



HITRUST CSF THREAT & MITIGATION ANALYSIS Q1 2026

Period Covered: 1/1/2026 to 3/31/2026

Static Security Programs Can't Keep Up. HITRUST Can.

As organizations enter a new era of AI-driven technology, the cyber threat landscape continues to evolve just as quickly. Security frameworks built primarily for consistency and conformity can struggle to keep pace with rapidly changing threats.

HITRUST addresses this challenge through its Cyber Threat Adaptive (CTA) program. By continuously analyzing real-world threat intelligence and applying those insights to the HITRUST CSF, HITRUST helps ensure its control requirements remain relevant and effective against emerging risks while also supporting organizations' ongoing compliance needs. As a result, organizations with HITRUST certifications are better prepared to address today's evolving threat environment.

Foreword

HITRUST has always known the threat landscape as dynamic. Attackers change tactics. New vulnerabilities surface. AI is accelerating parts of both discovery and exploitation. A static security framework cannot assume that yesterday's controls are enough.

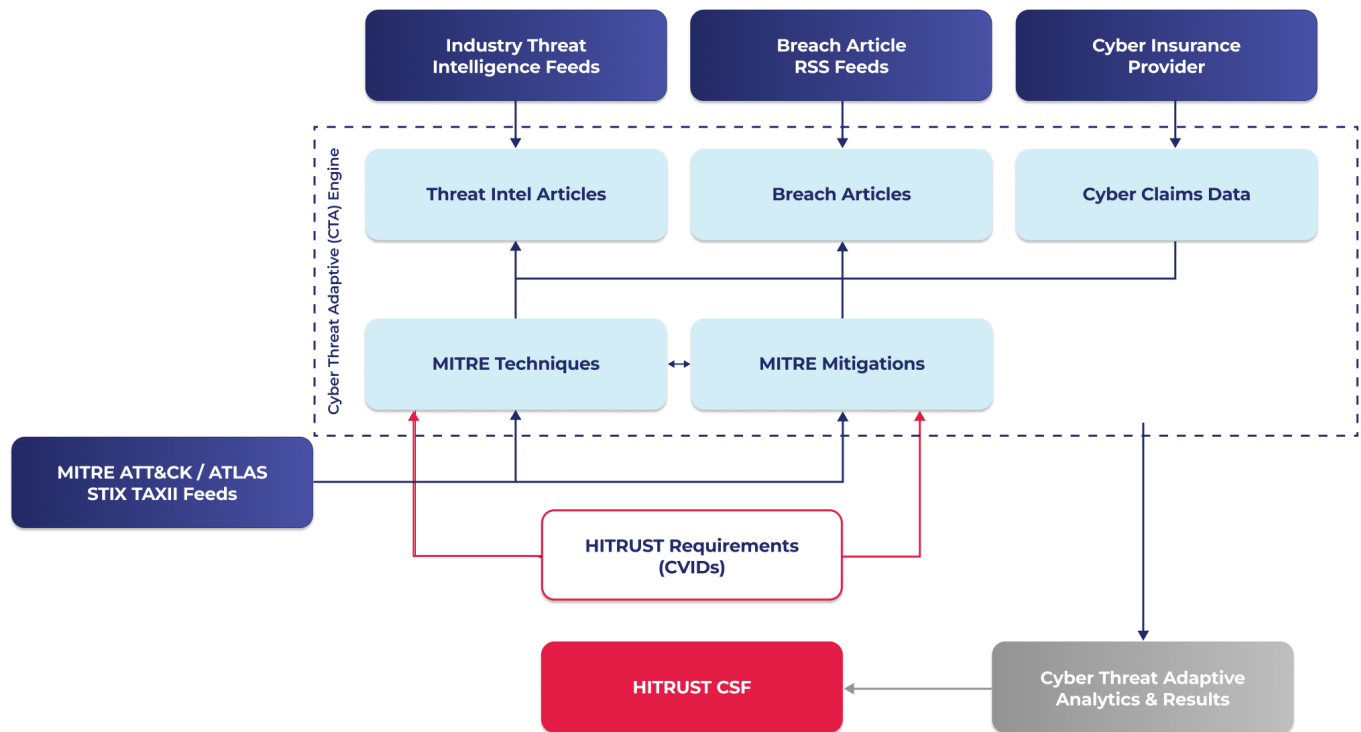
That is the reason HITRUST created the Cyber Threat Adaptive (CTA) program. CTA gives HITRUST a structured way to use threat intelligence, vulnerability research, and real-world attack data to keep the HITRUST CSF current. The purpose is to make sure HITRUST control requirements continue to address the risks organizations are actually facing, while preserving the assurance value of HITRUST assessments and certifications.

HITRUST has already started strengthening control guidance and expectations in areas where vulnerability discovery and exploit conditions are changing quickly. That work includes vulnerability identification and remediation, secure software development, dependency management, and detection and response. These updates are being handled through existing CSF, CTA, and MyCSF services, with a focus on improving threat relevance without creating unnecessary assessment burden.

The news of new frontier models like Anthropic's Mythos and the creation of Project Glasswing is a wakeup call to the industry. However, it is not unexpected from HITRUST's perspective. Time and again there is a disruptive technology that changes the game for defenders. It is the ability for defenders to respond, adapt and be resilient that is key. HITRUST is currently evaluating Project Glasswing and related information to determine the right CSF updates. As such, this CTA release does not include control changes specific to Project Glasswing. Future updates, including updates outside of cycle will be made.

Our Approach

HITRUST uses a comprehensive, continuous process to identify threats, align effective mitigations, and help organizations respond with confidence. Through this process, HITRUST regularly reviews and updates the framework to keep pace with the evolving threat landscape. The diagram below illustrates how we orchestrate this process and use it to progressively update the HITRUST framework:



The output of this process informs which requirements are included in HITRUST's e1 and i1 assessments and certifications. In general, the e1 is a streamlined baseline of security practices for organizations starting their HITRUST journey, with an important emphasis on cyber essentials, while the larger i1 offers a more comprehensive level of assurance and serves as the basis for our tailorable 2-year r2 assessment. The e1 addresses many of the most common entry points for attackers such as phishing, command and scripting, and process injection while the i1 introduces more robust controls around network traffic management, application allow listing, and more. The HITRUST AI Security Certification extends this approach to AI systems by enabling providers to demonstrate that they sufficiently mitigate cybersecurity threats to the AI technologies they have deployed. For more information on the e1, i1, r2, and AI Security Certification see:

- [HITRUST® 1-year \(e1\) Validated Assessment | HITRUST](#)
- [HITRUST® 1-year \(i1\) Validated Assessment | HITRUST®](#)
- [HITRUST® 1-year \(r2\) Validated Assessment | HITRUST®](#)
- [HITRUST® AI Security Assessment | HITRUST®](#)

Throughout this report, you will see references to MITRE ATT&CK and ATLAS techniques and mitigations which can be found here: [MITRE ATT&CK®](#) and [MITRE ATLAS®](#)

You will also see HITRUST Cross Version Identifiers (CVIDs) for HITRUST CSF requirements. To see the full text of the requirement statements, [Download the HITRUST CSF for free](#)

Summary of Findings

During our most recent analysis period, we expanded our evaluation beyond the threat techniques traditionally represented in MITRE ATT&CK by incorporating MITRE ATLAS, MITRE's knowledge base of adversary tactics and techniques targeting AI-enabled systems. Based on our review of 4,761 threat articles and 399,764 MITRE ATT&CK and MITRE ATLAS indicators, we confirmed that the requirements included in the HITRUST AI Security Certification remain responsive to the evolving AI threat landscape with over 97% coverage of adversarial techniques. Additionally, this review confirms that the e1, i1, r2 requirement selections continue to be responsive to the current threat landscape with 98.19% coverage of adversarial techniques for the e1 and 100% for the i1 and r2.

Top MITRE ATLAS Techniques and Mitigations

While the top traditional cyber attack methods have remained relatively consistent, AI-related tactics and techniques reflected in MITRE ATLAS have increased significantly. With the inclusion of these attacks in our analysis, we are seeing the volume of AI-specific techniques surpass many traditionally prevalent methods. The top three AI-related techniques observed in the first quarter of 2026 are:

Commonly Sighted AI Techniques in Q1 2026	MITRE Mitigations	HITRUST CVIDs
User Execution (AML.T0011) <i>Execution</i>		
Coming in as the new top attack technique seen in Q1 is the method of User Execution. In order to gain execution on a target system, adversaries rely on specific actions taken by a user, for example, opening a malicious document or file. The most popular sub-technique reported was the Malicious Package (AML.T0011.001). Creating malicious software packages has always been a successful, albeit time consuming, technique. With the use of AI, campaigns can be scaled exponentially. Other sub-techniques associated with User Execution include Unsafe AI Artifacts, Poisoned AI Agent Tool, and Malicious Link.	AML.M0011 - Restrict Library Loading AML.M0014 - Verify AI Artifacts AML.M0016 - Vulnerability Scanning AML.M0018 - User Training AML.M0023 - AI Bill of Materials	2870.0 (AISEC) 3048.0 (AISEC) 3040.0 (AISEC) 3066.0 (AISEC)
Phishing (AML.T0052) <i>Initial Access, Lateral Movement</i>		
Always a popular and successful technique, it is no surprise that adversaries are leveraging AI to phish their targets. Generative AI programs can be programmed with a meta prompt in order to phish for sensitive information and create complex schemes including the use of LLMs for synthetic text, visual deepfakes of faces, and audio deepfakes of speech. Using these tools in combination or as a tool to spear phish individuals, companies, and/or industries is becoming increasingly common.	AML.M0018 - User Training AML.M0034 – Deepfake Detection	3040.0 (AISEC)

Continued on the next page

Exfiltration via AI Agent Tool Invocation ([AML.T0086](#))

Exfiltration

When AI Agents can perform write operations, adversaries can prompt the agent to exfiltrate data. Targeted information can be encoded into input parameters and transmitted without raising alarm. Outside of exfiltrating sensitive data, variants can include sending email, creating or modifying existing documents, updating CRM records, or generating media such as images or videos.

AML.M0024 - AI Telemetry Logging
AML.M0026 - Privileged AI Agent Permissions Configuration
AML.M0027 - Single-User AI Agent Permissions Configuration
AML.M0028 - AI Agent Tools Permissions Configuration
AML.M0029 - Human In-the-Loop for AI Agent Actions
AML.M0030 - Restrict AI Agent Tool Invocation on Untrusted Data
AML.M0032 - Segmentation of AI Agent Components
AML.M0033 - Input and Output Validation for AI Agent Components

3054.0 (AISEC)
3039.0 (AISEC)
2948.0 (AISEC)

Suggested Actions

Based on the analysis of AI enabled attacks of this quarter, we recommend the following:

- Ensure AI specific security topics are included in **employee training** no less than annually for all teams involved in AI software and model creation and deployment, as well as data science and cybersecurity personnel.
 - See *HITRUST CVID 3040.0 (AISEC)*
- Prior to processing, **actively filter user inputs**, including attachments for suspicious and/or unexpected values or patterns which could be adversarial or malicious. Consider formulating filters in conjunction with a maintained, **documented inventory of data used to train, test, and validate AI models**.
 - See *HITRUST CVID 2948.0 (AISEC)*, *CVID 3066.0 (AISEC)* and *CVID 3048.0 (AISEC)*
- Practice **Human-in-the-loop** principle in the design of the AI application by allowing human operators the ability to evaluate model outputs before relying on them, as well as intervene if deemed necessary ahead of AI model-initiated actions.
 - See *HITRUST CVID 3039.0 (AISEC)*

Top MITRE ATT&CK Techniques and Mitigations

From this data, we identified the following as the top three techniques for this period:

Commonly Sighted Technique in Q1 2026	MITRE Mitigations	HITRUST CVIDs
Phishing (T1566) <i>Initial Access</i>		
Longstanding as the most common initial attack vector, this period was no exception and further validates the importance of anti-phishing training and email security. These six months saw an increase in spear phishing campaigns empowered by more advanced AI capabilities – enabling attackers to perform at a scale that was previously not available. These techniques were used in a blend of attacks aimed at either implanting persistent threats such as malware and ransomware, performing intelligence gathering, or achieving financial gain.	M1047 – Audit M1031 – Network Intrusion Prevention M1054 – Software Configuration M1017 – User Training	0599.0 (il) 0880.0 (il) 0886.1 (e1/il) 2316.0 (e1/il)
Drive-by Compromise (T1189) <i>Initial Access</i>		
An attacker may be able to gain access to a system through a user visiting a website over the normal course of browsing. Multiple methods of exploitation exist including compromising a legitimate website, malicious ads paid for and served through legitimate ad providers, and leveraging built-in application interfaces by inserting malicious scripts. Mitigation of a technique with such a broad attack surface is typically a piecemeal effort. Security best practices such as user training and ensuring browsers and plugins are up to date can go a long way. When combined with technical implementations of exploit protection applications, restricting certain web-based content as a whole, and browser sandboxing, many attacks using this technique can be thwarted.	M1048 – Application Isolation and Sandboxing M1050 – Exploit Protection M1021 – Restrict Web-Based Content M1051 – Update Software M1017 – User Training	0884.0 (il) 0946.0 (il) 0189.0 (il) 0873.0 (il) 1989.0 (il) 1369.0 (il) 2317.1 (e1/il) 2316.0 (e1/il) 0343.0 (il)
Exploit Public-Facing Application (T1190) <i>Initial Access</i>		
This technique relies on exploiting a weakness in an internet-facing host or system. This is unique from the Drive-by Compromise detailed above as the focus here is on compromising the host website/web server or database, versus the client endpoint visiting a website. There are several mitigating controls to implement to avoid falling victim to this attack technique, as well as limit impact in case exploitation does occur.	M1048 – Application Isolation and Sandboxing M1050 – Exploit Protection M1035 – Limit Access to Resource Over Network M1030 – Network Segmentation M1026 – Privileged Account Management M1051 – Update Software M1016 – Vulnerability Scanning	0884.0 (il) 0946.0 (il) 0117.0 (il) 0160.0 (e1/il) 0297.0 (il) 0035.0 (e1/il) 1369.0 (il) 2317.1 (e1/il) 2367.0 (e1/il)

Current Active Attacks

Additionally, we reviewed 259 real-world breaches reported during Q1 2026 and analyzed the techniques used to carry out those attacks. This analysis reconfirmed phishing-based attacks as the most prevalent trend observed in recent periods. In most cases, phishing techniques were used to facilitate data exfiltration or malware deployment. These findings remain consistent with the results of our quantitative threat analysis and continue to reinforce a major focus of our adaptive assessments.

Suggested Actions

Based on the results of this quarter, we recommend the following:

- Ensure that comprehensive **role-based security training** policies and procedures are developed, implemented effectively, and their adherence is measured and managed. Phishing and spear phishing attacks are continuously the most common and most effective attack vectors seen and are evolving rapidly with the increased application of large language model assisted attacks. Applying dedicated **phishing awareness training** is a vital step to protecting a workforce.
 - See HITRUST CVIDs 0343.0 (i1), 0599.0 (i1), 0880.0 (i1), 0886.1 (e1/i1), and 2316.9 (e1/i1)
- Ensure the **timely installation of anti-malware** protective measures and technologies, including regular updates, upgrades, and software scans. Additionally, **configure malicious code and spam protection** to detect known threats and low effort attacks.
 - See HITRUST CVIDs 0884.0 (i1) and 0873.0 (i1)
- Once a potential technical vulnerability has been identified, **identify associated risks and actions** to be taken. Perform these actions in a timely manner.
 - See HITRUST CVIDs 1369.0 (i1)
- Protect your network perimeter and key points within the network by implementing technical tools such as **intrusion detection systems (IDS)/intrusion prevention systems (IPS)**. Ensure these are implemented on the wireless side of the firewall, and update all engines, baselines, and signatures on a regular basis.
 - See HITRUST CVID 0946.0 (i1)

CONCLUSION



In an era of increasingly diverse cyberattacks, aligning your security assessment and certification to match your organization's risk profile is essential. HITRUST's e1, i1, and r2 certifications continue to be responsive to the most common techniques and their mitigations. HITRUST regularly reviews updated threat intelligence and breach data to continually refine the control selection in the e1, i1, and r2 assessments. This ensures that the assessments evolve in response to real-world adversarial tactics and remain aligned with emerging threats over time.

We are proposing to enhance a selection of e1, i1, and r2 requirements to better address context-dependent vulnerabilities and threats posed by AI-enabled adversaries.

HITRUST's AI Security Certification control set is similarly threat adaptive to ensure your AI-enabled services remains on the cutting edge of defense as new attack techniques materialize.

Whether you're looking to build confidence through a baseline "essentials" approach (e1), need a more comprehensive, implemented set of controls (i1), or achieve the highest level of assurance through the tailorable, risk-based 2-year certification (r2), each path is enhanced by HITRUST's continual threat monitoring to help safeguard against the most common and damaging attacks in today's rapidly changing threat landscape.