

PRODUCT GOVERNANCE

Responsible AI & AI Ethics Statement

Applies to the Synopsis platform, including rectifAI and synConnect.

VERSION 2.0 · EFFECTIVE 15 SEPTEMBER 2026

DOCUMENT CONTROL

DOCUMENT	Responsible AI & AI Ethics Statement
DOCUMENT REFERENCE	FX-POL-004
APPLIES TO	Synops platform, including rectifAI and synConnect
ISSUED BY	Foxtrot LLC, an Arizona limited liability company
VERSION	2.0
EFFECTIVE DATE	15 September 2026
SUPERSEDES	Version 1.0, effective 12 September 2026
ACCOUNTABLE PERSON	Farshad Mohammadi, Managing Member
REVIEW CYCLE	Annual, or on any material change to model behavior, scope, or deployment
RELATED DOCUMENTS	FX-POL-001 Code of Business Ethics & Conduct, FX-POL-003 Human Rights Policy Statement, FX-POL-011 Information Security & Data Protection Policy
CONTACT	compliance@foxtrotllc.com
CLASSIFICATION	Public. Published without gating so that customers, regulators and partners can hold us to it

Contents

1 Purpose.....	3
2 Products in scope	3
3 First principle: the AI advises, a qualified human decides	3
4 Approved data is not replaced	4
5 Human oversight in practice	4
6 Traceability and explainability	4
7 Data governance and sovereignty	4
8 Model change control	5
9 Known limitations, stated plainly	5
10 Safety culture and personnel fairness	6
11 Accessibility and working language.....	6
12 Security of the AI itself	6
13 Incident handling.....	6
14 Prohibited uses.....	7
15 Frameworks we align with	7
16 Accountability and contact.....	7
17 Approval.....	8
Revision history.....	8

1 Purpose

Synops builds AI into the Maintenance Control Center, which is an environment where decisions affect airworthiness, dispatch, and flight safety. That places obligations on us beyond ordinary software.

This statement sets out what our AI does, what it explicitly does not do, who remains accountable for a decision, and how we manage the risks. It is written to be read by maintenance control managers, quality and safety departments, IT security teams, and regulators, not only by procurement.

2 Products in scope

rectifAI is an AI maintenance controller providing three functions: verification, reliability analysis, and technical knowledge management. It uses machine learning and large language models over customer maintenance data, technical publications, and operational messaging.

synConnect handles SITA Type B and AIDX messaging. It is primarily deterministic message processing and integration. Where AI is applied to classification, routing or enrichment of a message, the commitments in this statement apply to those components, including traceability of the output and the human decision point that follows it.

Any future capability that uses AI to influence a maintenance, engineering, or dispatch decision falls under this statement automatically.

3 First principle: the AI advises, a qualified human decides

rectifAI does not make airworthiness determinations. It does not certify, approve, release, or defer anything.

Specifically, rectifAI does not and will not:

- issue or sign a certificate of release to service;
- make or apply a Minimum Equipment List or Configuration Deviation List decision as a final action;
- authorize a deferral, an extension, or a repair category;
- close a discrepancy, sign a task, or author a maintenance record entry in its own name;
- determine compliance with an Airworthiness Directive or Service Bulletin;
- dispatch or hold an aircraft.

Every one of those decisions remains with appropriately licensed and authorized personnel acting under the operator's or maintenance organization's own approvals. rectifAI produces analysis, options, references, and drafts. A human accepts, modifies, or rejects them, and that human's authority, not the system's output, is what makes a decision valid.

NO AUTONOMOUS RELEASE, AT ANY TIER

There is no configuration, licence tier, or customer request under which we will supply an autonomous airworthiness release.

4 Approved data is not replaced

Approved data remains the manufacturer and authority source material: the AMM, TSM, IPC, SRM, MEL and MMEL, service bulletins, airworthiness directives, and the operator's own approved procedures.

rectifAI locates, links to, and summarizes approved data. Its summaries and recommendations are not approved data and are never presented as such. Where an output paraphrases source material, the source task reference is shown so the user can and must verify against the original before acting. System messaging makes this distinction visible in the interface, not only in the contract.

5 Human oversight in practice

Oversight only counts if it is real, so we design against rubber-stamping:

- Outputs are presented as recommendations requiring action, never as completed decisions.
- Confidence, and the absence of confidence, is shown. Where the system does not have sufficient basis, it says so rather than producing a plausible answer.
- Users can override any output without friction, and overrides are recorded rather than discouraged.
- No safety-relevant action executes on a timer, a default, or an absence of response.
- Automation bias is addressed in user training, including guidance on when to distrust an output.

Customers are expected to reflect use of the tool in their own approved procedures and exposition documents, and to define which roles may rely on which outputs. We support that documentation and will provide the technical detail an authority requires, but we do not substitute our judgment for the customer's quality and safety organization.

6 Traceability and explainability

Every AI output carries an auditable record:

- the input data and time snapshot it was based on;
- source references for the technical content it relied on;
- model and system version identifiers;
- the user who received it and the action they took, including rejection or modification;
- an immutable, time-stamped log entry.

Logs are exportable in a form suitable for an internal audit, a regulatory audit, or an investigation. We do not present outputs whose basis cannot be reconstructed, and we treat "the model said so" as an unacceptable explanation.

7 Data governance and sovereignty

Customer data belongs to the customer. Our position is straightforward:

- We do not train, fine-tune, or improve models on customer data without specific written consent, recorded per customer.
- There is no cross-tenant learning. One airline's data never influences another airline's outputs.
- Tenants are logically segregated, with access controls and encryption in transit and at rest.

- rectifAI uses a data-sovereign architecture, and in-region deployment is available where residency is required, including deployment within the customer's jurisdiction.
- Subprocessors and any third-party model providers are disclosed, contractually bound, and listed on our trust page. Material changes are notified in advance.
- On termination, data is exported to the customer and deleted on the agreed schedule, with written confirmation.

Personal data of crew, engineers, and staff is handled under FX-POL-011 and applicable law, including the Saudi PDPL and the GDPR where relevant.

8 Model change control

Model behavior is part of the product, so it is managed like any other change to a safety-relevant system:

- models and prompts are versioned, and versions are recorded against every output;
- changes are evaluated against curated aviation test sets before release, including regression tests on previously correct behavior;
- material changes to model behavior, capability, or the underlying provider are notified to customers before deployment, with a description of the change and its expected effect;
- rollback to a prior version is possible and tested;
- no silent model substitution. A change of underlying foundation model is a material change and is communicated as one.

9 Known limitations, stated plainly

We would rather lose a deal than oversell the system. Users and reviewers should understand the following:

- Generative components can produce fluent output that is wrong. This is a property of the technology, not a defect we have eliminated. Verification against source data is mandatory, not advisory.
- Performance degrades on inputs unlike the training and evaluation data, including unfamiliar fleet types, unusual defect patterns, and non-standard local terminology.
- Quality depends on the customer's data. Incomplete histories, scanned or handwritten records, inconsistent defect coding, and stale publication revisions all reduce reliability.
- Reliability analysis identifies statistical patterns. It does not establish causation and is not an engineering determination.
- The platform is ground-based operational support software. It is not airborne software, not part of any type design, and no software assurance level under DO-178C or an equivalent standard is claimed or implied.
- We do not claim that the system is certified or approved by any aviation authority. No such approval exists for this category of tool, and any vendor claiming otherwise should be questioned.

10 Safety culture and personnel fairness

An AI system in maintenance control can easily become a surveillance tool. We do not build for that use.

- rectifAI is not designed or licensed as a tool for individual performance monitoring, ranking, discipline, or attributing blame to named engineers or controllers.
- Outputs must not be used to undermine just culture or to identify reporters within a confidential safety reporting system.
- Where a customer's use of the platform would involve monitoring of individual staff, that is a matter for the customer's law, agreements, and works or union consultation, and we will say so rather than quietly enabling it.
- We test for systematic bias in classification and prioritization outputs, including bias arising from fleet, base, shift, or language of the original record.

These commitments are part of how we meet FX-POL-003, which prohibits building or supplying capabilities intended for unlawful surveillance of individuals.

11 Accessibility and working language

Maintenance control is staffed by people working in a second or third language, often at night and under time pressure. The platform is built for that reality: interface language is plain, findings are legible without specialist AI knowledge, source references are shown rather than implied, and outputs do not depend on the user recognising a technical term in English only. Where a customer requires additional language support for safety-relevant text, we treat that as a product requirement rather than a preference.

12 Security of the AI itself

AI components introduce attack surface that conventional software does not:

- inputs from external sources, including Type B and AIDX messages and uploaded documents, are validated and treated as untrusted;
- we test against prompt injection and instruction-hijacking through ingested content;
- model access is authenticated, rate-limited, and logged;
- we protect against manipulation intended to produce a favorable but unsafe recommendation;
- AI-specific incidents are handled under the security incident process in FX-POL-011.

13 Incident handling

If an AI output contributes to an error, a missed defect, an incorrect record, or an operational event, we will:

1. notify the affected customer promptly, with what we know and what we do not;
2. preserve logs and the relevant model version for investigation;
3. support the customer's safety management system investigation and any regulatory reporting, and provide technical evidence to the authority on the customer's instruction;
4. establish root cause, state whether the behavior is reproducible, and give a remediation plan with dates;
5. disable or restrict the affected capability in the interim where safety requires it.

We do not contest a customer's decision to report an event to its authority, and we do not condition support on confidentiality about safety matters.

14 Prohibited uses

We will not supply, configure, or knowingly permit the platform to be used to:

- generate or alter maintenance records in a way that misrepresents who performed or certified work;
- fabricate, backdate, or reconstruct compliance evidence;
- replace a required inspection, test, or physical verification with a model output;
- justify circumventing an MEL limitation, AD compliance, or a mandatory instruction;
- produce a release to service or airworthiness determination without qualified human authorship;
- monitor or evaluate individual employees where that use is unlawful or not agreed.

Where we become aware of such use, we will raise it with the customer and, if unresolved, treat it as grounds for restriction or termination.

15 Frameworks we align with

We align our practice with the NIST AI Risk Management Framework, ISO/IEC 42001 as a management-system reference, the EASA guidance on trustworthy AI in aviation for human-centric, assistance-level AI, and the transparency and human-oversight principles of the EU AI Act, together with the Saudi Data and AI Authority's AI ethics principles for deployments in the Kingdom.

ALIGNMENT IS NOT CERTIFICATION

Alignment is a statement of how we work. We do not hold certification against these frameworks unless and until it is stated on our trust page with a certificate reference and date.

16 Accountability and contact

The Managing Member is accountable for AI risk, model change decisions, and this statement. Questions, concerns, and reports about AI behavior, output quality, or the use of this platform can be sent to:

CONTACT

Farshad Mohammadi, Managing Member
 Foxtrot LLC, 12112 N Rancho Vistoso Blvd, STE 100, Oro Valley, AZ 85755, United States
 Email: compliance@foxtrotllc.com

This statement is reviewed annually and republished with a revision date. It is published openly, without gating, so that customers, regulators, and partners can hold us to it.

17 Approval

This Responsible AI & AI Ethics Statement is adopted on behalf of Foxtrot LLC and takes effect on the date shown on the cover of this document. It supersedes version 1.0.

DocuSigned by:
Farshad Mohammadi
4620BECE0E724B4...

15 September 2026

FARSHAD MOHAMMADI, MANAGING MEMBER

DATE

Revision history

VERSION	DATE	SUMMARY OF CHANGE	APPROVED BY
2.0	15 September 2026	Added synConnect scope detail, accessibility and language commitments, human rights cross-reference, and named FX-POL-011 in place of the generic information security reference. Replaced placeholder contact details with compliance@foxtrotllc.com.	Farshad Mohammadi, Managing Member
1.0	12 September 2026	First issue.	Farshad Mohammadi, Managing Member

Controlled electronically. The authoritative version is the one published at foxtrotllc.com/policies.