

temper.ai

NIST AI Risk Management Framework (AI RMF 1.0)

Alignment & Mapping Document

Version	1.0 — Initial Draft
Date	April 2026
Owner	[Name] — [Title], temper.ai
Review Cycle	Annual, or upon material change to any AI system
Status	DRAFT — internal review required before external distribution
Classification	Confidential — available to customers under NDA on request

1. Purpose and Scope

This document records temper.ai's alignment with the NIST Artificial Intelligence Risk Management Framework (AI RMF 1.0), published by the National Institute of Standards and Technology in January 2023.

The AI RMF organises AI risk governance into four core functions — GOVERN, MAP, MEASURE, and MANAGE — and is referenced by US federal agencies, DoD contractors, and state/local governments as evidence of trustworthy AI practice. Aligning with this framework supports temper.ai's government-facing compliance posture and positions the company ahead of emerging AI-specific requirements in CMMC and DFARS.

This mapping covers all AI systems and features deployed within the temper.ai platform as of the document date. It should be updated whenever a new AI system or material model change is introduced, and reviewed in full at least annually.

2. AI Systems in Scope

The following AI systems and capabilities are covered by this mapping. Each should have a corresponding AI Transparency Card (see Appendix A).

System / Feature	Description	Deployment Context
[System Name 1]	[Brief description of what this system does and what model(s) it uses]	[Internal / Customer-facing / Government]
[System Name 2]	[Description]	[Context]
[Add rows as needed]		

Note: Government-facing AI systems or any system that may process Controlled Unclassified Information (CUI) require particular attention to GOVERN 1.1, MAP 5.1, MEASURE 2.7, and MANAGE 4.1.

3. GOVERN

The GOVERN function establishes organizational policies, accountability structures, and culture required to responsibly develop and deploy AI systems. It applies across the full AI lifecycle.

ID	Requirement	temper.ai Implementation / Evidence
GOVERN 1 — Policies, processes, and practices for AI risk management		
GV-1.1	Legal and regulatory requirements involving AI are understood, managed, and documented.	☐ Document: Map applicable regulations (CCPA, GDPR, DFARS 252.204-7012, state AI laws) against each system in scope. Owner: [Legal / GC] Evidence: Regulatory register
GV-1.2	Characteristics of trustworthy AI are integrated into policies and practice.	☐ Document: Define temper.ai's trustworthy AI principles (reliability, safety, fairness, privacy, transparency, accountability) and reflect them in product policy. Owner: [CPO / CTO]
GV-1.3	Organisational risk tolerance and risk management activities are established and communicated.	☐ Document: Define risk appetite for each AI system. Approved by exec team. Owner: [CEO / CTO]
GV-1.4	Organisational teams are committed to a risk management process.	☐ Document: Assign AI risk roles; establish a lightweight AI Risk Committee or review cadence. Owner: [CISO / CTO]

GV-1.5	AI risks are monitored on an ongoing basis.	☐ Document: Define monitoring cadence and escalation path for each production AI system. Owner: [Engineering Lead]
GV-1.6	Policies for AI system inventory are maintained.	☐ Document: Maintain an inventory of all AI systems (this document, Section 2, serves as a starting point). Owner: [CTO]
GV-1.7	Processes for safe decommissioning of AI systems are established.	☐ Document: Define decommission checklist: data retention/deletion, model deprecation, customer communication. Owner: [Engineering Lead]
GOVERN 2 — Accountability and team empowerment		
GV-2.1	Roles and responsibilities for AI risk management are documented and communicated.	☐ Document: RACI matrix for AI risk ownership across product, engineering, legal, and exec. Owner: [COO / CTO]
GV-2.2	Personnel receive AI risk management training appropriate to their role.	☐ Document: Define training requirements by role. Schedule initial training and annual refresh. Owner: [HR / CISO]

GV-2.3	Executive leadership takes responsibility for AI risk decisions.	☐ Document: Board/exec sign-off on AI risk policy and this mapping document. Owner: [CEO]
GOVERN 3 — Workforce diversity, equity, and inclusion		
GV-3.1	Decision-making teams are diverse in background and expertise.	☐ Document: Record AI system design team composition. Note any diversity gaps and mitigation steps. Owner: [CPO / HR]
GV-3.2	Policies for human-AI configuration reflect DEI considerations.	☐ Document: Confirm that human oversight design doesn't systematically disadvantage end users. Owner: [CPO]
GOVERN 4 — Risk-conscious culture		
GV-4.1	Safety-first policies are documented and followed.	☐ Document: Reference temper.ai's responsible AI policy and safety review gate in the development process. Owner: [CTO]
GV-4.2	Risk documentation and communication processes exist.	☐ Document: Define how AI risks are logged, escalated, and reported internally. Owner: [CISO]

GV-4.3	Testing and incident-sharing practices are in place.	☐ Document: Define pre-launch testing requirements and post-incident review process. Owner: [Engineering Lead]
GOVERN 5 — Stakeholder engagement		
GV-5.1	Feedback is collected from affected parties and external stakeholders.	☐ Document: Define customer feedback channels. Establish mechanism for AI-specific feedback. Owner: [CPO]
GV-5.2	Feedback is systematically integrated into AI risk management.	☐ Document: Process for triaging AI feedback into product roadmap and risk register. Owner: [CPO / CTO]
GOVERN 6 — Third-party and supply chain oversight		
GV-6.1	Third-party AI providers (model vendors, APIs) are assessed for risk.	☐ Document: Vendor risk register for all external AI models and APIs (e.g., OpenAI, Anthropic, etc.). Owner: [CISO / Legal]
GV-6.2	Contingency processes exist for third-party AI failures.	☐ Document: Fallback procedures and SLAs for each third-party AI dependency. Owner: [Engineering Lead / CTO]

4. MAP

The MAP function establishes the context of each AI system — its purpose, deployment environment, stakeholders, and risks. MAP activities should be completed for each system listed in Section 2.

ID	Requirement	temper.ai Implementation / Evidence
MAP 1 — Context establishment		
MP-1.1	Intended purposes, use cases, and deployment contexts are documented.	☐ Document: AI Transparency Card (Appendix A) per system, including intended vs. prohibited uses. Owner: [CPO / Product]
MP-1.2	Interdisciplinary input (legal, ethics, security, product) is included.	☐ Document: Record participants in AI system design reviews. Owner: [CPO]
MP-1.3	AI systems align to organisational mission and strategic goals.	☐ Document: Map each AI system to temper.ai's product mission and government-readiness objectives. Owner: [CEO / CPO]
MP-1.4	Business value of each AI system is defined.	☐ Document: Quantify expected value (efficiency, accuracy, revenue, risk reduction) per system. Owner: [CPO]

MP-1.5	Organisational risk tolerance is applied at the system level.	☐ Document: For each system, record whether risk profile is within tolerance; escalate if not. Owner: [CTO / CISO]
MP-1.6	System requirements include socio-technical considerations.	☐ Document: Address potential social and technical failure modes in system requirements. Owner: [CPO / Engineering]
MAP 2 — AI system categorisation		
MP-2.1	Tasks, methods, and datasets are defined and documented.	☐ Document: For each system, record task definition, model type, training/ fine-tuning data provenance. Owner: [ML Lead / Engineering]
MP-2.2	Knowledge limits and confidence boundaries are documented.	☐ Document: Record known edge cases, out-of-distribution failure modes, and confidence thresholds. Owner: [ML Lead]
MP-2.3	Scientific integrity and test/evaluation/ verification/validation (TEVV) considerations are addressed.	☐ Document: TEVV plan per system, including test set design, holdout data, and third-party review where relevant. Owner: [ML Lead / QA]
MAP 3 — Capabilities and expected outcomes		

MP-3.1	Potential benefits are examined and documented.	☐ Document: Expected benefits per system; update as evidence accumulates post-deployment. Owner: [CPO]
MP-3.2	Potential costs and harms are documented.	☐ Document: Risk register entry per system: failure modes, downstream harms, affected populations. Owner: [CPO / Legal]
MP-3.3	Scope and application are clearly specified.	☐ Document: Define what the system is NOT designed for (exclusions) alongside intended use. Owner: [CPO]
MP-3.4	Operator proficiency requirements are defined.	☐ Document: Training and competency requirements for internal users and customer admins. Owner: [CPO / Customer Success]
MP-3.5	Human oversight roles and responsibilities are defined.	☐ Document: Human-in-the-loop design, override mechanisms, and escalation paths per system. Owner: [CPO / Engineering]
MAP 4 — Component risk and benefit mapping		
MP-4.1	Technology and legal risks are identified and addressed.	☐ Document: Risk register covering IP, liability, third-party model terms, and export controls. Owner: [Legal / CTO]

MP-4.2	Internal risk control mechanisms are identified.	☐ Document: Technical and process controls mapped to each identified risk. Owner: [CISO / Engineering]
MAP 5 — Impact characterisation		
MP-5.1	Likelihood and magnitude of negative impacts are identified and documented.	☐ Document: For each system, score probability and severity of negative outcomes; document for government-facing systems in particular. Owner: [CPO / Legal]
MP-5.2	Regular engagement and feedback integration practices are established.	☐ Document: Feedback loop cadence; mechanism for affected-party input to reach AI risk function. Owner: [CPO]

5. MEASURE

The MEASURE function quantifies AI risks using appropriate metrics and evaluation methods. Measurement activities should be performed at launch and maintained continuously in production.

ID	Requirement	temper.ai Implementation / Evidence
MEASURE 1 — Methods and metrics identification		
MS-1.1	Risk measurement approaches are selected and documented.	[] Document: Per-system metrics plan — what is measured, how, at what frequency. Owner: [ML Lead / Engineering]
MS-1.2	Metric appropriateness is assessed regularly.	[] Document: Regular review of whether current metrics still capture the right risks (e.g., after model updates). Owner: [ML Lead]
MS-1.3	Internal and independent experts are involved in measurement design.	[] Document: Record internal and any external reviewers involved in defining AI evaluation criteria. Owner: [CTO / ML Lead]
MEASURE 2 — Trustworthy characteristics evaluation		
MS-2.1	Test sets and metrics are documented.	[] Document: Maintain test set registry per system; include benchmark datasets and custom evaluation sets. Owner: [ML Lead / QA]
MS-2.2	Human subject evaluation requirements are addressed.	[] Document: If user studies or human evaluation are conducted, record methodology and consent procedures. Owner: [CPO / Legal]

MS-2.3	Performance criteria are defined and measured.	[] Document: Quantitative acceptance thresholds (accuracy, latency, error rate) per system and use case. Owner: [ML Lead / Engineering]
MS-2.4	Production monitoring is in place.	[] Document: Real-time and batch monitoring dashboards; alerting on performance degradation. Owner: [Engineering / SRE]
MS-2.5	Validity and reliability of AI systems are demonstrated.	[] Document: Reproducibility testing; consistency across runs and environments. Owner: [ML Lead / QA]
MS-2.6	Safety risks are evaluated.	[] Document: Failure mode analysis; red-teaming or adversarial testing cadence. Owner: [Security / ML Lead]
MS-2.7	Security and resilience are evaluated.	[] Document: Adversarial input testing; prompt injection / model extraction testing for LLM-based systems. Owner: [Security Lead]
MS-2.8	Transparency and accountability risks are examined.	[] Document: Auditability of AI decisions; logging sufficient to reconstruct outputs on request. Owner: [Engineering / Legal]
MS-2.9	Model explanation and output validation are in place.	[] Document: Explainability mechanism per system (e.g., confidence scores, rationale output, citation sourcing). Owner: [ML Lead / CPO]

MS-2.10	Privacy risks are examined.	[] Document: PII exposure analysis per system; data minimisation and retention review. Owner: [CISO / Legal]
MS-2.11	Fairness and bias are evaluated.	[] Document: Bias evaluation methodology; demographic parity and disparate impact analysis where applicable. Owner: [ML Lead / CPO]
MS-2.12	Environmental impact is assessed.	[] Document: Estimate compute / carbon footprint of training and inference. Note any efficiency optimisations. Owner: [Engineering / CTO]
MS-2.13	TEVV metric effectiveness is evaluated.	[] Document: Post-deployment review: did pre-launch test metrics predict real-world performance? Owner: [ML Lead]
MEASURE 3 — Risk tracking		
MS-3.1	Existing and emergent risk identification approaches are defined.	[] Document: Mechanism for surfacing new risk types (e.g., LLM hallucination, data drift, supply chain model changes). Owner: [CISO / ML Lead]
MS-3.2	Difficult-to-assess risks are given special consideration.	[] Document: Record risks that are hard to quantify (e.g., long-term bias harms, compounding errors). Owner: [CPO / Legal]
MS-3.3	User and community feedback informs risk tracking.	[] Document: Process for incorporating customer feedback into risk register updates. Owner: [CPO / Customer Success]

MEASURE 4 — Measurement efficacy		
MS-4.1	Measurement approaches are informed by deployment context.	☐ Document: Confirm that production environment monitoring reflects actual usage patterns. Owner: [Engineering / SRE]
MS-4.2	Trustworthiness validation includes external consultation.	☐ Document: Record any third-party audits, red teams, or academic collaborations. Owner: [CTO]
MS-4.3	Performance improvements are identified and tracked.	☐ Document: Continuous improvement log — AI system changes and associated performance changes. Owner: [ML Lead]

6. MANAGE

The MANAGE function covers risk response, mitigation, monitoring, and incident handling across the AI lifecycle, including post-deployment.

ID	Requirement	temper.ai Implementation / Evidence
MANAGE 1 — Risk prioritisation and response		
MG-1.1	Whether AI systems are achieving their intended purpose is monitored.	[] Document: KPIs per system; review cadence; threshold for intervention. Owner: [CPO / Engineering]
MG-1.2	Risk treatments are prioritised by impact and likelihood.	[] Document: Risk register with priority scoring (High/Medium/Low). Reviewed quarterly. Owner: [CISO / CPO]
MG-1.3	Responses to high-priority risks are developed and maintained.	[] Document: Incident response playbooks per risk type (e.g., hallucination escalation, bias detection, data breach). Owner: [Engineering / Legal]
MG-1.4	Residual risks are documented.	[] Document: Risks accepted or deferred must be logged with rationale and owner. Owner: [CISO / CTO]
MANAGE 2 — Benefit maximisation and harm minimisation		

MG-2.1	Resources and alternative systems are considered to minimise risk.	☐ Document: Evaluate whether non-AI alternatives exist for higher-risk use cases. Owner: [CPO / CTO]
MG-2.2	Deployed systems are actively maintained to sustain value.	☐ Document: Model refresh schedule; performance degradation thresholds that trigger retraining or rollback. Owner: [ML Lead / Engineering]
MG-2.3	Procedures for unknown or unexpected risks are in place.	☐ Document: Incident triage process for novel failure modes not previously anticipated. Owner: [Engineering / CTO]
MG-2.4	System disengagement mechanisms are defined and tested.	☐ Document: Kill-switch or graceful degradation path for each AI system. Tested at least annually. Owner: [Engineering Lead]
MANAGE 3 — Third-party risk management		
MG-3.1	Third-party AI resources are monitored for changes.	☐ Document: Alerts for upstream model version changes, API deprecations, and vendor security notices. Owner: [Engineering / CISO]

MG-3.2	Pre-trained models are monitored for drift and integrity.	☐ Document: Process for evaluating base model updates before adoption in production. Owner: [ML Lead]
MANAGE 4 — Ongoing monitoring and improvement		
MG-4.1	A post-deployment monitoring plan is implemented.	☐ Document: Continuous monitoring stack — dashboards, alerting, on-call response. Review with this doc annually. Owner: [Engineering / SRE]
MG-4.2	Continual improvement is integrated into AI operations.	☐ Document: Feedback loop from monitoring and incidents into model/product improvement cycle. Owner: [ML Lead / CPO]
MG-4.3	Incidents are tracked, communicated, and fed back into risk management.	☐ Document: Incident register; post-mortems; communication to affected customers; updates to risk register. Owner: [CISO / CPO]

Appendix A: AI Transparency Card Template

Complete one of these cards for each AI system listed in Section 2. Transparency Cards are the primary evidence artefact for MAP 1.1 and GOVERN 1.6. They may be shared with government customers and procurement teams on request.

AI TRANSPARENCY CARD	
System Name	[Name of feature or model]
Version	[Version number or release date]
Last Updated	[Date]
Model Type	[e.g., LLM, classifier, retrieval, embedding, fine-tuned, RAG]
Underlying Model(s)	[e.g., Claude claude-sonnet-4-6 via Anthropic API; internal fine-tune on top of...]
PURPOSE & USE	
Intended Use	[What this system is designed to do; intended user population; deployment context]
Prohibited Use	[Explicitly prohibited use cases — e.g., not for autonomous decision-making without human review, not for processing classified data]
Deployment Environment	[Internal use only / Customer-facing SaaS / Government environment / Other]
INPUTS & OUTPUTS	
Input Data Types	[e.g., user text, documents, structured data, API responses]
Output Types	[e.g., generated text, classifications, ranked lists, scores]
Does output contain PII?	[Yes / No / Conditional — describe]
Is output used for automated decisions?	[Yes / No / Human-in-the-loop — describe oversight mechanism]
PERFORMANCE & TESTING	
Performance Metrics	[Key metrics used to evaluate this system: accuracy, precision/recall, latency, etc.]
Test Set Description	[How the evaluation dataset was constructed; size; diversity; date of last update]

Bias Evaluation	[Method used to detect bias or disparate impact; results; any known limitations]
Security Testing	[Adversarial testing performed: prompt injection, model extraction, data poisoning — describe scope and results]
RISKS & LIMITATIONS	
Known Limitations	[What the system cannot do well; edge cases; known failure modes]
Residual Risks	[Risks that are accepted or not yet fully mitigated — describe and reference POA&M if applicable]
Override / Feedback Mechanism	[How users or admins can flag erroneous outputs or request human review]
GOVERNANCE	
System Owner	[Name / Title at temper.ai]
Review Cadence	[How often this card is reviewed and updated; last reviewed date]
Third-Party Models Used	[List all external model APIs; note usage terms / data handling commitments]
Approved for Government Use?	[Yes / No / Conditional — list any restrictions for government-environment deployment]

Appendix B: Turing App Transparency Card

AI TRANSPARENCY CARD	
System Name	Turing Demonstrator
Version	v9.0.243
Last Updated	April 7, 2025
Model Type	Spiking Neural Network
Underlying Model(s)	Turing in-house model
PURPOSE & USE	
Intended Use	Simple demonstrator of a Spiking Neural Network
Prohibited Use	It is not intended for production environments, and it's intentionally not instrumented to accept real-world inputs (or real-world outputs).
Deployment Environment	No restrictions
INPUTS & OUTPUTS	
Input Data Types	User settings
Output Types	Screen depiction of neural firing patterns
Does output contain PII?	No
Is output used for automated decisions?	No
PERFORMANCE & TESTING	
Performance Metrics	User understanding of underlying principles of real-time spiking behavior of SNNs and how they differ from LLMs
Test Set Description	Varies with every run.
Bias Evaluation	Does not apply
Security Testing	Does not apply (no actual information is exchanged)
RISKS & LIMITATIONS	
Known Limitations	This is not a language model. The mapping of real-world data into spike streams is the hard problem to solve on a per-project basis.
Residual Risks	Not all problems are suitable to spiking decomposition and modularization

Override / Feedback Mechanism	Does not apply
GOVERNANCE	
System Owner	Michael Bilca, temper.ai
Review Cadence	Annual, or upon material version change. Last reviewed: April 2026.
Third-Party Models Used	None — the Turing Demonstrator runs an entirely in-house Spiking Neural Network. No external model APIs are used.
Approved for Government Use?	Yes — approved for demonstration purposes only (showing SNN firing behavior). Not approved for processing CUI or operational data.

Appendix C: Document Change Log

Version	Date	Changes	Author
1.0	April 2026	Initial draft — all items set to [] pending first internal review	Michael Bilca
1.01	April 7 2026	Added Turing app transparency card appendix	Michael Bilca

Instructions: Change [] to [x] as each item is completed. Where evidence artefacts (policies, registers, test reports) exist, add a reference to the file path or document name in the Implementation column.