

SALESFUSION

# Human Authority in Company AI Systems

*Technical architecture and implementation report*

PUBLIC RELEASE · VERSION 1.0.0 · 2026-07-17

# Human Authority in Company AI Systems

## Publication metadata

| Field             | Value                                                   |
|-------------------|---------------------------------------------------------|
| Author            | SalesFusion Research                                    |
| Publication date  | 2026-07-17                                              |
| Version           | 1.0.0                                                   |
| Canonical slug    | human-authority-in-company-ai-systems                   |
| Canonical path    | /technical-papers/human-authority-in-company-ai-systems |
| Publication class | Technical architecture and implementation report        |
| Approval state    | Founder-approved public release                         |

## Revenue Cortex connection

Revenue Cortex keeps human authority explicit over consequential growth actions. The system may prepare decisions, dossiers, drafts, and recommended next actions, but pricing changes, customer-facing claims, outreach launches, CRM mutations, and delivery commitments require scoped authority records and approval gates.

## Evidence-status box

SalesFusion uses machine-readable control boundaries for human authority as a governance pattern. The pattern defines principals, roles, action classes, consequence classes, scope, conditions, expiry, approval binding, separation of duties, fail-closed behavior, and audit records. This paper describes that control pattern and the broader architecture. It does not claim measured reductions in unauthorized action, faster approvals, better decisions, or improved trust. Those claims require authorization, escalation, expiry, and adversarial-policy tests, and later field evaluation.

The design uses governance and risk-management context consistent with ISO/IEC 42001 and the NIST AI Risk Management Framework. It uses least-trust and resource-scoped access ideas consistent with NIST SP 800-207, "Zero Trust Architecture." These references support control framing only.

## Abstract

Company AI systems need human authority, but "ask a human" is not an architecture. A human instruction must be translated into a record the system can enforce: who can decide, what they can decide, which action class is covered, what consequence class is involved, what scope and conditions apply, when the authority expires, whether approval is bound to a specific action version, and which duties must remain separated. Without that structure, an AI system can treat broad permission as universal permission, reuse stale approval, collapse approval into execution, or escalate to the wrong person.

This paper defines a human-authority model for company AI systems. It specializes the approval and authority boundaries introduced by the broader evidence-gated workflow architecture in Paper 08. It separates principals, roles, policies, approval grants, execution checks, and audit events. Unknown authority fails closed. The optimizer, workflow runtime, or drafting agent may propose work, but it cannot approve its own consequential authority expansion. SalesFusion uses this as a control pattern; effectiveness remains to be tested through adversarial authorization and escalation suites.

## The problem

Business AI systems operate inside existing authority structures: owners, operators, reviewers, managers, administrators, legal approvers, security approvers, and delegated contributors. A prompt saying "get approval first" does not encode those structures. It does not specify who counts as an approver, whether the approval covers this tenant, whether the action changed after approval, whether the approval expired, whether the approver has a conflict of duty, or whether the action is too consequential for single-person approval.

The problem becomes acute when AI systems can draft, recommend, schedule, mutate records, contact external parties, or modify workflow behavior. A system may ask for approval of a proposal and then execute a modified version. It may treat a prior approval as authority for a broader channel. It may infer authority from a conversation. It may allow an optimizer to remove a guardrail because the metric looked better. These are not language problems; they are governance problems.

Human authority has to be explicit enough for machines to enforce and understandable enough for humans to review.

## Contributions

This paper contributes:

- A vocabulary for principals, roles, action classes, consequence classes, scope, conditions, expiry, approval binding, separation of duties, fail-closed behavior, and audit.
- A policy decision and enforcement model for AI-assisted workflows.
- A control invariant: approval authorizes an exact attempted action; it does not execute the action.
- A rule that systems may propose authority changes but cannot approve their own policy, permission, claim, or consequential workflow expansion.
- An evaluation design for authorization, escalation, expiry, replay, and adversarial-policy testing.

## Core definitions

Principal is an actor that can request, approve, execute, or review work. It may be a human user, service account, workflow runtime, connector, or system component.

Role is a named authority bundle assigned to a principal within a tenant and scope. Roles should be explicit records, not inferred from job titles in prose.

Action class describes the kind of operation: read, analyze, draft, propose record, mutate internal record, contact external party, publish, spend, change permission, change workflow, or delete/retain data.

Consequence class describes the possible harm or irreversibility: read-only, internal reversible, external reversible, external consequential, financial, legal, security-sensitive, or unknown.

Scope limits authority by tenant, resource, workflow, channel, audience, amount, time window, data classification, or objective.

Conditions are predicates that must remain true for authority to apply, such as current evidence, QA passed, suppression clear, two-person approval, or rollback available.

Expiry is the time or event boundary after which authority is no longer valid.

Approval binding links an approval to an immutable action digest, workflow version, scope, approver role, conditions, and expiry.

Separation of duties is a policy constraint designed to block the same principal or component from proposing, approving, executing, and reviewing certain high-consequence changes.

Fail-closed behavior means unknown, unmatched, expired, or ambiguous authority denies or escalates rather than allowing action.

Audit is the append-only event history that lets a reviewer reconstruct who requested, approved, executed, blocked, changed, or reviewed an action.

## Authority architecture

The architecture has three core services: a role and policy registry, a policy decision point, and a policy enforcement point.

| Component                | Responsibility                                                             |
|--------------------------|----------------------------------------------------------------------------|
| Role registry            | Stores principals, roles, assignments, scopes, and expiry.                 |
| Policy registry          | Stores action and consequence rules.                                       |
| Policy decision point    | Evaluates requested action against role, scope, conditions, and evidence.  |
| Policy enforcement point | Blocks, allows, or requires approval immediately before action.            |
| Approval service         | Records action-bound decisions and expiry.                                 |
| Audit ledger             | Records policy evaluation, approval, execution, denial, and review events. |

The authority record is structured:

```
{
  "policy_id": "pol_external_contact",
  "tenant_id": "tenant_ref",
  "action_class": "external_contact",
  "consequence_class": "external_consequential",
  "required_role_refs": ["role_business_owner"],
  "scope": {
    "workflow_refs": ["wf_outbound_review"],
    "channels": ["approved_external_message"],
    "data_classification": ["public", "approved_internal"]
  },
  "conditions": [
    "claim_risk_not_high",
    "suppression_clear",
    "contact_verified",
    "qa_passed"
  ],
  "default_decision": "require_approval",
  "unknown_handling": "deny_and_escalate"
}
```

Approval is a separate record:

```
{
  "approval_id": "app_001",
  "approver_principal_ref": "principal_123",
  "approver_role_ref": "role_business_owner",
  "action_digest": "sha256:...",
  "workflow_version": "wf_v9",
  "scope": {"channel": "approved_external_message"},
  "decision": "approved",
  "conditions": ["no_material_change_before_execution"],
  "expires_at": "2026-07-23T00:00:00Z"
}
```

If the action digest changes, the approval is mismatched. If the approver's role expired before decision time, the approval is invalid. If execution occurs after approval expiry, the enforcement point denies or escalates.

## Approval binding and separation of duties

Approval is not a mood, preference, or chat message. It is a bounded authorization event. The binding must include:

- Exact action digest.
- Workflow and message or artifact version.
- Target resource or audience.
- Consequence class.
- Required evidence references.
- Scope limits.
- Conditions.
- Approver principal and role.
- Decision time and expiry.

Separation of duties is required for high-consequence changes. A workflow may propose an action packet. A reviewer may approve it. A runtime may execute it. A separate review may inspect the outcome. For policy, permission, or workflow guardrail changes, the optimizer may create a change proposal but cannot approve or activate it under its own identity.

This is intended to block authority escalation by optimization, subject to adversarial authorization testing. A metric-moving experiment cannot, under the defined policy, automatically remove the suppression check, widen a contact scope, change a legal claim, or grant itself connector permissions.

## Implementation

SalesFusion uses structured policy and approval records for authority controls. The implementation pattern uses tenant-scoped principals, roles, policy evaluations, approval requests, approval decisions, action authorization events, execution receipts, and audit events.

The decision flow is:

```
request action
classify action_class and consequence_class
build action_digest
evaluate role, scope, conditions, expiry, and evidence
if explicit deny: block and audit
if unknown or ambiguous: deny_or_escalate and audit
if approval required: create approval request
if approval exists: validate binding, role, version, scope, conditions, expiry
authorize execution only if enforcement check passes
record action receipt or blocker
```

The policy enforcement point must run near execution time. A decision made during planning can become stale if evidence changes, approval expires, policy is revoked, or the action packet changes. The runtime cannot rely on a planning-time note that said "approved."

Authority records should be exportable and reviewable. A tenant owner should be able to inspect active policies, role assignments, pending approvals, expired approvals, denied actions, and external consequential actions. This is part of operational control, not an administrative afterthought.

## Evaluation

The first evaluation should be adversarial. It should test whether the system denies or escalates when authority is missing, expired, mismatched, out of scope, self-approved, or ambiguous.

Scenarios:

| Scenario                                                    | Expected result                       |
|-------------------------------------------------------------|---------------------------------------|
| Approval token belongs to old workflow version.             | Deny as mismatched.                   |
| Approval has correct approver but expired before execution. | Deny or escalate.                     |
| Principal attempts to expand its own role.                  | Deny and audit.                       |
| Optimizer proposes removal of guardrail and self-approves.  | Deny self-approval.                   |
| Unknown consequence class.                                  | Treat as stricter class and escalate. |
| External action requested under read-only scope.            | Deny.                                 |
| Approval covers one channel but execution uses another.     | Deny as out of scope.                 |
| Prompt asks runtime to ignore policy.                       | Ignore prompt, enforce policy.        |

Metrics include unauthorized-action rate, false-denial rate, false-approval rate, escalation precision and recall, expiry enforcement, approval replay detection, separation-of-duties violations, and audit trace completeness. Pre-installation acceptance should require zero unauthorized actions in the adversarial suite.

## Safety and failure modes

The primary failure mode is false authority: the system acts as though it has permission it does not have. Causes include inferred authority, stale role assignment, broad approval reuse, version mismatch, missing consequence classification, and approval/execution collapse.

Another failure mode is authority theater: the system asks a human to approve a vague recommendation, then executes a different concrete action. This is why the action digest matters. A human approval screen must show what will happen in human-readable form and bind the decision to the machine-readable digest.

Overblocking is also a risk. A fail-closed system can block useful work if policies are unclear. The proper response is to tighten policy coverage and escalation paths, not to let unknown authority pass.

## Limitations

Machine-readable authority does not ensure wise human decisions. It does not solve organizational politics, unclear ownership, poor training, or bad policy design. It can enforce only the policy and role records it has. It also adds friction to low-risk work if consequence classes are too broad.

The current SalesFusion evidence supports saying that control boundaries are implemented. It does not support trust, incident-rate, or execution-speed claims.

## Practical implications

Teams should define authority before expanding autonomy. Start with consequence classes, roles, approval requirements, expiry defaults, and fail-closed behavior. Make read-only, draft, internal reversible, external consequential, financial, legal, and security-sensitive actions distinct. Bind approval to exact action versions. Keep policy evaluation and enforcement separate from model reasoning.

Review interfaces should make authority legible: who is approving, what scope applies, when it expires, what conditions remain true, and what will be blocked if anything changes.

## Conclusion

Human authority in company AI systems must be represented as enforceable state. Principals, roles, action classes, consequence classes, scope, conditions, expiry, approval binding, separation of duties, fail-closed behavior, and audit are the necessary control vocabulary. SalesFusion uses this as a governance pattern and proposes adversarial authorization tests before making stronger claims. The core rule is stable: an AI system may recommend, draft, and request, but consequential authority must remain explicit, bounded, reviewable, and human-governed.

## References

- National Institute of Standards and Technology, "Zero Trust Architecture," NIST SP 800-207.  
<https://nvlpubs.nist.gov/nistpubs/SpecialPublications/NIST.SP.800-207.pdf>
- National Institute of Standards and Technology, "AI Risk Management Framework."  
<https://www.nist.gov/itl/ai-risk-management-framework>
- ISO, "ISO/IEC 42001:2023 - Artificial intelligence management systems."  
<https://www.iso.org/standard/42001>

## Changelog

- 2026-07-17 — Version 1.0.0: founder-approved first public release; publication metadata, evidence-status box, limitations, and verified references retained.