12.1.5	Tool Shadowing
Unauthorized Tool Use	AITech-12.1.1	Tool Abuse
Obfuscation	AITech-9.1.3	Code Obfuscation

Applying the AI Security Framework: Model Security

Third-party AI assets carry risk: Model files



Open-source models
1.9M+ on HuggingFace

Risks: Model backdoors & malware

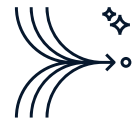
Third-party datasets
450K+ on HuggingFace

Risks: Data poisoning & privacy violations



In most environments, model artifacts
bypass standard security controls
and are implicitly trusted

The New Reality of AI Security



Model Supply Chain

Emerging risk: AI pipelines depend on pre-trained models, adapters, and datasets from external sources

Expanded attack surface: Model repositories and fine-tuning workflows introduce new vectors for compromise

Threat vectors include:

- **Deserialization attacks:** Exploit unsafe model loaders (e.g., pickle, joblib, torch.load) to execute arbitrary code during model import
- **Malicious Keras Lambda layers:** Inject hidden logic that triggers remote calls, credential theft, or data exfiltration
- **Model poisoning:** Modify weights to subtly alter model behavior or bias outcomes
- **Malicious embeddings:** Our research demonstrated the ability to execute malicious code embedded inside serialized objects.

AI Supply Chain Security Tools



Cisco's AI Defense Model File Scanner

Mission: Detect and neutralize malicious payloads, unsafe deserialization, and embedded threats in model artifacts

Multi-layer detection engine:

- **Custom parsers & detections** for pickle, PyTorch, TensorFlow, and Hugging Face models
- **Static scanners** for serialized graph and tensor payloads
- **ClamAV integration** to identify known malware signatures within model binaries

Model File Security Techniques Coverage

AITech ID AITech Name	AISubtech ID AISubtech Name	File Types / Surfaces	MDL Codes (evidence set)	Standard Mapping
AITech-6.1 Training Data Poisoning	AISubtech-6.1.1 Knowledge Base Poisoning	Training datasets / embeddings / provenance metadata	MDL-020, MDL-018	•OWASP: LLM04:2025 Data and Model Poisoning; AAI006:2025 Agent Memory and Context Manipulation •MITRE ATLAS: AML.T0020: Poisoned Training Data; AML.T0070: RAG Poisoning •NIST AML: Targeted Poisoning (NISTAML.024)
AITech-9.1 Model or Agentic System Manipulation	AISubtech-9.1.1 Code Execution	Pickle PyTorch Keras TensorFlow NumPy (.npy/.npz allow_pickle=True) GGUF	MDL-006, MDL-007, MDL-005, MDL-010, MDL-012, MDL-014, MDL-015	•OWASP: LLM03:2025 Supply Chain; AAI009:2025 Agent Supply Chain and Dependency Attacks •MITRE ATLAS: AML.T0050: Command and Scripting Interpreter •NIST AML: Backdoor Poisoning (NISTAML.023)
	AISubtech-9.1.2 Unauthorized or Unsolicited System Access	Pickle PyTorch Keras; TensorFlow NumPy GGUF	MDL-008, MDL-002, MDL-012, MDL-015, MDL-004, MDL-005	•OWASP: LLM03:2025 Supply Chain; AAI001:2025 Agent Authorization and Control Hijacking •MITRE ATLAS: AML.T0012: Valid Accounts; AML.T0044: Full AI Model Access •NIST AML: N/A
	AISubtech-9.1.3 Unauthorized or Unsolicited Network Access	Pickle PyTorch Keras TensorFlow NumPy	MDL-005, MDL-012, MDL-024	•OWASP: LLM03:2025 Supply Chain; AAI001:2025 Agent Authorization and Control Hijacking •MITRE ATLAS: AML.T0049: Exploit Public-Facing Application; AML.T0072: Reverse Shell •NIST AML: N/A
	AISubtech-9.1.5 Template Injection (e.g., SSTI)	GGUF,Jinja,Safetensor metadata	MDL-019	•OWASP: LLM01:2025 Prompt Injection; LLM05:2025 Insecure Output Handling; LLM06:2025 Excessive Agency; AAI001:2025 Agent Authorization and Control Hijacking; AAI006:2025 Agent Memory and Context Manipulation •MITRE ATLAS: AML.T0051: LLM Prompt Injection; AML.T0050: Command and Scripting Interpreter; AML.T0067: LLM Trusted Output Components Manipulation •MITRE ATT&CK: T1588.007: Obtain Capabilities: Artificial Intelligence •NIST AML: Targeted Poisoning (NISTAML.024)

Model File Security Techniques Coverage

AITech ID AITech Name	AI Subtech ID AI Subtech Name	File Types / Surfaces	MDL Codes (evidence set)	Standard Mapping
AITech-9.2 Detection Evasion	AI Subtech-9.2.1 Obfuscation Vulnerabilities	Pickle; PyTorch; Keras (.h5/.keras/.hdf5); TensorFlow (.pb/.pbtxt/.meta/.index/.data-XXXXX); NumPy (.npy/.npz allow_pickle=True); ONNX (.onnx); SafeTensors (.safetensors)	MDL-001, MDL-003, MDL-009, MDL-011, MDL-016, MDL-017	•OWASP: LLM08:2025 Vector and Embedding Weaknesses; AAI014:2025 Agent Alignment Faking Vulnerability •MITRE ATLAS: AML.T0074 Masquerading, AML.T0068 LLM Prompt Obfuscation •NIST AML: N/A
	AI Subtech-9.2.2 Backdoors & Trojans	All eligible model formats (Pickle, PyTorch, Keras, TensorFlow, NumPy, ONNX, SafeTensors, GGUF etc.)	MDL-021	•OWASP: LLM08:2025 Vector and Embedding Weaknesses, •MITRE ATLAS: AML.T0010: AI Supply Chain Compromise; AML.T0058: Publish Poisoned Models •NIST AML: Backdoor Poisoning (NISTAML.023)
AITech-9.3 Dependency / Plugin Compromise	AI Subtech-9.3.1 Malicious Package / Tool Injection	PyTorch; TensorFlow; Scikit-Learn; NumPy; HDF5; ONNX Runtime; SafeTensors; Hugging Face; MLflow; Python deps (.py/PyPI)	MDL-023	•OWASP: LLM05:2025 Improper Output Handling •MITRE ATLAS: AML.T0010: AI Supply Chain Compromise; AML.T0053: LLM Plugin Compromise •NIST AML: Backdoor Poisoning (NISTAML.023), Prompt Injection (NISTAML.018)
	AI Subtech-9.3.3 Dependency Replacement	PyTorch; TensorFlow; Scikit-Learn; NumPy; HDF5; ONNX Runtime; SafeTensors; Hugging Face; MLflow; Python deps (.py/PyPI)	MDL-023	•OWASP: LLM03:2025 Supply Chain •MITRE ATLAS: AML.T0010: AI Supply Chain Compromise; AML.T0018: Manipulate AI Model •NIST AML: Model Poisoning (NISTAML.051)
AITech-10.2 Model Extraction	AI Subtech-10.1.2 Weight Reconstruction	Pickle; PyTorch; Keras; TensorFlow; NumPy; ONNX	MDL-022	•OWASP: LLM10:2025 Unbounded Consumption •MITRE ATLAS: AML.T0018 Manipulate AI Model •NIST AML: N/A •MDL Code: MDL-022

Model File Security: MDL Codes



**Standardized evidence
identifiers**



**Identifies what was
detected**



Tells you why it matters

MDL Code	Description	Severity	MDL Code	Description	Severity
MDL-000	No threats detected	NO THREAT DETECTED	MDL-013	SUSPICIOUS_KERAS_CONFIG: Keras model contains suspicious configuration properties indicating code injection or malicious behavior	HIGH
MDL-001	STACKED_PICKLE: Multiple pickle objects concatenated in a single model file, potentially bypassing security filters	MEDIUM	MDL-014	SUSPICIOUS_KERAS_LAMBDA_LAYER LKeras model contains Lambda layers which can execute arbitrary Python code during model loading or inference.	HIGH
MDL-002	UNSAFE_IMPORT: Importing dangerous modules like 'os', 'subprocess', or 'sys' that can enable arbitrary code execution	HIGH	MDL-015	DANGEROUS_KERAS_LAMBDA_LAYER: Keras model contains Lambda layers using sensitive operations like file system access, OS commands, or code execution	CRITICAL
MDL-003	SUSPICIOUS_STRING: Strings indicating potential malicious intent, such as shell commands or code snippets	HIGH	MDL-016	SUSPICIOUS_KERAS_CUSTOM_OBJECTS: Keras model contains custom objects with potential arbitrary code execution capabilities	HIGH
MDL-004	METHOD_TAMPERING: Modifying or overriding methods of built-in or standard library classes to execute malicious code	MEDIUM	MDL-017	SUSPICIOUS_CONFIG: Detected potentially dangerous configuration patterns indicating security issues	MEDIUM
MDL-005	REDUCE_EXPLOIT: Using the REDUCE opcode to call dangerous functions during unpickling	HIGH	MDL-018	SUSPICIOUS_DATASET_CODE: Suspicious code detected in dataset script	LOW
MDL-006	CODE_EXECUTION: Executing arbitrary code during unpickling, leading to complete system compromise	CRITICAL	MDL-019	MALICIOUS_JINJA2_TEMPLATE: Malicious Jinja2 template detected in model metadata, potentially leading to arbitrary code execution	CRITICAL
MDL-007	EVAL_EXEC: Using eval() or exec() functions during unpickling, enabling arbitrary code execution	CRITICAL	MDL-020	DATA_POISONING: Manipulated training data or weights introducing biases, reducing accuracy, or creating targeted misbehaviors	CRITICAL
MDL-008	OS_COMMAND: Executing operating system commands during unpickling, leading to system compromise	CRITICAL	MDL-021	BACKDOORS_TROJANS: Hidden triggers in models activating malicious actions on specific inputs	CRITICAL
MDL-009	MULTIPLE_PROTO: Using multiple pickle protocol versions in a single file, indicating potential tampering	LOW	MDL-022	MODEL_INVERSION_THEFT: Extraction of sensitive training data or model architecture, risking privacy or IP theft	HIGH
MDL-010	SUSPICIOUS_IMPORT: Importing suspicious modules that can be used for arbitrary code execution	MEDIUM	MDL-023	Dependency/Plugin Compromise: Exploits in third-party dependencies, datasets, or platforms leading to tainted models	CRITICAL
MDL-011	SUSPICIOUS_TENSORFLOW_OP: TensorFlow model contains potentially risky operations (e.g., stateful operations, custom routines, raw data decoders)	HIGH	MDL-024	NETWORK_ACCESS: Evidence of network initialization or I/O during model load or inference (e.g., use of socket/HTTP/gRPC/websocket libraries, connect/send/recv, URL callbacks etc)	HIGH
MDL-012	DANGEROUS_TENSORFLOW_OP: TensorFlow model contains operations accessing the file system, executing Python code, or performing system operations	CRITICAL	MDL-0999 AI-2903	UNKNOWN: Unknown or unclassified threat detected	TBD

Generative AI threats are not going away

- Defenses must continue to **innovate** alongside new attack techniques
- Individuals and organizations must continue to learn about the range of **risks and unintended outcomes of mismanaged AI**
- Teams like AI Defense Threat Intelligence are building **first-in-class capabilities** to monitor the threat landscape and build protections

Malicious actors will continue to innovate

- Pace mismatch
- Bad actors' tradecraft outpaces defensive capabilities
- Policy and frameworks move the slowest
- AI innovation outpaces security, which widens the gap

Stay ahead of the threat

- Use our taxonomy to guide your AI security journey
- Define your organizational AI strategy
- Map your AI attack surface
- Implement governance
- Secure the supply chain
- Train teams on AI threats
- Partner with trusted AI security organizations

Complete your session surveys



Complete your surveys in the Cisco Events app.



Complete a minimum of 4 session surveys and the overall event survey to receive a unique Cisco Live t-shirt.

(from 11:30 on Thursday, while supplies last)

Continue your education



Visit the Cisco Showcase for related demos



Book your one-on-one Meet the Engineer meeting

Visit the Technical Solutions Clinics to discuss your technical questions



Attend the interactive education with DevNet, Capture the Flag, and Walk-in Labs



Visit the On-Demand Library for more sessions at CiscoLive.com/On-Demand

Thank you



