

Responsible AI Governance Policy

Version 1



Document Information

Name of the document	AI Governance
Release date	24-July -2026
Owned by	Joshua O'Brien (Head of AI)
Governed by	Udaya Bhaskar Reddy

Revision History

Version No	Version Date	Details of Change
1	20-Jun-2026-22 July 2026	Initially drafted. Rewritten to align with platform architecture: deterministic boundary defined; tool gate model replaces confidence thresholds; tenant-scoped fine-tuning clarified; infrastructure and provider updates. Aligned with a four-repository production codebase audit; model serving, fine-tuning, retrieval, voice, and logging practices documented in detail. Restructured as a high-level policy for executive, customer, and audit audiences. Implementation detail moved to the AI Inventory and AI System Cards. Product scope aligned with the refreshed positioning, including Finance service delivery. Regulatory monitoring expanded (EU AI Act; US state and municipal AI employment laws). Added: ISO/IEC 42001 alignment objective; AI Governance Review Forum; customer-facing AI transparency artefacts; individual recourse; AI quality KPIs. Editorial and consistency pass: version corrected to v5 across the cover and revision history; policy ownership reconciled to the Head of AI; Section 2 clarified that the model never itself calls a connected system; Glossary added as Appendix D (including a definition of "HR-adjacent"); minor typographical corrections.

Reviewer and Approver

Name	Title	Comments	Date Reviewed
Udaya Bhaskar Reddy	Co-Founder & CTO	Approved	24-July-2026

Contents

Contents

1. Purpose, Context, and Scope	4
Scope	4
2. How AI Operates in the Platform: The Deterministic Boundary	5
The role of the AI	5
What the AI can and cannot access	5
How a request is handled	5
How answers stay grounded	6
Where AI participates inside an automation	6
3. Governing Principles	6
4. Governance and Accountability	7
External alignment	7
Objectives and KPIs	8
5. AI Risk Classification	8
6. AI Inventory and Impact Assessment	9
7. AI Lifecycle Controls	9
Before production	10
At deployment	10
Change management	10
Validation and operation	10
8. Data Governance for AI	11
9. AI Security	12
10. Fairness, Explainability, and Transparency	12
Fairness	12
Explainability	13
Transparency and recourse	13
11. Resilience and Incident Management	13
12. Employee Use, Training, and Shadow AI	14
13. Regulatory Monitoring and Policy Review	15
Appendix A: AI Risk Classification Checklist	16
Appendix B: AI System Card Template	16
Appendix C: Approved Internal AI Tools	17
Appendix D: Glossary	20

1. Purpose, Context, and Scope

Rezolve.ai builds an agentic AI service desk for IT, HR, and Finance. The platform spans conversational and voice assistance (Sidekick and VoicelQ), AI-assisted ticket work (Agent Assist), governed automation that customers build and run themselves (Agent Studio, the agent Marketplace, and MCP tool integrations), service-desk intelligence (DeskIQ and conversational dashboards), AI knowledge management and enterprise search, and a full ITSM system of record that can run standalone or as the AI layer on a platform a customer already operates, such as ServiceNow — together with emerging capabilities in digital employee experience and AI-assisted IT operations. The authoritative list of AI components in production is the AI Inventory (Section 6); this policy governs all of them.

The platform runs on multi-cloud infrastructure in a multi-tenant architecture. It uses large language models from approved third-party providers — OpenAI, Anthropic, and Google — alongside Rezolve.ai's own tenant-scoped Otto models. Supporting AI services — retrieval, web search, model serving, fine-tuning, and voice — are provided by approved sub-processors, as published in Rezolve.ai's sub-processor list and cross-referenced in the AI Inventory. Which model powers which capability is a configuration decision, never a code change, and is governed by the change process in Section 7.

At its foundation, the platform is a ticketing system with a governed data layer. AI operates within a defined role (Section 2), and everything outside that role is conventional, deterministic software. AI-specific governance is nonetheless necessary because of three characteristics:

- **Non-determinism, confined to the AI layer:** within the AI layer, the same input can produce a different output, so that layer is monitored continuously, not just tested before release. Everything outside it behaves deterministically.
- **Autonomy is configured, never assumed:** agents can act in connected IT, HR, and Finance systems, but autonomy is a configuration decision, never a default. Every agent tool call carries an enforced gate — allow, ask, approve, or deny — with named approvers and an instant kill switch, and the AI can instead orchestrate work that a person reviews and clicks through.
- **Data proximity at inference:** Rezolve.ai does not train models on customer data, but models temporarily process the content provided to them during a request, which requires controls beyond those in existing security and privacy policies.

Scope

This policy applies to: AI and machine-learning components embedded in all Rezolve.ai products; internal AI tools used by employees; third-party AI providers and supporting AI services integrated into the platform; and all employees, contractors, and third parties building or operating AI on behalf of Rezolve.ai.

2. How AI Operates in the Platform: The Deterministic Boundary

This section makes the system's boundaries unambiguous — what the AI does, what it does not do, and what it could do if a future use case required it. Every other control in this policy assumes and protects this boundary.

The role of the AI

The language model performs two primary functions:

1. **Understanding and routing** — it interprets a person's request, compares it against the automations available on the platform, and triggers the appropriate one; and
2. **Knowledge-grounded answering** — it answers questions using the knowledge that has been approved and provided to it.

Beyond these functions, the model participates inside an automation only as an explicitly documented step. Everything else — the automations that act in connected systems, and the governance of data held in the ticketing

system — is conventional code executing defined logic: a given condition produces a given action. No model inference takes place in that layer.

What the AI can and cannot access

The model has access only to what is explicitly provided to it: knowledge content uploaded or synchronised into the platform, and what a person types — or, on a call, says. It has no access to connected enterprise systems, document repositories, or private networks. Content from those sources reaches the model only where a customer has deliberately brought it into the knowledge engine. The choice of model is a configuration decision, and that choice never expands what the model can access. Where an agent acts in a connected system, the model itself never makes that call: its constrained output selects an approved automation, and deterministic, gated code performs the action — so “the AI acts” always means the model selects and code executes.

How a request is handled

Request handling follows a consistent pattern. An employee asks a natural-language question — for example, “What was my pay stub?” The model determines whether a matching automation exists. If one does, it hands the request to that automation and is removed from the process entirely; a deterministic flow then makes the defined system calls, retrieves the information, and applies all data governance. The model does not see that data and is not invoked again for the remainder of the transaction. If no automation matches, the assistant asks a clarifying question rather than proceeding. Authentication, employee identifiers, and all sensitive data handling occur inside the deterministic flow — never within the model.

How answers stay grounded

When the right response is an answer rather than an action, the assistant answers only from approved knowledge and — where a customer has enabled it — from approved external sources via a governed web search service. An accuracy safeguard selects content only when it clearly addresses the question. When nothing fits, the assistant says so and offers the next step — a clarifying question, a ticket, or a person — rather than inventing an answer. How cautious the assistant is on a partial match is a customer-level setting, and it governs answering only; it is separate from action gating.

Where AI participates inside an automation

There are defined, documented cases where a model contributes a step within an automation — for example, ticket triage and categorisation, sentiment classification, or drafting a knowledge article for review. Each such step is recorded in the component's AI System Card together with the reason it is needed and the data it touches; the model produces structured, constrained outputs, and code — not the model — maps those outputs to the operations that execute. In Finance document processing, classifications carry a confidence score, and anything below the customer's configured threshold is routed to a person rather than processed automatically. Sensitive automations built in Agent Studio are reviewed by a person before publication.

3. Governing Principles

Principle	What It Requires
Responsible by design	Guardrails, enforced tool gates, approval steps, and human escalation paths are built in before deployment — not retrofitted after incidents.
Code governs; prompts guide	Hard constraints — routing, confirmations, validation, and high-risk actions — are enforced in software, never left to a model's judgement or a written instruction alone.
No training on customer data	Customer data is never used to train or improve shared AI models. Customers may choose to fine-tune their own tenant-scoped models on their own data; those

Principle	What It Requires
	models are isolated per tenant, the underlying base models are never modified, and cross-tenant leakage is not possible.
Humans stay in control	Every agent tool carries an enforced gate — allow, ask, approve, or deny — with named approvers and an instant kill switch. Nothing gated for review executes until a person confirms or approves it, and every AI interaction offers a path to a human.
Tenants stay isolated	Tenant isolation is an architectural requirement, enforced through logical segregation and security controls and tested on every release.
Transparent and explainable	Customers know which AI capabilities run on their tenant, what data flows where, and how to raise a concern. Every run is captured in a glass-box explainability record: who asked, what it did, which tools it used, and who approved.
Monitored on a risk basis	AI systems are monitored for output quality, operational performance, and anomalous behaviour in proportion to their risk.

4. Governance and Accountability

AI governance sits within Rezolve.ai's existing information security management structure, with an AI agenda item at the quarterly management review. In addition, an AI Governance Review Forum — a cross-functional group drawing on engineering, security, HR and compliance, and customer-facing leadership — convenes quarterly and before any Tier 1 or HR-adjacent deployment, providing a review path that does not depend on any single approver.

CTO / CISO	Approves this policy, sets risk appetite, and signs off on new Tier 1 components and material new AI sub-processors.
Head of AI	Owns this policy, maintains the AI Inventory, approves AI sub-processors, and can suspend any AI component. Reviews AI KPIs quarterly.
Head of Engineering	Ensures governance requirements are sprint deliverables and release gates.
AI System Owner	Named owner for each AI component, accountable for its inventory entry, risk assessment, monitoring, and post-incident review.
HR and Compliance	Covers AI acceptable use in onboarding and training, manages violations, and reviews HR-adjacent AI for fairness.
All employees	Use only approved AI tools, never place customer data in unapproved AI services, and report anomalous outputs immediately.

External alignment

The programme is informed by the NIST AI Risk Management Framework, and Rezolve.ai is aligning its AI management practices to ISO/IEC 42001, the international standard for AI management systems, with certification as a stated objective. Progress is reported at the management review. The scope of the AI management system is defined and maintained in a documented scope statement approved at the management review.

Objectives and KPIs

The objectives of this programme are to keep AI systems secure, reliable, and fit for purpose; protect the data they process; comply with legal, regulatory, and contractual obligations; prevent any unauthorised use of customer data

for training; preserve the deterministic boundary and human oversight; and improve continuously through monitoring and lessons learned.

AI Inventory complete — every deployed component has a current entry	100%
Tier 1 and Tier 2 components with a completed risk assessment	100%
P1 and P2 AI incidents with a completed post-incident review	100% within 10 business days
Shadow AI reports investigated and closed	100% within 5 business days
Annual AI training completion	100% of workforce
AI sub-processor agreements current	100%
Scheduled output-quality reviews completed	100%
Guardrail-trigger and approval-override trends reviewed	Quarterly

The AI governance programme is subject to internal audit at least annually under the internal audit programme. Findings from audits, monitoring, and reviews that are not incidents are recorded as nonconformities, assigned an owner and a corrective action, and tracked to closure; recurring nonconformities are escalated to the management review.

The management review assesses the programme's effectiveness at least annually — incidents and trends, risks and treatment, audit findings, KPI performance, regulatory developments, and customer concerns — with actions tracked to completion.

5. AI Risk Classification

Every AI component is classified before production using the checklist in Appendix A. The tier determines review depth, documentation, and oversight cadence. Risks that cannot be remediated to an acceptable level are documented, assigned an owner, and formally accepted before deployment proceeds.

Tier	Profile	Examples	Gate Before Deploy
1 — High	Autonomous action in a connected system with no human confirmation step, or processing of sensitive personal data with significant individual impact.	Password reset agent; access provisioning; HR decision-support.	CISO sign-off; impact assessment; isolation verified; gates and approvers configured; escalation tested; Review Forum consulted.
2 — Medium	A person sees the output before any system change, or the component only answers; personal data may be in context.	Sidekick answering; Agent Assist; VoicelQ responses.	System Owner sign-off; prompt-injection review; fairness check for HR-adjacent use.
3 — Low	Internal productivity or development tool; no customer data; a person reviews all outputs.	Internal AI tools.	Register in the AI Inventory; acceptable-use acknowledgement.

6. AI Inventory and Impact Assessment

Rezolve.ai maintains an inventory of every AI system and component in production or active development; nothing reaches production without an entry. Each entry records the component, its owner, risk tier, purpose, model and provider, data categories, third-party dependencies, human-oversight mechanism, and status. The inventory - not this policy - is the authoritative internal record of specific models and AI components, and it is reviewed quarterly. AI providers and supporting services that process customer data are governed as sub-processors under Rezolve.ai's privacy policy and customer agreements; the inventory cross-references that list rather than duplicating it.

Tier 1 components, and any component processing personal data or acting autonomously, require an AI Impact Assessment before production. The assessment covers privacy, security, fairness, operational resilience, human oversight, third-party dependencies, and potential impacts on individuals — including whether an automated action can be reversed. Assessments are refreshed whenever the component, its model, or its data scope materially changes.

7. AI Lifecycle Controls

Before production

- Every component is registered, classified, and assessed — no exceptions. Tier 1 and Tier 2 components pass a threat model and a deterministic-boundary review confirming that every hard constraint is enforced in code; both are release gates.
- Any model step inside an automation is documented in the design and the AI System Card, including why it is needed and what data it touches. Sensitive Agent Studio automations are human-reviewed before publication.
- New AI sub-processors require CTO approval, and their terms must be confirmed to cover no training on customer data, deletion after processing, data-residency commitments consistent with customer contracts, and prompt breach notification.
- Prompts are governed assets — versioned, peer-reviewed, and protected like production credentials — and Tier 1 systems must demonstrate a working human escalation path before go-live.

At deployment

- Gates are set for every tool an agent can call — allow, ask, approve, or deny — at a platform default each agent can tighten with conditions and named approvers. Anything gated for review does not execute until a person decides, the decision is logged, and a kill switch can suspend any agent or integration instantly.
- Agents run under dedicated, least-privilege identities that cannot escalate themselves, and each agent can call only its approved tools. Each step receives only the context it needs.
- Every AI interaction is captured in an auditable record — who asked, what it did, which tools it used, and who approved — retained for 12 months with restricted access. AI-generated responses are clearly identified as such (on voice, through the call greeting), and a person is always reachable.

Change management

- Model assignment is a per-tenant configuration, changed without code deployment and always within the model's existing data boundary. Tier 1 model changes require a change request and regression testing; urgent provider retirements follow an emergency path; and changes to supporting AI services, such as the web search provider, follow the same governance.
- Material model changes or retirements that affect customers are communicated to them.

Validation and operation

- AI outputs are treated as untrusted until validated — through human confirmation or approval, code-enforced gates and schema checks, or guardrail evaluation — and never trigger privileged actions unless explicitly approved by design and documented.
- Output quality is monitored continuously through automated checks, with sampled human review at a cadence defined per component in its System Card - at least weekly for Tier 1 and monthly for Tier 2; injection attempts are security events; AI components are in scope for the annual penetration test; and prompt or model changes ship behind staged rollout with evaluation gates that block on regressions.
- On decommissioning, prompts, tenant-scoped models, and logs are handled under the Data Retention Policy, related sub-processor relationships are formally closed, and the inventory is updated.

8. Data Governance for AI

Control	Requirement
No training on customer data	Customer data is never used to train or improve shared models. Optional tenant-scoped fine-tuning runs only through approved sub-processors; the resulting models are evaluated before serving, isolated and served per tenant, deleted when the relationship ends, and never alter the base model.
What the model sees	Only content explicitly provided: approved knowledge and what the person types or says. Data in connected systems is handled by deterministic automations after the model hands off — the model does not see it — and the model has no access to private networks or repositories.
Data minimisation	Only the fields a task requires are ever included in a model request; the fields in scope are documented per component, and broad data dumps into prompts are prohibited.
PII redaction	Sensitive-data detection and anonymisation is applied before content reaches a model, according to each tenant's configuration; extending coverage across all products is a tracked objective.
Tenant isolation	Retrieval and search are scoped to the requesting tenant and enforced at the database layer; cross-tenant access is architecturally prevented and tested every release.
Web search scope	Where a customer enables external answering, queries run through an approved web search service against approved sources only; it is off unless the customer turns it on.
Data residency	Processing regions match the residency commitments in the customer's agreement; routing outside a contracted region requires CISO approval.
Individual rights	Personal data processed through AI remains subject to applicable privacy rights — access, correction, deletion, restriction, and objection — and individuals affected by an AI-influenced outcome can request an explanation and human review (Section 10).
Employee use	Employees must not place customer data, personal data, or company IP into any unapproved AI tool, including personal accounts of approved tools.

9. AI Security

These controls extend Rezolve.ai's ISO 27001 and SOC 2 control set and integrate with the existing security policies. In summary:

- A model is never the enforcement mechanism for a hard constraint. Routing, confirmations, validation, error handling, and high-risk actions are enforced in code; models produce constrained, structured outputs that code maps to systems — never raw identifiers, credentials, or destinations.
- All user-provided input is sanitised and structurally separated from system instructions before it reaches a model, and content returned from tools, retrieval, or the web is treated as untrusted before it re-enters the model's context.
- A guardrail layer evaluates model-proposed actions before anything executes; guardrail decisions are logged with reasons, and unusual spikes alert the CISO. Known injection and jailbreak patterns are part of automated regression testing, and prompt injection is a mandatory case in every penetration test involving AI.
- Agents operate under dedicated, least-privilege identities with approved tool lists and no ability to escalate their own permissions. Platform access is governed by enterprise single sign-on with role-based access control and per-tenant isolation.
- Provider credentials and API keys are held in managed, access-controlled secret stores — never in source code or logs — and any exposure is treated as a security incident.
- Customer-connected (bring-your-own) MCP tools are governed, tested, and audited under exactly the same gate, logging, and review requirements as Rezolve.ai's own.
- Models, prompts, and AI configurations are under version control and change management, so every change is documented and auditable.

10. Fairness, Explainability, and Transparency

Fairness

Rezolve.ai's AI operates in IT, HR, and Finance contexts, where outputs affect how requests are categorised, prioritised, and resolved. Before deployment, Tier 1 and Tier 2 systems are tested for consistency across equivalently phrased requests; HR-adjacent systems undergo an additional fairness review completed with HR and Compliance; and voice capabilities are tested for language and accent sensitivity, with material degradation blocking that variant. Customer-reported fairness concerns are logged, reviewed monthly, and reported to the management review.

Explainability

Every run is captured in the platform's glass-box explainability record, in plain language: how the request was understood and routed, what knowledge was retrieved, which agents and tools ran, what was decided, and who approved. This record supports customer governance, audit, and Rezolve.ai's own oversight, and it underpins the recourse commitment below.

Transparency and recourse

Rezolve.ai maintains a Trust Center describing which AI capabilities are active, which providers power them, how customer data is protected, and how to raise a concern. To support customer due diligence and regulatory readiness, Rezolve.ai also maintains customer-facing AI transparency documentation — including a model overview for its own Otto models and per-capability descriptions of what each AI feature does, its limitations, and the data it uses.

- Every AI-generated response is identified as AI-generated; on voice, disclosure is delivered through the call greeting.
- A person is always reachable, and the path to a person is never harder than the AI path. If an AI capability is unavailable, the product degrades gracefully to a human-handled state rather than failing silently.
- Any individual affected by an AI-influenced outcome may request an explanation and human review; the explainability record makes both practical.

- Security-questionnaire responses describing AI capabilities are reviewed by the CTO or a designated security lead and must accurately reflect current practice.

11. Resilience and Incident Management

The platform operates multiple approved model providers with automatic failover, and products degrade gracefully to deterministic or human-handled behaviour when a model is unavailable. For Tier 1 components, fallback paths are documented and tested under the business continuity and disaster recovery programme.

AI incidents are handled under the Incident Management Policy and tagged for trend analysis:

Category	Description	Response
P1 — Critical	Prompt injection causing unauthorised action; cross-tenant data exposure; AI-driven breach.	CTO within 1 hour. Component suspended via kill switch; incident response activated; regulatory notification clocks assessed if personal data is involved.
P2 — High	An incorrect automated action with real-world consequence; guardrail failure; gate or approval bypass.	CTO and System Owner within 4 hours. Component suspended or gates tightened; root-cause analysis within 5 business days.
P3 — Medium	Repeated anomalous outputs; guardrail-trigger spike without breach.	System Owner within 24 hours; engineering review within 5 business days.
P4 — Low	Single anomalous response; output-quality complaint; shadow-AI self-report.	Reviewed within 5 business days; patterns tracked monthly.

Every P1 and P2 review must answer whether the design and gate configuration were adequate, whether the audit record could fully reconstruct the event, and what changes result.

12. Employee Use, Training, and Shadow AI

Employees use only approved AI tools (Appendix C); anything else requires written CISO approval before use. It is prohibited to place customer data, personal data, unreleased code, or commercial negotiation content into any AI tool including personal accounts of approved ones; to use AI-generated work in production or client-facing contexts without review and testing; to present AI-generated content as human-authored in compliance or legal artefacts; or to attempt to bypass guardrails, gates, or approvals in any AI system, including Rezolve.ai's own — the latter is treated as a security incident.

Because prohibition without detection is not a control, unapproved AI use is detected through monthly network-log review, quarterly managed-device review, and data-loss-prevention alerts reviewed within 24 hours. Employees who self-report inadvertent use before an incident is identified are handled under a lighter-touch process, encouraging early disclosure. Violations are managed under the Acceptable Use Policy and Code of Conduct.

All staff complete an annual AI training module covering approved tools, prohibited actions, and how to report anomalous output. Engineers on Tier 1 or Tier 2 components complete role-specific training — including prompt-injection defence and gate configuration — before assignment, and AI System Owners complete a structured handover before assuming ownership.

13. Regulatory Monitoring and Policy Review

The CTO, supported by the AI Governance Review Forum, monitors AI and data regulation applicable to Rezolve.ai and its customers, including: the EU AI Act — including the classification of employment-related capabilities under its

high-risk category and a documented assessment of workplace sentiment analysis against its prohibited-practice and transparency provisions; the GDPR, including its safeguards around automated decision-making; HIPAA; India's DPDP Act; the CCPA/CPRA and other US state privacy laws; emerging US state and municipal AI employment laws such as New York City Local Law 144 and the Colorado AI Act; and ISO standards, including ISO/IEC 42001. Material developments trigger an out-of-cycle review.

For each AI capability, Rezolve.ai documents its role and the applicable regulatory obligations, and records that determination in the component's AI System Card. Where customers deploy Rezolve.ai capabilities, Rezolve.ai provides the documentation they need to meet their own obligations - including descriptions of the capability, its human-oversight mechanisms, and access to the relevant audit records. This policy is reviewed annually, and out of cycle upon a new Tier 1 component or AI sub-processor, a P1 or P2 AI incident, a material change to the AI stack, or a relevant regulatory development. Updated versions are reviewed by the Head of AI, approved by the CTO and communicated to all staff. Rezolve.ai is committed to continually improving this programme as technology, business requirements, customer expectations, and regulation evolve.

Appendix A: AI Risk Classification Checklist

Complete before registering any AI component. If any of Q1 to Q5 is Yes, submit to the CISO for tier confirmation before development begins.

#	Question	Yes	No
Q1	Does this component take autonomous action in a connected system without a mandatory human confirmation step — that is, is any of its tools gated allow?		
Q2	Does this component process personal data during inference?		
Q3	Does it have any retrieval pathway that could surface data from more than one customer tenant?		
Q4	Does it call an external AI service not already on the approved sub-processor list?		
Q5	Could a successful prompt injection attack cause a state change in a production system or expose data?		
Q6	Does it touch any HR-adjacent workflow?		
Q7	Is this an internal tool only, with human review of all outputs and no customer data in scope?		

Scoring: Q1 Yes = Tier 1 minimum. Two or more of Q2–Q6 Yes = Tier 1. One of Q2–Q6 Yes = Tier 2. Q7 Yes and none of Q1–Q6 Yes = Tier 3.

Appendix B: AI System Card Template

Required for Tier 1 and Tier 2 components before production. The System Card is where implementation detail lives — models, providers, data fields, gate configuration, and documented AI steps — keeping this policy stable while the technology evolves. Reviewed at each component review cycle.

Field	Content
Component name, product, risk tier, owner	
What it does / cannot do	Plain-language capability and the out-of-scope behaviours enforced by code, gates, and guardrails.
Model / provider	Model, version, and endpoint per function, agent, or task, as applicable.
Data in scope at inference	The exact fields included in model requests — no more — and their provenance.
Tenant isolation	How cross-tenant access is prevented and tested.
Tool gate configuration	Gate posture per tool — allow / ask / approve / deny — with conditions and named approvers.
AI steps inside automations	Each step where a model participates, why it is required, and the data it touches — or “None.”

Field	Content
Human oversight	Escalation path and output-validation controls (guardrails, schema checks, approvals, human review).
Fairness testing	Method and outcome, or N/A with justification.
Resilience / exit plan (Tier 1)	Fallback and failover if a provider or model becomes unavailable.
Logging, references, status	Logging confirmed with retention; threat model and impact assessment references; deployment status and next review.

Appendix C: Approved Internal AI Tools

Maintained as part of the Approved Apps register. Any tool not listed requires CISO approval (standard: 5 business days; urgent with written justification: 24 hours). Do not use a tool until approval is received.

Approval applies to company-managed workspace; personal or free-tier accounts of listed tools are not approved. Each entry is reviewed at least annually.

This appendix covers AI tools used by employees for internal work; AI providers and supporting services embedded in Rezolve.ai products are governed as sub-processors and recorded in the AI Inventory and Privacy Policy.

Tool	Approved Use	Restrictions
General AI assistants		
OpenAI ChatGPT	Drafting, research, analysis, brainstorming	Company ID only; no customer data or PII; no unreleased code
Anthropic Claude	Drafting, analysis, document and code review	Company workspace only; no customer data or PII; no unreleased code
Perplexity	Research and web-sourced answers	Company account; queries must not contain customer data or confidential terms
Grammarly	Grammar and style checking	Business tier; disabled on documents containing customer data or negotiation content
Fyxr.AI	Email drafting and inbox triage	Approved mailboxes only; no customer contract or negotiation threads
Fireflies.ai	Meeting transcription and notes	Approved accounts only; external/customer calls require participant consent; recordings deleted per retention policy
Engineering and Development		
Cursor	AI assisted coding	Business plan, privacy mode on; approved repositories only; output

Tool	Approved Use	Restrictions
		peer-reviewed before merge
Windsurf	AI assisted coding	Same restrictions as Cursor
GitHub Copilot	Code completion and review	Company org seat; output peer-reviewed before merge
Warp.dev	AI assisted terminal	No production credentials or customer data in prompts
Sourcery	AI code review	Approved repositories only
Replit	Prototyping and experiments	Non-production, non-customer code only
Lovable	UI prototyping	Prototypes only; no customer data; production code follows standard SDLC review
Postman	API testing assistance	No production credentials in AI prompts
Marketing and Sales		
Canva	Design and marketing content	Public/marketing content only; no customer logos or data without consent
Gamma	Presentation drafting	Internal and marketing decks; no customer-confidential content
Adobe	Design and creative production	Licensed seats; generated assets reviewed for IP suitability before publication
Descript	Video editing	No voice cloning of any real person without written consent
Elevenlabs	Voice generation	Written consent required for any real-person voice; output labelled per marketing guidelines
Semrush AI	Content optimisation	AI-drafted content human-reviewed before publication
Supademo	Product demo creation	Demo tenants and synthetic data only — never customer tenants
Recurpost	Social scheduling	Human review before posting
Clay	Prospect enrichment	Business-contact data only; comply with applicable privacy law and source terms of service

Tool	Approved Use	Restrictions
Instantly	Email sequencing	Marketing-approved templates; suppression lists enforced
Phantombuster	Prospect automation	CISO-approved workflows only; no scraping in breach of platform terms; business contacts only
AI embedded in business applications		
AI features within approved apps — incl. HubSpot (Breeze), Zoom (AI Companion), Monday.com, Figma, Webflow, Whimsical, Eraser.io, Mailchimp, GetResponse, Zapier, Supabase	Per the host application's approved use	Governed by the host app's entry in the Approved Apps register; the data rules in Section 12 apply equally to the AI features; enabling a new AI feature that sends data to a new third-party model requires CISO notification

Appendix D: Glossary

Plain-language definitions of the terms of art used in this policy. Where a term is also defined in a customer agreement or another Rezolve.ai policy, that definition governs for the purposes of that document.

Term	Definition
HR-adjacent	Any AI use that touches an HR or employment-related workflow, or the support of an employment decision — for example onboarding and offboarding, leave, benefits, case handling, conduct or performance matters, or the triage and prioritisation that affects how an employee's request is treated. HR-adjacent components require an additional fairness review completed with HR and Compliance, are referred to the AI Governance Review Forum, and are escalated in tiering (Appendix A, Q6).
Deterministic boundary	The architectural line between the AI layer — non-deterministic model inference (understanding, routing, and knowledge-grounded answering) — and everything else, which is conventional, deterministic code that performs actions and governs data. The model never crosses this boundary to act directly in a connected system (Section 2).
Tool gate	The enforced control on every action an agent can take, set per tool to allow, ask, approve, or deny, with named approvers. Anything gated for review does not execute until a person decides, and the decision is logged.
Kill switch	A control that instantly suspends any agent or integration.
Glass-box explainability record	The per-run audit record capturing who asked, how the request was understood and routed, what knowledge was used, which tools ran, what was decided, and who approved. It underpins customer governance, audit, and the individual-recourse commitment (Section 10).
Guardrail layer	The control that evaluates model-proposed actions before anything executes, logging decisions with reasons and alerting on anomalies.
Tenant-scoped model / Otto	Rezolve.ai's own models, which a customer may optionally fine-tune on its own data. They are isolated

	and served per tenant; the underlying base model is never modified and cross-tenant leakage is not possible.
Sub-processor	An approved third party that processes customer data on Rezolve.ai's behalf — for example an LLM provider or a retrieval, model-serving, fine-tuning, or voice service. Sub-processors are governed under the privacy policy and customer agreements and are cross-referenced in the AI Inventory.
AI Inventory	The authoritative internal register of every AI component in production or active development. Nothing reaches production without an entry (Section 6).
AI System Card	The per-component record holding implementation detail — models, providers, data fields, gate configuration, and documented AI steps. Required for Tier 1 and Tier 2 components (Appendix B).
Risk Tier (1 / 2 / 3)	The classification (Appendix A) that determines review depth, documentation, and oversight cadence. Tier 1 (High): autonomous action in a connected system, or processing of sensitive personal data with significant individual impact. Tier 2 (Medium): a person sees the output before any system change, or the component only answers. Tier 3 (Low): internal productivity tools with no customer data and human review of all outputs.
Agent Studio	The environment in which customers build and run their own governed automations on the platform.
AI Governance Review Forum	The cross-functional group — drawing on engineering, security, HR and Compliance, and customer-facing leadership — that convenes quarterly and before any Tier 1 or HR-adjacent deployment, providing a review path that does not depend on any single approver (Section 4).

Review

Next review cycle: March 2027

Management may review and revise this policy at any time, based on organizational needs or circumstances.

Note: All documents related to policies and procedures—any reference to Actionable Science is as good as Rezolve.ai.