



ANCHOR CYBER SECURITY LLC

NIST AI RMF 1.0 – Field Reference

Same shape as GOVERN in CSF 2.0, purpose-built for AI systems. For any organization deploying, building on, or assessing AI in production.

\$ ai-rmf --glossary

Vocabulary

AI Actor	Anyone playing a role in an AI system's lifecycle — designer, developer, deployer, operator, evaluator, end user.
AI Actor Task	A concrete lifecycle activity (e.g. "TEVV" — test, evaluation, verification, validation) an AI Actor performs.
Trustworthy AI	An AI system judged against seven characteristics (below) rather than a single "safe/unsafe" label.
Profile	Same concept as a CSF Profile — which AI RMF outcomes apply, for a specific use case, sector, or system.
Risk Tolerance	How much AI risk an organization is willing to accept in pursuit of its objectives — set at the GOVERN level, applied everywhere else.

The Four Functions / 19 categories · 72 subcategories

GOVERN

Cross-cutting, not sequential

Policies, accountability structures, culture, and third-party/supply-chain risk for AI — sits above the other three the same way CSF's GOVERN does.

cat 6 sub 19

MAP

Context and categorization

Establish context, categorize the system, understand capabilities/costs/benefits, and characterize impacts to people and society.

cat 5 sub 18

MEASURE

Analysis and tracking

Apply methods/metrics, evaluate against the trustworthy characteristics, track risk over time, assess measurement itself.

cat 4 sub 22

MANAGE

Prioritize and respond

Act on what MAP and MEASURE found — prioritize risk response, manage third-party AI risk, document treatment and communication plans.

cat 4 sub 13

Category Map / NIST AI 100-1, Jan 2023

CODE	CATEGORY	SUBCATS
GOVERN		
GOVERN 1	Policies, processes, and practices for mapping, measuring, and managing AI risk are in place and effective	7
GOVERN 2	Accountability structures — teams and individuals empowered and trained for AI risk work	3
GOVERN 3	Workforce diversity, equity, inclusion, and accessibility in AI risk work	2
GOVERN 4	Organizational culture that considers and communicates AI risk	3
GOVERN 5	Robust engagement with relevant AI Actors	2
GOVERN 6	Policies for third-party software, data, and supply chain AI risk	2
MAP		
MAP 1	Context is established and understood	6
MAP 2	Categorization of the AI system is performed	3
MAP 3	AI capabilities, usage, goals, and expected costs/benefits are understood	5
MAP 4	Risks and benefits mapped for all components, including third-party software and data	2
MAP 5	Impacts to individuals, groups, communities, and society are characterized	2
MEASURE		
MEASURE 1	Appropriate methods and metrics identified and applied	3
MEASURE 2	AI systems evaluated for the seven trustworthy characteristics	13
MEASURE 3	Mechanisms for tracking identified risks over time are in place	3
MEASURE 4	Feedback on measurement efficacy is gathered and assessed	3
MANAGE		
MANAGE 1	Risk from MAP and MEASURE is prioritized, responded to, and managed	4
MANAGE 2	Strategies to maximize benefit and minimize harm are planned and documented	4
MANAGE 3	Third-party AI risks and benefits are managed	2
MANAGE 4	Risk treatment, response, recovery, and communication plans are documented and monitored	3

```
$ ai-rmf --trustworthy
```

Seven Trustworthy Characteristics / what MEASURE actually evaluates

Valid & Reliable

Accurate for its intended purpose, consistent under expected conditions.

Safe

Doesn't endanger human life, health, property, or the environment.

Secure & Resilient

Withstands adversarial manipulation and recovers from disruption.

Accountable & Transparent

Clear ownership of outcomes; appropriate disclosure about the system.

Explainable & Interpretable

A human can understand the output and why the system produced it.

Privacy-Enhanced

Anonymity, confidentiality, and user control over data are safeguarded.

Fair, Harmful Bias Managed

Equity considered across systemic, computational, and human-cognitive bias.

```
$ ai-rmf --apply
```

Where AI RMF Shows Up in Practice

01 AI-inspection or filtering tools.

A tool that inspects, filters, or makes decisions about AI traffic is itself an AI Actor — MAP 1/2 apply to it before they apply to a customer's downstream AI usage.

02 Client AI adoption reviews.

Use GOVERN 1/2 to check whether a client has any AI risk policy or accountable owner before they roll out an LLM tool company-wide.

03 Your own AI use policy.

Cross-check it against GOVERN 4 (culture) and GOVERN 6 (third-party AI risk) — most internal AI policies only cover acceptable use, not vendor risk.

04 Incident review involving an AI system.

MEASURE 3 and MANAGE 4 are the two categories that ask "are we tracking this risk over time," not just "did we fix it once."