

Platform overview

Data first, then models, not the other way around.

A platform overview for preparing enterprise data and models so AI workflows can move beyond demos.

The Data & Model Platform sits inside the same operating chain as Process Discovery, AxeStudio, and Continuum — so prepared data moves directly into production workflows without a hand-off to a separate vendor or a separate integration project.

ASSET PROMISE

elps organizations make data and models usable for real AI operations.

BEST AUDIENCE

Data leaders, ML/AI teams, CIOs, AI transformation teams

RECOMMENDED USE

Use before or after an enterprise briefing to help the buyer align internal stakeholders and define a practical next step.

METHODOLOGY

Data and model readiness follows a six-stage sequence. Each stage produces a working artifact — not a report. Skipping a stage means the model or workflow that follows is operating on data whose quality is unknown.

01

Inventory source systems

List every system, repository, and document store the target workflow depends on. For each source, record who owns it, how it is accessed, how frequently it changes, and whether it has ever been used for training or evaluation. Organizations routinely surface sources they did not know existed and find sources they assumed were clean that are not.

02

Cleanse and normalize

Remove duplicates, resolve conflicting records, standardize formats, and handle missing values. The goal is not perfection — it is a known, documented level of quality that the team can defend. Every cleansing decision should be logged: what changed, why, and who approved it.

03

Label and annotate

Assign structured signal to raw data so a model can learn from it or an evaluation set can measure against it. Labeling requires a schema, a labeling guide, inter-annotator agreement checks, and a review path for edge cases. Labels created without a guide produce noise, not signal.



04

Prepare structured datasets

Assemble training sets, fine-tuning corpora, retrieval indexes, and evaluation sets from the cleansed, labeled source material. Each dataset should carry provenance: which sources contributed, which version of the cleansing and labeling pipeline produced it, and what quality checks it passed.

05

Evaluate and validate

Run the evaluation dataset against the model or retrieval system before any production dependency is created. Record benchmark results, failure categories, and the conditions under which performance degrades. An evaluation set without a recorded baseline cannot tell you whether a future model version improved or regressed.

06

Connect to MLOps and LLMOps

Hand the validated data assets and evaluation baselines to the operations layer. Performance monitoring, drift detection, retraining triggers, and version management all require a stable reference point — which the preceding five stages produce. Without this connection, data preparation is a one-time event rather than an ongoing operational discipline.

01

Why data readiness decides AI success

Enterprise AI fails when the model is asked to operate on messy, incomplete, untrusted, or inaccessible data. The Data & Model Platform helps prepare the foundation that agentic workflows depend on.

02

What the platform supports

The platform is designed to support the practical data and model work required before and after AI deployment.

• Data cleansing

• Data labeling

• Dataset preparation

• Model training and fine-tuning support

• Evaluation datasets

• Document and knowledge preparation

• Quality checks

• MLOps and LLMOps alignment

03

How it connects to the rest of Axeron

Process Discovery identifies what data the workflow needs. The Data & Model Platform prepares that data. TeamAI and custom systems use it. Continuum governs behavior. AIOps/LLMOps monitors performance over time.

04

Buyer value

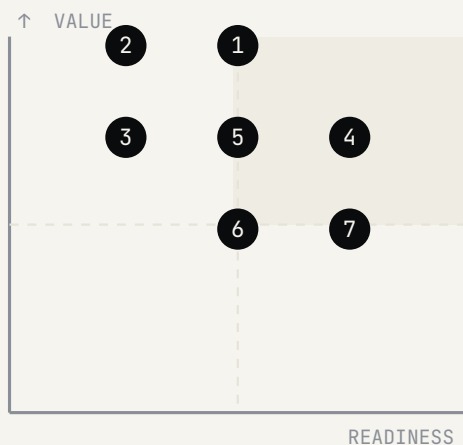
The platform reduces the gap between AI strategy and usable AI systems by turning fragmented information into structured, evaluated, and operationally useful assets.

Best-fit situations

This asset is most relevant when the buyer has scattered data, inconsistent documents, poor labeling, model performance concerns, or plans to train, fine-tune, evaluate, or operate models internally.

AI-OPPORTUNITY MATRIX

The Data Asset Priority Matrix helps teams decide which data sources to prepare first. Placement reflects two qualitative judgments: how much the target AI workflow depends on this data type, and how available and clean that data already is. Neither axis is a precise score — the purpose is to surface where preparation effort will produce the most immediate workflow impact.



Prepare now

High impact, high availability. The workflow depends on this data and it is already accessible in a known location. Cleansing and structuring this source unblocks the workflow immediately.

Invest to unlock

High impact, low availability. The workflow cannot operate without this data, but it is scattered, unlabeled, or inaccessible. Preparation here is non-optional — delay here delays the workflow.

Automate maintenance

Lower impact, high availability. The data is clean and accessible but not central to the first use case. Establish a routine quality process and return when the use case scope expands.

Deprioritize

Low impact, low availability. Expensive to prepare and the current workflow does not depend on it. Revisit after the first use case is in production.

- Unstructured internal documents
- Labeled training examples
- Evaluation and test sets
- Structured transaction records
- Historical case records
- External reference data
- Audit and compliance logs

WORKFLOW-READINESS SCORECARD

Run this rubric before committing to a data preparation workstream. Score each dimension as Weak, Workable, or Strong. A Weak on Data Ownership or Evaluation Baseline is a blocker — the platform cannot produce trustworthy outputs if those dimensions are unresolved. A mix of Workable and Strong across the remaining dimensions means the team can begin.

01 Source system inventory

Does the team have a complete list of every system, file store, and document repository the target workflow will draw from? Strong: every source is named, owned, and accessible with a defined query or export path. Workable: most sources are known; one or two need confirmation from system owners. Weak: the team is not certain which systems hold the relevant data.

**02 Data ownership**

Is there a named person or team responsible for each source dataset — someone who can approve changes, resolve conflicts, and handle access requests? Strong: every source has a named data owner with the authority to act. Workable: ownership is known but shared informally across two teams. Weak: data ownership is unclear or contested.

03 Labeling capacity

Does the organization have the domain expertise and bandwidth to produce accurate labels for the data the model will depend on? Strong: subject-matter experts are available, a labeling guide can be written, and inter-annotator review is feasible. Workable: expertise exists but is scarce; labeling will need to be prioritized and scoped tightly. Weak: the people who understand the data well enough to label it accurately are not available for this work.

04 Evaluation baseline

Is there an agreed set of test cases or benchmark questions the model or workflow must answer correctly before going to production? Strong: a representative evaluation set exists or can be assembled from historical cases within two weeks. Workable: the team can construct one, but it will take longer and require subject-matter input. Weak: no evaluation baseline exists and there is no clear path to creating one.

05 MLOps and LLMOps alignment

Is there a defined path for how prepared data and model versions will be versioned, monitored, and updated after initial deployment? Strong: an MLOps or LLMOps practice is in place with defined tooling, ownership, and retraining criteria. Workable: the team understands the need and has tooling candidates, but the process is not yet defined. Weak: production monitoring and retraining are not yet on anyone's roadmap.

06 Data governance and access control

Are there defined rules for who can access, modify, and export each data source — and are those rules enforced in the current systems? Strong: access controls are in place and documented; sensitive data handling rules are written and followed. Workable: controls exist but are inconsistently applied; a documented policy would close the gap within the project window. Weak: access to source data is informal and undocumented.

WHY AI PILOTS FAIL

Most data preparation failures are not technical. They follow a small set of organizational patterns that are recognizable before the work begins. Each one below is a condition the Data & Model Platform is designed to surface and resolve before any model depends on the underlying data.

WHY AI PILOTS FAIL

- × Do not begin model training or fine-tuning before the evaluation set exists. A model trained without a benchmark cannot be compared to its successor — the team loses the ability to know whether any future change improved or degraded performance. The evaluation set is not a late-stage deliverable; it is a prerequisite for the first training run.
- × Do not treat data cleansing as a one-time project. Enterprise data changes: systems are updated, records are added, formats shift. A cleansing pass that is not connected to an ongoing quality process produces a dataset that decays the moment it leaves the pipeline. Build the monitoring step before the model depends on the data.
- × Do not label data without a written guide. Labels produced from individual judgment rather than a shared schema introduce systematic inconsistency that compounds as the dataset grows. Every labeling task needs a guide, edge-case examples, and a review step before the labels enter a training or evaluation set.
- × Do not use the same data for both training and evaluation. A model evaluated on data it trained on cannot be trusted on data it has not seen. Partition the source data before labeling begins — not after — so the split is clean and the evaluation results are meaningful.

- Write the evaluation criteria before labeling begins. Knowing what the model must get right shapes which data to label, how to label it, and how many examples are needed. Evaluation-first discipline prevents the team from producing a large labeled dataset that does not measure what matters.
- Log every cleansing and transformation decision. When a data quality problem surfaces in production, the team needs to trace it back to a source record and a transformation step. Without a log, every investigation starts from scratch.
- Assign a named owner to each dataset before preparation begins. The owner approves the labeling guide, resolves conflicts, and accepts the dataset into production. Datasets without owners drift.
- Connect the first evaluation baseline to the pilot KPI defined in Process Discovery. The data team and the workflow team should be measuring the same thing — otherwise a model that scores well on internal benchmarks can still fail to move the operational metric the business cares about.

WHAT YOU WALK AWAY WITH

At the close of a data preparation engagement, the team has working artifacts — not a summary of findings. Each output is designed to be used by the implementation team immediately: the engineers building the agent workflow, the platform team configuring MLOps, and the governance team setting up Continuum.

- 01 Cleansed and normalized source datasets with a documented transformation log — recording what changed, why, and which version of the cleansing pipeline produced the output.
- 02 Labeled training and fine-tuning corpora with provenance metadata: source systems, labeling guide version, inter-annotator agreement results, and review sign-off.
- 03 Evaluation and test sets partitioned from the source data before labeling, covering representative cases and known edge conditions, with a recorded baseline score for the first model version.
- 04 Data Asset Priority Map showing which source types were prepared, their readiness and impact ratings, and the recommended sequencing for any remaining preparation work.
- 05 MLOps and LLMOps integration notes — source system connections, dataset versioning paths, retraining trigger criteria, and the monitoring baseline the operations team will track against.

WHERE IT FEEDS

The Data & Model Platform is not a standalone capability. Process Discovery identifies which data the target workflow depends on and where that data lives. The platform then prepares it — cleansing, labeling, structuring, and validating — so the agentic workflows built in AxeStudio have a trustworthy foundation to operate on. TeamAI and custom agent systems draw from the prepared data and evaluation sets at runtime. Continuum governs how the model behaves in production and gates self-improvement proposals before they affect live behavior. AIOps and LLMOps monitor performance over time, detect drift, and trigger retraining when the evaluation baseline indicates degradation. The design constraint throughout is that data preparation, model operations, workflow implementation, and governance run inside the same operating model. A team that prepares data with Axeron does not hand it to a different vendor to build the workflow — the prepared data moves directly into the implementation phase without reinterpretation.

BEFORE YOU START

- ☐ Identify source systems
- ☐ Define labeling needs
- ☐ Define model operations path
- ☐ Assess data quality
- ☐ Prepare evaluation sets
- ☐ Connect data readiness to use-case KPIs

NEXT STEP

Use the Data & Model Platform when data readiness is blocking AI execution.

Axeron turns resource-level education into a scoped conversation: one process, one measurable outcome, one deployment model, and one governance path.