

ENTERPRISE AI WORKFLOW RELIABILITY STANDARD · 2026

Reliable AI workflows do not hide uncertainty.

A public operating standard for source-backed, rule-tested, human-governed enterprise AI workflows.

Core principle. A workflow is not reliable because its output looks complete. It is reliable when it can show what it knows, what it does not know, where critical values came from, which rules were applied, who owns authority, and when the system must stop.

Four states. One explicit authority model.

Human oversight is not “review everything.” The workflow selects a defined state from evidence, risk and authority boundaries.

AUTO_PROCEED

Low-risk and verified

Source-backed inputs and deterministic checks pass inside a pre-authorized boundary.

HUMAN_REVIEW

Ambiguous or exceptional

Progress pauses while evidence, duplicate conditions or non-consequential exceptions are inspected.

HUMAN_APPROVAL

Consequential action

The system prepares the action, but an authorized human retains execution authority.

STOP_ESCALATE

Evidence or policy gap

Critical missing/conflicting evidence or lack of a safe rule blocks progression.

Human handoff is a designed control surface. It is not positioned as automation failure; it is how accountability remains explicit when risk or uncertainty rises.

Reliability is a system, not a model score.

Each control has a minimum evidence requirement and a fail condition.

R1**Source integrity & provenance**

Critical fields are traceable to an approved source or remain explicitly unknown.

FAIL: unsupported critical value is presented as fact.

R2**Unknown preservation**

Absence of evidence survives without plausible substitution.

FAIL: missing evidence becomes fact.

R3**Structured output & parse gate**

Machine-consumed output passes a declared schema and bounded retry policy.

FAIL: invalid output reaches downstream systems.

R4**Deterministic validation**

Rules that can be determined by deterministic rules.

FAIL: model judgment replaces analysis.

R5**Decision authority & human gate**

Consequential actions expose an owner and explicit approve / edit / reject / escalate control. External write remains unchanged and passes.

FAIL: consequential action bypasses the required authority boundary.

Control the capability, not just the prompt.

Reliable behavior depends on permission boundaries, logging, regression evidence and an explicit stop/recovery model.

R6

Least agency & tool boundary

AI components receive only the minimum functions, permissions and write scope required.

FAIL: unnecessary write/delete/broad capability is available.

R7

Observability & decision log

Source, extraction, rule result, state, human action and final disposition can be reconstructed.

FAIL: the team cannot explain why the workflow advanced or stopped.

R8

Regression & adversarial testing

Normal, malformed and adversarial cases are versioned with expected outcomes and visible failures.

FAIL: behavior changes without rerunnable evidence.

R9

Incident stop & recovery

Stop triggers, incident ownership, containment and recovery evidence are defined in advance.

FAIL: execution continues after a known critical control failure.

R10

Approved-result measurement

Success is measured at the approved business result: time, critical error, rework, escalation and cost.

FAIL: automation percentage is the only success metric.

Evidence is labeled by what it actually proves.

Current results are synthetic and reproducible. They demonstrate control behavior - not production accuracy, compliance or ROI.

10/10

deterministic cases

5/5

critical rule families

10/10

structured extraction

10/10

ingestion cases

50/50

provenance cases

100/100

regression cases

5/5

mutations detected

E2

highest current public proof level

EVIDENCE LADDER

No marketing shortcut across evidence levels.

LEVEL	EVIDENCE	ALLOWED PUBLIC WORDING
E0	Concept only	No result claim.
E1	Synthetic demonstration	"Synthetic example / demo".
E2	Reproducible public test	"Re-runnable public test" with scope and limitations.
E3	Controlled pilot	Measured pilot result with scope, period, sample and limitations.
E4	Production customer result with permission	Case result with methodology and permission.

Start with one workflow, not an AI transformation promise.

Expansion is earned after a single workflow passes both the reliability gate and the business gate.

01**One outcome**

A measurable approved result and baseline.

02**One source boundary**

Approved sources and explicit unknown behavior.

03**One authority boundary**

Who may review, approve, reject or escalate.

04**One stop rule**

A fail condition agreed before implementation.

Measure the approved result.

Primary metrics

Time to Approved Result · Critical Error Rate · Evidence Coverage · Rework Rate · Escalation Resolution Time · Cost per Approved Result.

What not to optimize first

Automation percentage, model output volume, token usage or “agent count” without an approved-result business metric.

Commercial principle: human review rate is a diagnostic, not a vanity metric. A workflow may become more valuable by escalating the right cases, not by escalating fewer cases.

Informed by recognized public frameworks. Not disguised certification.

External references are governance and risk inputs. Synapse does not claim ISO certification, NIST approval or OWASP endorsement.

NIST AI Risk Management Framework

Voluntary AI risk-management framework and trustworthiness context.

NIST AI 600-1 · Generative AI Profile

Cross-sectoral GenAI risk profile, including provenance and risk-management considerations.

ISO/IEC 42001:2023 · AI management systems

Management-system approach to AI governance, accountability, transparency and continual improvement.

OWASP GenAI / LLM Top 10

Security risks including improper output handling and excessive agency.

High-consequence boundary

In finance, health, legal and other high-consequence contexts, the workflow may support extraction, classification, anomaly flagging, evidence organization and decision preparation. Final regulated/professional judgment remains with the appropriate authorized human unless a separately validated and legally permitted control model applies.

SYNAPSE AUTOMATE

Automate the work. Keep authority human.

The goal is not to make AI look certain. The goal is to make the workflow accountable when certainty is not available.

Assess one workflow → synapseautomate.github.io/surec-analizi.html

Türkçe özet

Güvenilirlik, modelin düzgün cevap vermesinden fazlasıdır. Kritik bilgi kaynağa bağlı olmalı; deterministik kural model yorumundan ayrılmalı; sonuç doğuran eylemde yetki görünür olmalı; eksik veya çelişkili kanıtta akış durabilmelidir.

Public proof boundary

Current public evidence is synthetic/reproducible (E2). No production accuracy, customer ROI, certification or regulatory-compliance claim is implied.