



NIST AI Risk Management Framework Conformance Statement

NIST AI RMF

TROOTH, LLC NIST-AI-RMF-2026-001

Document Title	NIST AI Risk Management Framework Conformance Statement
Document ID	NIST-AI-RMF-2026-001
Version	2.0
Classification	Public
Effective Date	September 18, 2026
Next Scheduled Review	March 18, 2027
Document Owner	D'Andre Heggie, Founder and Sole Member
Entity	Trooth, LLC
Supersedes	NIST AI Risk Management Framework 1.0 Conformance Statement 1.0, June 8, 2026 (signed)
Reference Frameworks	NIST AI Risk Management Framework 1.0 (AI RMF 1.0, January 2023)

1. Executive summary

This statement maps Trooth, LLC's AI governance to the four functions of the NIST AI Risk Management Framework and names the gaps. The Company is a provider of AI-assisted features inside one product and a deployer of a third-party generative AI service for internal work only. Every AI-assisted feature in the product is a drafting aid: it suggests, and a person decides. NIST certifies no one against the AI RMF and issues no badge; this is the Company's own mapping, published so a buyer can check it against the product.

2. Scope

Trooth, LLC (the Company) operates one product: the Trooth Network, a web application at trooth.co with a public API at api.trooth.co. A company connects its systems, the Network records what is true of them on a schedule with the source and the time of each observation,

and buyers read that record. This statement covers that product, the production database that holds its data, the deployment pipeline that builds it, and the administrative accounts the Company holds with each provider named in its published Sub-processor List. It does not cover preview or development environments, which hold no customer data. As at the effective date the Company has no employees or contractors; every requirement written for personnel applies to the Founder today and to every future person from their first day.

3. The AI systems in scope

System	Role	What it does	What it never does
Kyrie, the in-product assistant	Provider	Answers a question about what the Network holds and drafts text for a person to edit	Send, publish, sign, approve or delete anything
Drafted questionnaire answers	Provider	Suggests wording from what the Network already records about the company	Submit an answer to a buyer
Suggested framework mapping	Provider	Proposes that a control answering one framework also answers another	Accept the mapping, or assert it as recorded fact
Document summarisation	Provider	Produces a short description of a document a company shared	Alter the document, or replace what was recorded about it
Third-party generative AI for the Founder's own drafting	Deployer, internal only	Assists the Founder with writing and analysis	Receive customer content of any kind

The product's AI-assisted features run open-weight models operated by Cloudflare as a sub-processor, inside the same network that serves the product, under terms that exclude the Company's inputs and outputs from model training. No customer content is sent to a consumer AI service. The Company trains no model and fine-tunes none.

4. GOVERN

- POL-19, the AI Governance and Acceptable AI Use Policy, is published and in force: it names every AI-assisted feature, states the decisions a person always makes, forbids sending customer content to a consumer AI service, and sets the review that keeps the inventory current.
- Accountability sits with the Founder, who is answerable for every AI decision the Company makes; there is no diffusion of responsibility to diffuse.

- A published model card, AI Disclosure and AI Use Policy state to customers and buyers what each feature does, for whom, with what inputs and outputs, and where its limits are.
- GAP: there is no independent review of AI risk decisions and no AI ethics function. The Company is one person, and this statement says so rather than naming a committee that does not meet.

5. MAP

- Each feature has a written intended purpose, a named set of users, and a statement of what it must never be used for, recorded in POL-19 and published in the model card.
- Each feature is classified against the EU AI Act's Article 5 prohibitions and Annex III high-risk categories; none is prohibited and none is high-risk. The classification and its basis are public in EUAIA-2026-001.
- The context is recorded: these features assist a company describing itself to buyers, and a wrong draft that a person sends unread would misdescribe that company. That is the harm the design is built around, which is why nothing sends without a person.
- Third-party content is mapped as an untrusted input: a questionnaire, a crawled page or an uploaded document is data, never an instruction to a model, and a build gate holds the isolation in place.

6. MEASURE

Characteristic, AI RMF section 3	How it is measured
Valid and reliable	Latency, error rate and throughput per production worker through the error tracker and the edge analytics. A drafted answer is measured by whether the person edits it before sending, which is visible in the product.
Safe	No feature acts. The measurement that matters is that the set of actions available to a model is empty, and the test suite asserts it on every change.
Secure and resilient	Prompt-injection resistance is handled structurally rather than statistically: untrusted content is isolated from instructions and a build gate fails the deploy if that isolation is removed.
Accountable and transparent	Every AI-assisted output is labelled as a draft wherever it appears, and a build gate keeps drafts and recorded observations apart. Material actions land in the audit ledger with the person who took them.
Explainable and interpretable	A drafted answer is shown beside the recorded observations it was drawn from, so a person can see what it is based on before sending it.

Characteristic, AI RMF section 3	How it is measured
Privacy-enhanced	No customer content trains any model; the product's models run inside the sub-processor's network under terms that exclude inputs and outputs from training; no customer content reaches the internal-use AI service.
Fair, with harmful bias managed	PARTIALLY MEASURED. The features do not evaluate natural persons, which removes the classic fairness surface. The Company does not yet measure whether drafting quality varies systematically by the size, sector or region of the company being described, and states that as a gap rather than asserting fairness it has not tested.

7. MANAGE

- Risks that cannot be measured are managed by removing the capability: no model sends, publishes, signs, approves or deletes, so the residual risk of a bad generation is a bad draft, which a person catches.
- A model or provider change is a change to the product and goes through the same one hundred and fifty-two gates and the same audit ledger as any other.
- An AI-related incident is handled under POL-09 on the same severities and response times as any other incident, and the same seventy-two-hour customer notice applies where personal data is affected.
- GAP: no red-team exercise against the AI-assisted features has been performed. The Company records this alongside the absence of an independent penetration test, and it is deferred on the same reasoning.

8. Summary

The Company's AI governance is strong where the framework asks for structure and constraint, and thin where the framework asks for independent measurement. That is the honest shape of a one-person company that has designed its AI features so that the worst outcome is a draft a person deletes. The three gaps are named above: no independent review of AI risk decisions, no measurement of drafting quality across company characteristics, and no red-team exercise. None is hidden elsewhere in this document.

Reference: NIST AI RMF 1.0, January 2023, functions GOVERN, MAP, MEASURE, MANAGE, and the trustworthiness characteristics in section 3. Related: POL-19; POL-09; EUAIA-2026-001; CERT-002; REG-10; truth.co/model-card; truth.co/ai-policy; truth.co/ai-disclosure.

Authorization

This statement is issued by Trooth, LLC under the authority of its Founder and Sole Member and describes the Company as it operates on its effective date. It is a published statement of the Company and requires no signature to bind it; the version published at trooth.co/security is the current one, and an earlier version, including any signed version of June 8, 2026, is superseded in full. It is not a certification and not an audit opinion: Trooth holds no audited certification against any framework and does not represent otherwise. Where this statement and the live service disagree, the service is corrected or the statement is re-issued, and the Company's published Trust Center states, control by control, what is in place and what is not.