

[AUTHORITY LIFECYCLE GOVERNANCE]

AUTHREX

TECHNOLOGY AND CAPABILITY CATALOG

Runtime authority, assurance, and verification for AI-enabled and autonomous systems. Seven governance frameworks, one executable governance kernel, two frozen research consoles, six software application profiles, twelve hardware reference designs and fourteen catalogued simulation records, documented at their actual standing with evidence, limitations, and a gated government evaluation path for each.

READ THE CORRECTIONS FIRST → PAGE 8

EVIDENCE CLASSES → PAGE 2

CLAIMS REGISTER → PAGE 57

7

GOVERNANCE
FRAMEWORKS

1

EXECUTABLE
KERNEL

12

HARDWARE DESIGNS
(NONE BUILT)

14

CATALOGUED
SIM RECORDS

3

PEER-REVIEWED
ARTICLES

8

PROVISIONALS
REPORTED AS FILED

0

INDEPENDENT
VALIDATIONS

[01] · WHAT EXISTS

Specifications, formal models checked at stated bounds, an executable governance kernel with a container test suite, deterministic seeded consoles with hash-chained records, bill-of-materials-level hardware designs, three peer-reviewed journal articles, and eight provisional applications reported as filed.

[02] · WHAT DOES NOT

No hardware built. No hardware-in-the-loop test. No field validation. No independent replication. No deployment, certification, accreditation, or government adoption. One favorable AFRL white-paper evaluation on record; not funded; not an endorsement.

[03] · THE ASK

Phase 1 technical review of models, code, and evidence (no hardware, no funding); Phase 2 government-defined scenarios on the seeded consoles; Phase 3 degradation and failure injection; then a sponsor testbed and hardware-in-the-loop.

INDEPENDENT RESEARCH PORTFOLIO. NOT A GOVERNMENT PUBLICATION; NOT AFFILIATED WITH, ENDORSED BY, OR PRODUCED FOR ANY GOVERNMENT AGENCY. NO CLAIM OF ADOPTION, CONTRACT, CERTIFICATION, ACCREDITATION, OR CLEARANCE IS MADE. EVERY RESULT IS SELF-REPORTED AND NOT INDEPENDENTLY VERIFIED. DISTRIBUTION: PUBLIC RELEASE (OPEN INFORMATION). LICENSE: CC BY 4.0. NO FORMAL DOD DISTRIBUTION STATEMENT IS ASSIGNED.

BURAK OKTENLI, MBA · PRINCIPAL INVESTIGATOR · INDEPENDENT RESEARCHER, AUTHREX RESEARCH PROGRAM, WASHINGTON, DC · ORCID 0009-0001-8573-1667
AUTHREX.SYSTEMS · BURAKOKTENLI.COM · INFO@BURAKOKTENLI.COM · VERSION 6.1 · SEPTEMBER 4, 2026 · SUPERSEDES AUTHREX-CAT-2026-002 (V5, 1 SEP 2026)

DOCUMENT CONTROL, DISTRIBUTION, AND READING RULES

DOCUMENT TITLE	AUTHREX Technology and Capability Catalog: Runtime Authority, Assurance, and Verification for AI-Enabled and Autonomous Systems
IDENTIFIER / VERSION	AUTHREX-CAT-2026-003 · Version 6.1 · September 4, 2026. Supersedes AUTHREX-CAT-2026-002 (V5, September 1, 2026), which is archived; V4 remains an off-site archive. Version 6.1 incorporates the corrections of a hostile pre-distribution red-team audit of the V6 draft (correction log CL-21 to CL-27); the V6 draft was not distributed.
AUTHOR	Burak Oktenli, MBA · Independent Researcher, AUTHREX Research Program, Washington, DC · ORCID 0009-0001-8573-1667
PURPOSE	To let a DoD technical evaluator, program manager, laboratory, warfare center, or innovation organization determine, quickly and accurately, what each capability is, what has been built and tested, what has not, and what government participation would advance it.
DISTRIBUTION	Public release. All content is drawn from publicly deposited, self-published, or self-reported research material. No CUI, no classified indicators, no third-party proprietary material, and no operational data are included. No formal DoD distribution statement is assigned because none is authorized; the document is open information under CC BY 4.0.
EXPORT CONTROL	No export-control determination has been made for the source artifacts (specifications, kernel, demonstrators). Counsel review is open. The catalog itself describes published open-access material only.
VERIFICATION STATUS	Every figure is self-reported by the author. No artifact in this catalog has been independently reproduced or technically verified by another party. DOIs and publicly deposited artifact identifiers are independently resolvable where linked, and published hashes can be recomputed from the packages that carry them. Provisional-application numbers, filing dates, entity status, and receipt numbers are supported by USPTO filing receipts held in the program evidence file; provisional applications are unexamined and may not be publicly searchable. Acceptance records for the conference items and the AFRL evaluation memorandum are held in the same file and are listed in the controlled evidence manifest that accompanies this catalog.
BASELINE SOURCES	Catalog V5 and its 44-row register (catalog.csv, 1 Sep 2026); authrex.systems deploy package (1 Sep 2026) including formal-coverage, kernel, ABORT and DA pages; burakoktenli.com release package (85 pages, site audit of 1 to 4 Sep 2026); per-platform Zenodo deposit files; the program white paper (July 2026, superseded where noted). Conflicts are logged in the correction log.

LANGUAGE DISCIPLINE

Model-checked properties are held at stated bounds, never verified or proven. Simulation results are demonstrated, never validated. Provisional applications are reported as filed. A withdrawn figure is withdrawn, not retracted. Standards are mapped by the author, architecturally relevant, or design targets; nothing is compliant or certified. Government correspondence is an evaluation record where a memorandum exists and is never sponsorship or endorsement. Readiness is self-assessed; no formal technology readiness assessment has been performed. These rules are issued as a controlled vocabulary standard (AUTHREX-CAT-2026-003-VOCAB) that governs every public statement made under the program name; the claims register at the back of this catalog is the source of truth, and a public claim without a register row is not published.

EVIDENCE CLASSES USED ON EVERY PAGE

Each material technical statement in this catalog carries one of seven classes. They are not interchangeable and no class substitutes for another.

Class	Meaning in this catalog	Example in the portfolio
VERIFIED	Directly supported by a documented verification process that a third party could repeat from the artifact: formal check at stated bounds, reproducible test, or inspection. Model-relative and bounded where formal.	HMAA automaton properties held at depth 9; AuthorityLapse v3 across 21 configurations
TESTED	Observed through documented execution: test suite, seeded simulation, benchmark, soak, or mutation campaign. Self-run unless stated.	Kernel container suite 83 passed; ABORT 1,000,000-run soak
IMPLEMENTED	Exists as functioning software, a codebase, a console, or an executable artifact.	Governance Kernel; standalone consoles
MODELED	Supported by a mathematical, formal, analytical, or simulation model but not physically demonstrated.	TLA+ surrogate countermodels; ERAM working paper
PROPOSED	Designed but not implemented: reference designs, specifications, planned tests.	BLADE hardware family; software application profiles
CONCEPTUAL	Early-stage concept requiring substantial further investigation.	Mission-authority continuity profile (in design; not catalogued)
NOT YET VALIDATED	A plausible claim or function for which adequate validation evidence does not yet exist.	MAIVA quorum repair; DA majority-capture resistance

AUTHREX AT A GLANCE

This page stands alone if forwarded without the catalog. A reader with thirty seconds should find who, what, why, proof, maturity, and the ask.

	7	1	12	14	3	8	0
	FRAMEWORKS	EXECUTABLE KERNEL	REFERENCE DESIGNS, 0 BUILT	CATALOGUED SIM RECORDS	PEER-REVIEWED ARTICLES	PROVISIONALS FILED	INDEPENDENT VALIDATIONS
WHO	A single principal investigator (Burak Oktenli, MBA; B.S. Computer Science, University of South Florida, 2020; M.B.A., Lynn University, 2026; M.P.S. Applied Intelligence, Georgetown University, in progress; prior industry experience in production data engineering). No standing engineering team, hardware laboratory, or independent test authority inside the program.						
WHAT	A runtime authority and assurance architecture that makes machine permissions explicit, enforceable, contractible, and auditable as operational evidence changes. Each of position, sensing, communications, and software integrity carries its own validity state; the system may execute only what every required condition admits; a strong source never compensates for a failed one; on failure the admissible set narrows at once and widens only by authorization. Every outcome (EXECUTE, DELAY, HANDOFF, ABORT) is written to a hash-chained record.						
WHY	The gap appears where two conditions hold together: a machine can select and execute an action faster than its supervisor can evaluate it, and the sensing that justifies the action can be degraded, jammed, or deceived. Uncrewed air systems under GNSS denial, counter-UAS engagement decisions, ground robotics, spacecraft under light delay, and agentic AI on networks all sit inside that gap.						
PROOF	HMAA authority automaton model checked in TLA+ (23,748 distinct states at depth 9; 5 invariants and 1 liveness held; 2 vacuous at the bound; discrete automaton only). AUTHREX-ABORT: two TLA+ modules across 21 and 35 configurations, 105 tests, 1,000,000-run soak. Governance Kernel: 83 tests passed in a container. Three peer-reviewed journal articles (Springer Nature x2, IJCSS); one favorable AFRL white-paper evaluation on record (CANVAS, June 2026; not funded; not an endorsement). All self-run on synthetic data.						
MATURITY	Research stage. Self-assessed TRL 2 to 4 by artifact; no formal TRL determination. No hardware-in-the-loop test, no field validation, no independent replication, no deployment, no certification, no government adoption. Twelve hardware designs exist at bill-of-materials level; none is built.						
ASK	Phase 1 technical review of models, code, and evidence (no hardware, no funding); Phase 2 government-defined scenarios on the seeded consoles; Phase 3 degradation and failure injection; then a government testbed and hardware-in-the-loop. Specific needs per capability: recorded sensor streams, a SiL flight or vehicle stack, a platform safety-interface specification, an independent evaluator, and a sponsor-stated mission use case.						
CONTACT	Burak Oktenli, MBA · AUTHREX Research Program · info@burakoktenli.com · authrex.systems · burakoktenli.com · ORCID 0009-0001-8573-1667						

Self-reported. Not independently verified. Not a government publication. No adoption, contract, certification, or endorsement is claimed. The evidence register and claims register at the back of the catalog are the record of what exists.

CONTENTS

Document control, distribution, and reading rules	2
Detachable executive brief	4
Executive overview	6
Portfolio at a glance	7
Corrections first: what changed since Version 5	8
Where the problem appears (mission relevance)	9
Program architecture: layer position, control loop, lifecycle, authority relationship	10
SATA · Sensor and Actor Trust Assessment	12
HMAA · Human-Machine Authority Allocation	15
ADARA · Adversarial Deception-Aware Reasoning	18
MAIVA · Multi-Agent Integrity Voting	21
FLAME · Forced Latency for Authority-Managed Escalation	23
CARA · Controlled Authority Recovery Architecture	26
ERAM · Escalation Risk Assessment Model	28
GOVERNANCE KERNEL · AUTHREX Governance Kernel (QA10)	31
AUTHREX-ABORT · Evidence Conditioned Authority Lapse Console (Terminal Authority Profile)	34
AUTHREX-DA · Data Admissibility Layer (Engineering Preview)	37
Software application profiles (SW-01 to SW-06)	39
Governor consoles and simulation records (SIM-01 to SIM-12)	40
BLADE hardware reference designs (BLADE-01 to BLADE-12)	42
Formal methods	44
Test and evaluation framework	46
AI and autonomy assurance	48
Cybersecurity and zero-trust relevance	48
Standards and policy alignment	49
Intellectual property and publications	50
Technology readiness register	51
Transition roadmap and government evaluation path	52
Government engagement matrix	53
Integration options	54
Portfolio maturity and mission-area matrices	55
Technical evidence index	56
Claims register	57
Correction log	58
Glossary and acronyms	62

EXECUTIVE OVERVIEW

What the program does, what it has built, and what it is asking for.

[WHAT THE PROGRAM DOES]

The AUTHREX Research Program develops runtime authority governance and assurance for autonomous and AI-enabled systems. It sits beside the mission autonomy rather than inside it and separates what a system can technically do from what it remains authorized to do at this moment, on this evidence. The program has produced seven governance frameworks that together form a seven-stage authority lifecycle, one executable governance kernel for agentic AI, two frozen research consoles (AUTHREX-ABORT for evidence-conditioned authority lapse and AUTHREX-DA for data admissibility), six software application profiles, twelve hardware reference designs, and a seeded, deterministic simulation portfolio, all supported by a public disclosure corpus with permanent identifiers.

[MAJOR TECHNICAL DOMAINS]

Autonomous systems (air, ground, maritime, space); AI assurance and agentic AI; mission assurance and command and control; cyber-physical systems and operational technology; formal methods; resilient PNT as an evidence problem; hardware roots of trust (reference designs only).

[CENTRAL ENGINEERING POSITION]

Authority is granted in advance by a human, bounded by an explicit envelope, and conditioned on evidence. When the evidence degrades, the admissible set of actions narrows at once; when it fails, authority lapses and irreversible actions are latched down. Authority widens only by an authorized transition, never by timeout and never by recovery of a signal alone. A confident model cannot compensate for invalid navigation.

[PORTFOLIO MATURITY, STATED PLAINLY]

- Frameworks: specified, simulated on synthetic data, and formally modeled at stated bounds where stated; HMAA_Refine is faithful on the discrete side; two of twenty TLA+ modules are faithful to implemented code. Self-assessed TRL 3.
- Governance Kernel: implemented, container-tested (83 passed, 1 skipped), not released, not independently validated. Self-assessed TRL 3 to 4.
- AUTHREX-ABORT: two TLA+ modules model checked at stated bounds, 105 implementation tests, open gates published. Self-assessed TRL 3.
- Hardware: twelve bill-of-materials-level reference designs; none built; no physical measurement. Self-assessed TRL 2, with BLADE-AV, AGENT-HSM and SWARM at 2 to 3 on simulation or emulator evidence.
- Independent validation: none. Government evaluation: one favorable AFRL white-paper evaluation memorandum (June 2026), not funded, not an endorsement. Deployment: none.

[GOVERNMENT ENGAGEMENT SOUGHT]

A technical review of the models, code, and evidence; a government-defined mission use case with an explicit confirm-or-abort tier; government-defined scenarios run on the seeded consoles; degradation and failure injection; then access to a software-in-the-loop stack, a hardware-in-the-loop environment, and an independent evaluator. Candidate vehicles that fit the maturity: research agreement or CRADA, BAA or SBIR/STTR topic response, testbed collaboration with an FFRDC, UARC, or laboratory, subcontract to a platform integrator, or a Phase 1 to 3 pilot. No mechanism is stated to be available or guaranteed.

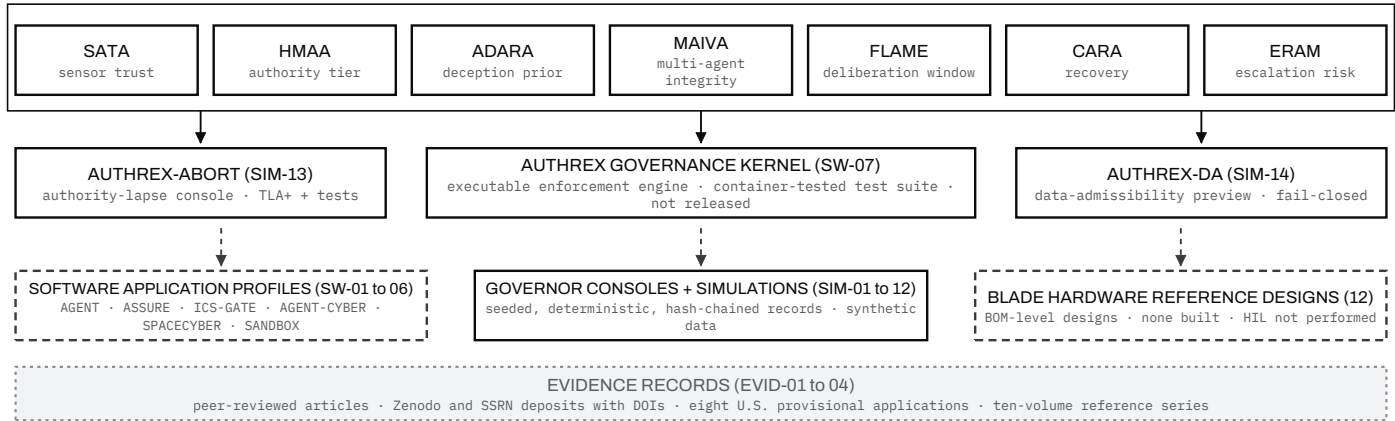
[WHAT THIS CATALOG DOES NOT CLAIM]

- That formal verification of a model proves any deployed implementation; bounded model checking of a configured finite model is not a proof of the physical system.
- That simulation equals field validation; every result here is simulation on synthetic data unless hardware is named, and no hardware is built.
- That any government organization has adopted, contracted, evaluated (beyond the AFRL record), certified, or endorsed the work.
- That agreement across evidence sources is truth; correlated false information is an open limitation (NEG-3 in AUTHREX-DA; identical-lie counterexample in ADARA).

PORTFOLIO AT A GLANCE

Forty-four register items in four families. What each is, what it addresses, how mature, what the evidence is, how it would connect, and what the government would need to provide.

GOVERNANCE FRAMEWORKS (7) • specified, simulated, formally modeled where stated



Encoding: solid = executable artifact exists; dashed = specification or reference design; gray = documentary record; arrows = the frameworks are instantiated by the artifact below.

Figure 1. Portfolio structure. The frameworks define the authority logic; the kernel, consoles, and simulations execute it in software; the hardware family specifies it for platforms; the evidence records document it.

Capability	Mission problem	Primary function	Current maturity	Evidence	Integration path	Government need
SATA	Sensor integrity cannot be assumed under jamming, spoofing, compromise	Per-source trust scalar from four diagnostics with attestation	TRL 3 (self); published protocol	JHSS 2026 article; Zenodo; simulation	Sensor-ingest boundary module	Recorded sensor streams; TPM node
HMAA	Binary permit/deny cannot express partial trust or degrade gracefully	Graded A3 to A0 tier with asymmetric hysteresis	TRL 3 (self); TLA+ at bounds	23,748 states depth 9; 98 tests	Enforcement point on the actuator path	Scenario with confirm-or-abort tier; SIL stack
ADARA	Manipulated perception inherited by governance	Deception prior modulating authority downward	TRL 3 (self)	Simulation; one faithful module; counterexample	Parallel tap on SATA output	Deception scenarios with ground truth; red team
MAIVA	One compromised agent or contaminated majority carries the outcome	Byzantine-resilient trimmed weighted median; CUSUM	TRL 3 (self); open defect	Faithful quorum model; 37 self-tests	Aggregator before the tier	Tolerable-compromise assumption; multi-vehicle SIL
FLAME	Veto not exercisable before an irreversible outcome	Consequence-scaled deliberation window; abort on timeout	TRL 3 (self)	Simulation; counterexample; requirement	Between tier and execution	Decision timeline; operators for a trial
CARA	Recovery short-circuited by one favorable signal	GREP recovery with non-compensatory terminal gate	TRL 3 (self)	Self-run enumeration; simulation	Invoked on ABORT or lockout	Recovery policy for a platform class
ERAM	Local actions cascade and compress strategic decision time	Escalation-risk model feeding the consequence tier	TRL 2 to 3 (self); working paper	SSRN 6176802; 600-run console	Cross-cutting monitor	Releasable exercise record
Governance Kernel	Agents hold whatever privilege their keys hold	Tier gating, approval queue, signed ledger at the tool boundary	TRL 3 to 4 (self); not released	83 tests passed in container	Proxy / runtime wrapper (MCP)	External reviewer team; tool inventory and policy
AUTHREX-ABORT	Fallbacks add capability when the system is least trustworthy	Evidence-conditioned authority lapse with latch	TRL 3 (self); gates open	2 TLA+ modules; 105 tests; soak	Abstract platform adapter	Safety-interface specification; T&E; adoption
AUTHREX-DA	Poisoned or repaired data admitted without an authority decision	Fail-closed admissibility verdicts with reason codes	TRL 2 to 3 (self); WP-0 blocked	101 tests; single-seed root; NEG-1 to 5	Pipeline stage (design)	Labeled poisoned records
Software profiles (6)	Same gap in agent, OT, cyber, space, and T&E; settings	Pipeline applied per setting; governor consoles	TRL 2 to 3 (self)	Specifications; consoles	Runtime wrapper or middleware	Select one profile for Phase 2
BLADE designs (12)	Hardware enforcement and roots of trust per domain	BOM-level governance nodes	TRL 2 (self); none built	Zenodo deposits; simulation; HSM emulator	Mission-system module (design)	Bench, HIL rig, range; integration partner

CORRECTIONS FIRST: WHAT CHANGED SINCE VERSION 5

The program applies withdrawals in place and records them beside the original. A reader who holds an earlier document should read this page before any capability page.

ID	Item	Earlier statement	Statement in this edition
CL-01	HMAA model-checking figure	Earlier figure of 48,751 states appears in the June 2026 two-page summary and older pages	Withdrawn; canonical figure is 23,748 distinct states at depth 9 with its scope clause, never stated without it
CL-02	SATA archived TLC run	formal-coverage.html still states a separate archived result of 523 distinct states at depth 9; the specification, configuration and log for that run could not be produced on request	Withdrawn from this catalog pending recovery of the artifacts; the as-published article figure (18,892 states) is cited only as published
CL-03	Rover 350-run and UAV 250-run counts	Withdrawn 7 Aug 2026 as unreproducible; still present on burakoktenli.com pages (evaluation.html, platform pages) and in the July white paper and two BAA white papers	Replaced by "seven fault scenarios" and "five adversarial scenarios" demonstrated in simulation; the aggregate "2,800+ runs" is not carried
CL-04	"Physically assembled" testbeds	July 2026 white paper describes the rover and UAV testbeds as physically assembled; the register records no platform as built	This edition states not built, physical build not yet documented
CL-05	Validation language on live pages	"Simulation Validated", "Validated in rover testbed", "Six Validated Scenarios", "ALL VALIDATIONS PASSED" appeared on framework and platform pages	Removed from the 1 Sep 2026 deploy; this edition uses demonstrated in simulation
CL-06	Simulation count: 24 versus 14 versus 37	V5 says twenty-four seeded simulations (counting governor consoles); the register has 14 SIM rows; the sites state 37 browser simulations (19 plus 18)	This edition uses 14 catalogued records (SIM-01 to 14) and states the other counts with their definitions; never interchangeable
CL-07	Name expansions	HMAA (Allocation, Architecture, Hierarchical Mission Authority Architecture); SATA (Attestation and Trust Anchoring); ADARA (Risk Architecture); MAIVA (Verification); FLAME (Flash War Latency Architecture); CARA (Control Authority Regulation; Cognitive Authority Recovery in the Zenodo title)	Canonical program names used; variants listed on each capability page; original artifact titles preserved
CL-08	Provisional receipts	V5 lists the three July 2026 filings as receipt pending; the patents page (verified 1 Sep 2026) carries receipts 78658472, 78659350, 78660113	Receipts carried in this edition
CL-09	Peer-reviewed article count	V5 states one peer-reviewed article; two further articles published 24 Aug and 31 Aug 2026	Three stated, with the caveat that one is not an AUTHREX technical paper
CL-10	BLADE-UUV and domain counts	The July white paper and the authrex.systems index list BLADE-UUV; the register has no such entry; domain counts vary (nine, ten, eleven) across properties	BLADE-UUV not carried (no deposit, no register row); the catalog counts twelve register platforms and does not state a domain count
CL-11	Standards-mapping overclaims	standards-mapping.html states "no unsafe state reachable from any initial state under any sensor input" and describes AUTHREX as "required by" MIL-HDBK-516C sections	Not carried; this edition uses mapped by the author, design target, architecturally relevant
CL-12	Superseded governance formulations	June 2026 wording; scores the evidence; grants authority in proportion; narrows in proportion; so the human stays in command; "narrows, and does not widen on its own"	Replaced by the per-source evidence-vector, non-compensatory-gate, admissible-set formulation; "narrows at once; widens only by authorization"
CL-13	Absolute hold rule	Earlier form: every irreversible action requires confirmation and cannot execute without it	Policy-conditioned form: where policy requires confirmation, execution occurs only after it; if absent at hold expiry the predeclared outcome is delay, handoff, or abort; pre-authorized autonomous execution must already exist in the envelope; timeout never creates authority

The full correction log, including items pending owner decision and pending page fixes on the public sites, is at the back of the catalog.

WHERE THE PROBLEM APPEARS

The problem is not autonomy in general. It appears where a machine can act faster than its supervisor can evaluate and the sensing that justifies the action can be degraded, jammed, or deceived. Either condition alone is covered by existing practice. Both together is the gap.

Where both conditions hold	Decision window	Degradable evidence	Relevant capabilities	What the portfolio offers today
Uncrewed air systems in GNSS-denied or contested electromagnetic conditions	Seconds	GNSS, RF, datalink	SATA, HMAA, ABORT	Trust falls; tier contracts at once; navigation failure is a hard veto; latch on lapse (simulation)
Counter-UAS engagement decisions	Seconds	Radar track, RF identification, EO/IR	ADARA, FLAME, HMAA, BLADE-CUAS	Deception prior extends the deliberation window; HANDOFF at A1 (simulation)
Ground robotic movement and convoy operations	Seconds to minutes	Lidar, vision, positioning	SATA, HMAA, CARA, Rover design	Speed and route constrained at A2 and A1; A0 (revoked) safe-stop; GREP recovery (simulation)
Spacecraft operations under light delay	Supervision impossible in real time	Star tracker, ranging, telemetry	FLAME, CARA, SPACECYBER, BLADE-SPACE	DELAY then ABORT to A0 if evidence is not refreshed; staged recovery (design and simulation)
Agentic AI acting on networks and infrastructure	Milliseconds	Tool output, retrieved content, provenance	Governance Kernel, AUTHREX-AGENT, DA	Tier enforcement and approval gating at the tool boundary; signed evidence (container-tested)
Operational technology and critical infrastructure	Seconds to minutes	Sensor plausibility, command provenance	ICS-GATE, BLADE-INFRA-OT, FLAME	Mandatory deliberation before control writes; fail-closed boundary adjudication (design and console)
Multi-platform and swarm operations	Seconds	Inter-agent attestations, RF	MAIVA, BLADE-SWARM	Trimmed weighted median; safe-halt by default under denied RF (simulation; open quorum defect)

OPERATIONAL PROBLEM TO CAPABILITY

Operational problem	Capability	How it helps (as demonstrated in simulation)
Sensor confidence degrades	SATA, HMAA	Trust scalar falls; authority tier contracts immediately; upgrade is delayed by hysteresis and never restored by timeout
Communications lost	FLAME, CARA, ABORT	Latency budget holds the decision (DELAY); loss of required evidence contracts the admissible set; pre-authorized contingencies only
Target evidence becomes stale	FLAME, HMAA	Deliberation window and evidence age gate the irreversible action; HANDOFF at A1 or ABORT to A0
Agent exceeds assigned permission	Governance Kernel, AUTHREX-AGENT	Tier enforcement and human approval gating at the agent runtime; signed audit evidence
Platform enters degraded mode	CARA	Deterministic GREP recovery with a non-compensatory terminal gate
Audit reconstruction required	All consoles	Hash-chained record with deterministic replay; a reviewer re-runs a result from the same packaged build
Data of uncertain provenance offered	AUTHREX-DA	ADMIT, ADMIT_FLAG, QUARANTINE, or REJECT with reason code; never repaired
Safety invariant must hold	HMAA, ABORT (TLA+)	Invariants held at stated bounds; bounded model checking is not a proof of the physical system

The applications above are potential mission applications. No adoption, contract, or evaluation by any government organization is claimed, and no specific DoD program is stated to be using or considering the work.

WHERE THE AUTHORITY LAYER SITS

AUTHREX governs the path between the autonomy software and the actuator or tool. It does not modify the mission autonomy and it is not physically isolated from it; separation is architectural, enforced through the authority interface the host platform exposes.

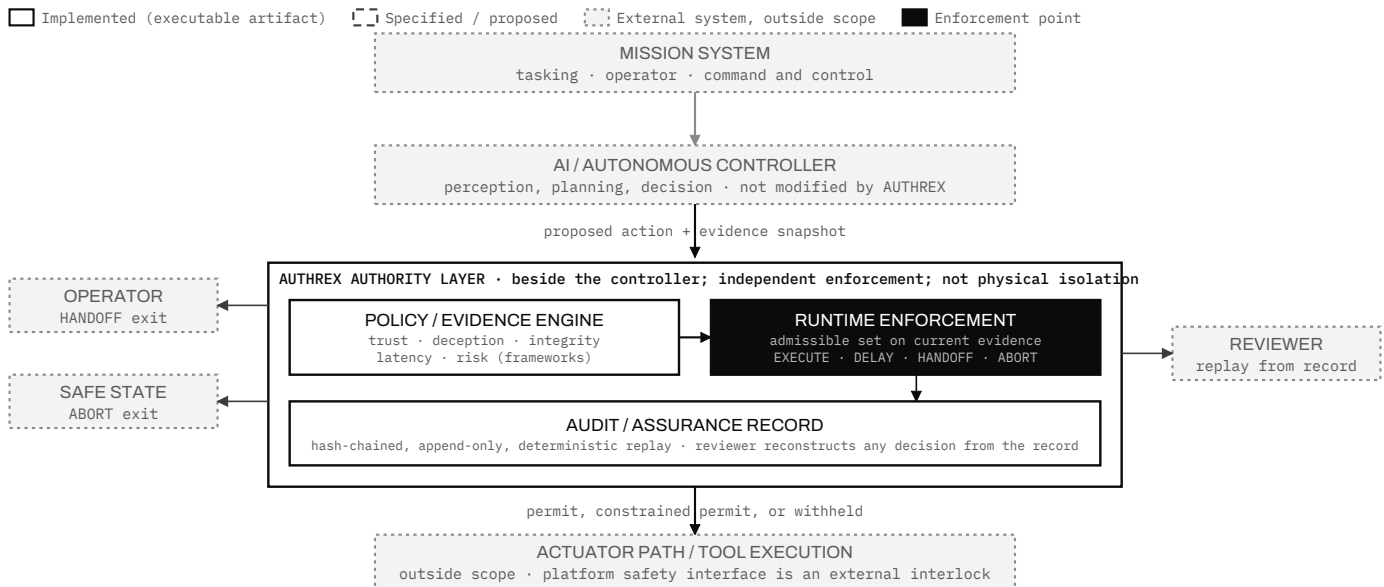
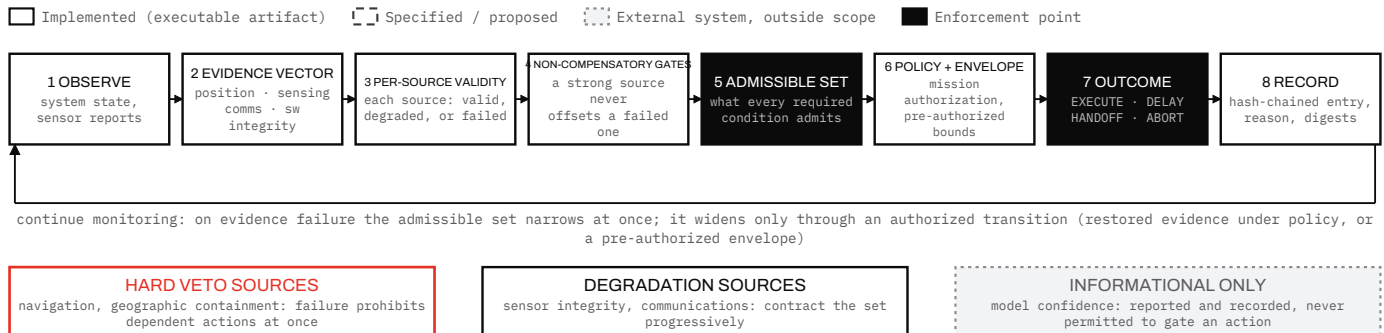


Figure 2. Layered view and trust boundary. Mission tasking, the autonomy stack, and the actuator are outside scope. Operator handoff, safe state, and the reviewer's replay are the three exits. The platform safety interface is an external interlock; AUTHREX supplies the governance and evidence layer that decides whether authority still holds and initiates no physical effect itself.

THE GOVERNING FORMULATION

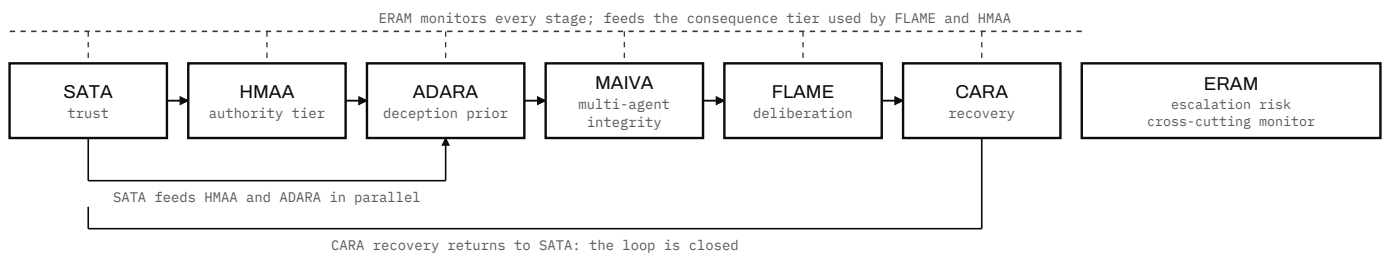
Each of position, sensing, communications, and software integrity carries its own validity state (an evidence vector). The system may execute only what every required condition admits (an admissible set). A strong source never compensates for a failed one (non-compensatory gates). On failure the admissible set narrows at once to what the remaining evidence admits; it widens only by an authorized transition, either restored evidence under policy or a pre-authorized envelope. Where mission policy requires human confirmation before an irreversible action, execution occurs only after that confirmation; if confirmation is absent when the bounded hold expires, the predeclared outcome is delay, handoff, or abort. Where policy explicitly authorizes autonomous execution after a bounded hold, that authority must already exist in the pre-authorized envelope and is never created by timeout alone. Permission is granted by mission authorization and applicable policy under assumptions about evidence validity.

THE AUTHORITY CONTROL LOOP AND THE SEVEN-STAGE LIFECYCLE



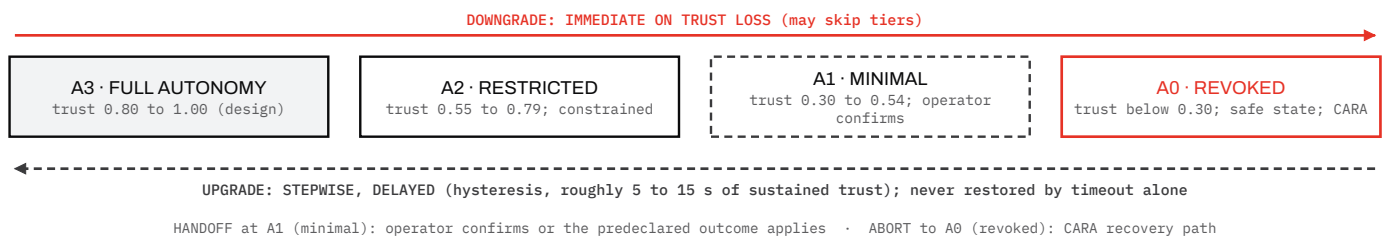
Evidence classes as implemented in AUTHREX-ABORT (SIM-13). The loop is the program-level design pattern; individual artifacts implement subsets of it (see each capability page).

Figure 3. Authority control loop. Steps 5 and 7 are enforcement points. The loop is the design pattern the artifacts implement in whole or in part; AUTHREX-ABORT implements the evidence classes shown; the Governance Kernel implements steps 6 to 8 for tool calls.



HMAA is the only component that issues an authority tier; every outcome (EXECUTE, DELAY, HANDOFF, ABORT) is issued from that one place.

Figure 4. Seven-stage lifecycle. SATA feeds HMAA and ADARA in parallel; MAIVA supplies multi-agent integrity; FLAME governs the latency budget; CARA performs recovery and returns to SATA; ERAM monitors escalation risk across the whole pipeline.



HMAA authority automaton. Model checked in TLA+ at stated bounds (23,748 distinct states at depth 9); the discrete automaton only, not continuous behaviour between decision instants.

Figure 5. HMAA graded authority with asymmetric hysteresis and the operator handoff path.

HUMAN-MACHINE AUTHORITY RELATIONSHIP AND THE DEGRADATION LADDER

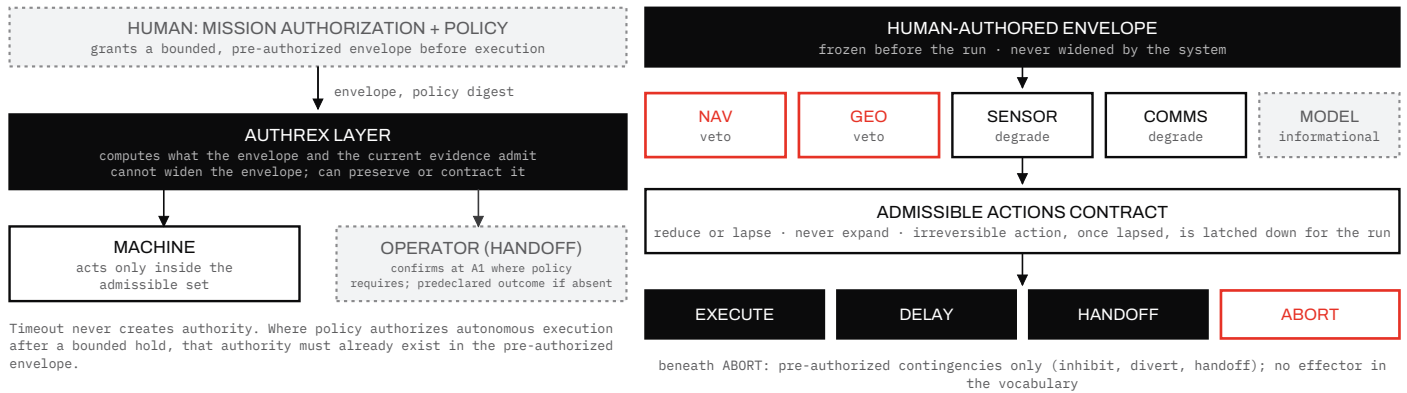


Figure 6 (left). The human grants a bounded envelope; the layer computes what it admits on current evidence; the machine acts only inside the admissible set; the operator confirms at A1 where policy requires. Figure 7 (right). Degradation ladder as implemented in AUTHREX-ABORT: hard vetoes, degradation sources, and an informational source that never gates.

AUTHORITY FLOW UNDER DEGRADATION

Step	State	What is recorded
1	Normal authority (A3): full envelope admissible	Envelope digest; evidence snapshot identity
2	Evidence degrades: one or more sources move to degraded or failed	Per-source validity; diagnostic vector
3	Trust and policy check: non-compensatory gates evaluated	Gate results; policy version
4	Authority contracts: admissible set narrows at once; irreversible actions may latch	Tier change; lapsed set; reason
5	Restricted action set (A2, A1): DELAY where the window applies; HANDOFF where policy requires confirmation	Window; operator decision or predeclared outcome
6	ABORT or safe state (A0): pre-authorized contingencies only; CARA recovery path	Contingency request and acknowledgment; recovery phases
7	Audit record: hash-chained entries with previous hash; deterministic replay	Ledger head, digest, findings, final state, run identity

Each step is a state in the model, not a policy statement; the record at the end is what makes the sequence reconstructible after the fact.

[CAPABILITY PROFILE · AUTHREX-01]

SATA · SENSOR AND ACTOR TRUST ASSESSMENT

Produce a continuous, per-source trust scalar from sensor evidence and hardware attestation so that authority is contracted before a corrupted input is acted on, not after an invariant breaks.

TESTED MODELED IMPLEMENTED

Name variants: Published article title: Sensor Attestation and Trust Anchoring (Journal of Hardware and Systems Security). Both expansions refer to the same protocol; the catalog uses the program name.

PRIMARY FUNCTION	Continuous sensor-trust scalar in [0,1] from weighted Dempster-Shafer fusion across four diagnostics (internal consistency, cross-sensor agreement, temporal stability, physical plausibility), anchored in an attestation chain design (TPM 2.0 measurements, signatures, monotonic counters).
MISSION AREAS	Cross-domain foundation layer: uncrewed air and ground systems under GNSS denial, counter-UAS track validation, maritime sensor fusion, spacecraft telemetry integrity.
EVIDENCE ANCHORS	Peer-reviewed article, J. Hardware and Systems Security (Springer Nature) 10, 17 (2026), DOI 10.1007/s41635-026-00190-4; Zenodo DOI 10.5281/zenodo.18936251 (specification, reference implementation, reproducibility artifacts); seeded browser simulation.
CURRENT STANDING	Specified and published; protocol state machine model checked as reported in the article; reference implementation and seeded simulation exist; not built on hardware; not independently reproduced.

SELF-ASSESSED TRL	TRL 3 (analytical and simulation proof of concept of the critical function). No formal TRL determination has been performed.
INTELLECTUAL PROPERTY	U.S. provisional application 64/002,453, filed March 11, 2026 (receipt 74808459), reported as filed; unexamined; confers no enforceable rights.

[MISSION PROBLEM]

Autonomy stacks act on sensor reports whose integrity cannot be assumed under jamming, spoofing, calibration drift, or hardware compromise. A control decision that is formally correct on corrupted inputs is still a wrong outcome, and the supervisor typically learns of the corruption only when an invariant is violated. Under GNSS denial or radio-frequency interference the corruption can be gradual and consistent across correlated sources, which defeats simple thresholding and majority voting. The mission consequence is an action executed with full authority on evidence that no longer deserves it.

[OPERATIONAL RELEVANCE]

- Uncrewed air and ground systems in GNSS-denied or contested electromagnetic conditions (decision windows of seconds).
- Counter-UAS track validation where radar, RF identification, and electro-optical reports must be weighed before an engagement recommendation.
- Spacecraft telemetry integrity under light delay, where a supervisor cannot re-check the sensor in real time.
- Resilient PNT: trust in position is treated as evidence with its own validity state rather than as an assumed input.

[CAPABILITY DESCRIPTION]

SATA is the foundation trust layer of the pipeline. For each source it computes a trust scalar in the interval zero to one by combining four diagnostics with weighted Dempster-Shafer belief functions, keeping ignorance (no evidence) distinct from distrust (contradicting evidence). The published protocol binds sensor reports to cryptographically committed attestation records so that a report can be checked for origin and replay, and specifies a formal adversary model under which the signature barrier is analyzed. Every downstream stage consumes the resulting scalar: HMAA converts it to an authority tier and ADARA tests it for manipulation. SATA does not decide what the platform may do; it states how much the evidence deserves to be believed.

[SATA • ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]

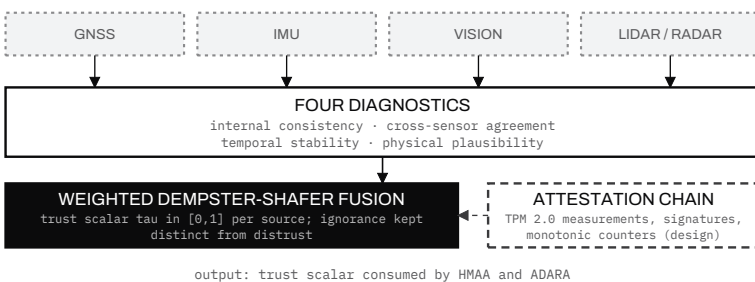


Figure 8. SATA technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[HOW IT WORKS]

1. Ingest raw and fused sensor reports at the host stack boundary, together with attestation records where the platform provides them.
2. Run the four diagnostics per source: internal consistency, cross-sensor agreement, temporal stability, physical plausibility.
3. Combine diagnostics with weighted Dempster-Shafer fusion into a trust scalar per source, preserving the ignorance-versus-distrust distinction.
4. Check the attestation chain (measurement, signature, monotonic counter) and reject replayed or unattested reports where the design is instantiated.
5. Emit the trust scalar over the authority interface to HMAA (tier computation) and ADARA (deception prior).
6. Record the scalar, diagnostics, and attestation outcome to the hash-chained ledger for later reconstruction.
7. Continue at the control-loop rate; a falling scalar contracts authority at once, and recovery of the scalar alone never restores authority.

[TECHNICAL DIFFERENTIATION]

Conventional runtime assurance switches to a safe controller when a monitored invariant is about to be violated and assumes the monitor's inputs are valid. Redundancy and voting mask a faulty channel but assume independent failure, which common-mode conditions such as spoofing defeat. SATA computes trust before authority is computed, so degraded inputs lower authority ahead of the violation rather than at the instant it occurs, and treats agreement across heterogeneous sources as evidence about trust rather than as truth. This is a complement to Simplex-style assurance and voting, not a replacement.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Ignorance kept distinct from distrust in fusion	Design property; exercised in simulation
Attestation chain: origin, replay resistance, monotonic counters	Design property (hardware anchoring not built); adversary model stated in the article
Signature-barrier unforgeability bound	Analytical bound reported in the published article under its stated adversary model
Protocol state-machine safety invariants	Reported in the article as model checked at stated bounds (see evidence table); not independently reproduced
Deterministic scalar for identical inputs	Design property; reproducible in the seeded console
Auditability of every trust decision	Implemented in the seeded console (hash-chained record); not implemented on hardware

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
Peer-reviewed publication of the protocol	JHSS 10, 17 (2026); received 4 Jun, accepted 14 Jul, published 21 Jul 2026	DOI 10.1007/s41635-026-00190-4
TLA+ bounded model check of the protocol state machine, as reported in the article	18,892 distinct reachable states (47,616 generated); zero violations of eight safety invariants (as published; not re-executed for this catalog)	Article and Zenodo 10.5281/zenodo.18936251
Monte Carlo conformance and sensitivity campaign, as reported in the article	10,000 runs; bit-exact reproducibility of the packaged artifacts	Article and Zenodo deposit
Seeded browser simulation	Deterministic replay from seed; synthetic data	burakoktenli.com/sata-simulation
Separately archived TLC run (523 distinct states, depth 9)	WITHDRAWN from this catalog: the site-check log records the run as archived (SATA_verification.zip), but the archive could not be produced for this edition; the figure is not carried	Correction log CL-02

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (reference implementation)
Executable artifact	Yes
Formal model	Yes
Simulation evidence	Yes
Hardware validation	No
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- Hardware anchoring (TPM 2.0 attestation chain) is a design; no platform carrying it has been built, so replay resistance and timing on real hardware are unmeasured.
- Sensor models are parameterized (Gaussian noise, step faults, drift ramps), not physics-based; adaptive adversaries are not modeled.
- The program's formal-coverage record lists no SATA module among the FAITHFUL models; the article-reported check is an as-published result.
- Simulation and formal-model evidence only; all data synthetic; no hardware-in-the-loop test, no field validation, no independent replication, no deployment.
- Not independently reproduced or technically verified by another party; second-environment reruns are confirmation by the same producer.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	The critical function (continuous trust from fused diagnostics with attestation) has an analytical treatment in a peer-reviewed article, a reference implementation, and a seeded simulation that exercises it: analytical and experimental proof of concept. It has not been validated as a component in a laboratory environment on representative hardware, which TRL 4 requires. Disclosure: the published article states TRL 4 on its own reading (simulation-evaluated, with explicit assurance gaps G1 to G4 and hardware not yet measured); this catalog assesses TRL 3 for the reason given, and the difference is recorded rather than reconciled.
EVIDENCE SUPPORTING	concept defined and published; analytical treatment (adversary model, unforgeability bound); software implementation (reference); simulation demonstration
EVIDENCE REQUIRED FOR NEXT TRL	(1) Component validation in a laboratory environment: SATA running on a representative compute node against recorded or injected sensor streams with attestation hardware present. (2) Independent replication of the published model check and Monte Carlo campaign from the packaged deposit. (3) Measured contraction latency and false-restriction rate under representative interference (HIL).

[INTEGRATION PROFILE]

INPUTS	Raw and fused sensor reports; attestation records (TPM measurements, signatures, counters) where available; platform clock.
OUTPUTS	Per-source trust scalar in [0,1]; diagnostic vector; attestation verdict; ledger entry.
INTERFACES	Authority interface exposed by the host stack (message bus or middleware, ROS 2 candidate); TPM 2.0 for attestation (design).
DEPLOYMENT	Software module beside the controller at the sensor-ingest boundary; standalone browser console for evaluation.
DEPENDENCIES	Host sensor reports; optional hardware root of trust; no network required for the console.
SECURITY BOUNDARY	Inside: fusion logic, ledger. Outside: sensors, attestation hardware, controller. A compromised attestation root defeats the anchoring.
DATA REQUIREMENTS	Synthetic today; government-furnished recorded sensor streams (unclassified or CUI as the sponsor determines) would enable Phase 2.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Technical peer review of the published protocol against the reviewer's own sensor-integrity threat models.
- Representative recorded sensor datasets with known degradation or spoofing episodes (unclassified is sufficient to start).
- Laboratory access with a ROS 2 platform and a TPM-equipped compute node for Phase 4 to 5 component validation.
- Independent replication of the packaged model-checking and Monte Carlo artifacts by a government laboratory or FFRDC.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE · AUTHREX-02]

HMAA · HUMAN-MACHINE AUTHORITY ALLOCATION

Convert trust and operational context into an explicit, graded, enforceable authority state so that human control remains meaningful when the machine acts faster than a supervisor can evaluate.

MODELED TESTED IMPLEMENTED

Name variants: Historical expansions: Human-Machine Authority Architecture (evidence site), Hierarchical Mission Authority Architecture (title of Zenodo record 18861653). Original artifact titles are preserved unchanged. Tier names: the current HMAA page uses Full Autonomy, Restricted, Minimal, Revoked; the archived V4 catalog used Full, Standard, Restricted, Safe, which is not carried.

PRIMARY FUNCTION	Graded four-level authority state (A3 full autonomy, A2 restricted, A1 minimal, A0 revoked; design thresholds on the fused trust scalar of 0.80, 0.55 and 0.30 on the HMAA page) computed deterministically from the trust scalar and operational context, with asymmetric hysteresis: downgrade is immediate on trust loss, upgrade is stepwise and delayed.
MISSION AREAS	Any platform where a human must retain meaningful control at machine speed: UAS engagement decisions, autonomous ground platforms, AI-enabled command-and-control recommendations.
EVIDENCE ANCHORS	TLA+ model checked with TLC (scope clause below); 42-file Python package with 98 tests and 7 adversarial experiments; seeded browser simulation; Zenodo DOI 10.5281/zenodo.18861653.
CURRENT STANDING	Formally analyzed at stated bounds and demonstrated in simulation; HMAA_Refine is FAITHFUL on the discrete side and BOUNDED on the continuous side in the formal-coverage record; not built on hardware; not independently reproduced.
SELF-ASSESSED TRL	TRL 3. No formal TRL determination has been performed.
INTELLECTUAL PROPERTY	U.S. provisional application 63/999,105, filed March 7, 2026 (receipt 74759595), reported as filed; unexamined; confers no enforceable rights.

[MISSION PROBLEM]

Human oversight of a machine is nominal when the machine acts faster than the human can evaluate. A binary permit-or-deny cannot express partial trust and cannot degrade gracefully, so operators either over-trust (leave full authority in place while evidence erodes) or over-restrict (lock the system out when a partial contraction would keep the mission alive). Oscillation is the second failure: a system that restores authority as soon as a signal recovers can flap between states at the sensor noise rate. The consequence is either an unsafe action at full authority or a mission lost to an unnecessary lockout.

[OPERATIONAL RELEVANCE]

- Engagement-authority chains for uncrewed aircraft where the confirm-or-abort decision must be explicit at a defined tier.
- Ground robotic movement and convoy operations under lidar or vision degradation (speed and route constrained at A2 and A1).

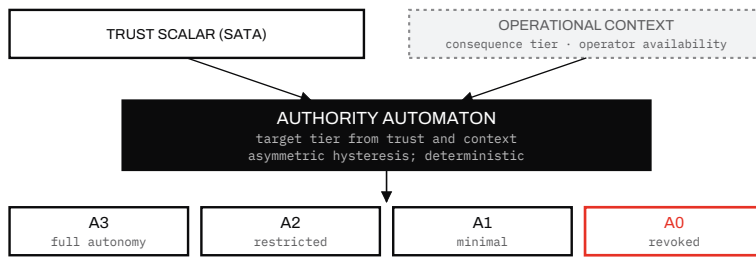
- AI-enabled command and control where a recommendation must carry the authority tier under which it was generated.
- Human-machine teaming and collaborative combat aircraft concepts, where the authority held by the machine must be legible to the pilot at all times.

[CAPABILITY DESCRIPTION]

HMAA is the single component of the pipeline that issues an authority tier, so a reviewer tracing any decision has one place to look. The automaton takes the trust scalar from SATA and the operational context (consequence tier, operator availability) and computes a target tier. Downgrade to the target is immediate and may skip tiers; upgrade is stepwise and delayed by a hysteresis interval of roughly 5 to 15 seconds of sustained trust so that a transient recovery cannot restore authority. At A1 (minimal) the operator confirms where policy requires it, or the predeclared outcome applies; at A0 (revoked) actuation is withheld and a safe state is commanded, which invokes CARA. The automaton is deterministic: identical inputs produce identical tiers, which is what makes deterministic replay of a recorded run possible.

[HMAA • ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]



tier issued to the enforcement point on the actuator path; HANDOFF at A1, ABORT to A0

Figure 9. HMAA technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[HOW IT WORKS]

1. Receive the per-source trust scalar from SATA and the current operational context.
2. Compute the target tier from trust and context using the configured, versioned policy thresholds.
3. If the target is below the current tier, downgrade immediately (tiers may be skipped) and record the transition with its reason.
4. If the target is above the current tier, hold; upgrade one tier only after the hysteresis interval has elapsed with trust sustained above threshold.
5. At A1 (minimal), route the pending action to the operator for confirmation where mission policy requires it; if confirmation is absent when the bounded hold expires, the predeclared outcome (delay, handoff, or abort) applies.
6. At A0 (revoked), withhold actuation, command the safe state, and hand control to CARA for structured recovery.
7. Write every tier change and outcome (EXECUTE, DELAY, HANDOFF, ABORT) to the hash-chained record.

[TECHNICAL DIFFERENTIATION]

A kill switch has two states and no memory; role-based access control grants permission by identity, not by the current evidence. HMAA makes authority a computed, time-varying, enforceable state with a formal model behind it, and it separates the rate of downgrade from the rate of upgrade so that safety and availability are governed by different rules. Against Simplex-style runtime assurance it adds a gradient between full autonomy and full intervention; against policy-as-code it supplies the runtime state that the policy acts on. It composes with both and replaces neither.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Tier confinement: no tier outside A3 to A0	Held in TLA+ at stated bounds (invariant)
Immediate downgrade on trust loss	Held in TLA+ at stated bounds (invariant)
Hysteresis on upgrade; no restoration by timeout	Held in TLA+ at stated bounds (invariant); timing values are design parameters
Stepwise upgrade path (UpgradelsStepwise)	Vacuous at this bound: not established
Trust-authority weak consistency	Vacuous at this bound: not established
Liveness: an eventual outcome is issued	One liveness property held at stated bounds
Model-to-code conformance	By construction and by the 98-test package; not by refinement proof (Obligation 6 open)
Human override and handoff at A1	Design property; exercised in the seeded console

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
TLC model check of the HMAA automaton	23,748 distinct reachable states at depth 9 (26,397,356 states generated); the result describes the discrete authority automaton under the assumption that instantaneous authority equals its target and does not cover continuous behaviour between decision instants; of 8 stated properties, 5 invariants and 1 liveness property held, 2 vacuous at this bound (UpgradelsStepwise, TrustAuthorityWeakConsistency)	HMAA_verification_v2_deposit.zip (deposit of record); Zenodo 10.5281/zenodo.18861653
Faithfulness of HMAA_Refine.tla	FAITHFUL on the discrete side; BOUNDED on the continuous side	authrex.systems/formal-coverage
Python package tests	42 files; 98 tests; 7 adversarial experiments; deterministic traces	Zenodo deposit
Seeded browser simulation	Deterministic replay from seed; synthetic data	burakoktenli.com/hmaa-simulation
Earlier figure of 48,751 states	WITHDRAWN; never to be reasserted	Correction log CL-01

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (Python package)
Executable artifact	Yes
Formal model	Yes
Simulation evidence	Yes
Hardware validation	No
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- Bounded model checking of the discrete automaton is not a proof of the physical system and does not cover continuous behaviour between decision instants.
- Two of eight stated properties are vacuous at the checked bound; they are not established.
- Obligation 6 (GLUE between discrete and continuous models) is open; initial-state correspondence was withdrawn as bound-dependent (held at MaxTick = 3, violated at 4).
- Hysteresis timing (5 to 15 s) is a design parameter tuned in simulation; not characterized against any real control loop.
- Simulation and formal-model evidence only; all data synthetic; no hardware-in-the-loop test, no field validation, no independent replication, no deployment.
- Not independently reproduced or technically verified by another party; second-environment reruns are confirmation by the same producer.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	The authority automaton has a formal model checked at stated bounds, a tested implementation, and seeded simulations exercising the critical function under trust loss and recovery. Laboratory validation on a representative controller loop (TRL 4) has not been performed.
EVIDENCE SUPPORTING	concept defined and specified; formal model (TLA+/TLC) at stated bounds; software implementation with test suite; simulation demonstration
EVIDENCE REQUIRED FOR NEXT TRL	(1) Software-in-the-loop with a representative controller or flight-stack interface (ArduPilot or ROS 2) in a laboratory environment. (2) Independent re-execution of the deposited TLC configuration and the 98-test package. (3) Closure or explicit scoping of the two vacuous properties at a larger bound, and a decision on Obligation 6.

[INTEGRATION PROFILE]

INPUTS	SATA trust scalar; consequence tier (from ERAM or mission policy); operator availability; mission clock.
OUTPUTS	Authority tier A3 to A0; outcome (EXECUTE, DELAY, HANDOFF, ABORT); ledger entry.
INTERFACES	Authority interface to the enforcement point on the actuator path; operator confirmation channel; ledger.
DEPLOYMENT	Software module beside the controller; the automaton is small enough for embedded targets but no embedded build exists.
COMPUTE REQUIREMENTS	Not characterized on target hardware.
SECURITY BOUNDARY	Inside: automaton, policy thresholds (versioned), ledger. Outside: trust inputs, operator channel, actuator interlock.
DATA REQUIREMENTS	Synthetic; government-defined scenarios can be run on the seeded console without data transfer.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Independent review of the TLA+ model and the 98-test package (Phase 1) by a formal-methods-capable government laboratory.
- A government-defined mission scenario with an explicit confirm-or-abort tier so that policy thresholds can be set by the sponsor rather than by the author.
- Access to a software-in-the-loop flight or vehicle stack for Phase 4.
- Guidance on operator-comprehension metrics for authority state, a success metric the program cannot evaluate alone.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE • AUTHREX-03]

ADARA • ADVERSARIAL DECEPTION-AWARE REASONING

Estimate the probability that the system's inputs are being manipulated and modulate authority downward before a deceptive input becomes an authorized action.

MODELED TESTED

Name variants: Deposit title: Adversarial Deception-Aware Risk Architecture (Zenodo record 19043924). Same framework.

PRIMARY FUNCTION	Deception prior $P(\text{adversarial})$ computed from input anomalies, temporal correlation, cross-sensor consistency, and Bayesian mission history; authority adjustment $A_{\text{adj}} = A_{\text{hmaa}} \times (1 - \lambda \times P_{\text{deception}})$; phantom-fleet module against fabricated hostile scenarios.
MISSION AREAS	Counter-UAS engagement under spoofing; maritime and air domain awareness against decoys; command and control against injected false tracks; electronic-warfare resilience of perception.
EVIDENCE ANCHORS	Zenodo DOI 10.5281/zenodo.19043924; seeded browser simulation; formal-coverage record: ADARA_XSI_EXACT.tla FAITHFUL; ADARA.tla and ADARA_IMPL_V2.tla surrogate countermodels; ADARA_IMPL.tla withdrawn.
CURRENT STANDING	Demonstrated in simulation; one faithful formal model of the exact cross-sensor inconsistency function; a strong-form counterexample is published and drives a stated requirement; not independently reproduced.
SELF-ASSESSED TRL	TRL 3. No formal TRL determination has been performed.
INTELLECTUAL PROPERTY	U.S. provisional application 64/110,218, filed July 13, 2026 (receipt 78658472), reported as filed; unexamined; confers no enforceable rights.

[MISSION PROBLEM]

Adversaries manipulate what an autonomous system perceives: spoofed positions, decoy tracks, blinded cameras, fabricated hostile scenarios. A governance layer that trusts the perception layer inherits its errors, and a trust score that looks only for disagreement can be defeated by an attack that makes every channel agree. Honest confusion (a degraded sensor) and hostile manipulation call for different responses, and an architecture that cannot tell them apart either over-restricts on every fault or under-restricts under attack. The consequence in counter-UAS is an engagement recommendation built on a manufactured track.

[OPERATIONAL RELEVANCE]

- Counter-UAS engagement decisions where the track feeding the recommendation may be spoofed or decoyed.
- Maritime and air domain awareness against decoys and AIS or ADS-B manipulation.
- Command and control against injected false tracks in a common operating picture.

- Contested electromagnetic environments in which perception attacks are expected rather than exceptional.

[CAPABILITY DESCRIPTION]

ADARA runs in parallel with HMAA on the SATA output. It computes a proactive deception prior from four evidence streams and uses it to modulate the authority tier downward as the probability of manipulation rises. The Phantom Fleet detection module identifies AI-hallucinated hostile forces from correlated anomalies across sensor modalities so that a fabricated fleet does not drive an engagement. Deception is a first-class input to authority, not a detection alarm handled elsewhere. The formal record adds a requirement learned from a counterexample: deception detection must weight agreement, not only disagreement, because identical lies across all channels read as perfect consistency to a spread-based score.

[ADARA • ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]

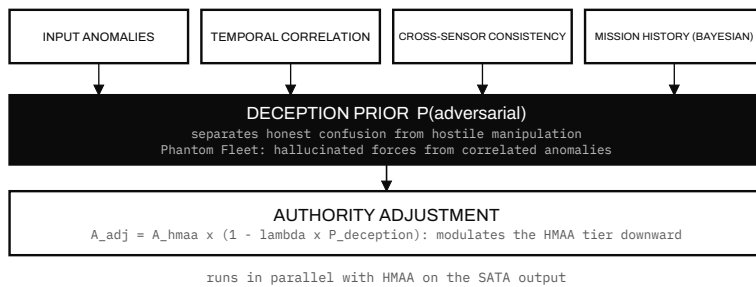


Figure 10. ADARA technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[HOW IT WORKS]

1. Receive the SATA per-source trust vector and the raw inconsistency features.
2. Compute input-anomaly, temporal-correlation, and cross-sensor-consistency scores; update the Bayesian mission-history prior.
3. Combine into $P(\text{adversarial})$, separating the honest-confusion hypothesis from the hostile-manipulation hypothesis.
4. Run the Phantom Fleet check (correlated anomalies across modalities) on any newly asserted hostile scenario.
5. Apply the authority adjustment: $A_{\text{adj}} = A_{\text{hmaa}} \times (1 - \lambda \times P_{\text{deception}})$, which can only reduce the HMAA tier.
6. Emit the adjusted authority and the deception prior to the enforcement point and to FLAME (longer deliberation under suspected deception).
7. Record the prior, its components, and the adjustment to the ledger.

[TECHNICAL DIFFERENTIATION]

Anomaly detection and runtime monitoring flag deviations and leave the response to someone else; ADARA converts the flag into an authority change with a recorded outcome. Redundancy and voting cannot see an attack that makes every channel agree; ADARA's stated requirement is that agreement itself be weighed as evidence. The comparison is structural and is supported by the published counterexample; no claim is made about the performance of any named product.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Adjustment can only reduce authority (monotone downward)	Design property; exercised in simulation
Exact cross-sensor inconsistency function	FAITHFUL formal model (ADARA_XSI_EXACT.tla imports and checks the implemented function)
Detection of identical-lie attacks	NOT established: published strong-form counterexample (all channels lie identically, conflict reads zero); requirement recorded
Deception prior recorded with components	Implemented in the seeded console
Near-minimum observation variant (ADARA_XSI_NEAR)	NOT FAITHFUL; result not relied upon

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
Seeded browser simulation	Deterministic replay from seed; synthetic deception scenarios	burakoktenli.com/adara-simulation
Faithful formal model of the inconsistency function	ADARA_XSI_EXACT.tla verdict FAITHFUL	authrex.systems/formal-coverage
Strong-form counterexample	Naive form held, strong form failed: identical readings read as agreement regardless of baselines	formal-coverage record
Withdrawn surrogate	ADARA_IMPL.tla withdrawn: discarded per-sensor baselines and the standard-deviation divisor	formal-coverage record
Zenodo deposit	Specification, source, simulation data	DOI 10.5281/zenodo.19043924

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (simulation engine)
Executable artifact	Yes
Formal model	Partial (one faithful module)
Simulation evidence	Yes
Hardware validation	No
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- The adversary is scripted (spoofing, jamming, drift injection); adaptive adversaries that respond to the system are not modeled.
- Identical-lie attacks across all channels are a published open failure mode; the requirement is stated but the repair is not implemented in the shipped engine.
- The lambda weighting is a design parameter tuned on synthetic data; no calibration against recorded deception episodes exists.
- Simulation and formal-model evidence only; all data synthetic; no hardware-in-the-loop test, no field validation, no independent replication, no deployment.
- Not independently reproduced or technically verified by another party; second-environment reruns are confirmation by the same producer.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	Critical function (deception prior modulating authority) is implemented in a seeded simulation and one component function has a faithful formal model; no laboratory validation against recorded or live deception has been performed.
EVIDENCE SUPPORTING	concept defined; partial formal model; simulation implementation
EVIDENCE REQUIRED FOR NEXT TRL	(1) Government-defined deception scenarios with ground truth (Phase 2) to measure false-restriction and miss rates. (2) Implementation and faithful-model check of the agreement-weighting repair. (3) Red-team evaluation against an adaptive adversary in a software-in-the-loop environment.

[INTEGRATION PROFILE]

INPUTS	SATA trust vector and inconsistency features; asserted hostile scenarios; mission history.
OUTPUTS	P(adversarial); adjusted authority; deception components; ledger entry.
INTERFACES	Parallel tap on the SATA output; modulation input to the HMAA tier and FLAME window.
DEPLOYMENT	Software module; browser console for evaluation.
SECURITY BOUNDARY	Inside: prior computation, phantom-fleet check. Outside: perception stack, track sources.
DATA REQUIREMENTS	Synthetic today; recorded spoofing or decoy episodes with ground truth would enable calibration.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Government-defined deception scenarios with ground truth for Phase 2 simulation evaluation.
- Red-team assessment by an electronic-warfare or counter-UAS test organization.
- Guidance on acceptable false-restriction rates for the target mission, which determines lambda.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE • AUTHREX-04]

MAIVA • MULTI-AGENT INTEGRITY VOTING

Aggregate authority across multiple agents so that a single compromised node, or a minority acting in concert, cannot carry the outcome, while treating agreement as evidence rather than truth.

MODELED

TESTED

NOT YET VALIDATED

Name variants: Deposit and patents-page title: Multi-Agent Integrity Verification. Same framework; MIT-licensed deposit v5.18.

PRIMARY FUNCTION	Byzantine-resilient aggregation using a trimmed weighted median resistant to f adversaries in a $3f+1$ roster, with CUSUM-augmented anomaly detection and an action-gate classification informed by DoDD 3000.09 categories (author mapping).
MISSION AREAS	Swarm and multi-platform operations; distributed sensing; federated model contributions; any decision aggregated across heterogeneous actors.
EVIDENCE ANCHORS	Zenodo DOI 10.5281/zenodo.19015517 (MIT); TLA+ specification; 37 self-tests; faithful model MAIVA_QUORUM.tla; seeded browser simulation; BLADE-SWARM refinement (5 safety, 3 liveness, model checked at stated bounds).
CURRENT STANDING	Demonstrated in simulation; one faithful formal model located a documented open defect in production source (quorum at $n = 1$); the repair is specified, not shipped; not independently reproduced.
SELF-ASSESSED TRL	TRL 3 for the aggregation logic; the quorum defect is open. No formal TRL determination has been performed.
INTELLECTUAL PROPERTY	U.S. provisional application 64/110,221, filed July 13, 2026 (receipt 78659350), reported as filed; unexamined; confers no enforceable rights. Deposit under MIT license.

[MISSION PROBLEM]

In multi-agent operations a single compromised agent, or a majority contaminated by the same false information, can carry the outcome. Consensus is not truth: majority voting assumes independent failure, and a quorum rule whose threshold is derived from the roster it is checking can only report on itself. The mission consequence is a swarm that commits to an action because its own bookkeeping says enough members agreed, when the agreement was manufactured.

[OPERATIONAL RELEVANCE]

- Attributable swarm autonomy (BLADE-SWARM reference design instantiates MAIVA with sub-quorum decomposition).
- Distributed sensing and multi-platform track fusion in contested communications.
- Federated learning contributions where a poisoned update must not be admitted on the strength of agreement alone (AUTHREX-DA extends this to data).
- Redundant computing where dissimilar channels must be weighed, not merely counted.

[CAPABILITY DESCRIPTION]

MAIVA aggregates authority-relevant evidence across agents before the tier is issued. Each agent runs local SATA and HMAA evaluation and contributes an attested vote to a Mission Authority Node, where trimmed weighted median aggregation discards the f most extreme contributions in a $3f+1$ roster and weights the remainder by trust; a three-layer CUSUM detector watches for slow drift in any member, and graduated escalation assigns per-level action permissions. The output is an integrity-weighted input to HMAA, not a decision. The faithful model MAIVA_QUORUM located a defect in the shipped quorum check: the fault bound f was derived from the roster size and the roster size was then compared against a threshold built from that bound, so at $n = 1$ a single agent constitutes a quorum. The repair requires the maximum compromised count to be an independent adversary assumption (MaxCompromised) stated beside the requirement, and the record notes that no runtime check can establish that bound: a system cannot verify its own compromise level.

[MAIVA • ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]

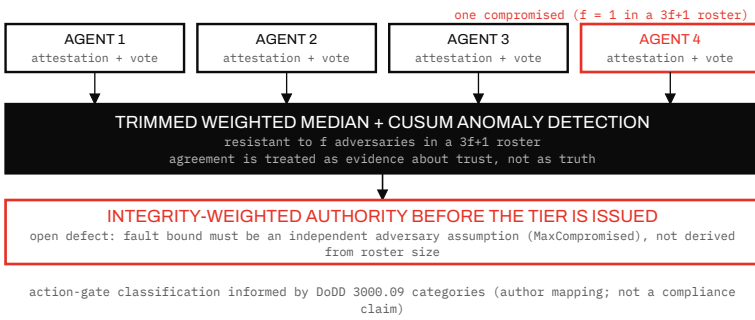


Figure 11. MAIVA technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[HOW IT WORKS]

1. Collect attested contributions from every agent in the roster with their current SATA trust.
2. Apply the fault-tolerance precondition: the number of compromised members must not exceed the declared f (an assumption, not a measurement).
3. Trim the f most extreme contributions; compute the trust-weighted median of the remainder.
4. Run CUSUM anomaly detection per member to flag drift below the trim threshold.
5. Classify the pending action by gate category (author mapping to DoDD 3000.09 categories) to decide whether integrity-weighted authority may proceed or must escalate.
6. Emit the integrity-weighted authority input to HMAA; flag members whose contributions were trimmed or whose CUSUM fired.
7. Record the roster, weights, trimmed members, and the declared f to the ledger.

[TECHNICAL DIFFERENTIATION]

Triple-modular redundancy and N-version voting mask a faulty channel by majority agreement and assume independent failure; common-mode conditions such as spoofing defeat them because every channel agrees and every channel is wrong. MAIVA weights agreement by trust and treats it as evidence about trust rather than as truth. Its own published defect shows the reverse risk: a quorum rule that is self-referential reports on itself, which is why the fault bound must be an external assumption.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Trimmed weighted median tolerates f adversaries in $3f+1$	Design property; exercised in simulation and self-tests
Q_NEVER_FAILS over the shipped quorum check	Held across 12,507,501 distinct states, in the narrowed sense that the self-referential check is satisfied for every roster size 1 to 5,000, which is the defect: at $n = 1$ one agent is a quorum
Repaired quorum (Quorum + MaxCompromised)	Holds at $f = 1$ and $f = 2$ in the repaired model; not shipped in the production engine
R1_ASSUMPTION: compromise bound is external	Stated requirement; no runtime check can establish it
BLADE-SWARM refinement S1 to S5, L1 to L3	Model checked at stated bounds (5 safety invariants, 3 liveness properties)
Strong form: attester compromise inside a met quorum	FAILED in the naive-versus-strong comparison; repaired by Quorum + MaxCompromised

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
Faithful model of the quorum decision	MAIVA_QUORUM.tla FAITHFUL; defect located at named source lines	authrex.systems/formal-coverage
Quorum-check state space	12,507,501 distinct states; check satisfied for $n = 1$ to 5,000 (the defect)	formal-coverage record
Self-tests	37 self-tests in the MIT deposit v5.18	Zenodo 10.5281/zenodo.19015517
Seeded browser simulation	Byzantine minority scenarios; deterministic replay	burakoktenli.com/maiva-simulation
BLADE-SWARM TLA+ refinement	5 safety, 3 liveness at stated bounds; $N = 10, 50, 500$ in simulation	Zenodo 10.5281/zenodo.20351198

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (MIT deposit)
Executable artifact	Yes
Formal model	Yes (one faithful)
Simulation evidence	Yes
Hardware validation	No
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- Documented open defect: the shipped quorum check is self-referential; the repair exists as a specification, not as released code.
- Multi-agent evaluation uses simulated agent populations; no physical multi-agent experiment has been performed.
- The fault bound f is an assumption that no runtime check can establish; correlated compromise beyond f is outside the guarantee.
- MAIVA.tla and MAIVA_REPAIRED.tla are surrogate countermodels: findings about the mechanism class, not about the implemented code.
- Simulation and formal-model evidence only; all data synthetic; no hardware-in-the-loop test, no field validation, no independent replication, no deployment.
- Not independently reproduced or technically verified by another party; second-environment reruns are confirmation by the same producer.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	Aggregation logic implemented and exercised in simulation with a formal model of one critical decision; the located defect means the shipped artifact does not yet satisfy its own stated requirement.
EVIDENCE SUPPORTING	concept and specification; faithful model of the quorum decision; self-tested implementation; simulation demonstration
EVIDENCE REQUIRED FOR NEXT TRL	(1) Ship the Quorum + MaxCompromised repair and re-check with the faithful model. (2) Multi-agent software-in-the-loop with independent agent processes and injected compromise (Phase 4). (3) Independent replication of the 37 self-tests and the TLC configurations.

[INTEGRATION PROFILE]

INPUTS	Attested contributions per agent; per-agent trust; declared fault bound f (MaxCompromised) from mission policy.
OUTPUTS	Integrity-weighted authority input; trimmed-member and CUSUM flags; ledger entry.
INTERFACES	Agent-to-aggregator message channel; HMAA input; BLADE-SWARM ICD-SWARM-001 (design).
DEPLOYMENT	Software module; requires more than one modality or actor to be meaningful.
SECURITY BOUNDARY	Inside: aggregation, CUSUM, roster bookkeeping. Outside: agents, their attestation roots, the channel.
DATA REQUIREMENTS	Synthetic; multi-platform recorded telemetry would enable Phase 2.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Adversarial evaluation of multi-agent agreement by a swarm or multi-platform test organization.
- A government-stated assumption on tolerable compromise (f) for the target mission, which the architecture cannot derive.
- Access to a multi-vehicle software-in-the-loop environment for Phase 4.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE • AUTHREX-05]

FLAME • FORCED LATENCY FOR AUTHORITY-MANAGED ESCALATION

Engineer the time between a proposed irreversible action and its execution so that a human veto is exercisable before the outcome, with abort as the only default on timeout.

MODELED TESTED

Name variants: Patent filing title: Flash War Latency Architecture (preserved as the source document name). Same framework; deposit v5.11.

PRIMARY FUNCTION	Mandatory deliberation window before an irreversible action, sized by a dynamic delay function $D(A, \text{tier}, \text{domain})$ over authority, consequence tier, and domain; five-state circuit breaker with cryptographically signed transitions and a keep-alive heartbeat; timeout defaults to abort, never execute.
MISSION AREAS	Engagement authorization chains; time-critical command and control; any decision with an irreversible effect, including counter-UAS mitigation and automated control writes in operational technology.

EVIDENCE ANCHORS	Zenodo DOI 10.5281/zenodo.19015618; seeded browser simulation; formal-coverage record: FLAME.tla surrogate countermodel with a published strong-form counterexample and a stated requirement.
CURRENT STANDING	Demonstrated in simulation; the deliberation mechanism is specified and simulated; the formal record shows the naive clock property can be satisfied with zero real time elapsed; not independently reproduced.
SELF-ASSESSED TRL	TRL 3. No formal TRL determination has been performed.
INTELLECTUAL PROPERTY	U.S. provisional application 64/005,607, filed March 14, 2026 (receipt 74858888), reported as filed; unexamined; confers no enforceable rights.

[MISSION PROBLEM]

A veto that cannot be exercised before an irreversible outcome is theoretical. Machine decision latency and human deliberation time are usually engineered separately, so a system can be formally supervised and practically unsupervisable. The failure is compounded when the fallback on a lost confirmation is to proceed. The consequence is an irreversible action executed inside a window no human could have used.

[OPERATIONAL RELEVANCE]

- Engagement authorization chains for uncrewed systems, where the confirm window must scale with consequence.
- Counter-UAS mitigation timing, where the decision window is seconds and the effect is irreversible.
- Operational technology: mandatory deliberation before automated load-shedding or control writes (BLADE-INFRA-OT, ICS-GATE).
- Spacecraft operations under light delay, where evidence age and budget gate the action (DELAY then ABORT if not refreshed).

[CAPABILITY DESCRIPTION]

FLAME treats latency as an engineered authority parameter. Before any irreversible action it imposes a deliberation window sized by $D(A, \text{tier}, \text{domain})$: higher consequence buys the human more time by construction. A five-state circuit breaker with cryptographically signed transitions enforces the window; a keep-alive heartbeat detects a stalled supervisor. If the window expires without the required confirmation the outcome is abort, not execute. Where mission policy explicitly authorizes autonomous execution after a bounded hold, that authority must already exist in the pre-authorized envelope; timeout never creates it. The formal record adds a requirement learned from a counterexample: the deliberation clock must advance only by elapsed time, because a clock that can be advanced by bookkeeping alone can be satisfied with zero real time elapsed.

[FLAME • ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]

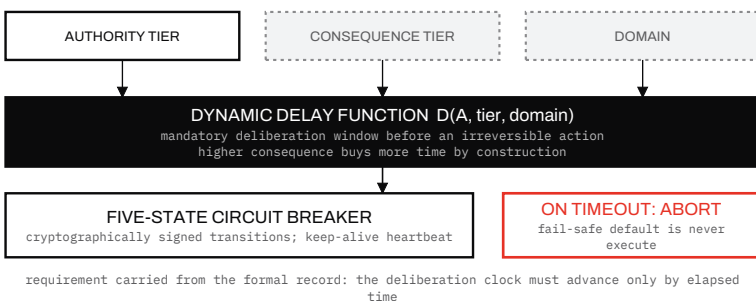


Figure 12. FLAME technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[HOW IT WORKS]

1. Receive the proposed action with its consequence tier, the current HMAA authority tier, and the domain.
2. Compute the deliberation window $D(A, \text{tier}, \text{domain})$ from the versioned policy table.
3. Arm the circuit breaker: transition to the deliberation state with a signed record and start the keep-alive heartbeat.
4. Hold the action (DELAY) while the window runs; surface it to the operator where policy requires confirmation.
5. On confirmation inside the window, transition to permit with a signed record; on timeout or lost heartbeat, transition to abort.
6. Where policy pre-authorizes autonomous execution after a bounded hold, permit only if that authority exists in the frozen envelope.
7. Record every transition, the window, and elapsed time to the ledger.

[TECHNICAL DIFFERENTIATION]

Human-on-the-loop supervision relies on an operator noticing in time; FLAME makes the time budget explicit and consequence-scaled, and it makes DELAY an outcome of the authority computation rather than an accident of workload. Static timers fail in both directions; a consequence-scaled window with abort as the default fails toward safety. The comparison is structural; no field measurement of any timer-based system is claimed.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Timeout defaults to abort, never execute	Design property; exercised in simulation
Window scales with consequence tier	Design property; policy table versioned
Signed circuit-breaker transitions	Implemented in the simulation engine; cryptographic substrate is software, not a hardware root
Deliberation clock advances only by elapsed time	REQUIREMENT stated from a counterexample; not established for the shipped engine
No authority created by timeout	Design rule (policy-conditioned hold); exercised in AUTHREX-ABORT and ENGAGE consoles

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
Seeded browser simulation	Deterministic replay; consequence-scaled windows on synthetic actions	burakoktenli.com/flame-simulation
Strong-form counterexample	Outcome EXECUTE with the deliberation clock satisfied and zero real time elapsed	authrex.systems/formal-coverage
Surrogate countermodel	FLAME.tla: finding about the mechanism class, not the implemented code	formal-coverage record
Zenodo deposit	Specification v5.11, source, simulation data	DOI 10.5281/zenodo.19015618

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (simulation engine)
Executable artifact	Yes
Formal model	Partial (surrogate)
Simulation evidence	Yes
Hardware validation	No
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- The elapsed-time requirement on the deliberation clock is stated but not established for the shipped engine; a bookkeeping clock is a published failure mode.
- Timing is idealized (about 100 Hz loop); jitter, communication delay, and processing bottlenecks on real hardware are not modeled.
- No human-factors measurement exists for whether the computed windows are usable by an operator.
- Simulation and formal-model evidence only; all data synthetic; no hardware-in-the-loop test, no field validation, no independent replication, no deployment.
- Not independently reproduced or technically verified by another party; second-environment reruns are confirmation by the same producer.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	Deliberation mechanism implemented in simulation with a specified breaker and signed transitions; timing has not been characterized on any real control loop or with human operators.
EVIDENCE SUPPORTING	concept and delay function specified; simulation implementation; surrogate formal analysis
EVIDENCE REQUIRED FOR NEXT TRL	(1) Implement and check the elapsed-time clock requirement in a faithful model. (2) Software-in-the-loop timing characterization with a representative controller (Phase 4). (3) Operator-in-the-loop trial of window usability against a government decision timeline.

[INTEGRATION PROFILE]

INPUTS	Proposed action with consequence tier; HMAA tier; domain; mission clock; operator confirmation channel.
OUTPUTS	DELAY, permit, or abort; signed breaker transitions; ledger entry.
INTERFACES	Between the authority tier and execution; operator channel; ledger.
DEPLOYMENT	Software module; browser console for evaluation.
SECURITY BOUNDARY	Inside: window computation, breaker, heartbeat. Outside: clock source, operator channel. A manipulable clock defeats the window.
DATA REQUIREMENTS	None beyond a government decision timeline for Phase 2 review.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Latency-budget review against a government decision timeline for a named mission (Phase 2).
- Access to a software-in-the-loop controller to measure end-to-end timing (Phase 4).
- Human-factors evaluation support to test window usability with operators.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations

and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE · AUTHREX-06]

CARA · CONTROLLED AUTHORITY RECOVERY ARCHITECTURE

Make recovery from authority lockout structured, deterministic, and impossible to short-circuit, because turning the system off is not always a safe or available option.

TESTED

MODELED

Name variants: Patents-page expansion: Control Authority Regulation Architecture; Zenodo record 18917790 title: Cognitive Authority Recovery Architecture. Same framework; original artifact titles preserved. GREP phases on the CARA page: Guard, Reduce, Evaluate, Promote (the BLADE-AV page uses a Restrict and Execute phrasing, not carried).

PRIMARY FUNCTION	Deterministic recovery from lockout through GREP phases (Guard, Reduce, Evaluate, Promote) with a terminal non-compensatory policy gate, so that a single favorable signal cannot restore authority; returns control to SATA, closing the loop.
MISSION AREAS	Spacecraft, uncrewed maritime and undersea systems, and industrial control systems where shutdown is not a safe state; any platform that must recover from A0 in a governed way.
EVIDENCE ANCHORS	Zenodo DOI 10.5281/zenodo.18917790; self-run 10-million-iteration enumeration and 68-state control-flow analysis (not independently reproduced); seeded browser simulation; formal-coverage record: CARA.tla surrogate countermodel with a published strong-form counterexample.
CURRENT STANDING	Demonstrated in simulation; enumeration artifacts are self-run; the formal record shows a replayed re-arm token defeats the naive form; requirement stated (token bound to a nonce); not independently reproduced.
SELF-ASSESSED TRL	TRL 3. No formal TRL determination has been performed.
INTELLECTUAL PROPERTY	U.S. provisional application 64/000,170, filed March 9, 2026 (receipt 74767602), reported as filed; unexamined; confers no enforceable rights.

[MISSION PROBLEM]

Recovery is where most authority architectures are silent. After a lockout, the practical pressure is to restore autonomy as soon as any signal looks good, and a recovery path that can be short-circuited by one favorable reading, or replayed by an attacker who captured a previous re-arm, becomes the widest hole in the system. For platforms where shutdown is itself hazardous (a spacecraft in a bad attitude, an undersea vehicle at depth, a process plant mid-cycle) the consequence of unstructured recovery is either a premature return to full authority or an authority cliff with no way back.

[OPERATIONAL RELEVANCE]

- Spacecraft safe-attitude hold and staged recovery on refreshed evidence (BLADE-SPACE, SPACECYBER reference designs).
- Uncrewed maritime and undersea vehicles where link loss must produce governed fallback rather than silence.
- Industrial control and utility operations where a controlled return to automatic mode must not be triggered by a single sensor.
- Ground robotic platforms recovering from an A0 (revoked) safe-stop.

[CAPABILITY DESCRIPTION]

CARA activates when HMAA authority enters A0 (revoked). Guard is an immediate safe-stop with all autonomous commands disabled; Reduce permits minimal safe behaviours only (return-safe or hover-hold on trusted sensors); Evaluate monitors trust recovery and requires sustained improvement before any restoration; Promote restores authority gradually, with constrained operation until full trust is confirmed. The recovery behaviours run safe-stop, return-safe, hover-hold, crawl-mode, degraded-teleop. The terminal gate is non-compensatory: every required condition must be independently satisfied, and a strong signal on one condition never offsets a failed one. Recovery returns control to SATA rather than directly to the controller, which closes the loop and means promotion is subject to the same trust assessment as initial authority. The formal record adds a requirement: the re-arm token must be bound to a nonce, because a replayed token satisfied the bookkeeping while ground truth recorded that it was never issued.

[CARA • ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]

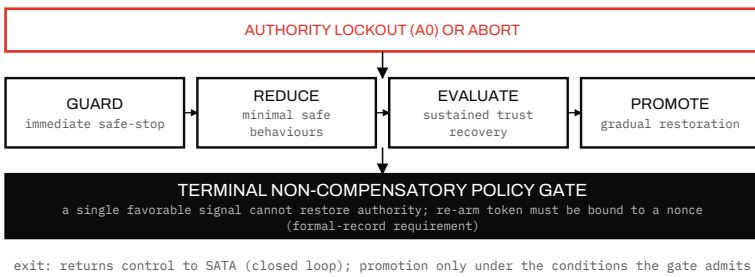


Figure 13. CARA technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[HOW IT WORKS]

1. Enter Guard on A0 (revoked): immediate safe-stop, all autonomous commands disabled, lockout cause recorded.
2. Reduce: minimal safe behaviours only (return-safe or hover-hold on trusted sensors); shed residual authority and pending actions.
3. Evaluate: monitor trust recovery through SATA and re-run any policy-specific conditions; each must pass independently and improvement must be sustained.
4. Apply the terminal non-compensatory gate; if any condition fails, remain in Guard or Reduce.
5. Promote: restore authority one tier at a time through HMAA, subject to hysteresis; a re-arm token bound to a nonce authorizes each promotion.
6. Return control to SATA so that the restored authority is continuously re-assessed.
7. Record each phase transition, the gate evaluation, and the token to the ledger.

[TECHNICAL DIFFERENTIATION]

A kill switch terminates and leaves recovery to procedure; conventional fail-safe designs specify the fall but not the climb. CARA makes recovery a first-class system function with its own gate, so a single favorable signal cannot restore authority and a replayed authorization cannot re-arm the system. The differentiation is structural; the enumeration evidence is self-run and no comparative measurement against another recovery scheme is claimed.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Non-compensatory terminal gate	Design property; exercised by the self-run enumeration and simulation
Deterministic phase sequence	Design property; 68-state control-flow analysis (self-run)
No promotion on a single favorable signal	Design property; exercised in simulation
Re-arm token bound to a nonce	REQUIREMENT from the published counterexample; not established for the shipped engine
Return to SATA (closed loop)	Implemented in the seeded consoles

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
Enumeration of the recovery state space	10,000,000 iterations, self-run; 68-state control-flow enumeration	Zenodo 10.5281/zenodo.18917790 (site claim; not independently reproduced)
Seeded browser simulation	GREP path exercised under injected fault clearance (scenario S7 in the evaluation protocol)	burakoktenli.com/cara-simulation
Strong-form counterexample	Replayed re-arm token: bookkeeping records a token present, ground truth records it never issued	authrex.systems/formal-coverage
Surrogate countermodel	CARA.tla: finding about the mechanism class	formal-coverage record

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (simulation engine)
Executable artifact	Yes
Formal model	Partial (surrogate)
Simulation evidence	Yes
Hardware validation	No
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- The 10-million-iteration enumeration is a self-run site claim; no independent reproduction exists.
- The nonce-bound re-arm token is a stated requirement, not a shipped implementation.
- Recovery timing (mean recovery time metric) has only been measured in idealized simulation.
- Simulation and formal-model evidence only; all data synthetic; no hardware-in-the-loop test, no field validation, no independent replication, no deployment.
- Not independently reproduced or technically verified by another party; second-environment reruns are confirmation by the same producer.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	Recovery logic implemented and exercised in simulation with self-run enumeration; no laboratory validation on a representative platform.
EVIDENCE SUPPORTING	concept and phases specified; simulation implementation; self-run enumeration
EVIDENCE REQUIRED FOR NEXT TRL	(1) Independent reproduction of the enumeration artifacts. (2) Implementation and check of the nonce-bound token requirement. (3) Recovery-path trial on a software-in-the-loop platform with injected lockout and fault clearance (Phase 4).

[INTEGRATION PROFILE]

INPUTS	ABORT or lockout event with cause; policy conditions for promotion; SATA re-assessment; re-arm authorization.
OUTPUTS	Phase state; gate verdict; promotion requests to HMAA; ledger entry.
INTERFACES	HMAA (A0 entry, promotion), SATA (re-assessment), platform safe-state interface (external interlock).
DEPLOYMENT	Software module; browser console for evaluation.
SECURITY BOUNDARY	Inside: phase logic, gate. Outside: safe-state actuation, token issuer. A replayable token defeats the gate until the nonce requirement ships.
DATA REQUIREMENTS	Synthetic; a government-defined safe-state and recovery policy for a named platform would enable Phase 2.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Recovery-path review by a platform safety authority (spacecraft, maritime, or ICS).
- A government-defined recovery policy (conditions for promotion) for a named platform class.
- Independent reproduction of the enumeration artifacts from the deposit.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE • AUTHREX-07]

ERAM • ESCALATION RISK ASSESSMENT MODEL

Quantify decision-time compression and cross-domain escalation cascades in AI-enabled command and control so that escalation risk constrains authority across the pipeline rather than only being reported.

MODELED TESTED

Name variants: No expansion variants. ERAM is a working paper (SSRN 6176802), not a peer-reviewed publication.

PRIMARY FUNCTION	Escalation-risk quantification and decision-time-compression modeling for AI-enabled C2; feeds the consequence tier consumed by FLAME and HMAA; implemented as the ERAM Strategic Command simulation (6 scenarios x 100 = 600 Monte Carlo runs) with Merkle provenance chains and a sealed decision ledger.
MISSION AREAS	AI-enabled command-and-control experimentation; strategic-level escalation studies; wargame support; joint all-domain C2 modeling.
EVIDENCE ANCHORS	SSRN working paper 6176802; ERAM Strategic Command console (SIM-01); formal-coverage record: ERAM.tla surrogate countermodel with a published strong-form counterexample (SilentDrop).

CURRENT STANDING	Working paper and seeded simulation; model-based with documented parameter assumptions; not validated against operational data; not peer reviewed; not independently reproduced.
SELF-ASSESSED TRL	TRL 2 to 3 (analytical model with a simulation demonstrator). No formal TRL determination has been performed.
INTELLECTUAL PROPERTY	U.S. provisional application 64/110,225, filed July 13, 2026 (receipt 78660113), reported as filed; unexamined; confers no enforceable rights.

[MISSION PROBLEM]

Local autonomous actions can cascade across domain boundaries and compress decision time at the strategic level: a defensive action in one domain becomes a cue in another, and each automated step removes time from the next human decision. Governance layers that look only at the local action miss the cascade. The consequence is escalation by diffusion, where no single decision was unauthorized but the aggregate outran every deliberate decision cycle.

[OPERATIONAL RELEVANCE]

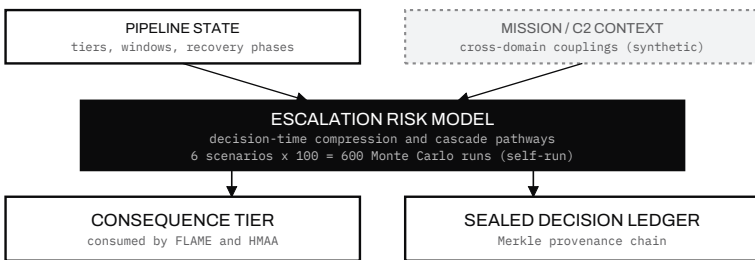
- AI-enabled command and control where recommendations must carry a consequence tier that reflects cross-domain coupling.
- Strategic-level escalation studies and wargame support with reproducible, seeded runs.
- Joint all-domain C2 experimentation environments (APEX JADC2 Theater and MDO Digital Twin simulations exercise the coupling).
- Audit reconstruction: the sealed ledger conventions used across the program originate here.

[CAPABILITY DESCRIPTION]

ERAM is the cross-cutting monitor of the pipeline. It consumes pipeline state (tiers, windows, recovery phases) and mission context, and models how local actions propagate across domains and compress decision time. Its output is a consequence tier that FLAME uses to size the deliberation window and HMAA uses in tier computation, so escalation risk constrains authority instead of only reporting it. The ERAM Strategic Command simulation runs six scenarios with 100 Monte Carlo runs each under a seeded generator, with formal properties checked on the model and a Merkle provenance chain over the run record. The formal record adds a requirement: ledger integrity needs an independent decision count to reconcile against, because a valid hash chain proves only that what was recorded was not altered, not that nothing was silently dropped.

[ERAM • ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]



formal record: a hash chain proves what was recorded was not altered, not what was never offered

Figure 14. ERAM technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[HOW IT WORKS]

1. Ingest pipeline state from every stage and the mission or C2 context (synthetic in the demonstrator).
2. Model cross-domain couplings and compute decision-time compression along each escalation pathway.
3. Derive the consequence tier for the pending action class and publish it to FLAME and HMAA.
4. Run the seeded Monte Carlo campaign for the scenario under evaluation and compute the escalation-risk distribution.
5. Seal the run record with a Merkle chain and reconcile the ledger against an independent decision count.
6. Report the distribution and the constraint applied; the outcome itself is issued by HMAA.

[TECHNICAL DIFFERENTIATION]

Escalation-risk analysis is usually a study product delivered after the fact; ERAM places it inside the runtime pipeline as a constraint on authority. Against anomaly detection it converts a signal into a consequence tier that changes how much time and authority the next action receives. The model is a working paper and the differentiation is architectural; no validation against operational escalation data exists.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Consequence tier constrains FLAME and HMAA	Design property; implemented in the integrated consoles
Seeded, reproducible Monte Carlo campaign	Implemented; bit-exact within the packaged build; not independently reproduced
Merkle provenance over the run record	Implemented in the console
Ledger completeness (nothing silently dropped)	NOT established: SilentDrop counterexample; requirement for an independent decision count recorded
Formal properties on the model	Checked in the ERAM Strategic Command console at stated bounds; ERAM.tla is a surrogate

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
Working paper	SSRN 6176802 (working paper, not peer reviewed)	SSRN
ERAM Strategic Command simulation	6 scenarios; 600 Monte Carlo runs; formal properties; Merkle chain; seeded PRNG	authrex.systems (SIM-01)
Strong-form counterexample	SilentDrop fires, ledger is empty, shadow is non-empty, hash chain valid	authrex.systems/formal-coverage
Earlier wording "validated across all 600 Monte Carlo runs"	Removed from the live site; results are demonstrated in simulation, not validated	Correction log CL-05

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (console)
Executable artifact	Yes
Formal model	Partial (surrogate)
Simulation evidence	Yes
Hardware validation	N/A
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- Model-based with documented parameter assumptions; not validated against operational or historical escalation data.
- Working paper status: not peer reviewed.
- Ledger completeness cannot be established from the chain alone; the independent decision count is a requirement, not a shipped mechanism.
- Six scenarios are author-defined; no government scenario has been run.
- Not independently reproduced or technically verified by another party; second-environment reruns are confirmation by the same producer.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 2 to 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	Analytical model with a reproducible simulation demonstrator; the critical function has not been validated against any real data, which limits the assessment to concept and analytical proof of concept.
EVIDENCE SUPPORTING	concept formulated (working paper); simulation demonstrator
EVIDENCE REQUIRED FOR NEXT TRL	(1) Peer-reviewed submission with sensitivity analysis. (2) Government-defined C2 scenarios on the console (Phase 2). (3) Comparison against a sponsor-provided escalation case study or wargame record.

[INTEGRATION PROFILE]

INPUTS	Pipeline state; mission and C2 context; scenario definitions; seed.
OUTPUTS	Consequence tier; escalation-risk distribution; sealed run record.
INTERFACES	FLAME and HMAA inputs; ledger; console export.
DEPLOYMENT	Standalone browser console (client-side, no network); analytical tool rather than an operational decision aid.
SECURITY BOUNDARY	Inside: model, chain. Outside: context sources. The chain does not prove completeness.
DATA REQUIREMENTS	Synthetic; wargame or exercise records under sponsor control would enable validation.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Working-paper review by a C2 or escalation-dynamics research organization.
- A government wargame or exercise dataset (releasable) for model validation.
- Integration into a sponsor's C2 experimentation environment as an analytical monitor, not a decision aid.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE · SW-07]

GOVERNANCE KERNEL · AUTHREX GOVERNANCE KERNEL (QA10)

Give LLM agents authority tiers instead of admin keys: gate every tool call by tier and human approval before it executes, and leave signed, tamper-evident evidence of what was allowed.

IMPLEMENTED	TESTED	NOT YET VALIDATED
PRIMARY FUNCTION	Runtime decision and enforcement engine for agentic AI: authority-tier enforcement (T3 to T0) at an agent-to-proxy-to-tool boundary, human approval gating with role-based access control, signed policy loading, checkpointed hash-chain audit evidence, reviewer packet export, loopback latency measurement, sustained concurrent-load testing, and ledger tamper detection.	
MISSION AREAS	Agentic AI acting on networks and infrastructure; MCP-based tool governance; software assurance for autonomous cyber and IT operations.	
EVIDENCE ANCHORS	Container test suite: 83 passed, 1 skipped (environment-dependent) across 53 test functions, python:3.11-slim, package v1.0-QA10, July 2026; Python 3.11 clean install pass; Docker Desktop build, demo and pytest pass; ledger tamper detection across clean, forged and truncated records; QA10 final artifact after eleven AI-assisted review rounds.	
CURRENT STANDING	Implemented in software and internally tested; a synthetic-data minimum viable product prepared for structured external technical review. Not released (public GitHub release pending counsel review; source access invited-review only). Not independently validated.	
SELF-ASSESSED TRL	TRL 3 to 4 (software prototype tested in a development environment with a test suite; no external environment). No formal TRL determination has been performed.	
INTELLECTUAL PROPERTY	License under review; evaluation or redistribution rights are not granted unless explicitly stated in the artifact package. Covered in part by the program's provisional applications; no per-artifact patent claim is made.	

[MISSION PROBLEM]

An LLM agent that can call tools holds, in practice, whatever privilege its credentials hold. Tool misuse and credential exposure are named risk categories in the CISA, NSA and Five Eyes guidance on agentic AI services (May 2026). A model guardrail constrains outputs; it does not decide whether an action is within the agent's assigned authority at this moment, and it leaves no independent evidence of what was permitted. The consequence is a privileged action executed by an agent that was never authorized to take it, with no record a reviewer can trust.

[OPERATIONAL RELEVANCE]

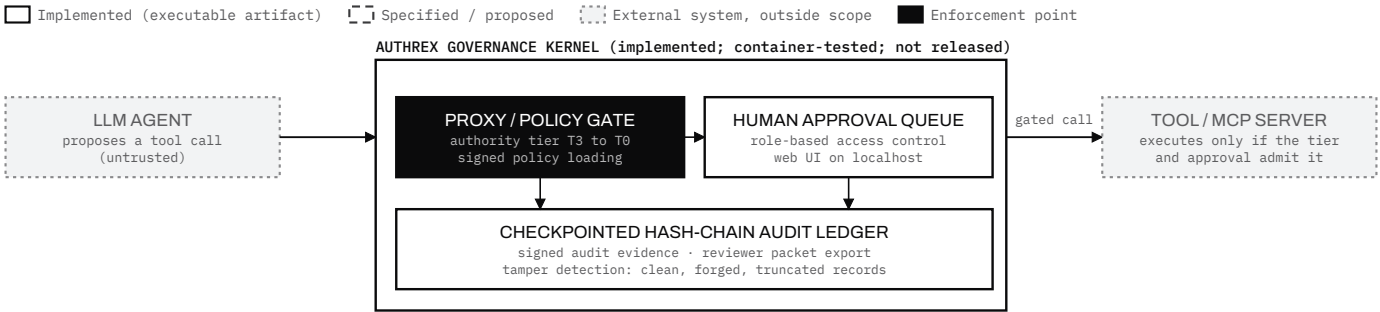
- Agentic AI operating on networks, infrastructure, and developer tooling (millisecond decision windows, tool output and retrieved content as degradable evidence).
- Software assurance and DevSecOps pipelines where an agent proposes changes to live systems.
- Zero-trust relevance: least privilege enforced per action rather than per session; auditability and provenance of every gated call.
- Cyber defense automation (AUTHREX-AGENT-CYBER profile) where an autonomous system may patch a live controller only under governed authority.

[CAPABILITY DESCRIPTION]

The Kernel is the executable enforcement engine of the program. An agent's tool call passes through a proxy where the policy gate checks the agent's authority tier against the tool's required tier; calls above the tier are routed to a human approval queue with role-based access control; approved or admitted calls are forwarded to the tool or MCP server. Every decision is written to a checkpointed hash-chain ledger with signed evidence, and a reviewer packet can be exported for external review. The demo suite exercises real HTTP agent-to-proxy-to-tool governance and real MCP SDK governance with the approval queue, a localhost web UI, loopback latency measurement, and sustained concurrent load. The Kernel calculates and enforces authority at its enforcement point; its guarantees depend on placement, platform integrity, protected keys, and the integrity of the downstream actuation path.

[GOVERNANCE KERNEL · ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]



Evidence: 83 passed, 1 skipped (environment-dependent) across 53 test functions in a python:3.11-slim container, package v1.0-QA10, July 2026; zero unauthorized privileged executions in synthetic scenarios. Guarantees depend on placement, platform integrity, protected keys, and the integrity of the downstream actuation path. Software-only containment can be bypassed without the hardware root of trust (BLADE-AGENT-HSM, not built).

Figure 15. Governance Kernel enforcement path. Encoding: solid = implemented; gray = external; filled = enforcement point.

[HOW IT WORKS]

1. Load signed policy (tool-to-tier map, roles, approval rules) at start; refuse unsigned policy.
2. Receive a tool call from the agent at the proxy; identify the agent and its current authority tier.
3. Compare the call against the required tier; admit, hold for human approval, or reject.
4. For held calls, present to the approval queue; an authorized role approves or denies under RBAC.
5. Forward admitted calls to the tool or MCP server; return the result to the agent.
6. Write the decision, approver, and result digest to the checkpointed hash-chain ledger with a signature.
7. On demand, verify the ledger (detecting forged records and tail truncation) and export the reviewer packet.

[TECHNICAL DIFFERENTIATION]

Role-based and attribute-based access control grant permission by identity or attribute, not by the current action and evidence; model guardrails constrain outputs regardless of which action they enable. The Kernel adds an evidence-conditioned, time-varying authority tier beneath the role and enforces it at the tool boundary with signed evidence. It composes with MCP tool servers rather than replacing them.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Zero unauthorized privileged executions in synthetic scenarios	Tested (self-generated suite passes its full set); not independently validated
Ledger tamper detection: clean, forged, truncated	Tested in the container suite
Human approval gating with RBAC	Implemented and demonstrated (MVP demo)
Signed policy loading	Implemented
Latency and concurrent load	Measured on loopback in the test environment; not benchmarked on any target
Software-only containment	Bypassable without a hardware root of trust; BLADE-AGENT-HSM (not built) is the designed companion

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
Container test suite	83 passed, 1 skipped (environment-dependent); 53 test functions; package v1.0-QA10; archive hash recorded in the internal QA log	authrex.systems/kernel (site claim pending external reproduction)
Install and build	Python 3.11 clean-environment install PASS; Docker Desktop (Mac) build PASS; containerized demo end to end as unprivileged user PASS	kernel.html
Alpha and MVP demos	HTTP agent-to-proxy-to-tool gating; MCP SDK governance with approval queue, RBAC web UI, signed evidence, packet export, latency, load	kernel.html
Review package contents	authrex_kernel/, tests/ (53 functions), harness/, docs/EXTERNAL-REVIEWE R-PACKET.md, demo_alpha.py, demo_mvp.py, Dockerfile, install.sh, VALIDATION-REPORT.md, CHANGELOG.md	by-request review package

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes
Executable artifact	Yes (container)
Formal model	No (not yet performed)
Simulation evidence	Yes (synthetic demos)
Hardware validation	No
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- Not released; source access is invited-review only while counsel review is pending; no independent security assessment; scalability not benchmarked.
- All results are from synthetic scenarios on self-generated test suites; a self-generated suite cannot establish that the policy semantics match any real mission policy.
- No formal model of the Kernel exists; the TLA+ work covers the frameworks and consoles, not this codebase.
- Software-only containment: a compromised host, leaked key, or bypassed proxy defeats enforcement.
- No SBOM, no formal vulnerability scan, no penetration test.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 3 to 4 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	A working software prototype with an executable test suite in a development environment satisfies the software proof-of-concept level and approaches laboratory validation; no validation outside the author's environment and no representative workload has been run, so TRL 4 is not claimed outright.
EVIDENCE SUPPORTING	software implementation; container test suite; demonstrations on loopback
EVIDENCE REQUIRED FOR NEXT TRL	(1) External reproduction of the container test suite from the review package by a government or FFRDC team. (2) Sandbox evaluation with government-defined tool calls and approval policies (Phase 2). (3) Independent security assessment and SBOM before any release.

[INTEGRATION PROFILE]

INPUTS	Agent tool-call requests; signed policy; approver actions; agent identity.
OUTPUTS	Admit / hold / reject decisions; forwarded calls; ledger; reviewer packet.
INTERFACES	HTTP proxy; MCP SDK; localhost web UI; file export.
DEPLOYMENT	Containerized (python:3.11-slim); runs as an unprivileged user; localhost only in the current package.
COMPUTE REQUIREMENTS	Not characterized beyond loopback latency in the test environment.
DEPENDENCIES	Python 3.11; Docker; MCP SDK; no external network required for the demo.
SECURITY BOUNDARY	Inside: proxy gate, approval queue, ledger, policy signature check. Outside: agent, tools, host OS, key storage. Key protection is the platform's responsibility.
DATA REQUIREMENTS	Synthetic; government-defined tool inventories and approval policies (unclassified) suffice for Phase 2.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Structured external technical review of the review package (the program's stated next evidence priority).
- Government-defined tool calls and approval policies for a sandbox evaluation.
- Independent security assessment (code review, dependency scan, penetration test) by a cyber test organization.
- Guidance on an accreditation-relevant evidence format (RMF artifacts) for a future release.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE • SIM-13]

AUTHREX-ABORT • EVIDENCE CONDITIONED AUTHORITY LAPSE CONSOLE (TERMINAL AUTHORITY PROFILE)

When the evidence a mission depended on fails, authority lapses: the system contracts the envelope a human authored, latches irreversible actions down, and falls back only to contingencies authorized in advance. It never acquires new authority.

MODELED	TESTED	IMPLEMENTED
PRIMARY FUNCTION	Deterministic synthetic console modeling evidence-conditioned authority lapse: a human-authored envelope frozen before the run; five evidence predicates (navigation and geographic containment as hard vetoes; sensor integrity and communications as degradation; model confidence as informational only) gate the available actions; irreversible actions, once lapsed, are latched down for the run; requests to the platform go through an abstract adapter with acknowledgments correlated back; every transition hash-chained and replayable.	
MISSION AREAS	Uncrewed air systems losing navigation or containment evidence; any autonomous platform whose fallback must not add capability at the moment it is least trustworthy; the program's contrast to kill-switch and flight-termination framings.	
EVIDENCE ANCHORS	Two TLA+ modules model checked at stated bounds: AuthorityLapse v3 across 21 configurations, InterfaceRecord v1 across 35 configurations, every property in isolation; Stage 1B reachable set 31,341 states and 130,869 transitions exported from TLC; Stage 1 availability relation 972 of 972 projected states (AvailRaw only, at the model bound); 105 implementation tests; 34 release checks; 1,000,000-run soak in a separate execution environment with zero failures; mutation campaign (7 valid mutants, 4 detected).	
CURRENT STANDING	Development research artifact, Build 0.6; formally model checked at stated bounds and implementation tested; not certified, accredited, integrated, operationally validated or independently verified; open gates published.	
SELF-ASSESSED TRL	TRL 3. No formal TRL determination has been performed.	
SCOPE STATEMENT	Governance and assurance only. No effector, no destructive mechanism, no target selection, no target tracking for engagement, no firing solution, no weapon release logic, no payload interface, no real coordinates, no real target data. In the page's own words, AUTHREX is not a kill switch and it does not detonate anything; no destructive action exists anywhere in the outcome vocabulary.	

[MISSION PROBLEM]

When an autonomous system loses the evidence its mission depended on, the common answer is that something takes over: a remote operator, a fallback autopilot, a terminal action. Each adds a new capability at the exact moment the system is least trustworthy. A fallback that assumes valid navigation to divert away from people fails precisely when navigation is what failed. The consequence is a platform executing a stronger irreversible authority on weaker evidence.

[OPERATIONAL RELEVANCE]

- Uncrewed air systems in GNSS-denied airspace: navigation is a hard veto, so loss of position prohibits dependent actions at once.
- Counter-UAS and launched-effects concepts where the contingency vocabulary must be fixed before launch and cannot be widened in flight.
- Test and evaluation of AI-enabled systems: a candidate evidence-lapse test pattern (predicate degradation, latch, replay) that a T&E; organization can run without hardware.
- Policy alignment where applicable: the artifact is not a weapon system; DoDD 3000.09 becomes relevant only if a governance layer of this kind is incorporated into a system within the directive's scope, and the correct language is alignment, not compliance.

[CAPABILITY DESCRIPTION]

A human authors a policy envelope before the run; it is frozen and can never be widened by the system. Each action depends on an explicitly identified evidence set. Navigation and geographic containment are hard vetoes: failure prohibits the dependent action immediately, with no score that compensates for not knowing where the system is. Sensor integrity and communications contract the set of available actions progressively. Model confidence is reported and recorded and never permitted to gate an action. When an irreversible action becomes unavailable it is latched down and cannot return during that run even if the evidence recovers. Only contingencies authorized in advance (inhibit, divert, handoff) may then be requested, through an abstract interface that produces no physical effect; an acknowledgment means the request was accepted in the model, not that anything happened in the world. One model finding changed the design: communications recovery violated irreversible non-reacquisition through a shared lease term, so a predicate is recoverable only if no irreversible action depends on it directly or through any shared term.

[AUTHREX-ABORT • ARCHITECTURE, ASSURANCE, EVIDENCE]
[TECHNICAL ARCHITECTURE]

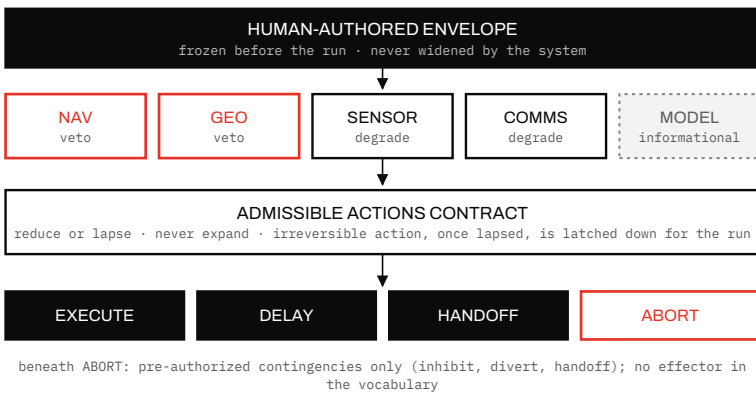


Figure 16. AUTHREX-ABORT technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[TECHNICAL DIFFERENTIATION]

Range-safety flight termination and kill switches add a terminal capability; this profile removes capability. ASTM F3269 run-time assurance is cited as design lineage (a monitor that constrains a complex function), never as compliance; the difference here is that the constraint is evidence-conditioned and latching, and that the contingency vocabulary is closed before launch. Results are model-relative and bounded, never proven for all executions.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
C1 envelope confinement: authority stays inside the human envelope	Formally checked at stated bounds and implementation tested
C2 adverse-step non-expansion: adverse evidence cannot expand authority	Formally checked and implementation tested
C3 irreversible non-reacquisition: lapsed irreversible authority cannot return	Formally checked and implementation tested
C4 execution conformity and C5 no unbacked irreversible execution	Formally checked (AuthorityLapse v3); implementation conformance to InterfaceRecord v1 is an open gate
Lapsed-set monotonicity and latch capture	Formally checked (AuthorityLapse v3)
Fifteen interface and record obligations plus type invariance	Formally checked (InterfaceRecord v1, 35 configurations)
Audit traceability	WITHDRAWN as a checked claim of the core (the core has no ledger); carried by Stage 1B InterfaceRecord instead
Behavioural negative controls	B12 closed; B14 and B2 open (by the property-specific standard)

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
AuthorityLapse.tla v3	21 configurations, every property checked in isolation; sha256 ea6c4fff...fbce2993	authrex.systems/authrex-abort; Build 0.6 package
InterfaceRecord.tla v1	35 configurations; 15 obligations plus type invariance; sha256 a3f5e2eb...78739c27	Build 0.6 package
Stage 1B reachable set (TLC export)	31,341 states; 130,869 transitions; state-set and transition-set hashes published	Build 0.6 package
Stage 1 availability relation	972 of 972 projected states, AvailRaw only, at the model bound (a projection, not the complete reachable graph)	Build 0.6 package
Implementation suite	105 tests; 34 release checks; engine sha256 ce3af3dd...509ebd7e	Build 0.6 package
Soak	1,000,000 runs, zero failures, sixteen-way parallelism, separate execution environment (same producer; not independent verification; not representative platform testing)	Build 0.6 package
Mutation campaign	7 valid mutants, 4 detected, 3 surviving, 1 invalid; mechanical score 0.57 is not the result; B12 closed, B14 and B2 open	authrex.systems/authrex-abort

[HOW IT WORKS]

1. Freeze the human-authored envelope and its policy digest before the run begins.
2. Evaluate the five evidence predicates each step; classify each as valid, degraded, or failed.
3. Compute the admissible action set: hard-veto failure removes dependent actions at once; degradation narrows progressively; model confidence never gates.
4. When an irreversible action leaves the admissible set, latch it down for the remainder of the run.
5. Issue the outcome (EXECUTE, DELAY, HANDOFF, ABORT); beneath ABORT, request only pre-authorized contingencies through the abstract adapter.
6. Correlate platform acknowledgments back to the issuing request; model loss, delay, duplication, stale epochs, and controller health (monotone downward within a run).
7. Append each transition to the hash-chained record with reason, scenario digest, and previous hash; a sealed run replays from its inputs and any tampering inside the integrity scope is detected.

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (standalone HTML console)
Executable artifact	Yes
Formal model	Yes (2 modules)
Simulation evidence	Yes
Hardware validation	No
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- Model checking establishes the properties within the declared models at stated bounds; it does not by itself establish that the implementation refines the model (frozen engine to InterfaceRecord v1 conformance is an open gate).
- Open gates, published: B12 issue-time correlation; B14 and B2 behavioural negative controls; engine-authored oracle applicability; replay and formal-bounds qualification; scenario program at 23 of a planned 42; generated-sequence campaign; browser and visual regression gates; requirements traceability; hazard package; cybersecurity package; schema-valid SBOM; interface control document; performance characterization; hardware-in-the-loop; external verification.
- The platform side of the interface is an abstract adapter; SafeState is deferred and not implemented; HOLD belongs to a different simulation and is out of scope here.
- No development package can self-issue NASA software classification concurrence or IV&V; determination, a DoD ATO or cATO decision, DoD residual-risk acceptance, DT or OT acceptance, legal review, DoDD 3000.09 approval, government certification, or third-party independent verification; all are recorded as not performed.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	Analytical (formal) and experimental (implementation-tested, seeded) proof of the critical function in a development environment; no laboratory validation against a representative platform interface.
EVIDENCE SUPPORTING	formal models at stated bounds; implementation tests and release checks; deterministic console with replay; soak in a second environment
EVIDENCE REQUIRED FOR NEXT TRL	(1) Close the implementation-to-InterfaceRecord conformance gate and the B14 and B2 negative controls. (2) Complete the scenario program (42) and the hazard, cybersecurity, SBOM, and ICD packages. (3) Hardware-in-the-loop, which the page lists among the open gates; it requires a representative platform safety-interface specification to replace the abstract adapter, and no bench demonstrator is stated on the record.

[INTEGRATION PROFILE]

INPUTS	Frozen envelope and policy digest; evidence predicate states; scenario inputs; platform acknowledgments.
OUTPUTS	Admissible set; outcome; contingency requests (inhibit, divert, handoff); hash-chained record; replay verdict.
INTERFACES	Abstract platform adapter (requests carry command, authority epoch, policy digest, evidence-snapshot identity; responses correlated).
DEPLOYMENT	Single standalone HTML file; no dependencies, no network requests, no browser storage; evaluation package verifies from a clean extraction with one gate command.
SECURITY BOUNDARY	Inside: envelope, predicates, latch, record. Outside: the platform safety interface and any effector (none exists here).
DATA REQUIREMENTS	Synthetic only; no real coordinates or target data anywhere in the artifact.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Independent re-execution of the 21 and 35 TLC configurations and the 105-test suite from the Build 0.6 package.
- A government T&E; organization to adopt the evidence-lapse pattern as a test case and run its own scenarios on the console.
- Access to a representative platform safety interface specification (unclassified) to replace the abstract adapter for the bench demonstrator.
- Hazard-analysis review support (MIL-STD-882E process) for the hazard package that is an open gate.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

[CAPABILITY PROFILE • SIM-14]

AUTHREX-DA • DATA ADMISSIBILITY LAYER (ENGINEERING PREVIEW)

Extend evidence-to-authority mapping from actions to data: every record or federated model contribution offered to a system receives a verdict on its admissibility, non-compensatory and fail-closed, before it is allowed to influence anything.

IMPLEMENTED	TESTED	PROPOSED	NOT YET VALIDATED
PRIMARY FUNCTION	Fail-closed admissibility gate: every record offered receives ADMIT, ADMIT_FLAG, QUARANTINE, or REJECT with a reason code and an evidence reference, written to an append-only record; the layer never repairs, imputes, or reconstructs a payload; deterministic under AUTHREX-DA-CANON-1 numeric canonicalization.		
MISSION AREAS	Federated learning and model-update admission; sensor and data-record integrity for AI-enabled C2 and ISR pipelines; data provenance and poisoning resistance.		
EVIDENCE ANCHORS	101 of 101 local tests; 12 ordered release gates over the whole shipped tree; 14 requirements at 12 TEST_PASS_LOCAL, 1 PARTIAL_LOCAL, 1 BLOCKED; archived canonical ledger root, engine-bound, seed 20260828, 120 rounds, 1,584 entries; 17-artifact assurance package under a three-part commitment; 19 audit passes (5 internal, 14 external), which is not independent reproduction; retained negative results NEG-1 to NEG-5.		
CURRENT STANDING	Proposed engineering preview, release 1.8, core engine 1.0; locally tested and gate-checked; not certified, accredited, authorized, operationally validated, or independently reproduced; WP-0 (quorum repair) blocked; no property DA-P1 to DA-P9 model checked.		
SELF-ASSESSED TRL	TRL 2 to 3. No formal TRL determination has been performed. The self-assessed rubric of 51 of 100 is expressly not a DoD score, a TRL, a certification, an authorization, or an endorsement.		
SCOPE STATEMENT	Governance and assurance only: no detection, identification, classification, selection, prioritization, or optimization of any effect.		

[MISSION PROBLEM]

AI-enabled systems ingest data records and model contributions whose provenance and integrity cannot be assumed, and most pipelines treat admission as a data-quality question answered by cleaning or imputation. A poisoned contribution admitted on the strength of agreement, or a record silently repaired into plausibility, changes what the system believes without any authority decision having been made. The consequence is a model or picture that was never authorized to change.

[OPERATIONAL RELEVANCE]

- Federated model contributions across platforms or partners, where a majority capture by contaminated contributors is the named open failure (NEG-3).
- Data records feeding a common operating picture or ISR product.
- Software supply chain and provenance: the same admissibility logic applies to artifacts as to records.
- Zero-trust data principles: never admit by default; verify before influence.

[CAPABILITY DESCRIPTION]

AUTHREX-DA is a variant configuration of the governance layer. Each offered record passes an admissibility gate that applies the evidence-to-authority mapping to data rather than to actions: required evidence conditions are evaluated independently, a strong condition never compensates for a failed one, and the default on any failure is to withhold admission. The four verdicts carry a reason code and an evidence reference and are written to an append-only hash chain with an external commitment check. Numeric canonicalization (AUTHREX-DA-CANON-1) makes the engine deterministic across platforms so that a reviewer can reproduce the archived ledger root from the seed. The layer records what it refused and why; it never alters a payload.

[AUTHREX-DA • ARCHITECTURE, ASSURANCE, EVIDENCE]

[TECHNICAL ARCHITECTURE]

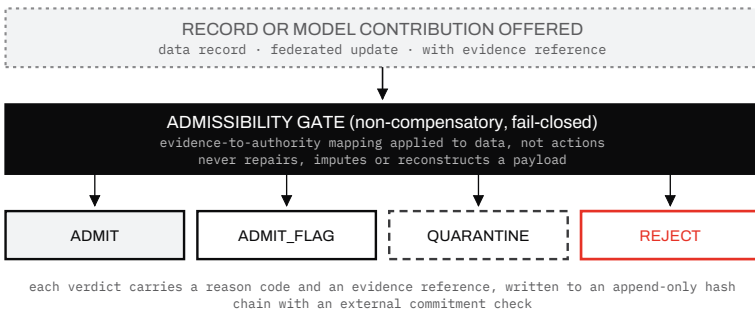


Figure 17. AUTHREX-DA technical architecture. Encoding: solid = implemented; dashed = specified; gray = external; filled = enforcement point.

[TECHNICAL DIFFERENTIATION]

Data-quality tooling cleans and imputes; DA refuses and records. Anomaly detection flags; DA converts the flag into an admissibility verdict that downstream consumers must honor. The architectural difference is that admission becomes an authority decision with an audit trail, not a preprocessing step.

[ASSURANCE MECHANISMS: DESIGN PROPERTY VERSUS ESTABLISHED PROPERTY]

Mechanism	Status
Fail-closed default on evaluation failure	Design property; locally tested
Non-compensatory admission	Design property; locally tested
Never repairs, imputes, or reconstructs a payload	Design property; locally tested
Deterministic canonical ledger root	Reproduced on one platform (engine-bound root archived); second-platform root pending
Properties DA-P1 to DA-P9	NOT model checked (WP-0 blocked; formal work surrogate only)
Resistance to majority capture	NOT established: NEG-3 majority-capture failure retained as an open requirement

[TECHNICAL EVIDENCE]

Evidence	Result	Source / artifact
Local test suite	101 of 101 tests; 12 ordered release gates over the whole shipped tree	SIM-14 release 1.8 package
Requirements status	14 requirements: 12 TEST_PASS_LOCAL, 1 PARTIAL_LOCAL, 1 BLOCKED	release 1.8 package
Canonical ledger root	seed 20260828, 120 rounds, 1,584 entries, engine-bound; single platform	release 1.8 package
Audit passes	19 on file (5 internal, 14 external); external audit review is not independent reproduction	release 1.8 package
Negative results	NEG-1 to NEG-5 retained, including NEG-3 majority-capture failure	release 1.8 package

[EVIDENCE SNAPSHOT]

EVIDENCE SNAPSHOT	
Implemented	Yes (preview engine)
Executable artifact	Yes
Formal model	No (surrogate only)
Simulation evidence	Yes (single seed)
Hardware validation	N/A
Field validation	No
Independent validation	No
Government evaluation	No
Operational deployment	No

[LIMITATIONS]

- WP-0 (quorum repair, closing the roster-derived fault-bound defect inherited from MAIVA) is blocked; it gates every later work package.
- One seed; no preregistered campaign; no confidence intervals; single-platform canonical root.
- Majority-capture failure (NEG-3) is an open requirement: a contaminated majority of contributors can carry admission.
- 38 numbered open items are published with the release; the engineering preview is not a release candidate.

[TECHNOLOGY READINESS (SELF-ASSESSED; NO FORMAL TRA)]

CURRENT ASSESSED TRL	TRL 2 to 3 (self-assessed against the DoD nine-level scale; no formal technology readiness assessment has been performed).
BASIS	Concept implemented as a locally tested preview; the critical property set is not formally established and the named open failure prevents a proof-of-concept claim for the multi-contributor case.
EVIDENCE SUPPORTING	concept and requirements; local implementation and tests; single-seed deterministic evidence
EVIDENCE REQUIRED FOR NEXT TRL	(1) Close WP-0 and model check DA-P1 to DA-P9 with a faithful model. (2) Preregistered multi-seed campaign with confidence intervals; second-platform canonical root. (3) Government-furnished records with known poisoned or malformed entries for a Phase 2 admission evaluation.

[INTEGRATION PROFILE]

INPUTS	Records or model contributions with evidence references; admissibility policy.
OUTPUTS	Verdict with reason code; append-only ledger; assurance package.
INTERFACES	Standalone console today; a pipeline stage in a data or federated-learning flow by design.
DEPLOYMENT	Client-side, deterministic browser build; no network.
SECURITY BOUNDARY	Inside: gate, canonicalization, chain. Outside: record sources, downstream consumers that must honor QUARANTINE and REJECT.
DATA REQUIREMENTS	Synthetic; government-furnished records with ground truth would enable evaluation.

[TRANSITION ROADMAP]

Phase	Activity for this capability	Who can execute it	Exit evidence
1 Technical review	Review of the artifacts listed in the evidence table	Program alone, with a government reviewer	Reviewer findings; correction log entries
2 Government-defined simulation	Sponsor scenarios on the seeded console	Program with sponsor scenarios	Deterministic, replayable results; hash-chained records
3 Degradation and failure injection	Evidence lapse, spoofing, delayed human response	Program with sponsor parameters	False-restriction, miss, and contraction-latency figures
4 Government testbed / SIL	Integration with a representative controller or interface	Requires sponsor interface and testbed	SIL integration report
5 Hardware-in-the-loop	Representative interference; timing on real hardware	Requires sponsor laboratory or rig	Measured latencies; TRL 4 to 5 evidence
6 Relevant-environment demonstration	Independent evaluation and red team	Requires independent evaluator	Independent report
7 Program integration	Only after earlier gates succeed	Requires transition partner	Not claimed to be available

[GOVERNMENT ENGAGEMENT REQUEST (SPECIFIC)]

- Review of the admissibility requirement set by a data-integrity or federated-learning program.
- A government dataset with labeled poisoned or malformed records for Phase 2.
- Independent reproduction of the canonical ledger root on a second platform.

[REVIEWER QUESTIONS THIS PROFILE ANSWERS]

What is it: see capability description. Why does DoD care: see operational relevance. Where would it fit: see integration profile. Does it exist: see evidence snapshot. What evidence supports it: see technical evidence. Has anyone independent tested it: no. Has it been tested in the real world: no. What does it rely on: see limitations and security boundary. What can cause it to fail and what does it do then: see assurance mechanisms and the published counterexamples. How mature: see the TRL block. What would the government need to provide: see the request above. Next measurable milestone: first item under evidence required for next TRL.

SOFTWARE APPLICATION PROFILES

Six reference architectures carry the same pipeline into settings that do not require a dedicated hardware node. All are single-author research; none is deployed in any production environment. The seventh software item is the Governance Kernel, profiled separately.

ID	Profile	Function	Evidence	Alignment (author mapping; no compliance)	Status and limits
SW-01	AUTHREX-AGENT	Authority-lifecycle shim for LLM-agent runtimes: four-level (T3 to T0) HMAA gating; per-tool HKDF authorization tokens; sub-agent spawn quorum; SATA input-trust scoring that collapses on credential patterns or prompt-injection signatures; FLAME bounded deliberation with fail-safe abort on timeout; hash-chained audit ledger; hardware-anchored by BLADE-AGENT-HSM when present.	15-section specification; working simulator with four reference use cases traced gate by gate; documented hardware companion (DOI 10.5281/zenodo.20299821).	Author mapping to the CISA, NSA and Five Eyes joint guidance Careful Adoption of Agentic AI Services (1 May 2026) and FY26 NDAA Sections 1513 and 6601	Specification plus simulator; not deployed; software-only containment bypassable without the hardware root of trust.
SW-02	AUTHREX-ASSURE	Pre-deployment authority governance: evaluates configured evidence conditions and records an internal release-or-hold recommendation; deployment decisions remain external to AUTHREX.	Specification plus governor console.	Author mapping to the Cyber Strategy for America (March 2026) and NDAA Section 1533	Reference architecture; not deployed.

ID	Profile	Function	Evidence	Alignment (author mapping; no compliance)	Status and limits
SW-03	AUTHREX-ICS-GATE	OT authority governance at the IT/OT boundary: authorizes the deterministic safety logic for AI in operational technology and never makes the safety decision itself.	Specification plus governor console.	Author mapping to the CISA and NSA Principles for Secure Integration of AI in OT (3 Dec 2025); NIST SP 800-82, IEC 62443 and NERC CIP referenced	Reference architecture; not deployed; the deterministic safety logic is the platform's.
SW-04	AUTHREX-AGENT-CYBER	Governs whether an autonomous cyber-reasoning system may patch a live controller; governance only, no offensive function; folds in the cited variants AUTHREX-ZTAGENT (zero trust for autonomous agents) and AUTHREX-MCPGOV (MCP server governance) as citation layers, not separate applications.	Specification plus governor console (AUTHREX-AGENT-SIM: full seven-gate pipeline over an autonomous cyber-defense campaign with ten injectable failure modes and a 55-test V&V; suite run in browser and headless Node).	DARPA AIXCC cited as lineage; Five Eyes agentive-AI guidance; NDAA Section 1513	Reference architecture; not deployed; no offensive capability by design.
SW-05	AUTHREX-SPACECYBER	Onboard authority for orbital autonomy: where signal delay removes the human from the loop, authority moves onboard under governed conditions.	Specification plus governor console.	Author mapping to NASA SBIR subtopic EXPAND.3.S26B and Space Policy Directive 5	Reference architecture; not deployed; no flight software.
SW-06	AUTHREX-SANDBOX	Test-and-evaluation authority governance: governs what a system under test may do inside an evaluation environment so evaluation is bounded, audited, and reversible; feeds ASSURE downstream.	Specification plus governor console.	Author mapping to NDAA Section 1534 (AI sandbox-environments task force, 1 Apr 2026)	Reference architecture; not deployed.

AUTHREX-PQC (post-quantum-ready signing as a property of how BLADE-AGENT-HSM signs; an interface model, real ML-DSA not bundled) and AUTHREX-AISBOM (an AI software bill of materials emitted by the ERAM ledger) are features of other artifacts, not standalone applications, and are not counted. The cited variants AUTHREX-ZTAGENT and AUTHREX-MCPGOV are folded into AUTHREX-AGENT-CYBER as citation layers and are not separate applications.

EVIDENCE CLASS	PROPOSED (specifications) with IMPLEMENTED governor consoles that exercise the profile logic on synthetic data; NOT YET VALIDATED.
SELF-ASSESSED TRL	2 to 3 by profile; no formal determination.
INTEGRATION MODEL	API integration or runtime wrapper at the agent, ICS, space-software, or test-environment boundary.
GOVERNMENT EVALUATION PATH	Select one profile with a sponsor use case; code review of the console; sandbox evaluation with government-defined inputs and approval policies (Phases 1 to 3).
WHAT THE GOVERNMENT WOULD NEED TO PROVIDE	A named setting (agent tool inventory, OT command set, orbital operations concept, or T&E; environment) and its policy; the profile cannot be evaluated against an abstract setting.

GOVERNOR CONSOLES AND SIMULATION RECORDS

Seeded, client-side, deterministic, synthetic data. The consoles are the primary current evidence for the governance logic. They demonstrate that the logic behaves as specified under seeded conditions; they do not prove field performance, timing on target hardware, or resistance to a live adversary.

GOVERNOR CONSOLE RECORDS WITH DOCUMENTED VERIFICATION

ID	Console	Function	Documented evidence (self-run)	Limits
SIM-10	AUTHREX-SAFING Safe-State Console v1.5.0	Safe-state governance: replay against five published reference roots; formal authority model check.	1061 checks pass (S1 to S59); formal authority model check enumerates 2,488,320 executions across 15 model states and 23 transitions to depth 7 with all nine modeled properties holding; 11,739 terminal decisions across three scenarios and fifty seeds with zero mismatches; 50 of 50 mandatory mutants killed; 43-check adversarial battery with signed results; standalone sha256 ef8b828e...92e6bd6.	Deterministic, seeded results on synthetic scenarios within the packaged build; not independent validation; no formal TRL determination. No dynamics, actuation, or targeting.

ID	Console	Function	Documented evidence (self-run)	Limits
SIM-11	AUTHREX-TEAM Team Authority Governance Console v0.16	Independent arbiter over an untrusted allocator: exactly one eligible authenticated holder per responsibility, safe-hold fallback, allowlisted supervisor override.	Boot self-test reports engine encapsulation; 20 of 20 scripted seeded scenarios contained with zero authority gaps and zero split-authority states; from seed 20260727 the build reproduces reference ledger root 5a937676...50a168c2; standalone sha256 797076b5...ea01399e is the only integrity claim.	Hardened interactive prototype, not a reproducible signed release; no release archive, detached signature, or signing key. Strong-form counterexample published: partition permits both sides to self-grant.
SIM-12	AUTHREX-REDLINE Positive Human Control Console v0.13	The critical authorization can be held only by two distinct authenticated humans; a machine is never eligible; failure safe state is NO-GO.	27 of 27 bundled deterministic scenarios contained with all six audits passing; seeded sweep of 250 seeds over 25,000 ticks with zero audit failures, zero machine key issuances, zero GO states without two distinct humans, and 843 fail-safe holds, none attributed to a human; from seed 20260728 over 400 ticks reproduces reference root 7bce87c4...59a2fbb72; standalone sha256 480459ef...43914164.	Hardened standalone prototype, not a reproducible signed release. Strong-form counterexample published: two credentials do not imply two humans.
SIM-08	AUTHREX-RTA Governor Console	Run-time assurance reference: places trust assessment upstream of a Simplex-style switch so authority grades down as inputs degrade.	Client-side, seeded PRNG; SHA-256 hash-chained decision ledger; replay determinism; in-browser V&V; five stress scenarios; synthetic data only.	Standalone research console; no vehicle dynamics, actuation, or live system. Strong-form counterexample published: detection latency leaves an untrusted command applied while the plant is unsafe.
SIM-09	AUTHREX-ENGAGE Governance Console	Engagement-decision governance with HOLD as an extension outcome; policy-conditioned hold before an irreversible action.	Client-side, seeded PRNG; SHA-256 hash-chained decision ledger; replay determinism; in-browser verification; five seeded scenarios; synthetic data only.	Standalone research console; no targeting, no kinetic effects, no live system. ENGAGE.tla is a surrogate countermodel.

Console integrity hashes and reference ledger roots are published on the console pages; a reviewer re-runs a published result from the same packaged build and compares roots.

[CONSOLES AND SIMULATIONS · CONTINUED]

FLAGSHIP SIMULATIONS

ID	Simulation	Frameworks exercised	Method	Size
SIM-01	ERAM Strategic Command	ERAM, SATA, FLAME, CARA	6 scenarios; 600 Monte Carlo runs; formal properties; Merkle chain	not stated
SIM-02	APEX JADC2 Theater	All seven	Full pipeline; PBFT; Merkle audit trail	about 28 KB
SIM-03	MDO Digital Twin	HMAA, CARA, MAIVA	Air, maritime, ground domains; Byzantine fault isolation	about 21 KB
SIM-04	Tactical COP (blue-on-blue)	HMAA, CARA, SATA, ADARA	Blue-on-blue interdiction via common operating picture	about 21 KB
SIM-05	AUTHREX OS	All seven	Distributed Web Worker nodes; COP; analytics; Merkle	not stated
SIM-06	Distributed Kernel	MAIVA, CARA	Actor-model mesh; PBFT consensus; emergent interlock	not stated
SIM-07	Tactical Core	HMAA, CARA	Zero-dependency kinematic engine; safe-abort handling	not stated

All fourteen SIM records are client-side with a seeded pseudo-random generator, bit-exact reproducible within the packaged browser build, and not independently reproduced. Console integrity hashes are published on authrex.systems. Count definitions, because four different public figures exist and each counts something different: 14 catalogued SIM records (this register, SIM-01 to SIM-14); 20 standalone consoles and 24 catalogued simulations in the authrex.systems simulation tree, which counts consoles and demonstrators rather than register rows; 22 catalogued simulations in the authrex.systems evidence tile, an earlier reconciliation of the same tree; and 37 browser simulations across the two public sites (19 on burakoktenli.com, 18 on authrex.systems), which counts launchable pages. The figures are never used interchangeably, and this catalog uses only the first.

WHAT THE SIMULATIONS DEMONSTRATE	Authority degrades with trust; deliberation windows hold irreversible actions; Byzantine minorities are outvoted within the modeled fault bound; recovery follows the defined staircase; every decision is reconstructible from the record.
WHAT THEY DO NOT DEMONSTRATE	Field performance; timing on target hardware; resistance to a live or adaptive adversary; behaviour of any physical sensor. Those claims await hardware-in-the-loop testing and independent red-team evaluation.
GOVERNMENT EVALUATION PATH	Run government-defined scenarios on the seeded consoles; re-run a published result from the same packaged build and compare ledger roots (Phase 2).

BLADE HARDWARE REFERENCE DESIGNS

Twelve domain-specific governance nodes that carry the pipeline to the bill-of-materials level. Each is published as a reproducible artifact package with a BOM, interface documentation, and simulation data. None has been built.

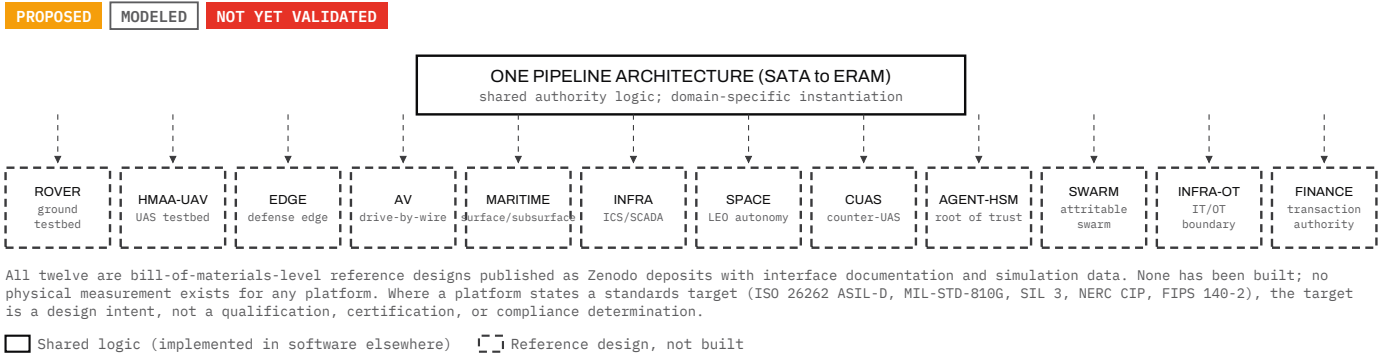


Figure 18. One pipeline architecture instantiated per domain; every platform is a reference design.

ID	Platform	Domain	Components	Ref. BOM cost	Status	Standards targets (design intent only)	Zenodo DOI
BLADE-01	Rover Testbed	Robotics testbed (ground)	37 components; 76 electrical connections	about \$484	Design complete; not built; physical build not yet documented	ROS 2; HMAA state machine model checked in TLA+ at stated bounds	10.5281/zenodo.19143190
BLADE-02	HMAA-UAV Platform	UAS testbed (air)	52 components	about \$4,200	Design complete; not built; no flight testing	MAVLink HIL bridge; ArduPilot	10.5281/zenodo.19128769
BLADE-03	BLADE-EDGE	Defense edge autonomy under denied and degraded communications	72 components	about \$139K	Design complete; not built	MIL-STD-810G (design target)	10.5281/zenodo.19177472
BLADE-04	BLADE-AV	Automotive drive-by-wire	62 components	about \$16.3K	Design complete; demonstrated in simulation; not built	ISO 26262 ASIL-D; SAE J3016 Level 4 (design targets)	10.5281/zenodo.19232130
BLADE-05	BLADE-MARITIME	Maritime surveillance (surface and subsurface)	84 components	about \$43K	Design complete; not built	IP68; MIL-STD-810G; MIL-STD-461G CE102 (design targets)	10.5281/zenodo.19246785
BLADE-06	BLADE-INFRA	Critical infrastructure ICS and SCADA	92 components	about \$11.6K	Design complete; not built	SIL 3; NERC CIP; FIPS 140-2 Level 3 (design targets)	10.5281/zenodo.19277887
BLADE-07	BLADE-SPACE	Orbital LEO autonomy	91 components; 6U-plus payload module	about \$505K	Preliminary design; 15-document package; V&V; specified, not executed	NASA SBIR EXPAND.3.S26B; SPD-5 (author alignment)	10.5281/zenodo.20183269
BLADE-08	BLADE-CUAS	Counter-UAS decision governance (passive; non-kinetic scope)	not stated	about \$43.5K	Design; simulation-based evaluation; not built	EO 14305; FY26 NDAA Safer Skies (author alignment as a research mapping)	10.5281/zenodo.20299604
BLADE-09	BLADE-AGENT-HSM	Agentic-AI hardware root of trust	USB-A and M.2 Key-E form factors	about \$199	Design; adversarial emulator v2.6; silicon not built	CC EAL6+ secure element; TPM 2.0 (component ratings, not a system certification)	10.5281/zenodo.20299821
BLADE-10	BLADE-SWARM	Attritable swarm autonomy	not stated	about \$1,333 per node	Simulated; TLA+ at stated bounds; testbed self-assessed TRL 2; not built	DoDD 3000.09 and NIST AI RMF (author alignment)	10.5281/zenodo.20351198
BLADE-11	BLADE-INFRA-OT	IT/OT boundary governance appliance	48 BOM line items; 1U	about 1U appliance (cost not stated)	Design; simulation-based evaluation; not built	NIST SP 800-82; IEC 62443; NERC CIP (author alignment)	10.5281/zenodo.20342067
BLADE-12	BLADE-FINANCE	Financial-sector AI decision authority	36 components	about \$9,228	Design; simulation on synthetic data; not built	Treasury FS AI RMF; EO 14179 (author alignment)	10.5281/zenodo.20374692

Reference costs are the author's hardware-only BOM estimates at deposit time, not program cost estimates.

[BLADE · PLATFORM PROFILES]

Platform profiles

BLADE-01 · Rover Testbed Robotics testbed (ground)

Design: Dual compute (Raspberry Pi 5 and ESP32) running the SATA-HMAA-CARA pipeline.

Evidence: Seven fault scenarios demonstrated in simulation (run counts withdrawn; see CL-03). Status: Design complete; not built; physical build not yet documented DOI: 10.5281/zenodo.19143190

BLADE-02 · HMAA-UAV Platform UAS testbed (air)

Design: Cube Orange+ flight controller with NVIDIA Jetson Orin NX companion and a MAVLink hardware-in-the-loop bridge.

Evidence: Five adversarial scenarios demonstrated in simulation (run counts withdrawn; see CL-03). Status: Design complete; not built; no flight testing DOI: 10.5281/zenodo.19128769

BLADE-03 · BLADE-EDGE Defense edge autonomy under denied and degraded communications

Design: Reference hardware instantiation of the full pipeline for high-consequence defense edge operations.

Evidence: Simulation and design analysis; portability of BLADE-AV shown against it in simulation. Status: Design complete; not built DOI: 10.5281/zenodo.19177472

BLADE-04 · BLADE-AV Automotive drive-by-wire

Design: Nine-module pipeline on Jetson AGX Orin plus Zynq UltraScale+; three-leg redundant KILOVAC LEV200 fail-safe relay; trust-driven forced driver handoff; brake-to-stop-in-lane degradation.

Evidence: Twelve attack scenarios demonstrated in simulation; no vehicle integration. Status: Design complete; demonstrated in simulation; not built DOI: 10.5281/zenodo.19232130

BLADE-05 · BLADE-MARITIME Maritime surveillance (surface and subsurface)

Design: Nine-module pipeline on Jetson AGX Orin 64GB plus Zynq UltraScale+ ZU7EV; hydroacoustic sonar, magnetic anomaly detection, AIS spoofing detection; four maritime extensions (hydroacoustic plus MAD fusion, recursive AIS deception risk, sea-state authority damping, acoustic-delay Byzantine consensus); dual-GPIO safety interlock with a 500 ms watchdog driving a normally-open relay.

Evidence: Thirteen fault-injection scenarios in simulation; no water trial. Status: Design complete; not built DOI: 10.5281/zenodo.19246785

BLADE-06 · BLADE-INFRA Critical infrastructure ICS and SCADA

Design: IEC 61850 GOOSE, Modbus TCP/RTU, PROFINET IO; Pilz PNOZ S7.1 SIL-3 safety relay.

Evidence: Simulation only. Status: Design complete; not built DOI: 10.5281/zenodo.19277887

BLADE-07 · BLADE-SPACE Orbital LEO autonomy

Design: Microchip RTG4 radiation-tolerant FPGA (primary and backup) plus Aitech S-A1760 Venus SBC (primary and backup) with a designed failover under 200 ms; three-fault-tolerant payload and thruster firing path (two independent normally-open solid-state relays plus pyrotechnic isolation); ECDSA P-256 audit-chain keypair in a rad-tolerant TPM; 30 krad TID behind 3 mm aluminium equivalent; 5-year design life (7-year stretch); \$505,440 reference BOM; 11.0 kg estimated mass.

Evidence: No platform-level simulation campaign completed; the least mature platform. Status: Preliminary design; 15-document package; V&V; specified, not executed DOI: 10.5281/zenodo.20183269

BLADE-08 · BLADE-CUAS Counter-UAS decision governance (passive; non-kinetic scope)

Design: Four-level HMAA federal-to-SLTT authority handoff; Dempster-Shafer consensus across radar, RF, EO/IR, Remote ID, LIDAR; ECDSA P-256 evidence record; about 75 percent design reuse from BLADE-EDGE.

Evidence: Demonstrated in simulation across the published BLADE family (the deposit uses the word validated; not carried here); no sensor integration. Status: Design; simulation-based evaluation; not built DOI: 10.5281/zenodo.20299604

BLADE-09 · BLADE-AGENT-HSM Agentic-AI hardware root of trust

Design: Non-exportable ECDSA P-256/P-384 keys in an NXP EdgeLock SE051; four-level authority state in an Infineon SLB 9670 TPM 2.0; HKDF per-tool tokens; multi-modal tamper cascade that zeroizes keys; five-opcode 64-byte ABI.

Evidence: Emulator: 275 of 275 deterministic checks, 3 reruns confirmed; browser-local, Web Crypto substrate, silicon timing modeled not measured. Status: Design; adversarial emulator v2.6; silicon not built DOI: 10.5281/zenodo.20299821

BLADE-10 · BLADE-SWARM Attributable swarm autonomy

Design: Byzantine-fault-tolerant two-phase sub-quorum consensus gated by SATA, HMAA, MAIVA across $N = 10, 50, 500$ agents, tolerating up to $f = (N-1)/3$ compromised agents per quorum; safe halt by default under denied or degraded RF; tamper-evident distributed audit ledger; X500 V2 quadrotor nodes; interface control document ICD-SWARM-001 v1.0. The platform page states TRL 3 to 4 for the simulator and formal specification and TRL 2 for the physical testbed design (self-assessed).

Evidence: TLA+ model checked at stated bounds (5 safety, 3 liveness); consensus results are simulation outcomes. Status: Simulated; TLA+ at stated bounds; testbed self-assessed TRL 2; not built DOI: 10.5281/zenodo.20351198

BLADE-11 · BLADE-INFRA-OT IT/OT boundary governance appliance

Design: Fail-closed, bump-in-the-wire appliance: each cross-boundary command is adjudicated to propagate, hold, or isolate across four OT regimes (nominal, elevated, lockdown, safe-halt); malformed input fails closed; Xilinx Kria K26 governance plane, x86 network plane; SHA-256 ledger.

Evidence: Simulation only; not deployed in any facility. Status: Design; simulation-based evaluation; not built DOI: 10.5281/zenodo.20342067

BLADE-12 · BLADE-FINANCE Financial-sector AI decision authority

Design: Eight-stage pipeline routing each transaction to autonomous clearance, supervised review, elevated confirmation, or manual hold; retrospective swarm review for low-and-slow activity; YubiHSM 2 in a FIPS 140-2 Level 3 enclosure.

Evidence: Synthetic data only; not deployed in any institution. Status: Design; simulation on synthetic data; not built DOI: 10.5281/zenodo.20374692

[LIMITATIONS COMMON TO THE FAMILY]

- No platform has been built. Evidence is design-level: bill of materials, interfaces, and simulation data. Physical sensor performance, RF propagation, real jamming and spoofing effects, timing on real hardware, and behaviour in a contested electromagnetic environment have not been measured.
- Standards named on the platform pages are design targets or author alignments, never qualifications, certifications, or compliance determinations. Component ratings (CC EAL6+ secure element, FIPS 140-2 Level 3 HSM enclosure) attach to purchased parts, not to the design as a system.
- BLADE-AGENT-HSM emulator results (275 of 275) are browser-local on a Web Crypto substrate; silicon timing is modeled, not measured; PQC is an interface model.
- The earlier statements that the rover and UAV testbeds were physically assembled, and the 350-run and 250-run counts, are withdrawn (correction log CL-03, CL-04).

[TECHNOLOGY READINESS AND NEXT VALIDATION STEP]

SELF-ASSESSED TRL	TRL 2 for the family (concept and design formulated; BOM-level). BLADE-AV (twelve attack scenarios in simulation), BLADE-AGENT-HSM (adversarial emulator) and BLADE-SWARM (TLA+ at stated bounds) reach 2 to 3 on simulation or emulator evidence. No formal determination.
EVIDENCE REQUIRED FOR TRL 3 TO 4	A first prototype assembly (BLADE-EDGE or the Rover testbed) with bench measurement of the governance path; then hardware-in-the-loop under representative interference. The program roadmap on authrex.systems lists BLADE-EDGE prototype assembly for Q3 to Q4 2026 as planned or in progress; no build is documented in the register.
GOVERNMENT NEED	Design review by a laboratory with hardware integration capacity; bench and HIL rig access; range access for RF effects; an integration partner (prime or platform provider) for any platform-specific instantiation.

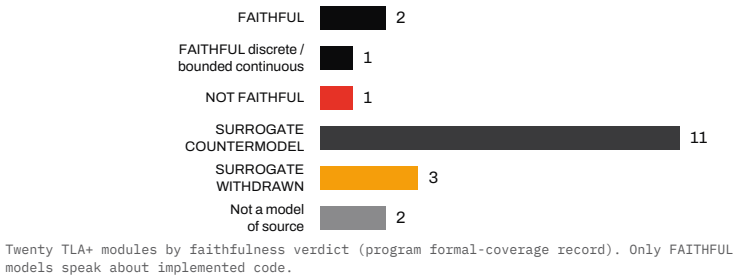
FORMAL METHODS: WHAT IS MODELED, WHAT WAS CHECKED, AND WHAT THAT DOES NOT ESTABLISH

FORMALISM	TLA+ specifications checked with the TLC model checker. No theorem proving. No refinement proof between models and code.
MODEL SCOPE	HMAA authority automaton (four tiers, hysteresis, handoff); AUTHREX-ABORT modules AuthorityLapse v3 and InterfaceRecord v1; faithful component models ADARA_XSI_EXACT and MAIVA_QUORUM; surrogate countermodels for ADARA, CARA, ENGAGE, ERAM, FLAME, MAIVA, REDLINE, RTA, TEAM; BLADE-SWARM refinement of the MAIVA module; SATA protocol state machine as reported in the published article.
STATE VARIABLES (HMAA)	Current tier, target tier, trust input, hysteresis timer, operator availability, pending outcome; the discrete automaton assumes instantaneous authority equals its target.
SAFETY PROPERTIES HELD (HMAA)	5 invariants at depth 9 (tier confinement, immediate downgrade, no timeout restoration, hysteresis on upgrade, outcome issued only from HMAA) plus 1 liveness property; 2 of 8 stated properties vacuous at this bound: UpgradelsStepwise and TrustAuthorityWeakConsistency.
SAFETY PROPERTIES HELD (ABORT)	C1 envelope confinement; C2 adverse-step non-expansion; C3 irreversible non-reacquisition; C4 execution conformity; C5 no unbacked irreversible execution; lapsed-set monotonicity; latch capture (AuthorityLapse v3, 21 configurations, each property in isolation). Fifteen interface and record obligations plus type invariance (InterfaceRecord v1, 35 configurations).
ASSUMPTIONS	Discrete time; declared fault bounds; instantaneous authority equals target (HMAA); abstract platform adapter (ABORT); no adaptive adversary; configuration bounds as stated per run.
STATE-SPACE EXPLORATION	HMAA: 23,748 distinct reachable states at depth 9 (26,397,356 states generated); the result describes the discrete authority automaton under the assumption that instantaneous authority equals its target and does not cover continuous behaviour between decision instants; of 8 stated properties, 5 invariants and 1 liveness property held, 2 vacuous at this bound (UpgradelsStepwise, TrustAuthorityWeakConsistency). ABORT Stage 1B: 31,341 states and 130,869 transitions exported from TLC; Stage 1 availability relation 972 of 972 projected states (AvailRaw only, a projection, not the complete reachable graph). MAIVA_QUORUM: Q_NEVER_FAILS held across 12,507,501 distinct states, in the narrowed sense described below. SAFING: 2,488,320 executions across 15 states and 23 transitions to depth 7, nine properties holding.

COUNTEREXAMPLES	No counterexample for the HMAA properties that held. Strong-form counterexamples are published for MAIVA, ERAM, ADARA, FLAME, CARA, REDLINE, TEAM, and RTA (table below). Obligation 6: the GLUE property is refuted in the bounded extension.
VERIFICATION INTEGRITY	62 configurations re-executed on a second operating system and architecture against sealed results: 62 distinct-state matches, 62 verdict matches, 0 differences. One configuration (MQ_Q_NOT_SINGLETON) has never reliably executed and is recorded as environment-specific and unresolved. Second-environment confirmation by the same producer; not independent verification.

[FORMAL METHODS • CONTINUED]

FAITHFULNESS TAXONOMY ACROSS TWENTY TLA+ MODULES



A model earns FAITHFUL only by importing and checking the implemented artifact. Two of twenty are FAITHFUL (ADARA_XSI_EXACT, MAIVA_QUORUM); HMAA_Refine is faithful on the discrete side and bounded on the continuous side; ADARA_XSI_NEAR is NOT FAITHFUL; eleven are SURROGATE COUNTERMODELS, whose results are findings about mechanism classes and not about implemented code; three surrogates are WITHDRAWN (ADARA_IMPL presented a surrogate as a code finding; OBL6 imported neither the archived specification nor a hybrid model; HMAA_Hybrid inverted the tier semantics); two files are not models of source (HMAA_Authority_FSM archived source; REQUIREMENTS a requirement specification). Each withdrawal is recorded beside its original. The page states that it corrects prior claims rather than adding to them.

THE CENTRAL FINDING: TWO PROPERTIES PER CLAIM

Each safety claim was written twice: a naive form over the system's own bookkeeping and a strong form over ground truth the system cannot access. Across independent domains nine naive forms held and nine strong forms failed. The gap is the finding: bookkeeping can be satisfied while ground truth is violated.

Architecture	Strong-form failure (class level)	Requirement recorded
MAIVA	Three attestations meet a quorum of three, but one attester is compromised, leaving two honest	Quorum + MaxCompromised; fault bound is an independent assumption (R1_ASSUMPTION: no runtime check can establish it)
ERAM	SilentDrop fires, the ledger is empty, the shadow is non-empty, the hash chain is valid	Ledger integrity needs an independent decision count
ADARA	All channels lie identically; conflict reads zero; verdict reads zero	Deception detection must weight agreement, not only disagreement
FLAME	EXECUTE with the deliberation clock satisfied and zero real time elapsed	A deliberation clock must advance only by elapsed time
CARA	Replayed re-arm token: bookkeeping records it present, ground truth records it never issued	A re-arm token must be bound to a nonce
REDLINE	Two credentials do not imply two humans	Recorded as a limit of credential-based positive control
TEAM	Partition permits both sides to self-grant	Recorded as a limit under partition
RTA	Detection latency leaves an untrusted command applied while the plant is unsafe	Recorded as a limit of monitor latency

THE MAIVA QUORUM DEFECT: FAITHFUL AND LOCATED IN PRODUCTION SOURCE

The shipped check derives the fault bound from the roster size ($f_{\text{Eff}} = \text{floor}((nA - 1) / 3)$), builds the quorum requirement from that bound ($2 \times f_{\text{Eff}} + 1$), and then compares the roster size against it. A test whose threshold is a function of its own input can only report on itself: the check is satisfied for every roster size from 1 to 5,000 and at $n = 1$ a single agent constitutes a quorum. The Byzantine formula $n \geq 3f + 1$ is correct in isolation; the defect is that f must be an independent assumption about the adversary. Scope is bounded to the quorum decision at the named lines; nothing downstream is covered. This is a documented open defect.

OBLIGATION 6 IS OPEN

Initial-state correspondence (ABS_INIT_CORRESP) was previously reported as holding and has been WITHDRAWN as bound-dependent: HELD at MaxTick = 3 (32,372 distinct states), VIOLATED at MaxTick = 4 (26,719 states) and at 6, 8, 12, 16. The mechanism is readable in the archived source: LockoutRelease increments the tick without the MaxTick guard that LockoutHold carries, so tick = 5 is reachable with MaxTick = 4. The shipped bound was 3 and the property fails at bound plus one. Step 4 of Obligation 6 is NOT DONE for two independent reasons: it checked TypeInvariant rather than correspondence, and its HOLDS was a bounded artifact.

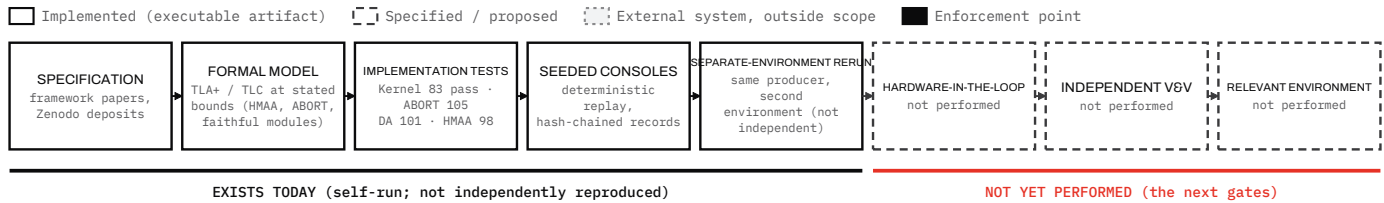
WHAT FORMAL VERIFICATION DOES ESTABLISH	That the stated properties hold within the declared models at the stated bounds and configurations; that a faithful model's finding applies to the imported implementation function; that a counterexample exhibits a reachable violation of the strong form.
---	---

WHAT IT DOES NOT ESTABLISH

Correctness of any implementation or deployment; behaviour beyond the checked bound (one bounded result was an artifact of its horizon); properties of the physical system; anything about modules labeled surrogate, beyond the mechanism class; anything for the vacuous properties; anything independent (no third party has re-executed any configuration).

TEST AND EVALUATION FRAMEWORK

Verification question, method, environment, result, and remaining gap for every testable artifact. Nothing here is asserted as completed testing beyond what the evidence register documents.



Reading rule: no stage substitutes for a later one. A model-checked automaton is not a proof of the physical system; a soak test in a second environment is not independent verification; a seeded console is not field validation. Conformance between the formal models and the implemented code is by construction and test, not by refinement proof (Obligation 6 is open).

Figure 19. Verification pipeline. Solid stages exist today and are self-run; dashed stages have not been performed.

Artifact	Verification question	Method / environment	Scenario count	Actual result	Repeatability	Remaining gap
HMAA	Does the automaton satisfy its 8 stated properties?	TLC on a laptop; depth 9	1 configuration	5 invariants + 1 liveness held; 2 vacuous	Deposit of record re-executable	Larger bound; refinement to code; Obligation 6
HMAA package	Does the code implement the automaton?	pytest; 42-file package	98 tests; 7 adversarial experiments	Pass (self-run)	Deterministic traces	Independent execution
Kernel	Zero unauthorized privileged executions?	pytest in python:3.11-slim container	53 test functions	83 passed, 1 skipped	Container build reproducible (archive hash logged)	External reproduction; security assessment; SBOM
ABORT	Do C1 to C5 and the record obligations hold; does the engine replay deterministically?	TLC (21 + 35 configurations); pytest; soak on a second machine	105 tests; 34 release checks; 1,000,000 runs; 23 of 42 scenarios	All held at bounds; zero soak failures; 7 mutants, 4 detected	Byte-identical harness output across extractions	Refinement gate; B14, B2; hazard and cyber packages; HIL
DA	Are verdicts deterministic and gates ordered?	Local pytest; 12 release gates	101 tests; 1 seed, 120 rounds	101 of 101; 12 TEST_PASS_LOCAL, 1 PARTIAL, 1 BLOCKED	Canonical root on one platform	WP-0; DA-P1 to P9; multi-seed; second platform
SAFING	Do the nine modeled properties hold; do decisions match reference roots?	Console model check to depth 7; 3 scenarios x 50 seeds	1061 checks; 11,739 decisions; 50 mutants; 43 adversarial checks	All pass; 50 of 50 mutants killed; zero mismatches	Five published reference roots	Independent replay
TEAM / REDLINE	Exactly one holder; two distinct humans for GO?	Scripted seeded scenarios; seeded sweep	20 of 20; 27 of 27; 250 seeds x 25,000 ticks	Contained; zero audit failures; 843 fail-safe holds	Reference ledger roots from stated seeds	Signed release; strong-form limits
ERAM console	Are the 6 scenarios reproducible; do formal properties hold?	Seeded PRNG; 100 runs per scenario	600 runs	Reproducible within the build	Bit-exact in the packaged build	Validation against any real record
Frameworks (SATA, ADARA, FLAME, CARA)	Does authority behave as specified under the seven adversarial scenarios (S1 to S7)?	Browser simulation; parameterized fault models; about 100 Hz idealized loop	7 scenarios (rover), 5 (UAV)	Demonstrated in simulation; run counts withdrawn	Seeded consoles	HIL; adaptive adversary; hardware timing
BLADE-AGENT-HSM	Do the emulated HSM checks hold under adversarial harnesses?	Node.js emulator v2.6; 7 suites	275 checks	275 of 275; 3 reruns confirmed	Deterministic	Silicon; measured timing
BLADE hardware	Any physical property	None performed	0	None	n/a	Everything physical

[TEST AND EVALUATION · CONTINUED]

STATUS BY TEST CLASS

Test class	Frameworks	Hardware	Software / Kernel	Consoles	Notes
Unit / functional	PARTIAL	NOT PERFORMED	COMPLETE (83 tests)	COMPLETE (per console)	Per artifact; no integrated system test
Integration	NOT PERFORMED	NOT PERFORMED	PARTIAL (demo chain)	PARTIAL (integrated sims)	No integration with any external system
Failure injection	PARTIAL	NOT PERFORMED	PARTIAL	PARTIAL	Trust loss and evidence lapse injected in simulation
Adversarial	PARTIAL	NOT PERFORMED	PARTIAL	PARTIAL	Scripted adversary; no independent red team
Degraded mode	PARTIAL	NOT PERFORMED	PLANNED	PARTIAL	Simulation only
Authority-state	COMPLETE (HMAA model)	NOT PERFORMED	PARTIAL	COMPLETE	Tier transitions and hysteresis checked in TLA+ and consoles
Recovery	PARTIAL	NOT PERFORMED	PLANNED	PARTIAL	CARA GREP path in simulation
Auditability	COMPLETE (replay)	NOT PERFORMED	COMPLETE (tamper detection)	COMPLETE	Deterministic replay and hash-chain checks
Monte Carlo	PARTIAL (SATA as published; ERAM)	NOT APPLICABLE	NOT PERFORMED	PARTIAL	Single-seed campaigns elsewhere; no confidence intervals
Mutation / negative control	NOT PERFORMED	NOT APPLICABLE	NOT PERFORMED	PARTIAL (ABORT, SAFING)	ABORT: B12 closed, B14 and B2 open
Formal (bounded)	COMPLETE at stated bounds (HMAA); PARTIAL (others)	NOT APPLICABLE	NOT PERFORMED	PARTIAL (ABORT, SAFING)	Not a proof of the physical system
SIL	NOT PERFORMED	NOT PERFORMED	NOT PERFORMED	NOT APPLICABLE	First government-facing gate
HIL	NOT PERFORMED	NOT PERFORMED	NOT PERFORMED	NOT APPLICABLE	Planned; no rig exists
Relevant environment	NOT PERFORMED	NOT PERFORMED	NOT PERFORMED	NOT APPLICABLE	Requires sponsor
Operational	NOT PERFORMED	NOT PERFORMED	NOT PERFORMED	NOT APPLICABLE	Not planned before earlier gates succeed

☐ Executed (self-run) ☐ Not yet performed

UNIT / IMPLEMENTATION pytest suites per artifact Kernel, HMAA, ABORT, DA, MAIVA self-tests	FORMAL (BOUNDED) TLC model checking properties isolated per configuration	SCENARIO / REPLAY seeded consoles; reference ledger roots; deterministic replay	NEGATIVE CONTROLS mutation campaigns; retained negative results (NEG-1 to 5)	SOAK / LOAD ABORT 1,000,000 runs; Kernel concurrent-load test	SIL / HIL / RANGE not performed; the first government-facing gate
---	--	--	--	--	---

All executed testing is self-run on synthetic data. Second-environment reruns (ABORT soak, formal-coverage 62-configuration rerun) are confirmation by the same producer, not independent verification.

Independent test and evaluation has not been performed by any government or third party for any artifact in this catalog.

Figure 20. Test architecture as it exists.

EVALUATION PROTOCOL ON RECORD (SIMULATION)

The published evaluation protocol defines seven adversarial scenarios (S1 camera occlusion, S2 lidar spoofing, S3 IMU drift, S4 RF jamming, S5 compound attack, S6 cross-sensor disagreement, S7 recovery), five metrics (unsafe action rate, false lockout rate, mean recovery time, authority transition correctness, detection latency), three baselines for comparison (binary threshold, ML anomaly detection, Simplex switching), and five stated assumptions (parameterized sensor models; idealized timing at about 100 Hz; scripted adversary; simulated agent populations; deterministic governance). The baselines are specified for reimplementation and no comparative result is claimed in this catalog. The per-scenario run counts that earlier accompanied the protocol are withdrawn (correction log CL-03).

AI AND AUTONOMY ASSURANCE

The program distinguishes AI performance (how well a model perceives or plans) from system assurance (whether the system stays within authorized bounds when the model is wrong). AUTHREX makes no claim about the former.

Assurance question	Program answer (design)	Evidence status
Human oversight	HANDOFF and ABORT are outcomes of the authority computation; the operator confirms at A1 where policy requires; a positive-control console (REDLINE) requires two distinct authenticated humans for the critical authorization.	Design; consoles; REDLINE limit published (two credentials are not two humans)
Authority assignment	Granted in advance by a human as a bounded envelope; computed at runtime as a tier; issued from one component (HMAA).	TLA+ at stated bounds; simulation
Degraded mode	Admissible set narrows at once; contingencies are pre-authorized; safe state and GREP recovery.	Simulation; ABORT formal modules
Uncertainty and confidence	Per-source trust scalar with ignorance kept distinct from distrust; model confidence is informational and never gates an action.	SATA design; ABORT implementation
Model failure	A wrong or overconfident model cannot widen authority; the machine keeps intelligence, the human keeps authority.	Design property (ABORT C2)
Sensor failure	Hard vetoes (navigation, containment) prohibit dependent actions at once; degradation sources contract progressively.	ABORT implementation; SATA diagnostics
Adversarial conditions	ADARA deception prior; MAIVA trimmed median; identical-lie and majority-capture failures are published as open.	Simulation; counterexamples
Distribution shift and out-of-distribution inputs	Treated as evidence degradation (temporal stability, physical plausibility diagnostics); not detected as such.	Partial; no OOD detector is claimed
Fallback state	A0 (revoked) safe state; never a new capability; SafeState in ABORT is deferred.	Design; partly implemented
Audit logs and traceability	Hash-chained, append-only records with deterministic replay; ledger tamper detection in the Kernel.	Implemented; completeness (nothing dropped) not established (ERAM counterexample)
Verification boundaries	Stated per artifact: bounds, configurations, assumptions, faithful versus surrogate.	Published formal-coverage record
Operator intervention	Explicit at A1; supervisor override allowlisted in TEAM; veto recorded as what it is.	Consoles
Fail-safe and abort behaviour	Timeout defaults to abort; abort is the last exit, with graded contraction and structured recovery before it.	FLAME design; ABORT formal modules

AI performance versus system assurance. Nothing in this catalog claims that any perception or planning model is accurate, robust, or trustworthy. The claims concern what the system remains authorized to do when the model or its inputs fail, and every such claim is bounded by the stated model, configuration, and synthetic data.

CYBERSECURITY AND ZERO-TRUST RELEVANCE

Evaluated against principles, not claimed as compliance. Wording throughout is "architecturally relevant to"; no zero-trust maturity assessment, RMF control assessment, or authorization decision exists.

Principle	Architectural relevance in the portfolio	What has been evaluated	What has not
Identity and authentication	Agent identity at the Kernel proxy; per-node ECDSA keys in BLADE-SWARM and AGENT-HSM designs; two distinct authenticated humans in REDLINE	Kernel demo; console scenarios	Any identity provider integration
Authorization and least privilege	Per-action tier gating (T3 to T0) beneath any role; per-tool HKDF tokens (AGENT design)	Kernel test suite (zero unauthorized privileged executions in synthetic scenarios)	Real workloads; policy semantics against a real mission
Segmentation	Proxy boundary between agent and tools; bump-in-the-wire OT boundary (BLADE-INFRA-OT design); ICS-GATE authorizes but never makes the safety decision	Design; console	Any network deployment
Auditability and provenance	Hash-chained, checkpointed ledgers; signed audit evidence; reviewer packet export; Merkle chains in ERAM	Tamper detection tested (clean, forged, truncated)	Ledger completeness; external time source; log retention policy
Data integrity	AUTHREX-DA admissibility verdicts; SATA attestation chain (design)	Local tests; single seed	Majority capture (NEG-3 open)

Principle	Architectural relevance in the portfolio	What has been evaluated	What has not
Software supply chain	Site serves a hash-pinned content security policy; no external scripts or third-party requests on authrex.systems	Site assessment of 1 Sep 2026	SBOM (none exists); vulnerability scan of vendored three.js; SSDF attestation
Secure update and secret protection	Signed policy loading (Kernel); non-exportable keys and tamper cascade (AGENT-HSM design)	Kernel signed policy; emulator	Key custody on any real platform; update mechanism
Tamper resistance	Multi-modal tamper cascade that zeroizes keys (design); ledger tamper detection (implemented)	Emulator; container tests	Physical tamper testing
Resilience and incident recovery	CARA recovery; fail-closed defaults; safe-halt under denied RF (SWARM design)	Simulation	Any live incident
Public software posture	Security headers live; no server-side logging; no authentication surface; client-side only	Assessed 1 Sep 2026	Independent penetration test; Section 508 conformance

Relationship to zero trust: zero trust verifies at the network layer (never trust, always verify); AUTHREX extends verification to sensor evidence and machine actions and supplies the runtime state that policy-as-code acts on. Relationship to RBAC and ABAC: AUTHREX adds an evidence-conditioned, time-varying authority beneath the role. These are complements. No product comparison is made.

STANDARDS AND POLICY ALIGNMENT

The relationship words are deliberate and graded: aligned in intent; informed by; architecturally relevant; mapped by the author; design target; evaluated against; compliant; certified. Nothing in this portfolio is at the last two levels. Governing versions should be verified against the issuing body before any publication that relies on them.

Standard, framework, or policy	Relationship	What it touches here	Not claimed
NIST AI RMF 1.0 (Govern, Map, Measure, Manage)	Informed by; mapped by the author	Runtime authority state as a managed risk control; auditability; documented limitations	No conformance assessment
DoD Responsible AI principles; DoD AI ethical principles	Aligned in intent	Traceable, governable behaviour; human override as a designed outcome	No DoD determination of any kind
DoDD 3000.09 (Autonomy in Weapon Systems)	Policy alignment where applicable	Action-gate categories used as an author mapping in MAIVA and BLADE-SWARM; the artifacts are not weapon systems	No approval; the directive applies to an incorporating system, not to this layer
DoD Zero Trust Strategy (2022) and NIST SP 800-207	Architecturally relevant	Per-action least privilege, continuous verification of evidence, auditability, provenance	No zero-trust maturity assessment
NIST SP 800-53 (AC, AU, SI, SC families); RMF (DoDI 8510.01)	Architecturally relevant	Access enforcement at a proxy, audit generation, integrity verification	No ATO, cATO, or control assessment
NIST SP 800-218 (SSDF); SBOM practice	Aspirational; gap recorded	A schema-valid SBOM is an open gate for AUTHREX-ABORT; none exists for the Kernel	No SSDF attestation
ASTM F3269-21, Standard Practice for Methods to Safely Bound Behavior of Aircraft Systems Containing Complex Functions Using Run-Time Assurance	Cited as design lineage	Monitor-plus-constrained-function pattern; AUTHREX-RTA console implements the reference model	Never compliance; no conformance to the practice is claimed
RTCA DO-178C and DO-333 (formal methods supplement)	Mapped by the author	Model checking as the kind of evidence DO-333 accepts; not produced to DO-333 objectives	No certification credit; no software level assigned
MIL-STD-882E (system safety)	Mapped by the author	Authority contraction as a candidate hazard-mitigation control; a hazard package is an open gate	No hazard analysis accepted by any authority
ISO 26262 (BLADE-AV)	Design target	Functional-safety design intent for the fail-safe relay, forced handoff, and brake-to-stop degradation; the deposit states ASIL-D as the target	Not certified; ASIL-D is a target pathway, not an achieved integrity level; no safety case exists
SAE J3016 (BLADE-AV)	Terminology reference	Taxonomy and definitions for levels of driving automation; used only to name the automation level the design addresses	Not a safety or assurance standard; no conformity, safety, or capability claim attaches to citing it
IEC 62443; NIST SP 800-82; NERC CIP (BLADE-INFRA, INFRA-OT, ICS-GATE)	Design targets; author mapping	Fail-closed boundary adjudication; OT regimes	No compliance determination
MOSA (10 U.S.C. 4401) and modular open systems practice	Conceptually aligned	Authority logic separated from the controller behind an authority interface; components replaceable	No MOSA conformance assessment performed
CISA, NSA and Five Eyes: Careful Adoption of Agentic AI Services (1 May 2026)	Author mapping	Tool misuse and credential exposure addressed by tier gating and per-tool tokens	No endorsement or evaluation by any issuing agency

Department naming: the statutory name Department of Defense is used in this catalog. Executive Order 14347 (September 5, 2025) authorizes Department of War as a secondary title in non-statutory contexts; where an instrument's own title uses the secondary title it is retained verbatim.

INTELLECTUAL PROPERTY

Eight U.S. provisional applications reported as filed. A provisional application establishes a filing date; it is not a granted patent, no examination has occurred, and no determination of novelty or patentability has been made. Nothing in this catalog is "patented".

Subject	Title as filed or program title	Application	Filed	Receipt	Status
HMAA	Human-Machine Authority Allocation	63/999,105	7 Mar 2026	74759595	Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending
CARA	Controlled Authority Recovery Architecture	64/000,170	9 Mar 2026	74767602	Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending
SATA	Sensor and Actor Trust Assessment	64/002,453	11 Mar 2026	74808459	Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending
FLAME	Forced Latency for Authority-Managed Escalation (filing title: Flash War Latency Architecture)	64/005,607	14 Mar 2026	74858888	Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending
ADARA	Adversarial Deception-Aware Reasoning	64/110,218	13 Jul 2026	78658472	Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending
MAIVA	Multi-Agent Integrity Voting	64/110,221	13 Jul 2026	78659350	Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending
ERAM	Escalation Risk Assessment Model	64/110,225	13 Jul 2026	78660113	Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending
EBCC	Evidence-Bounded Assertion Control Systems and Methods (program-level; not a governance architecture)	64/144,399	30 Aug 2026	confirmation 8448	Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending

All eight filed as micro entity, sole inventor. Patent-filing and priority deadlines, and any disclosure-bar analysis, are maintained in the controlled IP record and are to be confirmed with patent counsel; no deadline conclusion is stated here and no patent grant is claimed. Whether to file a non-provisional application claiming benefit of any of these filings is an open decision pending counsel. Deposits carry CC BY 4.0 (most) or MIT (MAIVA) licenses; the Governance Kernel license is under review and no evaluation or redistribution right is granted unless explicitly stated. Repository pages for six frameworks present proprietary all-rights-reserved terms (correction log: license matrix pending owner decision).

PUBLICATIONS AND TECHNICAL RECORD

Author	Title	Venue	Status (exact)	Identifier	Date
Oktenli, B.	SATA: Sensor Attestation and Trust Anchoring: A Formally Specified, Simulation-Evaluated Continuous-Trust Protocol for Autonomous-Systems Sensor Edges	Journal of Hardware and Systems Security (Springer Nature) 10, 17 (2026)	Published (peer reviewed)	10.1007/s41635-026-00190-4	Published 21 Jul 2026 (received 4 Jun; accepted 14 Jul)
Oktenli, B.	The Authority Gap: Why Meaningful Human Control Fails at the Tier Where It Is Signed	Digital War (Springer Nature, Palgrave Macmillan) 7, article 6 (2026)	Published (peer reviewed)	10.1057/s42984-026-00113-1	Published 31 Aug 2026 (accepted 21 Aug); Georgetown University affiliation on the record

Author	Title	Venue	Status (exact)	Identifier	Date
Oktenli, B.	The Ingratitude Paradox: A Causal Mechanism of Democratic Discounting of Successful Protection	International Journal of Contemporary Security Studies 2(2), 55 to 68 (2026)	Published (peer reviewed; not an AUTHREX technical paper)	DOI to be registered when the issue is finalized (December 2026)	Published 24 Aug 2026
Oktenli, B.	Trust-Proportional Authority Governance for Uncrewed Autonomous Systems: Formally Verified Runtime Authority Management	AIAA SciTech Forum 2027, session UAS-16 (Certification for Autonomous Systems), Orlando	Accepted for presentation following peer review; proceedings publication pending	DOI assigned on proceedings publication	Presentation 15 Jan 2027; final manuscript due 1 Dec 2026
Oktenli, B.	Escalation Risk Assessment Model (ERAM)	SSRN working paper 6176802	Working paper (not peer reviewed)	SSRN 6176802	2026
Oktenli, B.	AUTHREX Research Program: presentation	NDIA 2026 Emerging Technologies for Defense Conference and Exhibition, Washington, DC	Abstract accepted; selected as presenter; registration confirmed (not a peer-reviewed paper)	n/a	9 to 10 Sep 2026
Oktenli, B.	Open-access corpus	Zenodo (31 deposits) and SSRN (17 papers), 48 DOI-registered works	Self-archived; not peer reviewed	ORCID 0009-0001-8573-1667	as of 1 Sep 2026
Oktenli, B.	The Authority Equation (3 vols), The Authority Discipline (4 vols), Autonomous Authority (3 vols)	Authority & Architecture imprint; ten-volume technical reference series	Published print and ebook; not peer reviewed	LCCN 2026912260 (Vol. I); Bowker ISBNs	2026

External record (verified on the evidence site as of 1 to 2 September 2026): NHTSA cited the author by name in Federal Register 91 FR 30789, footnote 10. External sources have cited, quoted, republished or discussed this work; one is a peer-reviewed article (Middle East Policy, Wiley, DOI 10.1111/mepo.70106) whose endnote cites a Geopolitical Monitor commentary by the author, not a technical artifact. No count is stated here: the citation tally on the evidence site changes as entries are added, and it is peripheral to the engineering case. Commentary and analysis have appeared in Modern War Institute at West Point, RUSI, The Defense Post, RealClearDefense, The Space Review and Military AI. Commentary and citations are recognition of writing; they are not technical validation of any artifact.

EXTERNAL EVALUATION RECORD

A program white paper, "Conditions-Based Authority Governance and Assurance for the Virtual C2 Layer: Verified Delegation, Authority Tradeoff Analysis, and Intent-Constraint Validation," submitted under Air Force Research Laboratory Broad Agency Announcement FA8750-24-S-7003 (CANVAS), was favorably evaluated and assessed as of interest to the Air Force; budget constraints prevented funding at that time (memorandum dated 24 June 2026, AFRL/RFC, Rome Research Site). An evaluation record is not funding, validation, certification, sponsorship, or endorsement. No other government evaluation of any artifact exists.

TECHNOLOGY READINESS REGISTER

Self-assessed against the DoD nine-level TRL scale (DoD Technology Readiness Assessment guidance; software TRL descriptions). No formal technology readiness assessment has been performed by any authority. A simulation is not TRL 5 to 6; a formal model is not an operational prototype; a reference design is not a breadboard.

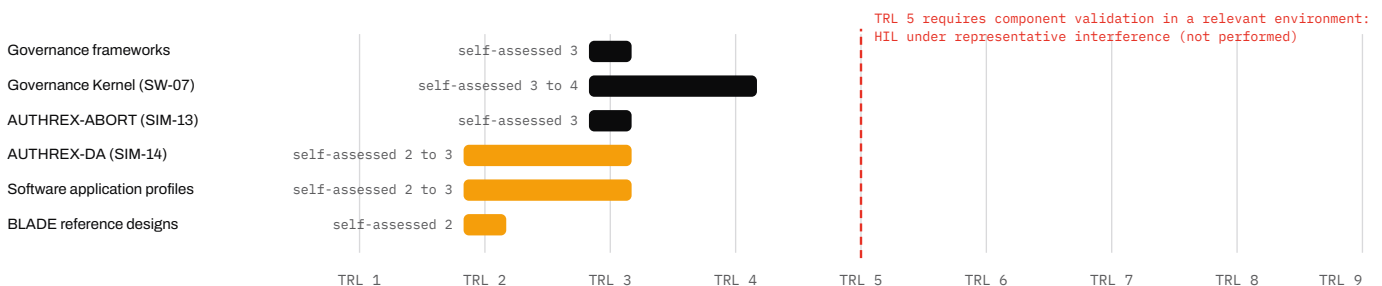


Figure 21. Self-assessed readiness by family. Near-term maturity gates differ by artifact and are listed in the register below. Hardware-in-the-loop under representative interference is a later transition gate for the platform-facing capabilities, not the next gate for every family.

Item	TRL (self)	Rationale	Missing evidence	Next milestone
SATA	3 (article states 4)	Published protocol; reference implementation; seeded simulation; formal check as published. The article self-assesses TRL 4; this register holds TRL 3 pending laboratory component validation	Laboratory component validation with attestation hardware; independent replication	HIL on ROS 2 with injected degradation

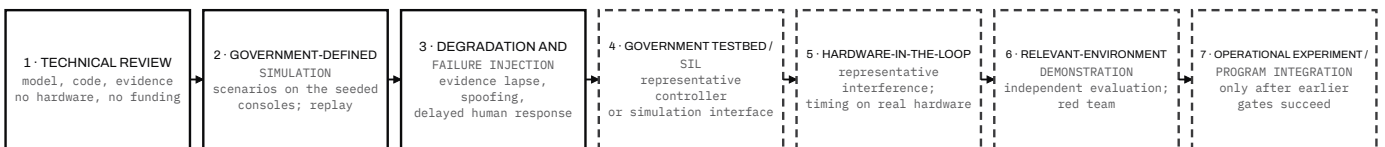
Item	TRL (self)	Rationale	Missing evidence	Next milestone
HMAA	3	TLA+ at stated bounds; 98-test package; simulation	SIL with a representative controller; closure of two vacuous properties; Obligation 6	SIL integration with a flight or vehicle stack
ADARA	3	Simulation; one faithful formal module; published counterexample	Repair shipped and checked; scenarios with ground truth	Government-defined deception scenarios
MAIVA	3	Self-tested implementation; faithful quorum model located an open defect	Repair shipped; multi-agent SIL	Repaired quorum re-check
FLAME	3	Simulation; surrogate formal analysis; stated clock requirement	Timing on a real loop; operator trial	Elapsed-time clock requirement checked in a faithful model
CARA	3	Simulation; self-run enumeration	Independent reproduction; nonce-bound token shipped	Recovery trial on SIL platform
ERAM	2 to 3	Working paper; seeded demonstrator	Validation against any real data; peer review	Peer-reviewed submission
Governance Kernel (SW-07)	3 to 4	Executable prototype with container test suite	External reproduction; security assessment; representative workload	Structured external review
Software profiles (SW-01 to 06)	2 to 3	Specifications with governor consoles	Any implementation beyond the console	Select one profile for Phase 2
AUTHREX-ABORT (SIM-13)	3	Two TLA+ modules; 105 tests; soak; mutation campaign	Refinement gate; negative controls; hazard and cyber packages	Engine-to-InterfaceRecord conformance gate and the B14 and B2 negative controls closed
AUTHREX-DA (SIM-14)	2 to 3	Local tests; single-seed root	WP-0; DA-P1 to P9; preregistered campaign	WP-0 closure
Governor consoles (SIM-08 to 12)	3 (as tools)	Deterministic consoles with reference roots	Signed releases; independent replay	Signed release of TEAM and REDLINE
BLADE-AV	2 to 3	BOM-level design; 12 attack scenarios in simulation	Bench integration of the fail-safe relay; drive-by-wire rig	Bench relay integration
BLADE-AGENT-HSM	2 to 3	Design plus adversarial emulator (275/275)	Silicon; measured timing	Prototype board
BLADE-SWARM	2 to 3	TLA+ at stated bounds; simulation at N up to 500	Any physical node	Multi-node SIL
BLADE-EDGE, AV, MARITIME, INFRA, INFRA-OT, CUAS, FINANCE, Rover, UAV	2	BOM-level reference designs; simulation data	Any physical build or measurement	First prototype assembly (BLADE-EDGE or Rover)
BLADE-SPACE	2	Preliminary design; V&V; specified, not executed	Executed V&V; orbital simulation campaign	Execute specified V&V;

Reading rule for TRL 3 versus 4 in this register: TRL 3 requires analytical and experimental proof of concept of the critical function, which simulation on synthetic data with a formal model satisfies; TRL 4 requires validation of a component in a laboratory environment, which for software means a prototype validated beyond the author's development environment against a representative interface or workload. Only the Governance Kernel approaches that threshold, and it is stated as 3 to 4 rather than 4.

TRANSITION ROADMAP

Current state, integration prototype, relevant environment, mission experiment, program integration: seven gated phases. Phases 1 to 3 require no hardware and no deployment funding.

☐ Executable now ☐ Sponsor-dependent



CAN START NOW: no hardware and no deployment funding required

REQUIRES A SPONSOR: use case, interface, testbed or range access, independent evaluator

Each phase is a gate. A phase proceeds only after the prior one succeeds and its evidence is published to the register. Solid boxes are phases the program can execute unaided today; dashed phases depend on government-furnished scenarios, interfaces, test environments, or an independent evaluator. No phase implies that any organization has agreed to sponsor it.

Figure 22. Gated evaluation path.

30 / 60 / 90 DAY PILOT STRUCTURE (PHASES 1 TO 3)

First 30 days	Days 31 to 60	Days 61 to 90
Architecture review; interface definition; government use-case selection; baseline scenario construction; integration planning	Prototype integration with a representative controller or simulation interface; government-defined tests; failure injection; metrics collection	Demonstration; independent evaluation; technical report; transition recommendation

WHAT EACH SIDE PROVIDES

GOVERNMENT, REQUIRED	Mission scenario; representative controller or simulation interface; evaluation metrics.	AUTHREX PROVIDES	Software and reference architectures; TLA+ models and enumeration artifacts; seeded consoles and test harness; assurance logic and integration support; evidence reports; source access under invited review.
GOVERNMENT, OPTIONAL	Test data; operational constraints; access to a testbed or range; an independent evaluator.	NOT OFFERED	Deployed hardware, certified software, or operational support, none of which exist.

SUCCESS METRICS A SPONSOR COULD ADOPT

Unauthorized transition prevention; authority-contraction latency; deterministic state behaviour under replay; failure-response correctness; false-restriction rate (authority reduced in a benign environment); false-negative rate; recovery behaviour; audit reconstruction success; computational overhead and share of the latency budget consumed; integration burden; operator comprehension of authority state.

TRANSITION PATHWAYS THAT FIT THE CURRENT MATURITY

Research agreement or CRADA; university research collaboration; BAA or SBIR/STTR topic response; testbed collaboration with an FFRDC, UARC, or laboratory; subcontract or prime integration for a specific platform; government pilot at Phases 1 to 3; consortium membership through an OTA. No mechanism is stated to be available or guaranteed; each depends on a sponsor identifying a use case. Eligibility for any vehicle is not asserted.

ACQUISITION-READY INFORMATION

Field	Value	Field	Value
Legal entity	Formation in progress; no entity identifiers available	UEI / CAGE	Not yet issued
SAM status	Not registered	NAICS	Not assigned
Small-business status	Not determined	Contract vehicles	None
Security clearance	Not provided	Facility clearance	None
Program staffing	One principal investigator	Teaming on award	Hardware integration, HIL, red team, certification evidence, and external formal-methods review would be teamed; named as gaps, not claimed as capabilities

GOVERNMENT ENGAGEMENT MATRIX

For every capability: the kind of organization that fits, the specific contribution needed, the evaluation proposed, and the transition objective. "We are looking for DoD partners" is not a request; these are.

Capability	Desired organization type	Government contribution needed	Proposed evaluation	Transition objective
SATA	Service laboratory (AFRL, ARL, NRL); FFRDC	Recorded sensor streams with degradation episodes; ROS 2 platform with TPM node	Phase 1 to 3 then HIL	Component validated in a laboratory environment (TRL 4)
HMAA	Formal-methods-capable laboratory; T&E; organization	Mission scenario with explicit confirm-or-abort tier; SIL flight or vehicle stack	Independent review of TLA+ and tests; SIL integration	TRL 4; policy thresholds set by a sponsor
ADARA	EW or counter-UAS test organization	Deception scenarios with ground truth; red team	Phase 2 simulation; adaptive red team	Calibrated false-restriction and miss rates
MAIVA	Swarm or multi-platform program office	Tolerable-compromise assumption; multi-vehicle SIL	Adversarial agreement evaluation	Repaired quorum validated in SIL
FLAME	C2 or engagement-chain program office; human-factors laboratory	Government decision timeline; operators for a usability trial	Latency-budget review; operator-in-the-loop	Windows usable against a named timeline
CARA	Platform safety authority (space, maritime, ICS)	Recovery policy for a named platform class	Recovery-path review; SIL lockout trial	Recovery validated in SIL
ERAM	C2 research organization; wargaming center	Releasable exercise or wargame record	Model validation against the record	Validated analytical monitor

Capability	Desired organization type	Government contribution needed	Proposed evaluation	Transition objective
Governance Kernel	Cyber test organization; DIU or service software factory	Tool inventory and approval policy; reviewer team	External reproduction; sandbox evaluation; security assessment	Independently reproduced; assessed for release
AUTHREX-ABORT	DT&E; or OT&E; organization; UAS program office	Platform safety interface specification; hazard-analysis review	Adopt the evidence-lapse pattern as a test case	Bench demonstrator with a real safety interface
AUTHREX-DA	Data-integrity or federated-learning program	Labeled records with poisoned entries	Admission evaluation	WP-0 closed; multi-seed campaign
BLADE family	Laboratory with hardware integration capacity; prime contractor	Bench, HIL rig, range access; interface specifications	Design review then prototype assembly	First platform built and measured

SEVEN LEVELS, FROM A CONVERSATION TO A PILOT

Level	What it involves	What the government receives	What is acquired
Technical exchange	30 to 60 minute architecture discussion	Answers on scope, evidence, limitations	Nothing
Use-case review	Apply the architecture to a specific mission scenario	Written mapping of the scenario to the authority lifecycle	Nothing
Independent technical assessment	Code, formal artifacts, test evidence provided for review	Source access under invited review; TLA+ models; enumeration artifacts	Access for review
Simulation evaluation	Government-defined scenarios on the seeded consoles	Deterministic, replayable results; hash-chained records	Access for review
Pilot integration	Integrate with a representative platform or simulation interface	Phases 1 to 3; technical report	Research and prototype integration engineering, not a product
Research collaboration	Advance formal methods, assurance, or mission applications	Joint artifacts; shared replication packages	Research
Industry teaming	Integrate through a prime or platform provider	Reference architectures and integration support	Reference architectures under license

INTEGRATION OPTIONS

How the capability connects. No interoperability with any named platform is claimed; the patterns describe the interface the authority logic expects.

Pattern	Applies to	What connects	Deployment	Dependencies
Policy enforcement layer	Frameworks	Between the host controller and the actuator path over the authority interface the platform exposes	Software module beside the controller	Host sensor reports; platform safety interface
Runtime wrapper / proxy	Governance Kernel, AUTHREX-AGENT	Wraps an LLM-agent runtime or MCP tool boundary; approval gating; signed evidence	Containerized; localhost in the current package	Python 3.11; Docker; MCP SDK
Middleware / boundary node	ICS-GATE, ASSURE, BLADE-INFRA-OT	Boundary between IT and OT or between assurance domains; fail-closed adjudication	Appliance (design) or service	OT protocol stack (design)
Mission-system module	BLADE platforms	Reference hardware node with defined interfaces and BOM	Embedded hardware (design only)	Platform integration partner
Hardware root of trust	SATA attestation, BLADE-AGENT-HSM	TPM 2.0 measurements and attestation chain; non-exportable keys	Hardware (design; emulator only)	TPM or secure element
Digital twin / simulation input	All consoles	Seeded scenarios with deterministic replay for T&E;	Standalone HTML; no installation; no network	A browser
Test and evaluation tool	Consoles, harnesses	Government-defined scenarios; failure injection; ledger comparison	Standalone	Sponsor scenarios
Data pipeline stage	AUTHREX-DA	Admissibility verdicts honored by downstream consumers	Standalone console (design as a stage)	Records with evidence references

MODULAR OPEN SYSTEMS

The authority logic is platform-agnostic and separated from the controller it governs; components are replaceable behind the authority interface and no retraining of the mission autonomy is required. This is conceptually aligned with the modular open systems approach; no MOSA conformance assessment has been performed and none is claimed.

DATA REQUIREMENTS

Everything in the portfolio runs on synthetic data. Phase 2 can proceed with government-defined scenarios and no data transfer. Recorded sensor streams, tool inventories, OT command sets, or exercise records at the unclassified or CUI level, as the sponsor determines, would enable Phases 3 and 4. No classified data is held or requested; no facility clearance exists.

PORTFOLIO MATURITY MATRIX

Marks: DEMONSTRATED (self-run evidence exists) · PARTIAL · PLANNED · n/a (not applicable). Only evidence in the register counts.

Capability	TRL (self)	Software	Formal model	Simulation	HIL	Relevant env.	Operational test
SATA	3	DEMONSTRATED	DEMONSTRATED (as published)	DEMONSTRATED	PLANNED	PLANNED	PLANNED
HMAA	3	DEMONSTRATED	DEMONSTRATED	DEMONSTRATED	PLANNED	PLANNED	PLANNED
ADARA	3	DEMONSTRATED	PARTIAL	DEMONSTRATED	PLANNED	PLANNED	PLANNED
MAIVA	3	DEMONSTRATED	PARTIAL (open defect)	DEMONSTRATED	PLANNED	PLANNED	PLANNED
FLAME	3	DEMONSTRATED	PARTIAL	DEMONSTRATED	PLANNED	PLANNED	PLANNED
CARA	3	DEMONSTRATED	PARTIAL	DEMONSTRATED	PLANNED	PLANNED	PLANNED
ERAM	2 to 3	DEMONSTRATED	PARTIAL	DEMONSTRATED	n/a	PLANNED	PLANNED
Governance Kernel	3 to 4	DEMONSTRATED	PLANNED	DEMONSTRATED	n/a	PLANNED	PLANNED
AUTHREX-ABORT	3	DEMONSTRATED	DEMONSTRATED	DEMONSTRATED	PLANNED	PLANNED	PLANNED
AUTHREX-DA	2 to 3	DEMONSTRATED	PLANNED	PARTIAL (1 seed)	n/a	PLANNED	PLANNED
Software profiles	2 to 3	PARTIAL (consoles)	PARTIAL (surrogates)	DEMONSTRATED	n/a	PLANNED	PLANNED
Governor consoles	3 (tools)	DEMONSTRATED	PARTIAL	DEMONSTRATED	n/a	n/a	n/a
BLADE designs	2	PARTIAL (sim data)	PARTIAL (SWARM)	DEMONSTRATED (design sims)	PLANNED	PLANNED	PLANNED

MISSION-AREA MATRIX

PRIMARY: designed for the domain and a reference design or console exists. APPLICABLE: the logic transfers with configuration. POTENTIAL: plausible, not instantiated. No entry implies deployment.

Capability	Autonomy (air/ground/sea)	PNT (as evidence)	C2	Cyber	Space	ISR	Critical infrastructure	Logistics
SATA	PRIMARY	PRIMARY	APPLICABLE	APPLICABLE	PRIMARY	APPLICABLE	APPLICABLE	POTENTIAL
HMAA	PRIMARY	APPLICABLE	PRIMARY	APPLICABLE	APPLICABLE	APPLICABLE	APPLICABLE	POTENTIAL
ADARA	PRIMARY	APPLICABLE	APPLICABLE	APPLICABLE	APPLICABLE	PRIMARY	POTENTIAL	n/a
MAIVA	APPLICABLE	n/a	APPLICABLE	APPLICABLE	APPLICABLE	APPLICABLE	POTENTIAL	n/a
FLAME	PRIMARY	n/a	PRIMARY	APPLICABLE	APPLICABLE	n/a	POTENTIAL	n/a
CARA	APPLICABLE	n/a	APPLICABLE	POTENTIAL	PRIMARY	n/a	APPLICABLE	n/a
ERAM	POTENTIAL	n/a	PRIMARY	APPLICABLE	POTENTIAL	n/a	POTENTIAL	n/a
Governance Kernel	n/a	n/a	APPLICABLE	PRIMARY	n/a	n/a	APPLICABLE	n/a
AUTHREX-ABORT	PRIMARY	PRIMARY	APPLICABLE	n/a	POTENTIAL	n/a	n/a	n/a
AUTHREX-DA	APPLICABLE	n/a	APPLICABLE	PRIMARY	n/a	APPLICABLE	APPLICABLE	n/a
BLADE designs	PRIMARY	APPLICABLE	APPLICABLE	n/a	PRIMARY	n/a	PRIMARY	n/a

TECHNICAL EVIDENCE INDEX

Capabilities mapped to supporting artifacts, with version, evidence function, and public availability. No artifact has been independently verified. DOIs and deposited identifiers resolve where linked; records that are not publicly resolvable are listed in the controlled evidence manifest at the foot of this table.

Capability	Artifact	Version / date	Evidence function	Public availability
SATA	JHSS article	10, 17 (2026)	Peer-reviewed protocol; as-published formal and Monte Carlo results	DOI 10.1007/s41635-026-00190-4
SATA	Zenodo deposit (spec, reference implementation, artifacts)	2026	Reproducibility package	DOI 10.5281/zenodo.18936251
SATA	Provisional application	64/002,453, 11 Mar 2026	Filing date	USPTO
HMAA	HMAA_verification_v2_deposit.zip	v2 (deposit of record)	TLC configuration and log for the 23,748-state result	Zenodo 10.5281/zenodo.18861653
HMAA	Python package (42 files, 98 tests)	2026	Implementation tests; deterministic traces	Zenodo deposit
HMAA	Provisional application	63/999,105, 7 Mar 2026	Filing date	USPTO
ADARA	Zenodo deposit; ADARA_XSI_EXACT.tla	2026	Faithful model of the inconsistency function; counterexample	DOI 10.5281/zenodo.19043924; authrex.systems/formal-coverage
MAIVA	MIT deposit v5.18; MAIVA_QUORUM.tla; 37 self-tests	v5.18	Faithful model; located open defect	DOI 10.5281/zenodo.19015517
FLAME	Zenodo deposit v5.11; FLAME.tla	v5.11	Specification; surrogate counterexample	DOI 10.5281/zenodo.19015618
CARA	Zenodo deposit; enumeration artifacts	2026	Self-run enumeration (10,000,000 iterations; 68-state control flow)	DOI 10.5281/zenodo.18917790
ERAM	SSRN working paper; ERAM Strategic Command console	6176802	Analytical model; 600-run demonstrator	SSRN; authrex.systems
Governance Kernel	Review package (authrex_kernel/, tests/, harness/, docs, Dockerfile, VALIDATION-REPORT.md)	v1.0-QA10, Jul 2026	Container test suite; demos; ledger tamper detection	By request, invited review only (counsel review pending)
AUTHREX-ABORT	Build 0.6 package: console, AuthorityLapse.tla v3, InterfaceRecord.tla v1, Stage 1B sets, harnesses	Build 0.6; engine sha256 ce3af3dd...	Formal checks; 105 tests; soak; mutation campaign	authrex.systems (console); package hashes published
AUTHREX-DA	Release 1.8 package; 17-artifact assurance package	release 1.8, engine 1.0	101 tests; gates; canonical root; NEG-1 to 5	authrex.systems
SAFING	AUTHREX-SAFING-Release-v1.5.0.zip	v1.5.0; sha256 ef8b828e...	1061 checks; model check; reference roots	authrex.systems
TEAM / REDLINE	Standalone consoles	v0.16 / v0.13	Scenario suites; reference ledger roots	authrex.systems (no signed release)
RTA / ENGAGE	Governor consoles	2026	Five stress / five seeded scenarios	authrex.systems
SIM-01 to 07	Flagship browser simulations	2026	Integrated pipeline demonstrations	authrex.systems
BLADE-01 to 12	Zenodo deposits with BOM, CONFIG, ELECTRICAL, MECHANICAL, GUIDE, schematic, paper	2026	Reference designs; simulation data; HSM test-report.json; SWARM TLA+ and verification report	DOIs per platform (BLADE register)
Formal coverage	Twenty TLA+ modules; sealed results; 62-configuration rerun	2026	Faithfulness taxonomy; counterexamples; withdrawals	authrex.systems/formal-coverage
Evaluation protocol	evaluation.html	2026	Scenarios S1 to S7; metrics; baselines; assumptions	burakoktenli.com/evaluation
Program record	Catalog V5 and catalog.csv (44 rows)	1 Sep 2026	Register of record (superseded by this edition)	authrex.systems/catalog
EBCC	Provisional application	64/144,399, 30 Aug 2026	Filing date (program-level)	USPTO
External evaluation	AFRL/RFC memorandum	24 Jun 2026	Favorable white-paper evaluation under FA8750-24-S-7003	On file with the author; not public (EVID-GOV-01)
Controlled evidence manifest	AUTHREX-Catalog-V6.1-Evidence-Manifest.md	4 Sep 2026	Lists every record that supports a claim in this catalog but is not publicly resolvable (USPTO receipts, conference acceptances, the AFRL memorandum), with its status	Accompanies this catalog; the underlying records are held in the program evidence file and released on request
Site synchronization worksheet	AUTHREX-Site-Synchronization-Worksheet.md (Revision B)	4 Sep 2026	Site-wide string sweep and seventeen exact replacements for the public statements this catalog has withdrawn or reworded	Accompanies this catalog

Capability	Artifact	Version / date	Evidence function	Public availability
Controlled vocabulary standard	AUTHREX-Controlled-Vocabulary-Standard.md	4 Sep 2026	The word rules, canonical figures and safe sentence patterns derived from the claims register; governs every public statement made under the program name	Accompanies this catalog

CLAIMS REGISTER

Every important quantitative claim in this catalog, its supporting artifact, and its evidence class. A claim not in this register is not asserted.

ID	Capability	Claim	Supporting artifact	Class	Caveat
C-01	HMAA	23,748 distinct reachable states at depth 9; 26,397,356 generated	HMAA_verification_v2_deposit.zip; formal-coverage; catalog.csv	VERIFIED (bounded, model-relative)	Never stated without the scope clause
C-02	HMAA	5 invariants and 1 liveness held; 2 vacuous	same	VERIFIED (bounded)	"held", never "verified" for individual properties
C-03	HMAA	42-file package; 98 tests; 7 adversarial experiments	Zenodo 18861653; patents page	TESTED (self-run)	Not independently executed
C-04	SATA	18,892 states; zero violations of 8 invariants; 10,000-run Monte Carlo	JHSS article; Zenodo 18936251	TESTED / MODELED (as published)	Program formal-coverage lists no faithful SATA module
C-05	SATA	523 distinct states at depth 9 (archived run)	formal-coverage.html	WITHDRAWN	Artifacts not producible; CL-02
C-06	MAIVA	Q_NEVER_FAILS held across 12,507,501 distinct states	formal-coverage	VERIFIED (narrowed sense)	Only beside the n = 1 quorum finding
C-07	MAIVA	37 self-tests; v5.18; MIT	Zenodo 19015517	TESTED	
C-08	CARA	10,000,000-iteration enumeration; 68-state control flow	Zenodo 18917790	TESTED (self-run site claim)	Not independently reproduced
C-09	ERAM	6 scenarios; 600 Monte Carlo runs	ERAM Strategic Command console	TESTED	Not validated against real data
C-10	Kernel	83 passed, 1 skipped; 53 test functions; python:3.11-slim; v1.0-QA10	kernel.html; internal QA log	TESTED	Site claim pending external reproduction
C-11	ABORT	AuthorityLapse v3 across 21 configurations; InterfaceRecord v1 across 35	Build 0.6 package	VERIFIED (bounded)	Every property in isolation
C-12	ABORT	Stage 1B 31,341 states; 130,869 transitions	TLC export; hashes published	VERIFIED (bounded)	
C-13	ABORT	Stage 1 availability relation 972 of 972 projected states, AvailRaw only	Build 0.6 package	VERIFIED (projection)	Never described as the complete reachable graph
C-14	ABORT	105 implementation tests; 34 release checks	Build 0.6 package	TESTED	
C-15	ABORT	1,000,000-run soak; zero failures; separate environment	Build 0.6 package	TESTED	Not independent verification
C-16	ABORT	7 valid mutants; 4 detected; B12 closed; B14, B2 open	authrex-abort page	TESTED (negative control)	Mechanical score 0.57 is not the result
C-17	DA	101 of 101 tests; 12 gates; 14 requirements (12/1/1)	release 1.8	TESTED	Local only
C-18	DA	Canonical root: seed 20260828, 120 rounds, 1,584 entries	release 1.8	TESTED	Single platform
C-19	DA	19 audit passes (5 internal, 14 external)	release 1.8	IMPLEMENTED (review record)	Not independent reproduction
C-20	SAFING	1061 checks; 2,488,320 executions, 15 states, 23 transitions, depth 7, 9 properties; 11,739 decisions; 50/50 mutants; 43 adversarial checks	SAFING v1.5.0 release	TESTED / VERIFIED (bounded)	
C-21	TEAM	20 of 20 scenarios; root 5a937676...	console v0.16	TESTED	No signed release
C-22	REDLINE	27 of 27; 250 seeds x 25,000 ticks; 843 fail-safe holds	console v0.13	TESTED	No signed release
C-23	Formal coverage	2 of 20 modules FAITHFUL; 11 surrogate; 3 withdrawn	formal-coverage	VERIFIED (taxonomy)	Surrogates say nothing about code
C-24	Formal coverage	62 configurations re-executed; 0 differences	formal-coverage	TESTED	Same producer, second environment
C-25	BLADE-AGENT-H SM	275 of 275 emulator checks; 3 reruns	test-report.json v2.6	TESTED (emulator)	Silicon not built

ID	Capability	Claim	Supporting artifact	Class	Caveat
C-26	BLADE-SWARM	TLA+ 5 safety, 3 liveness; N = 10, 50, 500	verification report v1.0; deposit	VERIFIED (bounded) / TESTED	Simulation outcomes
C-27	BLADE-AV	12 attack scenarios in simulation; 62 components; about \$16.3K	Zenodo 19232130	TESTED (sim) / PROPOSED	No vehicle integration
C-28	Rover / UAV	7 fault scenarios; 5 adversarial scenarios	deposits	TESTED (sim)	Run counts 350 / 250 withdrawn
C-29	Program	8 U.S. provisionals reported as filed	patents page; receipts	IMPLEMENTED (record)	Unexamined
C-30	Program	3 peer-reviewed journal articles; 48 DOI-registered works (31 Zenodo, 17 SSRN)	publications page (1 Sep 2026)	IMPLEMENTED (record)	One article is not an AUTHREX technical paper
C-31	Program	AFRL CANVAS favorable evaluation, 24 Jun 2026	memorandum on file	IMPLEMENTED (record)	Not funding or endorsement
C-32	Program	Hysteresis 5 to 15 s; about 100 Hz loop; sub-200 ms failover (SPACE)	framework pages; deposits	PROPOSED (design parameters)	Not measured on hardware
C-33	Program	Public-estate simulation counts at the recorded baseline: 14 controlled SIM records; 20 standalone consoles of 24 site-catalogued simulations; 37 launchable browser simulation pages across the two sites	Site synchronization worksheet Revision B; frozen public-estate inventory with the deploy identifier for each property	IMPLEMENTED (record)	Each figure counts a different object; snapshot-specific and revised whenever the estate changes; not performance evidence. The 14-record register is the controlled technical count. The 22-figure tile on the site index is superseded and is replaced under FIX-14.

CORRECTION LOG

Contradictions, stale claims, unsupported statements, and inconsistencies discovered while preparing this edition, with their disposition. Live pages are corrected in place; DOI-registered deposits keep the original verbatim with a dated erratum beside it.

ID	Item	Finding	Disposition in this edition	Status
CL-01	HMAA model-checking figure	Earlier figure of 48,751 states appears in the June 2026 two-page summary and older pages	Withdrawn; canonical figure is 23,748 distinct states at depth 9 with its scope clause, never stated without it	Applied throughout this edition
CL-02	SATA archived TLC run	formal-coverage.html still states a separate archived result of 523 distinct states at depth 9; the specification, configuration and log for that run could not be produced on request	Withdrawn from this catalog pending recovery of the artifacts; the as-published article figure (18,892 states) is cited only as published	Page fix pending on authrex.systems
CL-03	Rover 350-run and UAV 250-run counts	Withdrawn 7 Aug 2026 as unreproducible; still present on burakoktenli.com pages (evaluation.html, platform pages) and in the July white paper and two BAA white papers	Replaced by "seven fault scenarios" and "five adversarial scenarios" demonstrated in simulation; the aggregate "2,800+ runs" is not carried	Site pages pending; this edition clean
CL-04	"Physically assembled" testbeds	July 2026 white paper describes the rover and UAV testbeds as physically assembled; the register records no platform as built	This edition states not built, physical build not yet documented	White paper superseded
CL-05	Validation language on live pages	"Simulation Validated", "Validated in rover testbed", "Six Validated Scenarios", "ALL VALIDATIONS PASSED" appeared on framework and platform pages	Removed from the 1 Sep 2026 deploy; this edition uses demonstrated in simulation	Closed on site; carried here for the record
CL-06	Simulation count: 24 versus 14 versus 37	V5 says twenty-four seeded simulations (counting governor consoles); the register has 14 SIM rows; the sites state 37 browser simulations (19 plus 18)	This edition uses 14 catalogued records (SIM-01 to 14) and states the other counts with their definitions; never interchangeable	Applied
CL-07	Name expansions	HMAA (Allocation, Architecture, Hierarchical Mission Authority Architecture); SATA (Attestation and Trust Anchoring); ADARA (Risk Architecture); MAIVA (Verification); FLAME (Flash War Latency Architecture); CARA (Control Authority Regulation; Cognitive Authority Recovery in the Zenodo title)	Canonical program names used; variants listed on each capability page; original artifact titles preserved	Applied

ID	Item	Finding	Disposition in this edition	Status
CL-08	Provisional receipts	V5 lists the three July 2026 filings as receipt pending; the patents page (verified 1 Sep 2026) carries receipts 78658472, 78659350, 78660113	Receipts carried in this edition	Applied
CL-09	Peer-reviewed article count	V5 states one peer-reviewed article; two further articles published 24 Aug and 31 Aug 2026	Three stated, with the caveat that one is not an AUTHREX technical paper	Applied
CL-10	BLADE-UUV and domain counts	The July white paper and the authrex.systems index list BLADE-UUV; the register has no such entry; domain counts vary (nine, ten, eleven) across properties	BLADE-UUV not carried (no deposit, no register row); the catalog counts twelve register platforms and does not state a domain count	Owner decision on BLADE-UUV pending
CL-11	Standards-mapping overclaims	standards-mapping.html states "no unsafe state reachable from any initial state under any sensor input" and describes AUTHREX as "required by" MIL-HDBK-516C sections	Not carried; this edition uses mapped by the author, design target, architecturally relevant	Page fix pending on burakoktenli.com
CL-12	Superseded governance formulations	June 2026 wording: scores the evidence; grants authority in proportion; narrows in proportion; so the human stays in command; "narrows, and does not widen on its own"	Replaced by the per-source evidence-vector, non-compensatory-gate, admissible-set formulation; "narrows at once; widens only by authorization"	Applied
CL-13	Absolute hold rule	Earlier form: every irreversible action requires confirmation and cannot execute without it	Policy-conditioned form: where policy requires confirmation, execution occurs only after it; if absent at hold expiry the predeclared outcome is delay, handoff, or abort; pre-authorized autonomous execution must already exist in the envelope; timeout never creates authority	Applied
CL-14	First-in-field claims in repository READMEs	repo-flame and repo-adara reproduce README claims of "the first technical architecture" and "the first proactive deception prior"	Not carried; unfalsifiable positives are excluded	Owner to soften at source
CL-15	FPGA present-tense hardware claims	heilmeier.html states sensor trust is computed in hardware and that an FPGA implementation holds 5 ms end to end; no hardware exists	Not carried	Owner decision pending (CONFLICT-LOG CL-01)
CL-16	ABORT open-gate summary wording	V6 draft entry claimed authrex-abort.html overstates the mutation register; on re-reading, the page scopes the sentence to two properties ("for two properties, B14 and B2, the test suite currently cannot distinguish the accepted engine from a deliberately broken one")	Entry withdrawn; no page fix required; this edition states B12 closed, B14 and B2 open	Closed (withdrawn)
CL-17	Kernel register wording	catalog.html register JSON reads "1 liveness verified"	Canonical wording "held" used throughout	Page fix pending
CL-18	Program white paper figures	July 2026 white paper: seven provisionals, 40 DOI-registered outputs, nine domains, mutation testing 41 mutants against a 386-test suite, international consultation entries	Superseded counts updated (eight; 48); items without a register artifact are not carried	White paper marked superseded
CL-19	External evaluation records	Earlier program materials referred to government submissions in general terms that could be read as evaluation support	Only the AFRL CANVAS favorable evaluation memorandum (24 Jun 2026) is recorded, in its own terms; no other submission, review, or scoring is cited as evaluation support	Applied
CL-20	Georgetown affiliation form	Some locked simulation strings carry an unqualified "MPS Applied Intelligence" byline	This edition uses the in-progress form only	Applied
CL-21	HMAA tier names	V6 draft carried the archived V4 tier names (full, standard, restricted, safe); the current HMAA page defines A3 Full Autonomy, A2 Restricted, A1 Minimal, A0 Revoked	Current names used throughout; V4 names recorded as historical	Applied (red team)
CL-22	Simulation size misattributed	V6 draft gave ERAM Strategic Command a size of about 86 KB; the index attaches 86 KB to AUTHREX-TEAM	Size removed for SIM-01; sizes retained only where the index states them	Applied (red team)
CL-23	Unsupported statements removed	V6 draft stated about 104 media placements, a TLC version of 2.19, ABORT gate identifiers G02-C and G09, an AFRL debriefing statement, an NDIA poster slot, and three AUTHREX-AGENT reference use cases; none is in the source record (the agent page traces four)	All removed or corrected to the source wording	Applied (red team)
CL-24	Hash fragment typos	V6 draft wrote the TEAM standalone sha256 as 79707655... and the AuthorityLapse v3 sha256 as ...bfce2993; the published values are 797076b5...ea01399e and ea6c4fff...fbce2993	Corrected against the published hashes	Applied (red team)
CL-25	Software profile count on the cover	V6 draft cover stated seven software application profiles while counting the Governance Kernel separately (six profiles plus the Kernel is the register)	Cover states six profiles and one kernel	Applied (red team)
CL-26	SATA readiness discrepancy	The published SATA article states TRL 4 on its own reading; this catalog assesses TRL 3	Difference disclosed in the SATA profile and the TRL register rather than reconciled silently	Applied (red team)
CL-27	Figure numbering and document identifiers	V6 draft numbered figures out of document order and reused AUTHREX-CAT-2026-003 unsuffixed on companion documents	Figures numbered in order of appearance; companions carry -EXEC, -MATRIX, -SHEET- and -REG suffixes	Applied (red team)

ID	Item	Finding	Disposition in this edition	Status
CL-28	Package configuration control	The V6.1 package still carried V6 filenames, PDF Subject metadata and workbook headings while the pages read V6.1	Every deliverable now reads V6.1 in filename, document identifier, PDF Title, Subject and Keywords, page footer and workbook README	Applied (second-opinion pass)
CL-29	Identifier resolvability	Page 2 stated that DOIs, application numbers and hashes are resolvable public records, which overstates the status of provisional-application numbers	DOIs and deposited identifiers are resolvable where linked and hashes recomputable; provisional numbers, dates, entity status and receipts are supported by filing receipts in the program evidence file and may not be publicly searchable	Applied (second-opinion pass)
CL-30	Provisional status wording	Seven rows read "Provisional; pending USPTO action; unexamined", which implies substantive prosecution that a provisional does not undergo	All eight rows read "Provisional application reported as filed under 35 U.S.C. 111(b); unexamined; no patent granted; follow-on filing decision pending"	Applied (second-opinion pass)
CL-31	TRL figure caption	The readiness figure caption said hardware-in-the-loop is the next gate for every family, contradicting the register beneath it (peer review for ERAM, external review for the Kernel, WP-0 for DA, SIL for HMAA, quorum repair for MAIVA)	Caption states that near-term gates differ by artifact and that HIL is a later transition gate for the platform-facing capabilities	Applied (second-opinion pass)
CL-32	External citation count	The catalog carried the site figure of twenty-two external sources across fourteen countries; the site record shows that figure has moved (13 across nine, then 21 across nine, then 22 across fourteen)	No count is stated; the sentence records that external sources have cited or discussed the work and that this is recognition of writing, not technical validation	Applied (second-opinion pass)
CL-33	Standards rows split	ISO 26262 and SAE J3016 were grouped as design targets, and ASTM F3269-21 was cited by a shortened description	ISO 26262 is a design target; SAE J3016 is listed separately as a terminology reference carrying no safety or conformity claim; ASTM F3269-21 is cited by its official title	Applied (second-opinion pass)
CL-34	Non-public evidence traceability	USPTO receipts, the AIAA and NDIA acceptance records and the AFRL memorandum support claims in this catalog but are not in the public package, so no reviewer can certify them from the catalog alone	A controlled evidence manifest is issued with the catalog listing each record, the claim it supports and its status; the records themselves remain in the program evidence file and are released on request	Applied (second-opinion pass); the records themselves are an owner action
CL-35	Patent deadline statement	The IP page stated approximate twelve-month non-provisional windows and referred to one-year disclosure bars preceding two filings, which is a legal conclusion that depends on each filing and disclosure history	Replaced: filing and priority deadlines and any disclosure-bar analysis are maintained in the controlled IP record and are to be confirmed with patent counsel; no deadline conclusion is stated and no patent grant is claimed	Applied (third-pass review)
CL-36	Simulation-count definitions	Four public figures exist (14 register records; 20 consoles of 24 catalogued in the site simulation tree; 22 catalogued simulations in the site evidence tile; 37 browser simulations across both sites) and the catalog reconciled only three of them	All four are stated with what each counts; the catalog uses the 14-record register only	Applied (third-pass review); the site should state the same reconciliation
CL-37	Site-wide synchronization scope	The synchronization worksheet covered seven statements; a third-pass review of the live sites found further contradictions on heilmeier.html (mathematical-proof and regardless-of-inputs wording, rover and UAV at TRL 4, DoDD 3000.09 compliance framing, 37 simulations "show the approach works", drop-in integration), testbed.html and index pages ("implemented experimental system", "physical research platforms"), and standards-mapping.html (direct DO-333 evidence, technical compliance mechanism)	The worksheet is reissued as a site-wide sweep with a string checklist and exact replacement wording for each item; every occurrence is to be reconciled against this catalog's claims register before distribution	Open: owner action on the public sites
CL-38	Controlled vocabulary for future public statements	The catalog holds the strictest language in the estate, but nothing governed what a new site page or repository README could say, so the estate could drift back to validated, proved, TRL 4 or hardware-present-tense wording after a synchronization	A controlled vocabulary standard (AUTHREX-CAT-2026-003-VOCAB) is issued with this catalog: banned and fixed terms, canonical figures, required companions, safe sentence patterns, and a pre-publish check requiring a claims-register row before any public claim	Applied (fourth-pass review); adoption on the sites is an owner action
CL-39	Companion document metadata	The ten capability sheets carried a document-management title that did not name the capability or the version in PDF metadata	Each sheet carries its code, full capability name and version in Title, its identifier and evidence context in Subject, and its code, identifier and self-assessed TRL in Keywords	Applied (fourth-pass review)
CL-40	Audit document staleness	The red-team audit stated 62 pages, ended its page table at 62 with the back cover, described the synchronization worksheet as a six-step verification and referred to six site fixes, all of which the 63-page build and Revision B had superseded	The page-by-page section is regenerated from the built PDF rather than patched, the step count and the seventeen worked replacements are stated correctly, and the audit carries its own version and build reference	Applied (fifth-pass review)
CL-41	Freeze baseline incomplete	Revision B step 7 froze seventeen artifacts and omitted the red-team audit, which is distributed with the package and had itself gone stale	Step 7 names all eighteen artifacts including the audit, and records a deploy identifier for each public property rather than a single site commit	Applied (fifth-pass review)

ID	Item	Finding	Disposition in this edition	Status
CL-42	Public counts without a register row	The controlled vocabulary standard requires a claims-register row for any public claim, while the synchronization worksheet instructed the sites to publish simulation counts that had no row, making the worksheet an exception to the rule it was meant to enforce	Claim C-33 records the public-estate counts with their supporting artifact, evidence class and caveats; FIX-13 and FIX-14 now comply with the rule rather than sitting outside it	Applied (fifth-pass review)
CL-43	Banned-word rule too broad	Section 1.1 of the vocabulary standard banned words including verified and deployed and instructed that every hit be replaced or removed, which would have deleted the negative disclaimers the catalog depends on and the defined uppercase evidence-class label VERIFIED	The prohibition is scoped to unqualified affirmative claims about an AUTHREX artifact; negative disclaimers, quoted historical language in a correction record and the defined uppercase label are permitted in their bounded meanings; the pre-publish step reads review every hit rather than remove every hit	Applied (fifth-pass review)
CL-44	Overlap-scan claim not reproducible	The audit stated that a programmatic scan of every word box reported zero overlaps, without naming the extractor or the intersection definition, so another reviewer using different word rectangles could reach a different result	The audit states the scanner, the extraction method and the intersection threshold, separates the source-layout scan result from rendered inspection of all thirteen PDFs, and describes bounding-box intersections from adjacent glyphs as artifacts rather than defects	Applied (fifth-pass review)
CL-45	Audit freeze-baseline wording	The red-team audit described the audited build as accompanied by twelve companion artifacts listed in the freeze baseline; twelve is the count of companion PDFs, while the controlled baseline is eighteen artifacts	The audit states the seventeen other artifacts in the eighteen-artifact freeze baseline	Applied (sixth-pass review)
CL-46	Stale six-statement framing	The audit executive verdict still said the public sites carry six statements the catalog has withdrawn, a count superseded when the worksheet became a site-wide sweep with seventeen worked replacements and an instruction to resolve every hit	The sentence states that the estate contains multiple conflicting or superseded statements, counts and status formulations, with no fixed number	Applied (sixth-pass review)
CL-47	Score arithmetic	The audit said three points were withheld beneath a table scoring 90 of 100; ten points are withheld across five dimensions (technical credibility 2, evidence discipline 1, mission relevance 2, acquisition relevance 2, test and evaluation maturity 3)	The paragraph states ten points across the five named dimensions and what evidence each requires	Applied (sixth-pass review)
CL-48	Page-map range boundary	The regenerated page map ran the document-control range to the page where the next section begins, even where that section starts at the top of the page and no earlier content appears on it	The generator now ends a range at the page before the next section when that section begins at the top of a page, and carries the shared page only where a section genuinely continues onto it	Applied (sixth-pass review)
CL-49	Workbook README claim count	The registers workbook README described the claims sheet as 32 quantitative claims after C-33 was added, while the formula-driven summary correctly reported 33	The README reads 33 registered claims, a wording that also fits C-33, which is a configuration baseline rather than a technical measurement	Applied (seventh-pass review)
CL-50	Audit closing-section pass count	The red-team audit front matter recorded six review rounds while its closing section still said five passes had been run and that CL-40 to CL-44 were the last	The closing section states seven passes and names the fifth-pass, sixth-pass and seventh-pass findings; it is the same defect class CL-40 exists to prevent	Applied (seventh-pass review)
CL-51	Page-map boundary rule incomplete	The corrected range rule caught the document-control boundary but missed pages where a capability profile begins under its own bracket label, so ADARA, ERAM, the Governance Kernel, AUTHREX-ABORT and the glossary each ran one page into a section that owns its page outright	The generator skips leading bracket labels when deciding whether a section starts at the top of a page, so the rule applies to every boundary rather than only unlabeled ones	Applied (seventh-pass review)
CL-52	Orphaned heading and near-empty page	The BLADE platform-profiles label and heading sat alone on a near-empty page while the profiles began on the next, contradicting the catalog build rule that a heading is kept with its first block	The layout rule no longer nests a keep-together block inside another, which had made the group unmeasurable and pushed it to a fresh page; a bracket label above a heading is now carried with it. The heading, the caption and the first profiles share the page	Applied (seventh-pass review)
CL-53	Second stranded block found by a new check	A sparse-page check added during the seventh pass found the language-discipline block alone on an otherwise blank page, the same class of defect as the BLADE heading orphan but in the front matter	The document-control section is reordered so the language rules precede the evidence-class table, and the definition table is kept whole rather than splitting; the sparse-page check is now part of the build verification and reports no page under twelve body lines	Applied (seventh-pass review)
CL-54	Correction-log self-consistency	After the seventh-pass rebuild, the disposition recorded for CL-50 still said the audit closing section contained six passes and named only the fifth-pass and sixth-pass findings, while the rebuilt report correctly contained seven and named all three; the finding column was historical, but the disposition column was one revision behind	CL-50 carries the current seven-pass wording, and the audit, catalog and workbook are regenerated from one frozen baseline. A cross-document check now compares the pass count in the audit closing section against the disposition recorded here on every rebuild, because a disposition that quotes a moving number is the one part of a correction log that can silently go stale	Applied (final mechanical verification)

GLOSSARY OF PROGRAM TERMINOLOGY

Term	Canonical usage in this catalog
Authority (runtime)	The set of actions a system is permitted to execute at this moment, on this evidence, under its mission authorization and policy. Distinct from capability, which is what it can technically do.
Admissible set	The actions that every required evidence condition currently admits. Narrows at once on failure; widens only by an authorized transition.
Non-compensatory gate	A condition that must be independently satisfied; a strong result on another condition never offsets its failure.
Pre-authorized envelope	The bounded authority a human grants before execution, frozen for the run; the system may preserve or contract it and cannot widen it.
Hard veto	An evidence source whose failure prohibits dependent actions immediately (navigation, geographic containment).
Latch	Once an irreversible action lapses it cannot return during the run, even if the evidence recovers.
Held (property)	Satisfied by the model checker across the stated configurations and bounds. Never "verified" or "proven" in this catalog.
Vacuous (property)	Not exercised at the checked bound; not established.
FAITHFUL / SURROGATE	A formal model is FAITHFUL only if it imports and checks the implemented artifact; a SURROGATE COUNTERMODEL is a finding about a mechanism class, not about the code.
Naive form / strong form	Each safety claim written twice: over the system's own bookkeeping and over ground truth the system cannot access; the gap is the finding.
Demonstrated	Shown in simulation on synthetic data unless hardware is named. Not validation.
Second-environment confirmation	Re-execution by the same producer on a different machine; a real but weaker claim than independent verification.
Independent	Reproduced or technically verified by another party. No artifact in this catalog has this status.
Reported as filed	A provisional application whose receipt exists; unexamined; no enforceable rights.
Withdrawn	A figure no longer asserted as current; recorded beside its original so it stays discoverable. Not a publication retraction.
TRL (self-assessed)	The author's reading of the DoD nine-level scale against the evidence; no formal technology readiness assessment has been performed.
EXECUTE / DELAY / HANDOFF / ABORT	The canonical outcome set, issued by HMAA; HOLD is an ENGAGE extension; inhibit and divert are contingencies beneath ABORT.
A3 to A0; T3 to T0	Authority tiers for platforms (A: full autonomy, restricted, minimal, revoked, with design thresholds 0.80, 0.55, 0.30 on the fused trust scalar) and agent runtimes (T3 to T0). The archived V4 names (full, standard, restricted, safe) are not used.
GREP	CARA recovery phases: Guard, Reduce, Evaluate, Promote.
BLADE	Designator for the twelve hardware reference designs; letter expansion under owner review.
AUTHREX	Program name; no letter expansion.

ACRONYMS

Acronym	Expansion	Acronym	Expansion
AAL	Army Applications Laboratory	JADC2	joint all-domain command and control
AFRL	Air Force Research Laboratory	LLM	large language model
ATO / cATO	authorization (continuous authorization) to operate	MCP	Model Context Protocol
BAA	broad agency announcement	MOSA	modular open systems approach
BOM	bill of materials	OTA	other transaction agreement
C2	command and control	PBFT	practical Byzantine fault tolerance
COP	common operating picture	PNT	positioning, navigation and timing
CRADA	cooperative research and development agreement	PRNG	pseudo-random number generator
CUI	controlled unclassified information	RBAC	role-based access control
CUSUM	cumulative sum control chart	RMF	Risk Management Framework
C-UAS / C-sUAS	counter uncrewed (small uncrewed) aircraft systems	RTA	run-time assurance
DIU	Defense Innovation Unit	SBIR / STTR	Small Business Innovation Research / Technology Transfer
DT&E; / OT&E;	developmental / operational test and evaluation	SBOM	software bill of materials
FFRDC	federally funded research and development center	SLTT	state, local, tribal and territorial
GNSS	global navigation satellite system	TLA+ / TLC	Temporal Logic of Actions specification language / its model checker
HIL / SIL	hardware- / software-in-the-loop	TPM	trusted platform module
HKDF	HMAC-based key derivation function	TRL / TRA	technology readiness level / assessment
HSM	hardware security module	UARC	university affiliated research center
ICS / OT	industrial control systems / operational technology	UAS	uncrewed aircraft system
ISR	intelligence, surveillance and reconnaissance	V&V;	verification and validation
IV&V;	independent verification and validation		

[ENGAGE]

GOVERN AUTONOMY AT RUNTIME.

Burak Oktenli, MBA · Principal Investigator

Independent Researcher, AUTHREX Research Program, Washington, DC

authrex.systems · burakoktenli.com · info@burakoktenli.com · ORCID 0009-0001-8573-1667

TECHNICAL EXCHANGE → SCENARIO MAPPING → SOURCE ACCESS UNDER INVITED REVIEW → SEEDDED CONSOLES →

[DISCLAIMERS AND STATUS LABELS]

This catalog presents public-source, self-authored research artifacts. Every result is self-reported and has not been independently reproduced or technically verified by another party. Evaluation is simulation on synthetic data unless a physical test is specifically identified; none is. Hardware references are design specifications; no platform has been built. Formal results are bounded and model-relative; model checking of a configured finite model is not a proof of any implementation or of the physical system. Provisional patent applications are unexamined filings that confer no enforceable rights. Standards are mapped by the author or named as design targets; no compliance, certification, accreditation, or authorization to operate is claimed or implied.

The catalog does not represent official government endorsement, sponsorship, evaluation (beyond the AFRL memorandum described inside), or affiliation. It is not a government publication and carries no DoD distribution statement. Department of Defense is the statutory name used here. Distribution: public. License: CC BY 4.0.

Best read alongside the per-artifact deposits, which contain the full bills of materials, specifications, simulation data, TLA+ modules, and source code, and alongside the correction log, which is the record of what this edition changed and what remains pending.

[STANDING OFFER]

A 30 to 60 minute technical exchange; a written mapping of a government scenario to the authority lifecycle; source access under invited review; government-defined scenarios on the seeded consoles. No mechanism is stated to be available or guaranteed.