

AI Model Card: DiagAI Autofill v1.1

Compliance Framework: ISO/IEC 42001:2023

1. General Information

1.1. Basic Information

AI System Name	DiagAI Autofill
Model Identifier	AIMC-06
Model Version	1.1
Release Date	Planned 28 October 2026 (SeqOne Platform 2.2)
Model Type	AI-Powered Clinical Data Extraction and Intelligent Form Completion System
Development Team	GAA
Maintenance Status	Active
Related Products	Integrates with entire SeqOne Platform; complements DiagAI HPO
EU AI Act Classification	High-risk AI system (Article 6 & Annex 3.5b) – Ai intended to be used as a safety component in the management and operation of critical digital infrastructure, or as medical device

1.2. Licensing and Access

License Type	Proprietary
Access Level	Restricted to authorized clinical and research users
Terms of Use	Subject to SeqOne Platform Terms of Service and Clinical Use Agreements

2. Intended Use & Constraints

2.1. Intended Purpose

DiagAI Autofill is an AI-powered clinical decision support tool designed to reduce manual data entry burden, minimize transcription errors, and accelerate clinical workflows for genetic testing and rare disease diagnosis.

It automatically extracts meaningful information from structured and unstructured data such as SeqOne annotations and scientific publications and uses Large Language Models to generate a comment that summarizes key information about the variant and why it is appropriate to report the variant for the patient. It also includes information about recommended therapies in somatic context.

Primary Functions

- Extract clinically meaningful information from a genetic variation
- Summarize the information into a comment ready to integrate in a clinical report
- Extend or improve a draft comment from a user
- Provide clinical assistance in variant reporting

2.2. Target Users

Primary Users

Healthcare professionals responsible for routine clinical genomic analysis: geneticists, bioinformaticians, biologists, engineers, and NGS platform technicians.

Indirect Users

- Patients receiving diagnostic information
- Healthcare systems managing patient care
- Regulatory authorities overseeing device performance

2.3. Out-of-Scope Uses

- Primary medical record creation or modification
- Direct clinical diagnosis or treatment decisions
- Automated clinical reporting without mandatory human review
- Replacement of the informed consent process
- Billing code assignment for reimbursement (ICD, CPT codes)
- Legal or forensic documentation
- Prescription writing or medication ordering
- Direct patient-facing information collection
- Population health screening without clinical context
- Automated prior authorization submissions
- Quality of care assessments or provider evaluation
- Clinical trial eligibility determination without expert review
- Medical-legal documentation
- Compliance reporting or regulatory submissions without review

2.4 Deployment Constraints

Runtime Environment

Component	Detail
Interface	API calls and MCP calls
LLM Backend	Anthropic Haiku 4.5 (provided pre-trained; no custom training performed)
Input Format	Plain text (.txt) and structured data (.json)
Output Format	JSON as expected by the SeqOne Platform frontend
Upstream Integration	Entire SeqOne Platform; feeds tech report
Model Storage	AWS Bedrock

i No training was performed for DiagAI Autofill. The system uses the Anthropic Haiku model, which is already trained by Anthropic. SeqOne applies augmentation and few-shot prompting on top of this base model.

3. Model Architecture

Attribute	Value
System Type	Multi-source retrieval and augmentation pipeline with multi-pass LLM inference
Base Language Model	Claude Haiku 4.5, pre-trained by Anthropic, served through AWS Bedrock in the region of the application deployment.
Inference Approach	Retrieval-augmented generation with few-shot prompting. Deterministic decoding, temperature fixed at 0.
Technology Stack	REST API calls to internal and external knowledge sources, MCP calls for literature retrieval, managed LLM inference on AWS Bedrock
Hyperparameter Optimisation	Not applicable. No custom training is performed.
Prompt Management	Versioned prompt registry with aliased releases and a fallback to embedded templates
Explainability	Source provenance in the generated text, PMID citations, and a per-run execution trace identifier returned to the caller
Human Oversight	The biologist's own draft is an input, and the generated text is a proposal the biologist reviews and edits before validation

Overview

DiagAI Autofill generates a draft interpretation comment for a single genetic variant in a single analysis. It is an assistance function: it produces text that a biologist reads, corrects, and validates. It makes no classification decision of its own and writes nothing back to the patient record.

The system is organised in three layers. An orchestration layer receives the request and governs the sequence. A context assembly layer retrieves the clinical and scientific evidence relevant to the variant. An inference layer runs one or two language models, passes over that evidence and returns the assembled comment.

Orchestration

The orchestration layer exposes one endpoint, scoped to an analysis, a variant, and a transcript. It performs five governance functions before any generation occurs.

It authenticates the caller. Every downstream call to an internal service carries that caller's identity, so the evidence retrieved is restricted to the data the caller is entitled to see.

It validates the request. The optional inputs are the existing classification, the biologist's draft comment, the target language, the patient's HPO terms, and the patient indication. An invalid request is rejected before any model call.

It resolves the assay type of the analysis, germline or somatic, from the project that owns it. This value selects the clinical prompt. A failure to resolve it is surfaced as an error rather than defaulted, because guessing here would route a case to the wrong clinical reasoning.

It routes the request to the alteration path, small variant or copy number variant. Somatic copy number variants are declared not implemented and are refused explicitly.

It opens an execution trace for the run and returns the trace identifier with the response.

Core Pipeline Components

DiagAI Autofill employs a modular architecture where Large Language Models handle different information extraction tasks, coordinated by an orchestration layer that manages document processing, field routing, and result integration.

Component 1 — Variant Description and Classification Generation

- Architecture: LLM inference call with augmentation and few-shot prompt
- Functions: Variant annotation interpretation; scientific literature retrieval; variant classification and details augmentation
- Technology: API calls, MCP calls + LLM inference
- Input: txt + json | Output: json expected by the platform frontend
- Overview: This component produces the body of the comment: what the variant is, what it does, how the evidence supports its classification, and which publications support that reading. Its input is the complete evidence package, rendered into a structured prompt with few-shot examples, together with the biologist's draft and the existing classification when supplied. Its output is prose in the requested language. Prompt selection depends on the assay type, so germline and somatic cases are reasoned about under different clinical instructions.

Component 2 — Disease-Related Context Generation (Germline Analyses)

- Architecture: LLM inference call with augmentation and few-shot prompt
- Functions: Retrieves disease-related clinical context to support germline variant reporting
- Input: txt + json | Output: json expected by the platform frontend
- Overview: For germline analyses where disease associations were retrieved, a second inference pass produces the clinical context paragraph: the disease associated with the gene, its presentation, and its relevance to the patient's reported phenotype. This pass runs concurrently with Component 1 and is appended to its output.

Component 3 — Therapy recommendation (Somatic Analyses)

- Architecture: LLM inference call with augmentation and few-shot prompt
- Functions: Patient indications formatting, retrieval of recommendation from Somatic DiagAI, LLM inference with few shots
- Input: txt + json | Output: json expected by the platform frontend

Component 4 — Copy Number Variation comment (Germline analysis)

- Architecture: LLM inference call with augmentation and few-shot prompt
- Functions: Retrieves disease-related clinical context to support germline variant reporting
- Input: txt + json | Output: json expected by the platform frontend
- Overview : Copy number variants follow a separate path, restricted to germline analyses. The system retrieves the large variant record, the genes it overlaps with their associated diseases, and the clinical regions it intersects. It adds the patient's sex, the genome build, and the identified proband sample, and matches the region against a reference catalogue of known clinical intervals. A single inference pass then produces the comment. Retrieval failures on the gene and clinical-region sub-resources degrade the evidence available rather than failing the request; the aggregate dosage information already present on the variant record is used instead.

4. Training Data & Provenance

⚠ No custom training was performed for DiagAI Autofill. The system relies on Anthropic Haiku 4.5, a pre-trained foundational Large Language Model provided by Anthropic. SeqOne does not possess the training data for this base model. Augmentation and few-shot prompting are applied at inference time.

4.1. Training Data

Not applicable. DiagAI Autofill uses a pre-trained foundational model (Mistral Large). No proprietary training data was used or generated by SeqOne for model training.

Few-shot examples used for prompt augmentation are not training data in the traditional sense; they are curated prompt artefacts managed as configuration, not datasets.

4.2. Validation Datasets

DiagAI Autofill performs no custom training, so there is no training/validation split in the machine-learning sense. Validation is conducted as a **retrospective performance qualification** of the assembled pipeline against expert-authored reference comments. Two datasets are in scope, one per variant class.

Attribute	DS-17 (SNV / indel)	DS-18 (CNV)
Dataset file	snv_germline_dataset.tsv	cnv_germline_dataset.tsv
MLflow dataset digest	4e673fda	a08db7a4
Analysis context	Germline	Germline
Rows in source file	62	32
Rows with an expert reference	62	17
Reference standard	ClinGen expert-curated variant comments	Comments authored by a SeqOne clinical scientist
Independence	Independent of SeqOne — external expert curation	Internal expert reference; independent of the development team but not of the organisation

Scope and exclusions. Both datasets cover germline analyses only. The somatic therapy-recommendation component (Component 3) and the germline disease-context component (Component 2) are **not** covered by these two runs; they require a separate qualification before they can be declared validated.

4.3. Data Preprocessing

Annotated variants were directly extracted from the SeqOne Platform for both Short variants and large variants. Comments are generated using the autofill application directly.

5. Performance Metrics & Evaluation

5.1. Performance Metrics

Evaluation Process

Evaluation is automated and executed by the versioned `judge` harness. Every run records the code revision, the prompt hash, the judge configuration and the dataset digest, so any reported figure can be traced back to a specific execution.

Parameter	Value (both runs)
Judge model	<code>gpt-5.4-nano</code> , temperature 0
Judge repeats / retries	1 repeat, up to 3 retries on transient failure
Sampling seed	0
Deterministic pre-check	Enabled (literal matching before LLM judging)
Judge rubric	Logged per run as artefact <code>prompts/judge_prompts.txt</code>

The scheme has two tiers.

Tier 1 — finding taxonomy. Each generated comment is examined for twelve pre-declared defect types, identical across both runs: fabricated value, distorted value, unsupported claim, invented citation, missing ACMG tag, dropped field, partial field, missing caveat, missing category, irrelevant content, contradicts expert, misses expert point. These roll up into the reported metrics — fabrication and citation defects into *groundedness*, value and claim defects into *correctness*, omission defects into *completeness*, and the two expert-disagreement defects into *agreement precision and recall* against the reference comment.

Tier 2 — deterministic gates. A set of rule-based checks is applied per variant class and rolls up into *instruction following*. The gate sets are deliberately different because the two output formats have different clinical and editorial constraints, and they are of different size:

- Short variants (11 gates): no fabrication, no wrong field, no hallucination, no repeated sentence, disease wording, PMID grounded, language, loss-of-function mechanism, ACMG classification, plain prose only, missense predictors.
- Large variants (21 gates): paragraph count, paragraph 1 and 2 word counts, overall word count, no repeated sentence, language, objective paragraph, copy-gain wording, no missing-data wording, no clause dash, zygosity restricted to paragraph 2, no time interval, `no_achropuce_when_absent`, no safety violation, no fabrication, no wrong field, no hallucination, VUS not used for counselling, DGV discordance flagged, no haploinsufficiency claim for a gain, carrier statement requires a heterozygous call.

Because the gate sets differ in both size and severity, **instruction-following scores must not be compared between the two variant classes.**

Overall Field-Level Performance

All values are percentages, from the two MLflow runs cited above.

Metric	Short variants (SNV, n = 62)	Large variants (CNV, n = 32)
Groundedness	98.8	98.4
Correctness	96.8	98.5
Completeness	100.0	99.2
Agreement — precision	95.5	96.2
Agreement — recall	95.8	99.1
Instruction following	53.2	60.6
Execution errors	0	0

Interpretation

The clinical substance of the generated comments is strong and consistent across both variant classes. Groundedness and correctness sit between 96.9 and 98.8, meaning the judge found almost no fabricated values, distorted annotation values, unsupported claims or invented citations. Completeness is at or near ceiling (100.0 and 99.2): the required evidence categories are essentially always present. Agreement with the expert reference is high in both directions — precision 95.5–96.2 (what the system asserts is what the expert asserted) and recall 95.8–99.1 (the system rarely omits a point the expert made).

The deficiency is **instruction following** — 53.2 for short variants and 60.6 for large variants. This measures compliance with the editorial and formatting gates, not clinical accuracy: paragraph and word-count constraints, prose style, wording conventions, and placement rules. Roughly half of the gate checks on short-variant comments, and two in five on large-variant comments, are not satisfied.

The consequence is bounded by the system's operating model. Autofill produces a proposal that a biologist reads, edits and validates before anything reaches a report; it makes no classification decision and writes nothing back to the patient record. A gate failure therefore manifests as additional editing effort for the biologist, not as an unreviewed clinical error. It remains a material quality finding: it degrades the usefulness of the assistance function and should be treated as the priority improvement target for the next release.

Known limitations of the validation

1. **Sample size.** 62 short-variant and 32 large-variant cases, single site. Confidence intervals around each figure are wide; no interval is currently computed or reported.
2. **Automated reference judging.** Features are extracted by a LLM ([gpt-5.4-nano](#)), then aggregated into a score. No human-versus-judge concordance study has been performed, so the judge itself is not yet qualified as a measuring method.

- 3. **No reproducibility estimate.** Each run used a single judge pass (`judge_repeats = 1`). Temperature 0 and a fixed seed reduce but do not eliminate run-to-run variance, and no variance figure is available.
- 4. **Partial reference independence.** The large-variant reference comments were authored by a SeqOne clinical scientist rather than an external source.
- 5. **Coverage.** Germline only; the somatic component is unvalidated.

Planned corrective and preventive actions

#	Action	Rationale
1	Run a human-versus-judge concordance study on a sampled subset	Qualifies the LLM judge as a measuring method
2	Extend validation to the somatic therapy-recommendation component and the germline disease-context component	Validation datasets only contain Germline comments.

6. AI Risk Assessment & Impact

6.1. Identified Risks

ID	Hazard Description	Consequence / Harm	P	S	Risk Level	Mitigation Strategy	Residual Risk
RSK-76	Explainability Transparency - User interface is not understood / Misunderstanding of scores, rankings, or evidence. - AI-powered "black box" logic provides clinical scores without clear biological reasoning.	Incorrect results - Misinterpretation of data	2	2	Low	Mitigated by comprehensive training and clear UI design (in-product tooltips and contextual help on each score component; user manual section dedicated to the interpretation of evidence and rankings.) The system shall display a DiagAI score summary, with the different score sub-components : UP2, Phenogenius, Family Pattern, Inheritance compatibility, Quality Checks, PanelApp. (SR-616)	Low
RSK-84	Over-reliance on AI / clinical misinterpretation Predictions used without appropriate clinical context or expert review	Human fails to catch AI errors — incorrect decisions acted upon uncritically; missed pathogenic variant; erroneous clinical decision; potential patient harm	3	2	Medium	Mandatory human-in-the-loop validation; clear UI warnings; system labeled as decision support; predictions accompanied by confidence scores	Low
RSK-88	Data Management & Data Quality Reference database error propagated into DiagAI ClinVar, gnomAD, COSMIC or other reference databases used by DiagAI may contain erroneous, outdated, or conflicting variant annotations. These errors propagate into DiagAI pathogenicity classifications without the model detecting the source error, potentially causing systematic misclassification of variants across all patients.	Incorrect results - Misinterpretation of data	3	2	Medium	Independent validation step before database release to production.	Low
RSK-91	Model Drift / Knowledge Base Obsolescence Performance degradation as new genes, diseases, and variants are discovered after model training, and	Reduced accuracy; outdated pathogenicity predictions; missed clinically relevant variants	2	2	Low	Major knowledge base updates in our bioresources. Some DB (ClinVar, OMIM) are also plugged to directly use the latest information from the pipeline, minimizing the risk of	Low

ID	Hazard Description	Consequence / Harm	P	S	Risk Level	Mitigation Strategy	Residual Risk
	as genomic databases and evidence evolve post-training					missing an important annotation. Models are retrained at least annually	
RSK-94	Out-of-distribution (OOD) input / scope creep AI model used outside validated use cases	Unreliable outputs on unseen data patterns without detection	2	2	Low	Models are gatekept in variant specific pipelines.	Low
RSK-104	Data management & data quality Training data used to build or retrain a model is maliciously altered (data poisoning), a model artifact/weights file is substituted or tampered with without authorization, or crafted ("adversarial") input files are used to deliberately manipulate a model's output.	Incorrect results - Misinterpretation of data	3	2	Medium	Model artifacts are version-controlled and integrity-checked before deployment; access to training pipelines and datasets is restricted and logged in accordance with the ISMS (ISO/IEC 27001:2022/A1:2024); the release process requires validation of model provenance before a new version reaches production (see RSK-089).	Low
RSK-105	Generation of text responses from uncontrolled external sources (LitFinder / Autofill) LitFinder and Autofill use external literature and document sources that are not curated, version-controlled, or quality-checked by SeqOne, yet are generally reputed reliable. The generation step itself (LLM-based drafting/summarization) can also introduce content not actually supported by the retrieved source (confabulation), or misattribute a claim to a source that does not state it. Because the resulting text is presented as a polished, ready-to-use comment, the user may not systematically verify it against the original source before including it in a clinical report.	Incorrect results - Misinterpretation of data	3	2	Medium	Every AI-generated comment/autofilled field must be clearly labelled as AI-generated and traceable to its source document(s)/reference(s), so the user can verify it before use; mandatory human review and validation is required before any AI-generated comment or autofilled field is included in a clinical report (as already implemented for the AI-generated comment feature, SR-338 / SR-582). Every cited article has its id verified and is accessible in the interface Date and other metadata are displayed Retracted articles have explicit mentions in the associated pubmed links.	Low
RSK-108	Post-production monitoring &	Incorrect results - Misinterpretation of data	2	2	Low	Post-market surveillance activities are documented in the PMS Procedure, and will	Low

ID	Hazard Description	Consequence / Harm	P	S	Risk Level	Mitigation Strategy	Residual Risk
	performance degradation over time No structured post-market monitoring of ML-specific performance indicators (data drift, concept drift) is in place beyond general PMS activities. Gradual degradation of classification/scoring accuracy in the field is not detected in a timely manner					be followed on a quarterly basis.	

P = Probability (1=Improbable — 1 in 10 000 000; 2=Occasional — 1 in 1 000; 3=Probable — 1 in 100; 4=Frequent — 1 in 1) | S = Severity (1=Negligible — no/negligible patient risk; 2=Minor — temporary impairment, incorrect/missing variant; 3=Major — injury requiring medical intervention, false diagnostic; 4=Critical — death, permanent impairment) | Risk Level derived from SEQ-011 Risk Acceptability Matrix.

6.2. Impact Assessment

Question	Answer
High-impact individual decisions?	Yes — AI-assisted clinical variant reporting and diagnostic support
Sensitive data processed?	Yes — genomic data and clinical health information (PHI)
Overall risk level	High (prior to controls) — Low/Medium (post mitigation)
Referenced AIIA document	AIIA-001_DiagAI

7. Ethical Considerations

7.1. Fairness and Bias

Identified Bias Areas

Bias Type	Issue	Impact	Mitigation
Healthcare Setting	Training data primarily from academic medical centres with comprehensive documentation	Potentially reduced performance for community practices	Partner with community hospitals; balance academic vs. community sources; target ≥30% community setting training data
Document Format	Better performance on common EHR formats (Epic, Cerner) vs. non-standard systems	Institutions using less common systems may see reduced benefit	Format-agnostic extraction; include non-standard and legacy formats in evaluation
Language	System optimised for English-language documents	Reduced or no support for non-English documents	Integrate a Large Language model with good performances on multi language benchmarks.
Clinical Specialty	Optimised for genetic testing workflows	May not generalise to other specialties without adaptation	Performance validated in genetics context only — clearly documented as deployment constraint
Socioeconomic	Better documentation quality from well-resourced institutions	Potential disparity in benefit for under-resourced settings	Active data collection from diverse settings; quality-independent feature engineering

7.2. Transparency and Explainability

Interpretability Methods

Feature	Description	Availability	Limitation
Source Document	Original text that led to each extraction is linked as a citation source; click any field to see source	Standard — all users	Limited for multi-document merged extractions
Audit Trail	Complete immutable log: timestamp, user, action, original and changed values	Administrators and QA personnel	Accessible to authorised administrators only

Human Oversight Architecture

Oversight Principle	Implementation
Oversight Model	Human-in-the-loop: all AI outputs must be explicitly reviewed and confirmed by a qualified operator before any data is committed
Decision Authority	The AI system provides extracted data with confidence scores. Final decision and accountability remain with the qualified human operator at all times
Override Mechanism	Operators can accept, modify, or reject any AI output. All decisions — including overrides — are logged with timestamp, user ID, and optional rationale
Uncertainty Communication	Low-confidence extractions are visually flagged (red/yellow) and require explicit confirmation before proceeding

Oversight Principle	Implementation
Audit Trail	All AI outputs, confidence scores, and subsequent human decisions are logged for regulatory retention periods in the SeqOne Platform
Failure Mode Handling	If the AI system is unavailable or returns an error, the process falls back to manual entry — no silent degraded mode

7.3. Privacy & Data Protection

Legal Framework	GDPR (EU); ISO 27001; ISO 27701; HIPAA (US);
Encryption	AES-256 at rest; TLS 1.3 in transit
Security Certification	ISO 27001 certified; Annual penetration testing
Access Control	Role-based access control (RBAC); multi-factor authentication (MFA) required ; principle of least privilege
Data Pseudonymization	All patient data pseudonymized before use in model development and validation
Audit Logging	All data access events logged (HIPAA compliant)
Data Retention	GDPR Article 17 right to erasure implemented; automated deletion after retention period; secure disposal procedures

7.4. Clinical Ethics

- System designed to augment — not replace — clinical judgment; no automated clinical decisions without human oversight
- Clear delineation between AI suggestions and clinical determinations
- Informed consent required: clear communication of AI involvement, explanation of limitations,
- Genetic counseling must be involved in result disclosure to patients

8. Monitoring and Maintenance

8.1. Known Limitations

Limitation	Impact	Communication Required?
No custom training — relies on Claude Haikubase model	Performance bound by base model capabilities; no domain-specific fine-tuning	Yes — document in user guide
No Handwriting recognition	raw text must be provided, not handwritten text photo.	Yes — flagged in UI
Novel or non-standard document formats	Reduced accuracy for formats not validated; unusual layouts challenge section detection	Yes — manual review triggered
Implicit or inferred information	Difficulty extracting information stated indirectly; inference limitations	Yes — completeness checks flag gaps
Ambiguous medical abbreviations	Abbreviations with multiple meanings may be disambiguated incorrectly	Yes — expansion shown for user verification
Multilingual documents	Reduced accuracy for non-English text; code-switching challenging	Yes — language detection warning displayed
Temporal ambiguity	"Last year" without current date context; relative dates difficult without anchor	Yes — temporal uncertainty flagged
English-only validation	Performance not validated for non-English input languages	Yes — explicitly documented as out-of-scope

8.2. Documentation Update Triggers

Document	Update Trigger
Model Card (this document)	Within 30 days of any model changes
User Manual	Within 14 days of UI or functionality changes
Risk Assessment	Quarterly review; updated as needed
Performance Reports	Monthly generation

8.3. Decommissioning Plan

- Formal decommissioning decision approved by AI System Owner and top management
- All model artifacts (weights, training data, configurations) archived securely per data retention policy
- All integrations and API endpoints disabled; operators notified with 5 days' notice
- Final version of Model Card and all AIMS records archived in the document management system
- Post-decommissioning review: lessons learned captured for future model development

9. Governance & Accountability

9.1. Regulatory Compliance

AI Management	ISO/IEC 42001:2023 — AI Management System
Medical Device	EU IVDR (2017/746) — In Vitro Diagnostic Regulation
Data Protection	GDPR (EU); HIPAA (US); PIPEDA (Canada); national data protection laws as applicable
Clinical Standards	ACMG/AMP Guidelines for Sequence Variant Interpretation; ACMG Recommendations for Reporting Secondary Findings
Clinical AI	EU AI Act — High-risk AI system classification assessment ongoing
Digital Health	NICE Evidence Standards Framework for Digital Health Technologies
US FDA	FDA guidance on Clinical Decision Support Software

9.2. Approval & Accountability

Role	Responsibility
Model Developer / AI Team	Technical development, matrix construction, internal validation, maintenance
Data Scientist	Algorithm design, data source integration, performance analysis
Product Manager	Minor update approval; user requirement management
Compliance Officer	Regulatory compliance oversight; GDPR/IVDR coordination
Medical/Scientific Advisor	Clinical performance interpretation; guideline alignment

10. Version History

Version	Description of changes	Approval date
1	Initial Model Card creation	2025-12-05
2	Update following the Fall26 (2.2) release	See approval date

Table of signatures	Name / Function	Date	Signature
Written by	Nicolas Duforet Datascience manager	2026-09-25	Nicolas Duforet
Reviewed by	Nicolas Philippe Chief Scientific Officer	2026-09-25	Nicolas PHILIPPE
Approved by	Amélie Martinez AIGO	2026-09-25	Amélie Martinez