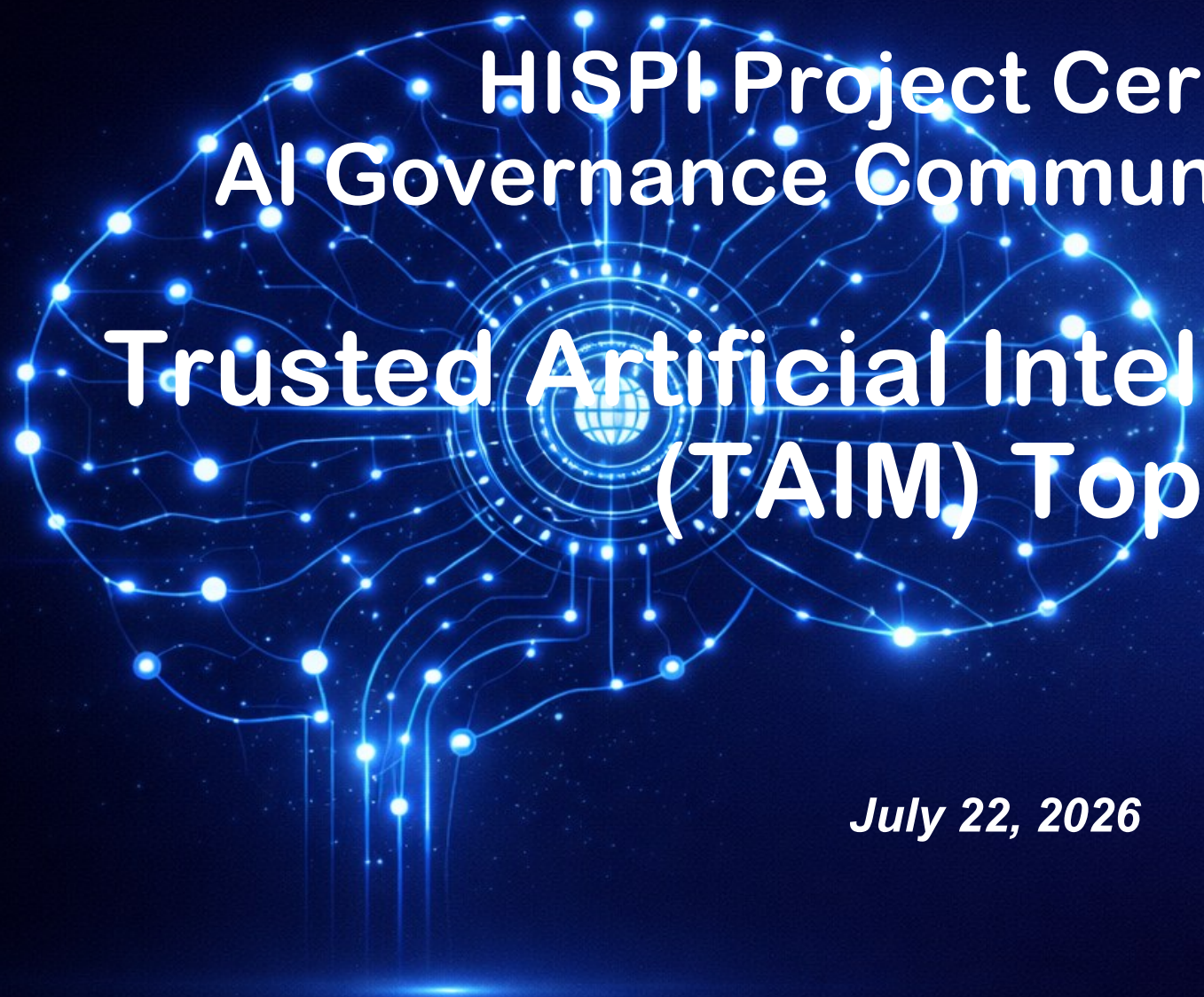




HISPI Project Cerebellum AI Governance Community of Practice

Trusted Artificial Intelligence Model (TAIM) Top 20

July 22, 2026



Project Cerebellum

AI should cause no harm but enhance the quality of human life

Trusted AI starts with clear guardrails, accountable decisions, and measured risk.

Vision

Provide effective guardrails for AI use cases.

Mission

Harmonize best practices, standards, and frameworks for AI ecosystems.

Trusted AI Model (TAIM)



Risk of AI: Production impact from poor guardrails

A broadly scoped task can become a production incident when access, testing, and rollback are not governed.

What

Agent deleted production data.

Why

Broad token allowed unsafe action.

Lesson

Govern access, testing, and rollback.

Risk control takeaway

AI agents need the same basic safeguards as other high-impact automation: least privilege, separation of environments, change control, logging, rollback, and tested recovery.

AI Incidents: The risk landscape

- **Fraud and deepfakes** Voice clones, fake endorsements, scam ads
- **Bias and access** Unfair outcomes, language barriers, exclusion
- **Privacy and identity** Misuse of names, likeness, personal data
- **Safety and reliability** Bad alerts, unsafe automation, failures
- **Security and supply chain** Third-party systems, agent access, data exposure

Even a narrowly scoped task can become a production incident when access, testing, and rollback are not governed.

Trusted AI Model (TAIM) Domains

(NIST AI RMF functions)

TAIM helps organizations introduce AI deliberately and mitigate risk across the AI lifecycle.

GOVERN

Set direction and
accountability

MAP

Understand context
and risks

MEASURE

Test and evaluate
trustworthiness

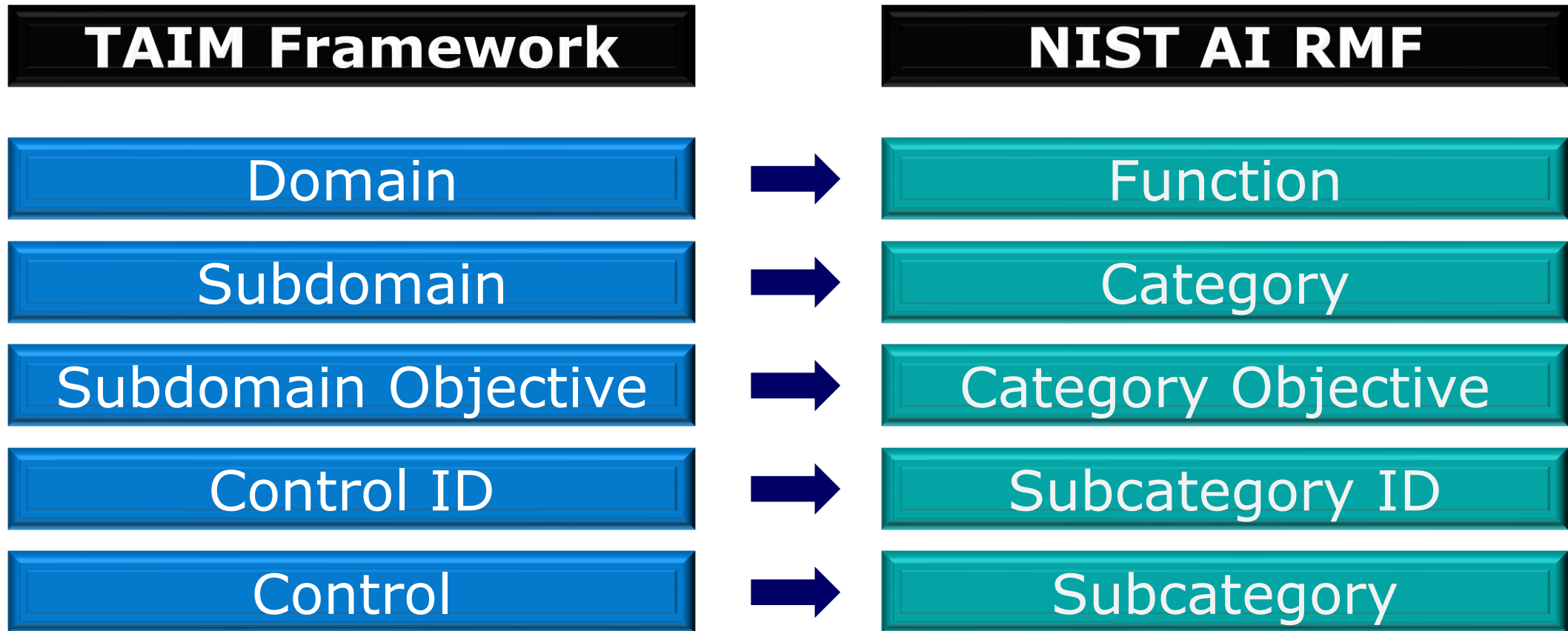
MANAGE

Respond, improve,
and report

The domains create a repeatable pathway: decide what AI should do, understand where it can go wrong, test the trust factors, and manage what happens next.

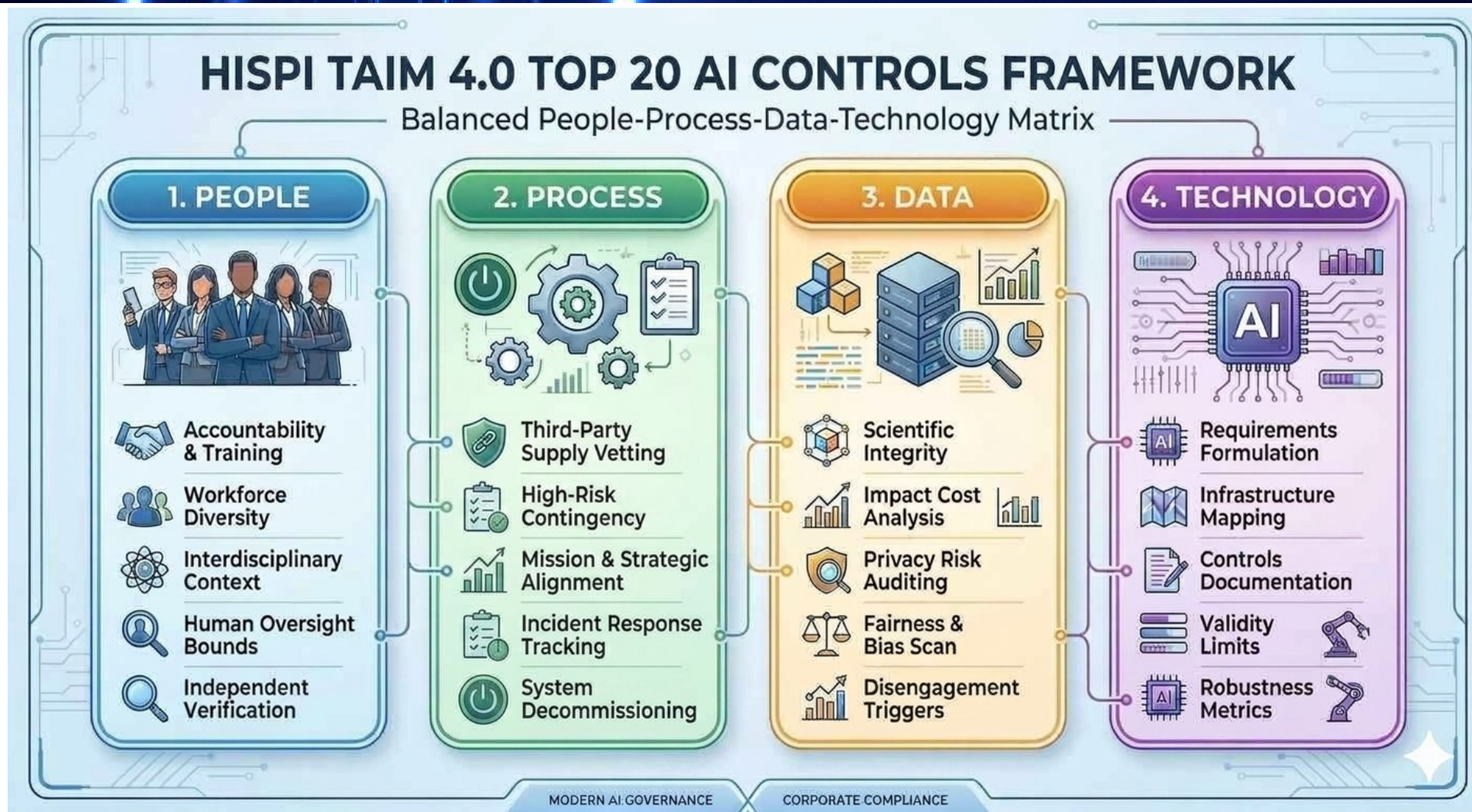
Trusted AI Model (TAIM) Taxonomy

(relationship to NIST AI RMF)



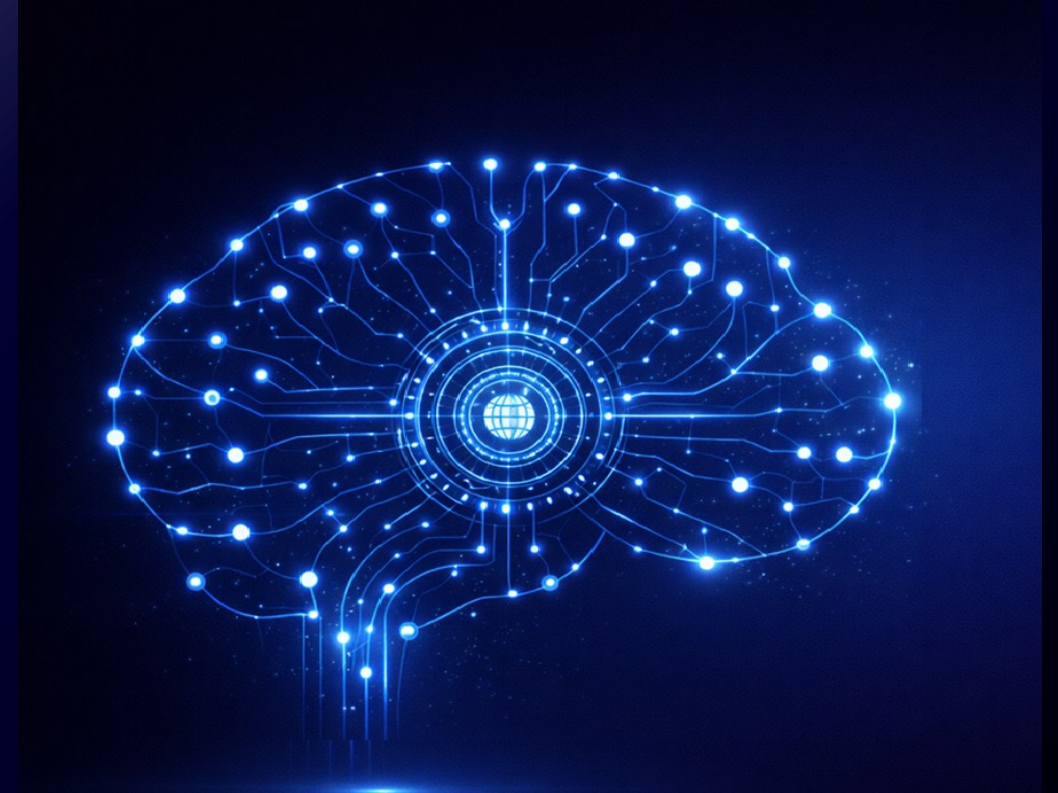
Purpose: give organizations a practical control language while preserving alignment to the NIST Artificial Intelligence Risk Management Framework.

TAIM Controls



TAIM 4.0

Top 20 Controls



TAIM Snapshot

A single model view connects domains, objectives, controls, mappings, principles, and supporting evidence.

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P
	Domain (Function)	Subdomain (Category ID)	Subdomain Objective (Category Objective)	Control ID (Subcategory ID)	Control (Subcategory)	Control Question	Recommendation	Top 20 Control	Related AI Incidents	PEOPLE	PROCESS	DATA	TECHNOLOGY	ISO/IEC TR 24028:2020 Information technology - Artificial intelligence - Overview of trustworthiness in artificial intelligence	BS ISO/IEC 23053:2022 Framework for Artificial Intelligence (AI) Systems Using Machine Learning (ML) (British Standard)	ISO/IEC 23894:2023 Information technology - Artificial intelligence - Guidance on risk management
1																
10	GOVERN	GOVERN 2	managing AI risks.	GOVERN 2.2	The organization's personnel and partners receive AI risk management training to enable them to perform their duties and responsibilities consistent with related policies, procedures, and agreements.	Do the organization's personnel and partners receive AI risk management training to enable them to perform their duties and responsibilities consistent with related policies, procedures, and agreements?	Ensure that the organization's personnel and partners receive AI risk management training to enable them to perform their duties and responsibilities consistent with related policies, procedures, and agreements.	✓	https://ai-incidents.projectcerebellum.com/taim/govern-2-2	✓	✓	✓	✓	5.4		5.4.1 - ISO 31000:2018
				GOVERN 2.3	Executive leadership of the organization takes responsibility for decisions about risks associated with AI system development and deployment.	Does executive leadership of the organization take responsibility for decisions about risks associated with AI system development and deployment?	Ensure that executive leadership of the organization takes responsibility for decisions about risks associated with AI system development and deployment.			✓	✓	✓	✓		A.2.2.3	5.5
11			Workforce diversity, equity, inclusion, and accessibility processes are specified in the	GOVERN 3.1	Decision-making related to mapping, measuring, and managing AI risks throughout the lifecycle is informed by a diverse team (e.g., diversity of demographics	Are decision-making related to mapping, measuring, and managing AI risks throughout the lifecycle informed by a diverse team (e.g., diversity of demographics, disciplines	Ensure that decision-making related to mapping, measuring, and managing AI risks throughout the lifecycle is informed by a diverse team (e.g., diversity of demographics, disciplines	✓	https://ai-incidents.projectcerebellum.com/taim/govern-3-1	✓	✓	✓	✓	7.1	7.2	6.2 - ISO 31000:2018
	TRUSTED-AI-MODEL-(TAIM)-V4.0															

Incidents

<https://ai-incidents.projectcerebellum.com/taim/govern-1-7>

Applicable Regulation

PPDT Governance Framework

PEOPLE

PROCESS

DATA

TECHNOLOGY

Control ID

GOVERN 1.7

Risk Statement

To ensure a purpose-driven culture focused on risk understanding and management.

Subdomain Objective

Processes and procedures are in place for decommissioning and phasing out AI systems safely and in a manner that does not increase risks or decrease the organization's trustworthiness.

Control

Processes and procedures are in place for decommissioning and phasing out AI systems safely and in a manner that does not increase risks or decrease the organization's trustworthiness.

Mappings/Crosswalk

BS ISO/IEC 23053:2022 8.1
ISO/IEC DIS 5338:2022 6.3.3.3, 6.4.8.2, 6.4.8.3, 6.4.17.1
ISO/IEC 38507:2022 4.3, 5.3
ISO/IEC 22989:2022 6.1, 6.2.9
ISO/IEC 42001:2023 B.6.2.6 - AI system operation and monitoring
Texas: Responsible Artificial Intelligence Governance Act (TRAIGA)

Core Principles

TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓					✓	

©2026 HISPI



11

Incidents		Control ID	Subdomain Objective	Control	Mappings/Crosswalk	
https://ai-incidents.projectcerebellum.com/taim/govern-2-2		GOVERN 2.2	Accountability structures are in place so that the appropriate teams and individuals are empowered, responsible, and trained for mapping, measuring, and managing AI risks.	The organization’s personnel and partners receive AI risk management training to enable them to perform their duties and responsibilities consistent with related policies, procedures, and agreements.	ISO/IEC TR 24028:2020 5.4 ISO/IEC 31000:2018 5.4.1 ISO/IEC DIS 5338:2022 6.4.12.1 ISO/IEC 38507:2022 4.1, 4.3, 5.1, 6.7.2, 6.7.5, A.1.3 ISO/IEC 22989:2022 6.2.2 NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 7.2 – Competence Texas: Responsible Artificial Intelligence Governance Act (TRAIGA)	
Applicable Regulation		Risk Statement				
To ensure a purpose-driven culture focused on risk understanding and management.						
HIPAA		PCI-DSS				
SOC2		EU AI ACT				
EU GDPR		WHITE HOUSE AI EO				
PPDT Governance Framework						
PEOPLE		PROCESS				
DATA		TECHNOLOGY				
Core Principles						
TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓			✓		

©2026 HISPI



12



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk
https://ai-incidents.projectcerebellum.com/taim/govern-3-1	GOVERN 3.1	Workforce diversity, equity, inclusion, and accessibility processes are prioritized in the mapping, measuring, and managing of AI risks throughout the lifecycle.	Decision-making related to mapping, measuring, and managing AI risks throughout the lifecycle is informed by a diverse team (e.g., diversity of demographics, disciplines, experience, expertise, and backgrounds).	ISO/IEC TR 24028:2020 7.1 BS ISO/IEC 23053:2022 7.2 ISO/IEC 31000:2018 6.2 ISO/IEC DIS 5338:2022 6.3.3 ISO/IEC 38507:2022 6.1, 6.3, 6.5, 6.7.1, 6.7.2, 6.7.3, Annex A ISO/IEC 22989:2022 6.2.2 ISO/IEC 42001:2023 B.4.6 - Human resources B.5.4 - Assessing AI system impact on individuals and groups of individuals Texas: Responsible Artificial Intelligence Governance Act (TRAIGA)
	Risk Statement			
	To ensure a purpose-driven culture focused on risk understanding and management.			
Applicable Regulation				
PPDT Governance Framework				
PEOPLE	PROCESS			
DATA	TECHNOLOGY			



Incidents		Control ID	Subdomain Objective	Control	Mappings/Crosswalk		
https://ai-incidents.projectcerebellum.com/taim/govern-6-1		GOVERN 6.1	Policies and procedures are in place to address AI risks and benefits arising from third-party software and data and other supply chain issues.	Policies and procedures are in place that address AI risks associated with third-party entities, including risks of infringement of a third-party’s intellectual property or other rights.	ISO/IEC TR 24028:2020 5.5 ISO/IEC 31000:2018 5.4.1 ISO/IEC DIS 5338:2022 6.4.3.3 ISO/IEC 38507:2022 5.5, 6.4, A.3 ISO/IEC 22989:2022 8.6.1 NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 B.10.2 - Allocating responsibilities B.10.3 – Suppliers Texas: Responsible Artificial Intelligence Governance Act (TRAIGA)		
Applicable Regulation		Risk Statement					
To ensure a purpose-driven culture focused on risk understanding and management.							
HIPAA		PCI-DSS					
SOC2		EU AI ACT					
EU GDPR		WHITE HOUSE AI EO					
PPDT Governance Framework							
PEOPLE		PROCESS					
DATA		TECHNOLOGY					
Core Principles							
TRANSPARENCY (EXPLAINABLE)		ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓		✓		✓	✓	✓	

©2026 HISPI



14



<https://ai-incidents.projectcerebellum.com/taim/govern-6-2>

Applicable Regulation

HIPAA	PCI-DSS
SOC2	EU AI ACT
EU GDPR	WHITE HOUSE AI EO

PPDT Governance Framework

PEOPLE	PROCESS
DATA	TECHNOLOGY

PEOPLE	PROCESS
DATA	TECHNOLOGY

GOVERN 6.2

To ensure a purpose-driven culture focused on risk understanding and management.

Policies and procedures are in place to address AI risks and benefits arising from third-party software and data and other supply chain issues.

Contingency processes are in place to handle failures or incidents in third-party data or AI systems deemed to be high-risk.

ISO/IEC TR 24028:2020 7.1
BS ISO/IEC 23053:2022 8.3
ISO/IEC 31000:2018 5.4.1
ISO/IEC DIS 5338:2022
6.4.15.3
ISO/IEC 38507:2022 6.7.5
ISO/IEC 22989:2022 5.15.4
NIST CSF 2.0
COBIT 5
ISO/IEC 42001:2023
B.10.2 - Allocating
responsibilities
B.10.3 – Suppliers
Texas: Responsible Artificial
Intelligence Governance Act
(TRAIGA)

TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓		✓	✓	✓	



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk				
https://ai-incidents.projectcerebellum.com/taim/map-1-2	MAP 1.2	Context is established and understood.	Interdisciplinary AI actors, competencies, skills, and capacities for establishing context reflect demographic diversity and broad domain and user experience expertise, and their participation is documented. Opportunities for interdisciplinary collaboration are prioritized.	ISO/IEC TR 24028:2020 9.8 ISO/IEC DIS 5338:2022 5.1, 6.2.3.3, 6.2.4.1, 6.2.4.3, 6.2.6.1, 6.2.6.3, 6.4.9.4 ISO/IEC 38507:2022 A.1.3 ISO/IEC 22989:2022 5.1, 6.2.2 NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 B.4.6 - Human resources 7.2 – Competence Utah: AI Policy Act (2024)				
Applicable Regulation	Risk Statement	Core Principles						
HIPAA PCI-DSS SOC2 EU AI ACT EU GDPR WHITE HOUSE AI EO	To ensure sufficient contextual knowledge about AI system impacts to inform an initial go/no-go decision about whether to design, develop, or deploy an AI system.							
PPDT Governance Framework								
PEOPLE PROCESS DATA								
		TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
		✓	✓	✓	✓	✓	✓	✓

©2026 HISPI



16



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk				
https://ai-incidents.projectcerebellum.com/taim/map-1-3	MAP 1.3	Context is established and understood.	The organization’s mission and relevant goals for AI technology are understood and documented.	ISO/IEC TR 24028:2020 9.10.2.1 BS ISO/IEC 23053:2022 6.2.1, 6.3, 7.6, 8.1, 8.8 ISO/IEC DIS 5338:2022 6.4.1, 6.4.13.1, A.2.5 ISO/IEC 38507:2022 4.1, A.1.1, A.1.4, A.2 ISO/IEC 22989:2022 6.2.2 ISO/IEC 42001:2023 B.4.6 - Human resources 7.2 – Competence				
Applicable Regulation	Risk Statement							
	To ensure sufficient contextual knowledge about AI system impacts to inform an initial go/no-go decision about whether to design, develop, or deploy an AI system.							
PPDT Governance Framework		Core Principles						
PEOPLE	PROCESS							
DATA	TECHNOLOGY							
		TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
		✓	✓		✓			

©2026 HISPI



17



Incidents		Control ID	Subdomain Objective	Control	Mappings/Crosswalk	
https://ai-incidents.projectcerebellum.com/taim/map-1-6		MAP 1.6	Context is established and understood.	System requirements (e.g., “the system shall respect the privacy of its users”) are elicited from and understood by relevant AI actors.	ISO/IEC TR 24028:2020 9.8 BS ISO/IEC 23053:2022 8.3 ISO/IEC 31000:2018 5.4.1 ISO/IEC DIS 5338:2022 6.3.4.3, 6.4.1.3, 6.4.3.3 6.4.8.3, 6.4.17.3 ISO/IEC 38507:2022 5.3, 6.4, 6.6.2, 6.7.2, 6.7.3, 6.7.5, A.3 ISO/IEC 22989:2022 3.2.11, 3.5.14, 3.5.15, 3.5.16, 5.15.8, 5.19, 6.2, 6.1, 6.2.2, 8.6.2, Annex A NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 B.6.2.2 - AI system requirements and specification B.5.4 - Assessing AI system impact on individuals and groups of individuals B.5.5 - Assessing societal impacts of AI systems	
Applicable Regulation		Risk Statement				
To ensure sufficient contextual knowledge about AI system impacts to inform an initial go/no-go decision about whether to design, develop, or deploy an AI system.						
HIPAA		PCI-DSS				
SOC2		EU AI ACT				
EU GDPR		WHITE HOUSE AI EO				
PPDT Governance Framework						
PEOPLE		PROCESS				
DATA		TECHNOLOGY				
Core Principles						
TRANSPARENCY (EXPLAINABLE)		ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)
ROBUSTNESS (EXPLAINABLE)						
✓		✓		✓	✓	

©2026 HISPI



18



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk
https://ai-incidents.projectcerebellum.com/taim/map-2-3	MAP 2.3 Risk Statement	Categorization of the AI system is performed.	Scientific integrity and Test & Evaluation, Validation & Verification (TEVV) considerations are identified and documented, including those related to experimental design, data collection and selection (e.g., availability, representativeness, suitability), system trustworthiness, and construct validation.	ISO/IEC TR 24028:2020 5.3 BS ISO/IEC 23053:2022 6.4, 6.5.3.1, 8.2, 8.4, 8.5, 8.6, 8.8 ISO/IEC DIS 5338:2022 5.3, 6.3.1.3, 6.4.13, 6.4.14, A.2.5, A.3 ISO/IEC 38507:2022 4.3, 6.6.1, 6.6.2 ISO/IEC 22989:2022 3.5.17, 5.10, 5.11.9.2, 5.16, 6.1, 6.2.3, 6.2.4 ISO/IEC 42001:2023 B.6.1.3 - Processes for responsible design and development of AI systems B.6.2.7 - AI system technical documentation B.7.2 - Data for development and enhancement of AI system B.7.3 - Acquisition of data B.7.4 - Quality of data for AI systems B.7.5 - Data provenance B.7.6 - Data preparation B.6.2.4 - AI system verification and validation
Applicable Regulation	To ensure sufficient contextual knowledge about AI system impacts to inform an initial go/no-go decision about whether to design, develop, or deploy an AI system.			
PPDT Governance Framework				
PEOPLE DATA	PROCESS TECHNOLOGY			

Core Principles

TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓				✓	



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk		
https://ai-incidents.projectcerebellum.com/taim/map-3-2	MAP 3.2	AI capabilities, targeted usage, goals, and expected benefits and costs compared with appropriate benchmarks are understood.	Potential costs, including non-monetary costs, which result from expected or realized AI errors or system functionality and trustworthiness – as connected to organizational risk tolerance – are examined and documented.	ISO/IEC TR 24028:2020 8.10, 9.8 BS ISO/IEC 23053:2022 6.5.4.1, 6.5.4.2, 7.4 ISO/IEC 23894:2023 6.4.2.6 ISO/IEC DIS 5338:2022 6.4.3.3, 6.4.12.3, 6.4.13.3, 6.4.15.3, 6.4.16.3 ISO/IEC 38507:2022 4.2, 6.4 ISO/IEC 22989:2022 6.2.2, 6.2.6, 7.4.1 ISO/IEC 42001:2023 B.5.2 - AI system impact assessment process B.5.3 - Documentation of AI system impact assessments B.5.4 - Assessing AI system impact on individuals and groups of individuals B.5.5 - Assessing societal impacts of AI systems 8.2 - AI risk assessment 8.3 - AI risk treatment 8.4 - AI system impact assessment		
Applicable Regulation	Risk Statement					
	To ensure sufficient contextual knowledge about AI system impacts to inform an initial go/no-go decision about whether to design, develop, or deploy an AI system.					
PPDT Governance Framework						
PEOPLE	PROCESS					
DATA	TECHNOLOGY					
Core Principles						
TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓				✓	✓

©2026 HISPI



20



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk		
https://ai-incidents.projectcerebellum.com/taim/map-3-5	MAP 3.5	AI capabilities, targeted usage, goals, and expected benefits and costs compared with appropriate benchmarks are understood.	Processes for human oversight are defined, assessed, and documented in accordance with organizational policies from the GOVERN function.	BS ISO/IEC 23053:2022 8.1 ISO/IEC 31000:2018 5.4.3 ISO/IEC DIS 5338:2022 6.4.16.3, 6.3.4.3, 6.4.3.3, 6.4.14.1 ISO/IEC 38507:2022 3.2.1, 4.3, 5.1, 5.3, 5.4, 5.5, 6.2, 6.3, 6.5, 6.7.3, A.1.1, A.1.3, A.1.4 ISO/IEC 22989:2022 5.1 ISO/IEC 42001:2023 B.6.1.3 - Processes for responsible design and development of AI systems B.6.2.7 - AI system technical documentation B.8.2 - System documentation and information for users		
Applicable Regulation	Risk Statement					
	To ensure sufficient contextual knowledge about AI system impacts to inform an initial go/no-go decision about whether to design, develop, or deploy an AI system.					
PPDT Governance Framework						
PEOPLE	PROCESS					
DATA	TECHNOLOGY					
Core Principles						
TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓				✓	✓

©2026 HISPI



21



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk			
https://ai-incidents.projectcerebellum.com/taim/map-4-1	MAP 4.1	Risks and benefits are mapped for all components of the AI system including third-party software and data.	Approaches for mapping AI technology and legal risks of its components – including the use of third-party data or software – are in place, followed, and documented, as are risks of infringement of a third party’s intellectual property or other rights.	ISO/IEC TR 24028:2020 7.1 ISO/IEC 31000:2018 5.4.1, 5.6 ISO/IEC DIS 5338:2022 6.1.1.3, 6.4.1.3, 6.4.3.3, 6.4.8.2, 6.4.8.3, 6.4.17.1 ISO/IEC 38507:2022 5.5, 6.4, A.3 ISO/IEC 22989:2022 5.15.1, 5.17 NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 4.1 - Understanding the organization and its context B.2.2 - AI policy B.9.2 - Processes for responsible use of AI systems B.9.4 - Intended use of the AI system TEXAS: RESPONSIBLE ARTIFICIAL INTELLIGENCE GOVERNANCE ACT (TRAIGA)			
Applicable Regulation	Risk Statement						
HIPAA PCI-DSS SOC2 EU AI ACT EU GDPR WHITE HOUSE AI EO							
PPDT Governance Framework	To ensure sufficient contextual knowledge about AI system impacts to inform an initial go/no-go decision about whether to design, develop, or deploy an AI system.						
PEOPLE PROCESS DATA TECHNOLOGY							
Core Principles							
	TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
	✓	✓			✓		

©2026 HISPI



22



Incidents		Control ID	Subdomain Objective	Control	Mappings/Crosswalk	
https://ai-incidents.projectcerebellum.com/taim/map-4-2		MAP 4.2	Risks and benefits are mapped for all components of the AI system including third-party software and data.	Internal risk controls for components of the AI system, including third-party AI technologies, are identified and documented.	ISO/IEC TR 24028:2020 5.4 ISO/IEC 31000:2018 5.4.1 ISO/IEC DIS 5338:2022 6.4.8.3 ISO/IEC 38507:2022 6.7.2 ISO/IEC 22989:2022 6.1 NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 B.6.2.7 - AI system technical documentation B.8.2 - System documentation and information for users B.10.3 - Suppliers	
Applicable Regulation		Risk Statement				
To ensure sufficient contextual knowledge about AI system impacts to inform an initial go/no-go decision about whether to design, develop, or deploy an AI system.						
HIPAA PCI-DSS SOC2 EU AI ACT EU GDPR WHITE HOUSE AI EO						
PPDT Governance Framework						
PEOPLE PROCESS DATA TECHNOLOGY						
Core Principles						
TRANSPARENCY (EXPLAINABLE)		ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)
ROBUSTNESS (EXPLAINABLE)						
✓		✓			✓	

©2026 HISPI



23



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk
https://ai-incidents.projectcerebellum.com/taim/measure-1-3	MEASURE 1.3 Risk Statement To ensure objective, repeatable, or scalable test, evaluation, verification, and validation (TEVV) processes including metrics, methods, and methodologies are in place, followed, and documented.	Appropriate methods and metrics are identified and applied.	Internal experts who did not serve as front-line developers for the system and/or independent assessors are involved in regular assessments and updates. Domain experts, users, AI actors external to the team that developed or deployed the AI system, and affected communities are consulted in support of assessments as necessary per organizational risk tolerance.	ISO/IEC TR 24028:2020 9.2 ISO/IEC 31000:2018 5.4.1, 6.4.1 ISO/IEC DIS 5338:2022 6.4.7.1, 6.4.7.3, 6.4.9.4 ISO/IEC 38507:2022 4.3, 5.2.3, 6.4 ISO/IEC 22989:2022 5.19.1, 5.19.5.1, 5.19.5.4, 6.2.8 NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 6.1.2 - AI risk assessment 9.2.2 - Internal audit programme B.5.2 - AI system impact assessment process B.5.4 - Assessing AI system impact on individuals and groups of individuals B.5.5 - Assessing societal impacts of AI systems
Applicable Regulation HIPAA PCI-DSS SOC2 EU AI ACT EU GDPR WHITE HOUSE AI EO				
PPDT Governance Framework PEOPLE PROCESS DATA TECHNOLOGY				

Core Principles

TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓	✓	✓	✓	✓	



Incidents		Control ID	Subdomain Objective	Control	Mappings/Crosswalk	
https://ai-incidents.projectcerebellum.com/taim/measure-2-5		MEASURE 2.5	AI systems are evaluated for trustworthy characteristics.	The AI system to be deployed is demonstrated to be valid and reliable.	ISO/IEC TR 24028:2020 9.7, 9.11.2 ISO/IEC DIS 5338:2022 5.1, 6.3.4.3, 6.4.11, 6.4.13, 6.4.14, A.2.5, A.3 ISO/IEC 38507:2022 6.6.1 ISO/IEC 22989:2022 3.5.9, 3.5.16, 5.15.2, 5.15.3, 5.15.4 NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 B.6.2.4 - AI system verification and validation B.6.2.5 - AI system deployment B.6.2.7 - AI system technical documentation B.8.2 - System documentation and information for users	
Risk Statement						
To ensure objective, repeatable, or scalable test, evaluation, verification, and validation (TEVV) processes including metrics, methods, and methodologies are in place, followed, and documented.						
Applicable Regulation						
HIPAA	PCI-DSS					
SOC2	EU AI ACT					
EU GDPR	WHITE HOUSE AI EO					
PPDT Governance Framework						
PEOPLE	PROCESS					
DATA	TECHNOLOGY					
			Core Principles			
TRANSPARENCY (EXPLAINABLE)		ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)
✓		✓	✓	✓	✓	✓

©2026 HISPI



25



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk				
https://ai-incidents.projectcerebellum.com/taim/measure-2-6	MEASURE 2.6	AI systems are evaluated for trustworthy characteristics.	The AI system is evaluated regularly for safety risks – as identified in the MAP function. The AI system to be deployed is demonstrated to be safe, its residual negative risk does not exceed the risk tolerance, and it can fail safely, particularly if made to operate beyond its knowledge limits. Safety metrics reflect system reliability and robustness, real-time monitoring, and response times for AI system failures.	ISO/IEC TR 24028:2020 7.2, 9.9, 9.11.3 BS ISO/IEC 23053:2022 8.1 ISO/IEC 31000:2018 5.6 ISO/IEC DIS 5338:2022 6.3.4.3 ISO/IEC 38507:2022 6.7.2, 6.7.3 ISO/IEC 22989:2022 Annex A NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 B.6.2.8 - AI system recording of event logs B.6.2.6 - AI system operation and monitoring B.6.2.4 - AI system verification and validation 8.2 - AI risk assessment Colorado: Colorado AI Act (2024) Tennessee: ELVIS Act (2024) Utah: AI Policy Act (2024) Washington: Biometric consent laws (2017) TEXAS: RESPONSIBLE ARTIFICIAL INTELLIGENCE GOVERNANCE ACT				
Applicable Regulation	Risk Statement							
HIPAA PCI-DSS SOC2 EU AI ACT EU GDPR WHITE HOUSE AI EO								
PPDT Governance Framework	To ensure objective, repeatable, or scalable test, evaluation, verification, and validation (TEVV) processes including metrics, methods, and methodologies are in place, followed, and documented.	Core Principles						
PEOPLE PROCESS DATA TECHNOLOGY								
		TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
		✓	✓	✓	✓	✓	✓	✓

©2026 HISPI



26



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk			
https://ai-incidents.projectcerebellum.com/taim/measure-2-10	MEASURE 2.10	AI systems are evaluated for trustworthy characteristics.	Privacy risk of the AI system – as identified in the MAP function – is examined and documented.	ISO/IEC TR 24028:2020 9.6, 9.10.4 BS ISO/IEC 23053:2022 8.1 ISO/IEC 31000:2018 5.4.1 ISO/IEC DIS 5338:2022 5.1, 6.3.4.3, 6.4.1.3, 6.4.3.3, 6.4.8.3, 6.4.17.3 ISO/IEC 38507:2022 5.3, 6.4, 6.6.2, 6.7.2, 6.7.3, 6.7.5, A.3 ISO/IEC 22989:2022 5.19.6.2, 6.1, 6.2.2, 8.6.2, Annex A NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 B.5.2 - AI system impact assessment process B.7.2 - Data for development and enhancement of AI system B.7.3 - Acquisition of data B.2.3 - Alignment with other organizational policies TEXAS: RESPONSIBLE ARTIFICIAL INTELLIGENCE GOVERNANCE ACT (TRAIGA)			
Applicable Regulation	Risk Statement						
HIPAA PCI-DSS SOC2 EU AI ACT EU GDPR WHITE HOUSE AI EO	To ensure objective, repeatable, or scalable test, evaluation, verification, and validation (TEVV) processes including metrics, methods, and methodologies are in place, followed, and documented.						
PPDT Governance Framework	Core Principles						
PEOPLE PROCESS DATA TECHNOLOGY	TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
	✓	✓	✓	✓	✓	✓	✓

©2026 HISPI



27



<https://ai-incidents.projectcerebellum.com/taim/measure-2-11>

Applicable Regulation	
HIPAA	PCI-DSS
SOC2	EU AI ACT
EU GDPR	WHITE HOUSE AI EO

PPDT Governance Framework

PEOPLE	PROCESS
DATA	TECHNOLOGY

©2026 HISPI

MEASURE 2.11

Risk Statement

To ensure objective, repeatable, or scalable test, evaluation, verification, and validation (TEVV) processes including metrics, methods, and methodologies are in place, followed, and documented.

Risk Statement
To ensure objective, repeatable, or scalable test, evaluation, verification, and validation (TEVV) processes including metrics, methods, and methodologies are in place, followed, and documented.

AI systems are evaluated for trustworthy characteristics.

Fairness and bias – as identified in the MAP function – are evaluated and results are documented.

ISO/IEC TR 24028:2020 8.4
BS ISO/IEC 23053:2022
 6.3, 8.1, 8.3, A.2.2.2
ISO/IEC 31000:2018 5.4.1
ISO/IEC DIS 5338:2022
 5.1, 6.3.4.3, 6.3.8.3, 6.4.3.3, 6.4.14.3
ISO/IEC 38507:2022
 4.2, 5.4, 6.4, 6.5, 6.7.2, 6.7.4
ISO/IEC 22989:2022
 3.2.2, 3.5.4, 5.10, 5.15.9, 7.4.1, Annex
 A
NIST CSF 2.0
COBIT 5
ISO/IEC 42001:2023
 B.5.5 - Assessing societal impacts of
 AI systems
 B.5.4 - Assessing AI system impact on
 individuals and groups of individuals
Illinois: Biometric Information
Privacy Act (2008)
TEXAS: RESPONSIBLE ARTIFICIAL
INTELLIGENCE GOVERNANCE ACT
(TRAIGA)

TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓	✓	✓	✓	✓	✓



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk
https://ai-incidents.projectcerebellum.com/taim/manage-2-4	MANAGE 2.4 Risk Statement To ensure that plans for prioritizing risk and regular monitoring and improvement are in place.	Strategies to maximize AI benefits and minimize negative impacts are planned, prepared, implemented, documented, and informed by input from relevant AI actors.	Mechanisms are in place and applied, and responsibilities are assigned and understood, to supersede, disengage, or deactivate AI systems that demonstrate performance or outcomes inconsistent with intended use.	ISO/IEC 23894:2023 6.4.2.4, 6.4.2.5, 6.4.2.6 ISO/IEC 31000:2018 6.3.4, 6.5.2 ISO/IEC DIS 5338:2022 6.4.17.1 ISO/IEC 38507:2022 3.2.3, 6.7 ISO/IEC 22989:2022 3.1.22, 5.17, 6.1, 6.2.9, Annex A NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 B.9.4 - Intended use of the AI system B.8.2 - System documentation and information for users B.6.2.7 - AI system technical documentation B.6.1.3 - Processes for responsible design and development of AI systems TEXAS: RESPONSIBLE ARTIFICIAL INTELLIGENCE GOVERNANCE ACT (TRAIGA)
Applicable Regulation HIPAA PCI-DSS SOC2 EU AI ACT EU GDPR WHITE HOUSE AI EO				
PPDT Governance Framework PEOPLE PROCESS DATA TECHNOLOGY				

Core Principles						
TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓	✓	✓	✓	✓	✓



Incidents	Control ID	Subdomain Objective	Control	Mappings/Crosswalk		
https://ai-incidents.projectcerebellum.com/taim/manage-4-3	MANAGE 4.3	Risk treatments, including response and recovery, and communication plans for the identified and measured AI risks are documented and monitored regularly.	Incidents and errors are communicated to relevant AI actors, including affected communities.	BS ISO/IEC 23053:2022 3.1.4, A.2.2.2 ISO/IEC 31000:2018 6.5.2, 6.7 ISO/IEC DIS 5338:2022 6.4.3.3, 6.4.15.3 ISO/IEC 38507:2022 4.3, 6.1, 6.4, 6.7.3 ISO/IEC 22989:2022 NIST CSF 2.0 COBIT 5 ISO/IEC 42001:2023 9.3.2 - Management review inputs B.8.5 - Information for interested parties B.6.2.6 - AI system operation and monitoring TEXAS: RESPONSIBLE ARTIFICIAL INTELLIGENCE GOVERNANCE ACT (TRAIGA)		
	Risk Statement		Processes for tracking, responding to, and recovering from incidents and errors are followed and documented.			
	To ensure that plans for prioritizing risk and regular monitoring and improvement are in place.					
Applicable Regulation						
HIPAA	PCI-DSS					
SOC2	EU AI ACT					
EU GDPR	WHITE HOUSE AI EO					
PPDT Governance Framework						
PEOPLE	PROCESS					
DATA	TECHNOLOGY					
Core Principles						
TRANSPARENCY (EXPLAINABLE)	ACCOUNTABILITY (ETHICAL)	IMPARTIALITY/FAIRNESS (EXPLAINABLE)	INCLUSION (ETHICAL)	SECURITY AND PRIVACY (EXPLAINABLE)	RELIABILITY AND SAFETY (ETHICAL AND EXPLAINABLE)	ROBUSTNESS (EXPLAINABLE)
✓	✓			✓	✓	

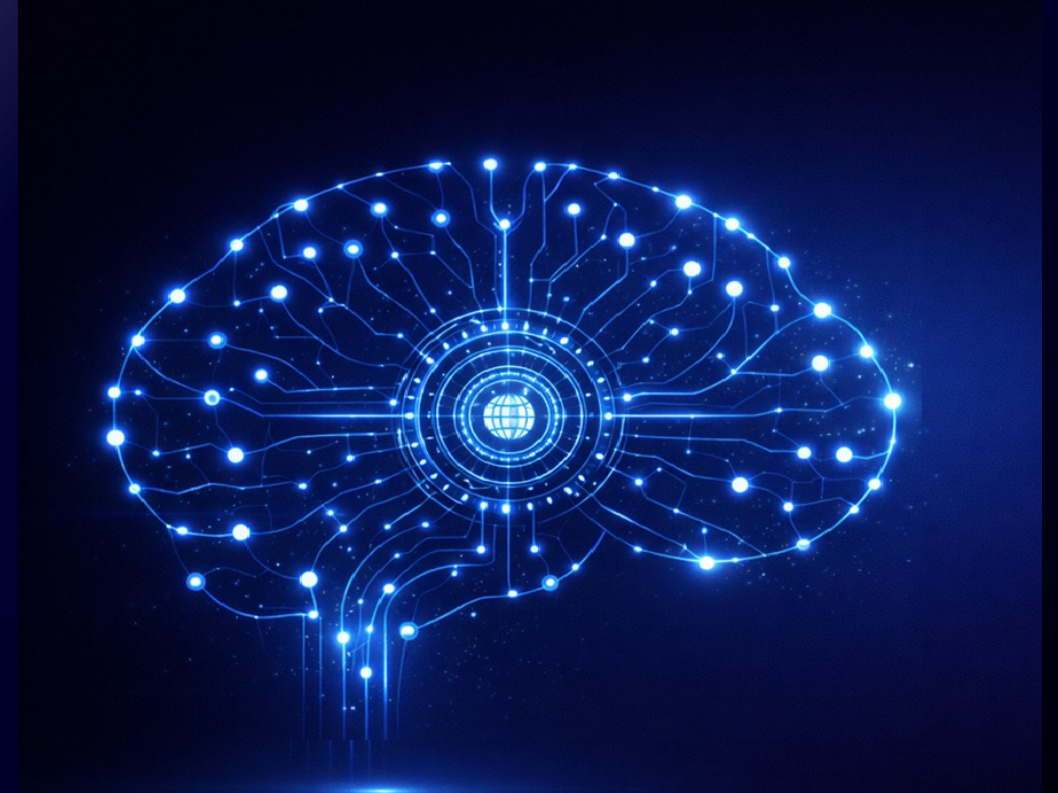
©2026 HISPI



30

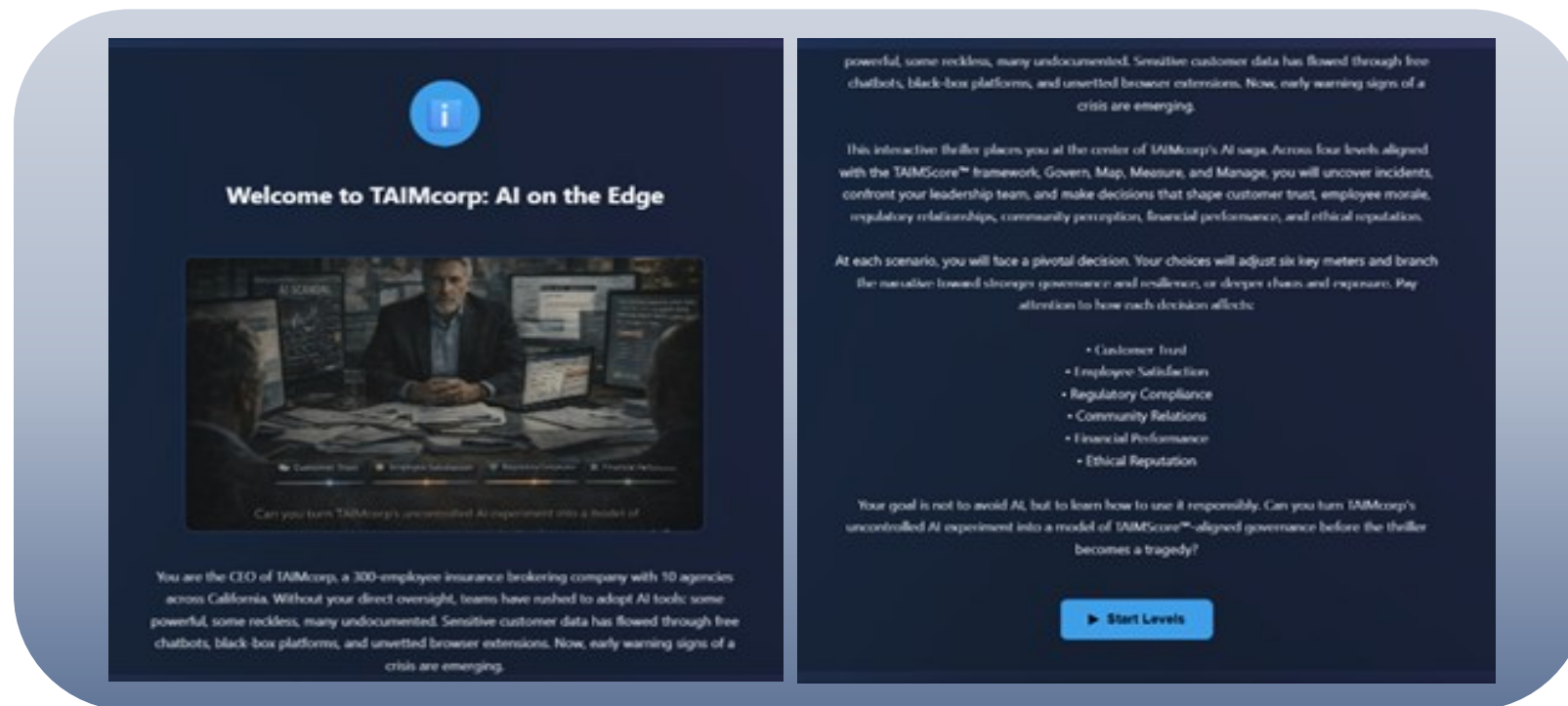


TAIMScore™ Platform



TAIMScore™ Gamification

A guided learning and assessment experience designed to keep users engaged while they move through AI trust questions. Progressive questions, feedback, and prompts help teams turn AI risk concepts into repeatable decisions.



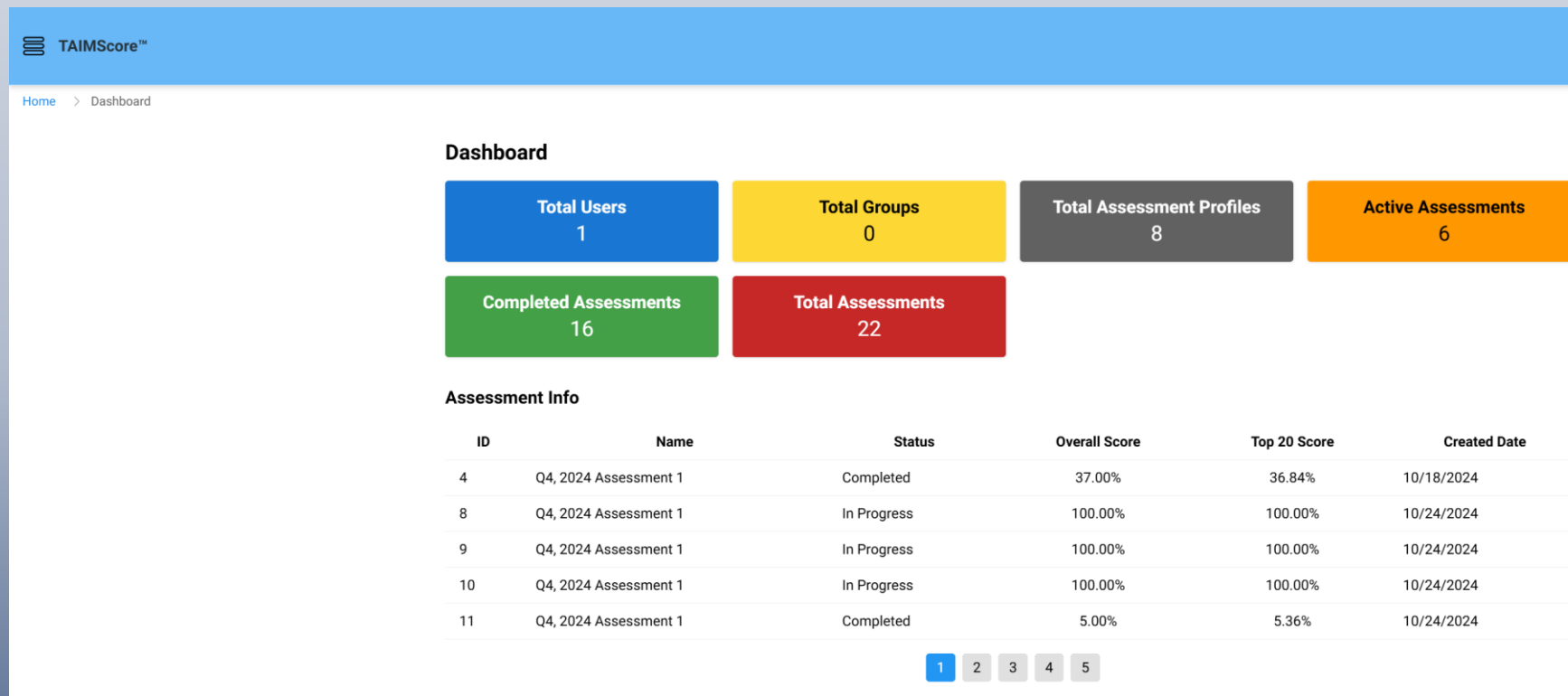
TAIMScore™ Interface

A structured interface gives users a clear path through profile data, assessment questions, scoring, and recommendations.



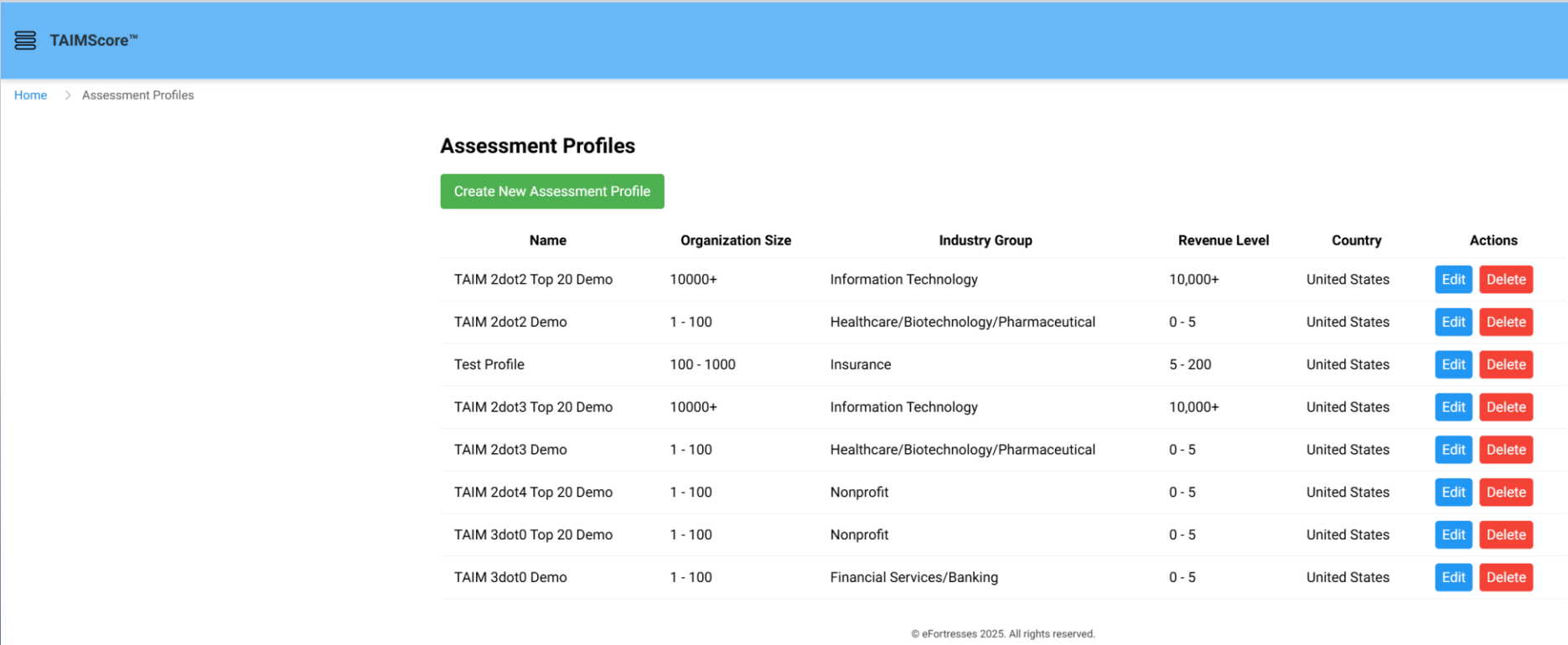
TAIMScore™ Dashboard

The dashboard summarizes maturity, risk themes, and areas that need attention.



TAIMScore™ Profile Screen

Profiles establish organizational context before scoring so results are tied to use case and environment.



TAIMScore™

Home > Assessment Profiles

Assessment Profiles

Create New Assessment Profile

Name	Organization Size	Industry Group	Revenue Level	Country	Actions
TAIM 2dot2 Top 20 Demo	10000+	Information Technology	10,000+	United States	Edit Delete
TAIM 2dot2 Demo	1 - 100	Healthcare/Biotechnology/Pharmaceutical	0 - 5	United States	Edit Delete
Test Profile	100 - 1000	Insurance	5 - 200	United States	Edit Delete
TAIM 2dot3 Top 20 Demo	10000+	Information Technology	10,000+	United States	Edit Delete
TAIM 2dot3 Demo	1 - 100	Healthcare/Biotechnology/Pharmaceutical	0 - 5	United States	Edit Delete
TAIM 2dot4 Top 20 Demo	1 - 100	Nonprofit	0 - 5	United States	Edit Delete
TAIM 3dot0 Top 20 Demo	1 - 100	Nonprofit	0 - 5	United States	Edit Delete
TAIM 3dot0 Demo	1 - 100	Financial Services/Banking	0 - 5	United States	Edit Delete

© eFortresses 2025. All rights reserved.

TAIMScore™ Results

Scoring outputs help users identify strengths, gaps, and areas for follow-up.

Group Name	Is Complete	Score
GOVERN 2	Yes	0.00%
GOVERN 6	Yes	0.00%
MAP 1	Yes	0.00%
MAP 4	Yes	0.00%
MAP 5	Yes	0.00%
MEASURE 2	Yes	0.00%
MEASURE 3	Yes	0.00%
MEASURE 4	Yes	0.00%
MANAGE 2	Yes	0.00%
MANAGE 4	Yes	0.00%

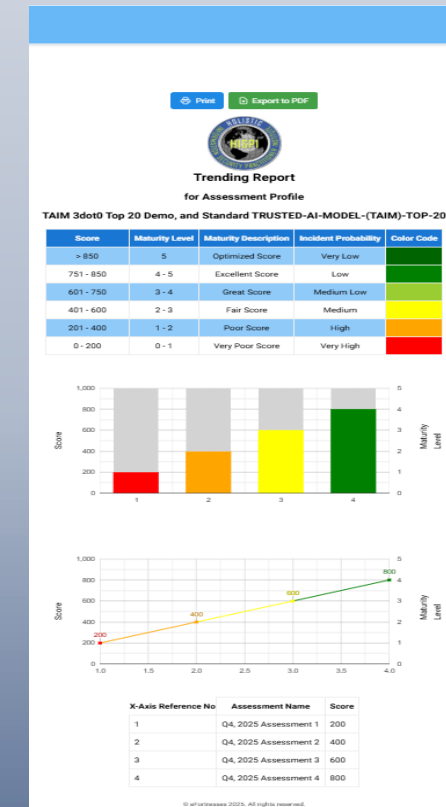
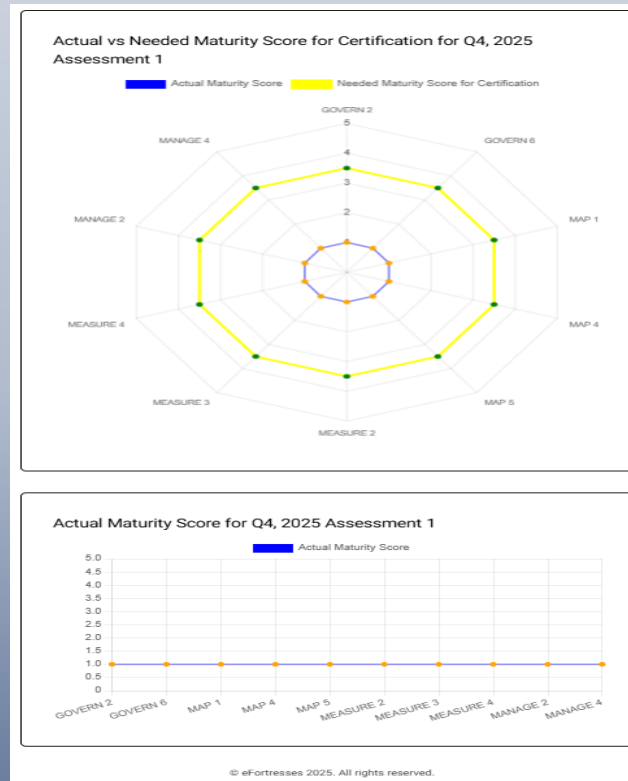
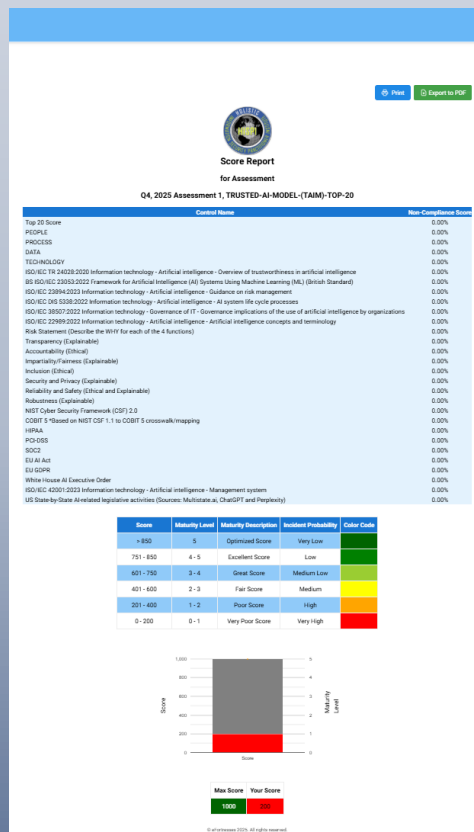
Sub Group Name	Sub Group Objective	Is Complete	Questions	Score
MANAGE 4.1	Risk treatments, including response and recovery, and communication plans for the identified and measured AI risks are documented and monitored regularly.	Yes	1	0.00%
MANAGE 4.3	Risk treatments, including response and recovery, and communication plans for the identified and measured AI risks are documented and monitored regularly.	Yes	1	0.00%

No	Question	Related To	Related AI Incident	Top 20 Score	Response	Maturity Level	Control Evidence
1	Are incidents and errors communicated to relevant AI actors, including affected communities? Are processes for tracking, responding to, and recovering from incidents and errors followed and documented?	PEOPLE: X PROCESS: X DATA: X TECHNOLOGY: X ISO/IEC TR 24028:2020 BS ISO/IEC 23053:2022 A.2.2.2 ISO/IEC 23894:2023 Inf 6.7 - ISO 31000:2018 ISO/IEC DIS 5338:2022 6.4.15.3 ISO/IEC 38507:2022 Inf 6.1 6.4	A deepfake bot is being - MIT Technology Review https://www.technology.com/AI/Undressing Incident An app called DeepNude Risk AI can be used to invade	3	Yes	1 - Initial	

Save Responses

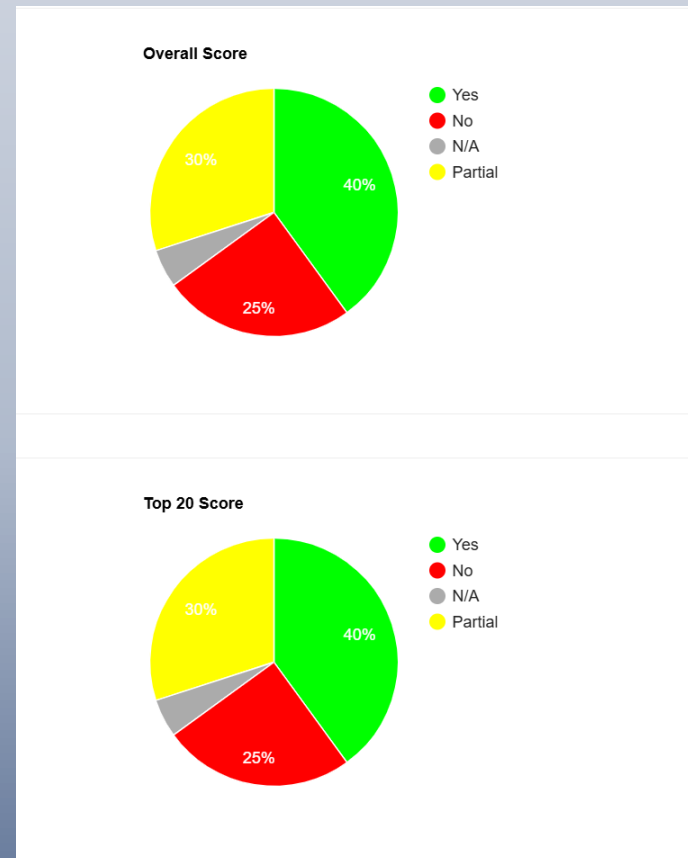
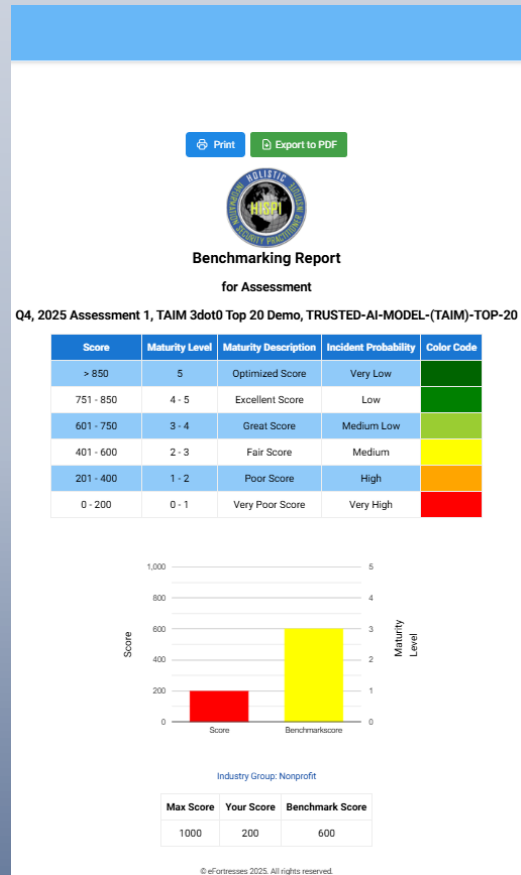
TAIMScore™ Reports

Reports translate assessment results into evidence, trends, and artifacts that can be shared with decision-makers.



TAIMScore™ Benchmark Report

Benchmarking helps organizations compare posture and prioritize improvement.



TAIMScore™ Recommendations

Recommendation content connects assessment results to practical next steps.

Recommendations				
Group Name	Objective	Weakness	Recommendation	Related AI Incident
MAP 1	Context is established and understood.	Are system requirements (e.g., "the system shall respect the privacy of its users") elicited from and understood by relevant AI actors? Does design decisions take socio-technical implications into account to address AI risks?	Ensure that system requirements (e.g., "the system shall respect the privacy of its users") are elicited from and understood by relevant AI actors. Ensure that design decisions take socio-technical implications into account to address AI risks.	OpenAI faces new class action lawsuit over data privacy - Washington Post https://www.washingtonpost.com/technology/2023/06/28/openai-chatgpt-lawsuit-class-action/ AI Data Aggregation Incident As generative AI needs a lot of data, OpenAI's solution was to scrap the internet for said data. The data is alleged to contain personal and copyrighted information. Risk Generative AI encourages sourcing a lot of data. If a company has the resources to shrug off their legal liability regarding the data they take, they can continue to take the data, even if it belongs to people.
MEASURE 2	AI systems are evaluated for trustworthy characteristics.	Is the AI system evaluated regularly for safety risks – as identified in the MAP function? Is the AI system to be deployed demonstrated to be safe, its residual negative risk does not exceed the risk tolerance, and it can fail safely, particularly if made to operate beyond its knowledge limits? Does safety metrics reflect system reliability and robustness, real-time monitoring, and response times for AI system failures?	Ensure that the AI system is evaluated regularly for safety risks – as identified in the MAP function. Ensure that the AI system to be deployed is demonstrated to be safe, its residual negative risk does not exceed the risk tolerance, and it can fail safely, particularly if made to operate beyond its knowledge limits. Ensure that safety metrics reflect system reliability and robustness, real-time monitoring, and response times for AI system failures.	Twitter taught Microsoft's AI chatbot to be racist in less than a day - Business Insider https://www.businessinsider.com/microsoft-deletes-racist-genocidal-tweets-from-ai-chatbot-tay-2016-3 AI can be changed. Incident Microsoft's "TAY" was released to the public as an experiment, with TAY being able to learn from prior conversations. Because TAY was open to input without sanitation, it outputted racial slurs, defended white-supremacist propaganda, and called for a genocide before being taken off the internet. Risk AI can be manipulated into producing undesirable output. Depending on the output presents risk to an organization such as public relations issues or unauthorized disclosure, depending on what it has access to.
MANAGE 2	Strategies to maximize AI benefits and minimize negative impacts are planned, prepared, implemented, documented, and informed by input from relevant AI actors.	Are mechanisms in place and applied, and responsibilities assigned and understood, to supersede, disengage, or deactivate AI systems that demonstrate performance or outcomes inconsistent with intended use?	Ensure that mechanisms are in place and applied, and responsibilities are assigned and understood, to supersede, disengage, or deactivate AI systems that demonstrate performance or outcomes inconsistent with intended use.	Malicious Actors Manipulating Photos and Videos for Sextortion Schemes - FBI IC3.org https://www.ic3.gov/Media/Y2023/PSA230605 AI Malicious actors and Sextortion Incident The FBI issued a Public Service Announcement (PSA), warning of synthetic content known as deepfakes, which comes after reports of victims, ranging from minors and adults, whose pictures and videos were altered to be explicit and shared around the internet to sexually harass or do sextortion schemes. Risk Risk regarding AI manipulation of media can range from public relation issues to personal ones, enabling controversies that might hurt revenue to public scrutiny of fabricated events.
MANAGE 4	Risk treatments, including response and recovery, and communication plans for the identified and measured AI risks are documented and monitored regularly.	Are incidents and errors communicated to relevant AI actors, including affected communities? Are processes for tracking, responding to, and recovering from incidents and errors followed and documented?	Ensure that incidents and errors are communicated to relevant AI actors, including affected communities. Ensure that processes for tracking, responding to, and recovering from incidents and errors are followed and documented.	A deepfake bot is being used to "undress" underage girls - MIT Technology Review https://www.technologyreview.com/2020/10/20/1010789/ai-deepfake-bot-undresses-women-and-underage-girls/ AI Undressing Incident An app called DeepNude was being used to undress women. After public backlash, the creators took down the application, later re-establishing on Telegram. Risk AI can be used to invade the privacy of others and publicly humiliate people.
© eFortresses 2025. All rights reserved.				

Questions

Let's continue the conversation!

Project Cerebellum resources and contact information

LinkedIn group <https://www.linkedin.com/groups/6624427/>

Email projectcerebellum@hispi.org

Project Cerebellum <https://projectcerebellum.com>

HISPI <https://hispi.org>

TAIMScore™

Trusted AI decisions made
visible.

