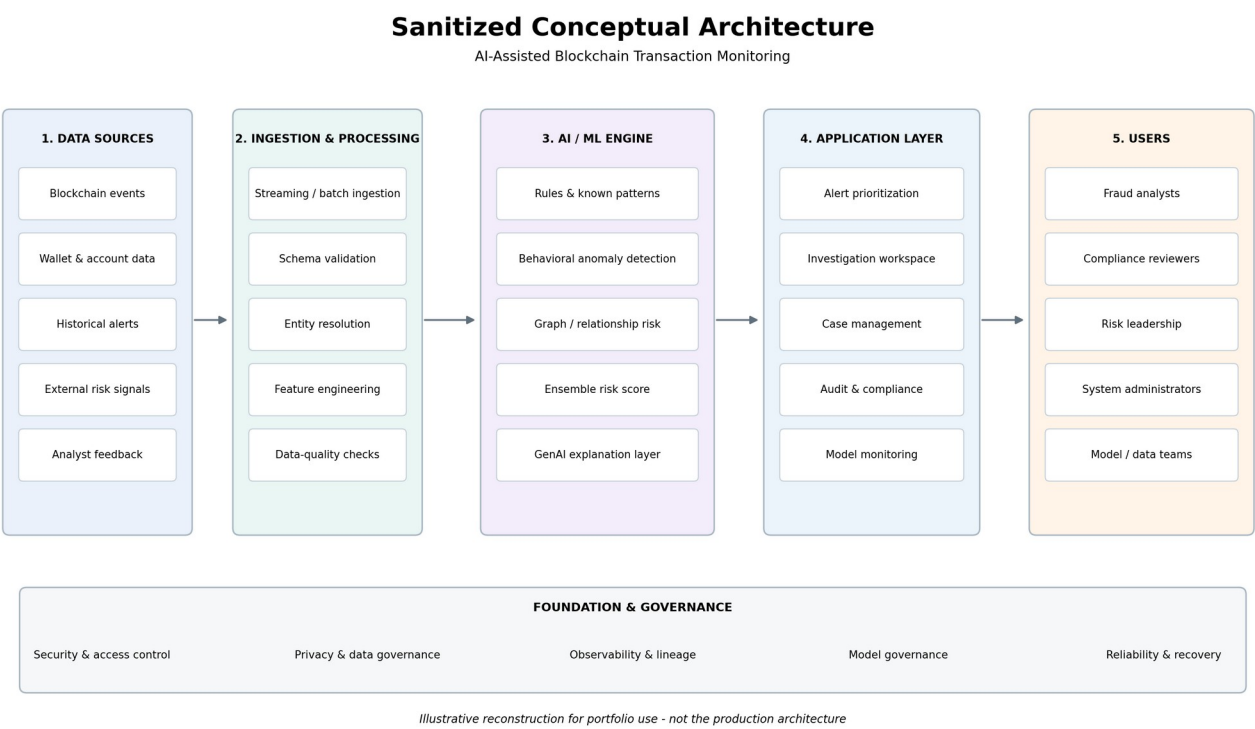


AI Fraud Detection Portfolio Package

Essential product management documents for a confidential blockchain transaction-monitoring project



Sanitized conceptual architecture reconstructed from memory; not the exact production design.

DOCUMENT STATUS Final Portfolio Edition	VERSION 1.0	AUTHOR Min Thu Kyaw	DATE July 2026
ROLE AI Product Manager & Solution Architect	COMPANY Confidential Blockchain Company	DOCUMENT CODE AFM-00	PUBLICATION Portfolio-safe reconstruction
PORTFOLIO PURPOSE Explain the case study boundary, document map, evidence model, and portfolio presentation approach.	EVIDENCE BASIS User-provided high-level architecture and user journey	CONFIDENTIALITY No code, data, PRD, metrics, vendors, thresholds, or internal screenshots	USE Hiring and portfolio review only

Document Control

Purpose	Provide a recruiter-friendly overview of the sanitized artifact set.
Primary audience	Hiring managers, AI product leaders, solution architects, risk leaders, and portfolio reviewers.
Evidence classification	Reconstructed fact / Illustrative design / Target / Hypothesis / Recommendation
Source limitation	Only high-level architecture and a representative user journey were supplied. No confidential artifacts are reproduced.
Related documents	AFM-01 through AFM-05
Change control	Update only when the sanitized scope, evidence basis, or portfolio narrative changes.

How to Read the Evidence Labels

Label	Meaning
RECONSTRUCTED	High-level information recalled and generalized by the author; exact implementation details are intentionally excluded.
ILLUSTRATIVE	A portfolio artifact created to demonstrate product-management method; not an internal company artifact.
TARGET	A proposed threshold to validate; never presented as an achieved result.
HYPOTHESIS	A belief that requires a defined test and decision rule.
RECOMMENDATION	A portfolio recommendation derived from the reconstructed problem and architecture.

Confidentiality and Publication Boundary

This is a sanitized portfolio reconstruction, not a reproduction of the confidential company product documentation.

The document demonstrates product thinking, AI product design, requirements discipline, system-level reasoning, and lifecycle governance without revealing proprietary implementation details.

Excluded content includes source code, production data, customer or wallet identifiers, fraud rules, thresholds, model parameters, vendor configuration, internal screenshots, exact architecture, security controls, incident history, operational volume, production performance, and commercial outcomes.

Architecture diagrams and workflows use generic terminology and are labeled as reconstructed or illustrative. Any metrics, roadmap items, release gates, or model-quality thresholds are targets for validation, not claims about the employer's production results.

Portfolio Narrative

I led the product and solution-level design of an AI-assisted fraud detection and transaction-monitoring capability in a confidential blockchain environment. The public portfolio version focuses on the problem, workflow, architecture, decisions, requirements, AI evaluation strategy, and measurement plan while intentionally excluding proprietary details.

Artifact Set

Code	Artifact	What it demonstrates
AFM-01	Product Concept Document	Problem framing, users, value proposition, product principles, scope, risks, and validation strategy.
AFM-02	Product Requirements Document	Testable functional, AI, data, security, analytics, and operational requirements with traceability.
AFM-03	AI Strategy & Evaluation Plan	Hybrid model strategy, feature categories, offline/online evaluation, HITL, explainability, drift, and GenAI controls.
AFM-04	As-Built Architecture & Traceability Report	Sanitized architecture, system boundaries, design decisions, requirement-to-component traceability, and limitations.
AFM-05	Pilot Measurement & Roadmap Plan	Pilot design, metric tree, instrumentation, gates, roadmap, backlog, and decision rules.

Recommended Portfolio Website Presentation

- Project card label: Confidential enterprise project - sanitized public case study.
- Primary CTA: View case study. Secondary CTA: Download portfolio document.
- Replace GitHub and internal PRD links with: Private company artifacts unavailable due to confidentiality.
- Show only recreated diagrams, representative workflows, and selected requirement excerpts.
- Keep exact metrics, thresholds, vendors, model parameters, transaction volumes, and security controls private.

End-to-End Detection and Investigation Workflow

Hybrid detection + human review + feedback loop

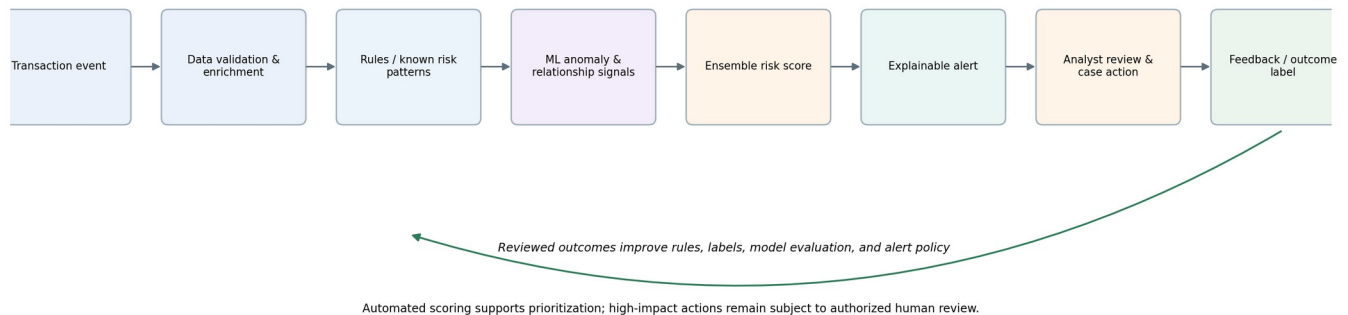


Figure AFM-00.1 - Sanitized detection and investigation workflow.

Product Manager Reflection

The product is the controlled decision workflow, not only the model

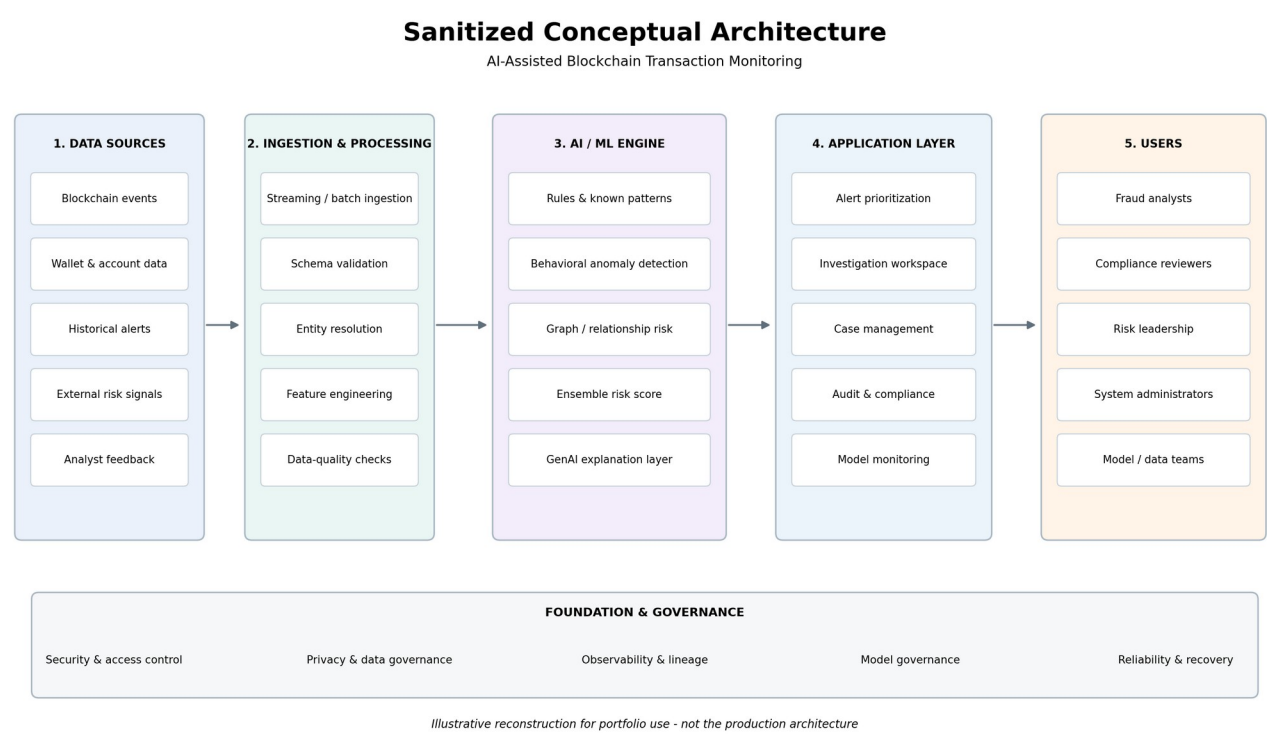
The architecture makes the AI models visible, but the customer outcome depends on the full operating loop: reliable data, calibrated prioritization, explainable evidence, analyst judgment, feedback capture, monitoring, and governance. A technically strong model can still fail as a product when alerts are unactionable, explanations are generic, workflows are fragmented, or outcomes are not captured for learning.

What this case study demonstrates

- Converting a high-risk business problem into measurable user and system requirements.
- Balancing fraud detection, false-positive cost, explainability, privacy, reliability, and analyst capacity.
- Designing hybrid AI where rules, machine learning, relationship analysis, and Generative AI have distinct responsibilities.
- Using evidence labels and phase gates to avoid presenting assumptions or targets as achieved outcomes.
- Working across product management and solution architecture while preserving confidentiality.

AI-Assisted Blockchain Transaction Monitoring System

Product Concept Document - sanitized confidential-project portfolio edition



Sanitized conceptual architecture reconstructed from memory; not the exact production design.

DOCUMENT STATUS Final Portfolio Edition	VERSION 1.0	AUTHOR Min Thu Kyaw	DATE July 2026
ROLE AI Product Manager & Solution Architect	COMPANY Confidential Blockchain Company	DOCUMENT CODE AFM-01	PUBLICATION Portfolio-safe reconstruction
PORTFOLIO PURPOSE Present the product opportunity, customer problem, solution thesis, AI role, scope, risks, and validation plan.	EVIDENCE BASIS User-provided high-level architecture and user journey	CONFIDENTIALITY No code, data, PRD, metrics, vendors, thresholds, or internal screenshots	USE Hiring and portfolio review only

Document Control

Purpose	Define the concept-stage product narrative and the evidence required to justify continued investment.
Primary audience	Hiring managers, product leaders, AI/ML leaders, risk and compliance stakeholders.
Evidence classification	Reconstructed fact / Illustrative design / Target / Hypothesis / Recommendation
Source limitation	Only high-level architecture and a representative user journey were supplied. No confidential artifacts are reproduced.
Related documents	AFM-02, AFM-03, AFM-04, AFM-05
Change control	Update only when the sanitized scope, evidence basis, or portfolio narrative changes.

How to Read the Evidence Labels

Label	Meaning
RECONSTRUCTED	High-level information recalled and generalized by the author; exact implementation details are intentionally excluded.
ILLUSTRATIVE	A portfolio artifact created to demonstrate product-management method; not an internal company artifact.
TARGET	A proposed threshold to validate; never presented as an achieved result.
HYPOTHESIS	A belief that requires a defined test and decision rule.
RECOMMENDATION	A portfolio recommendation derived from the reconstructed problem and architecture.

Contents

- 01 Executive Summary**
- 02 Problem and Opportunity**
- 03 Target Users and Jobs to Be Done**
- 04 User Journey
- 05 Product Vision and Principles
- 06 Value Proposition
- 07 Solution Concept
- 08 AI Opportunity Assessment
- 09 Scope and Non-Goals
- 10 Success Criteria
- 11 Assumptions and Risks
- 12 Validation Plan
- 13 Product Manager Reflection

01

Executive Summary

A high-level product concept for prioritizing suspicious blockchain transactions, supporting analyst investigations, and preserving auditable human decisions.

The product combines deterministic rules, behavioral and relationship-based machine-learning signals, an ensemble risk score, and an analyst-centered investigation workflow. Generative AI is positioned as an assistance layer for evidence-linked summaries, investigation guidance, and draft case narratives - not as the authoritative fraud decision-maker.

Concept statement

Transform fragmented blockchain activity and risk signals into prioritized, explainable, reviewable alerts that help authorized analysts investigate faster and make more consistent decisions.

Why this matters

- Rule-only systems can be transparent but brittle and may produce high alert volumes.
- Model-only systems can detect complex behavior but may be difficult to explain, govern, or calibrate operationally.
- Analysts need one coherent evidence trail rather than disconnected tools and raw signals.
- The product must optimize both detection quality and investigation capacity, not one metric in isolation.

02

Problem and Opportunity

The core product problem is not simply predicting fraud. It is deciding which events deserve attention, explaining why, and enabling a defensible human outcome.

Observed / reconstructed problem	Customer impact	Product opportunity
Large and changing alert volume	Analysts spend time triaging low-value alerts.	Risk-based prioritization and workload-aware alert policy.
Fragmented transaction, entity, and case context	Longer investigation time and inconsistent evidence assembly.	Unified investigation workspace and evidence timeline.
New fraud patterns evolve beyond static rules	Known-pattern coverage degrades over time.	Behavioral, graph, and anomaly signals with continuous evaluation.
Opaque scores reduce trust	Analysts may ignore or over-trust automation.	Reason codes, evidence provenance, confidence, and mandatory review.
Weak outcome capture limits learning	Rules and models do not improve from investigator decisions.	Structured disposition labels and feedback governance.

Opportunity statement

Build an AI-assisted transaction-monitoring product that improves the quality and speed of investigation decisions while maintaining human authority, auditability, privacy, and operational resilience.

03

Target Users and Jobs to Be Done

The primary user is an experienced fraud or risk analyst responsible for converting uncertain signals into an evidence-backed disposition.

Persona	Primary job	Success looks like	Key concern
Senior Fraud Analyst	When suspicious activity is detected, help me understand the highest-risk cases and assemble evidence so I can decide quickly and accurately.	High-risk cases are reviewed first, evidence is clear, and decisions are documented without repetitive work.	False positives, missing context, and explanation quality.
Compliance Reviewer	When a case requires oversight, help me verify evidence, rationale, policy alignment, and audit history.	Review is consistent, traceable, and reproducible.	Completeness, governance, and escalation policy.
Risk Operations Lead	When alert volume changes, help me manage capacity and quality across the team.	Workload, backlog, alert yield, and review quality are visible.	Operational efficiency and missed-risk exposure.
Model / Data Team	When model behavior changes, help me diagnose quality, drift, and feedback signals.	Performance can be evaluated by segment and version.	Data quality, label reliability, and monitoring.

04

User Journey

The experience must reduce context switching from the moment an alert is created through review, investigation, escalation, closure, and feedback.

Senior Fraud Analyst Journey

Representative persona: Emma Chen | Goal: Identify suspicious blockchain activity quickly while minimizing false positives and preserving investigation evidence.

Stage	User goal	Key actions	Primary friction	Product opportunity	AI support
1. Monitor alerts	Identify suspicious activity	Review queue, severity, account an	High volume and weak prioritiz	Rank alerts by business risk and	Risk score, reason codes, pattern sun
2. Review alert	Understand why the alert	Inspect activity timeline, connect	Fragmented context and incons	Create one explainable alert rec	Evidence-linked summary and confid
3. Investigate	Collect evidence and valid	Trace fund flow, compare history, d	Multiple tools and manual evide	Consolidate investigation conte	Relationship analysis, similar-case ret
4. Decide & report	Take action and preserve	Escalate, close, request review, doc	Repetitive reporting and weak f	Standardize decisions and outco	Draft case narrative with mandatory I

Illustrative user journey reconstructed for portfolio demonstration; not a production workflow or approved operating procedure.

Figure AFM-01.1 - Representative analyst journey. All persona details are illustrative.

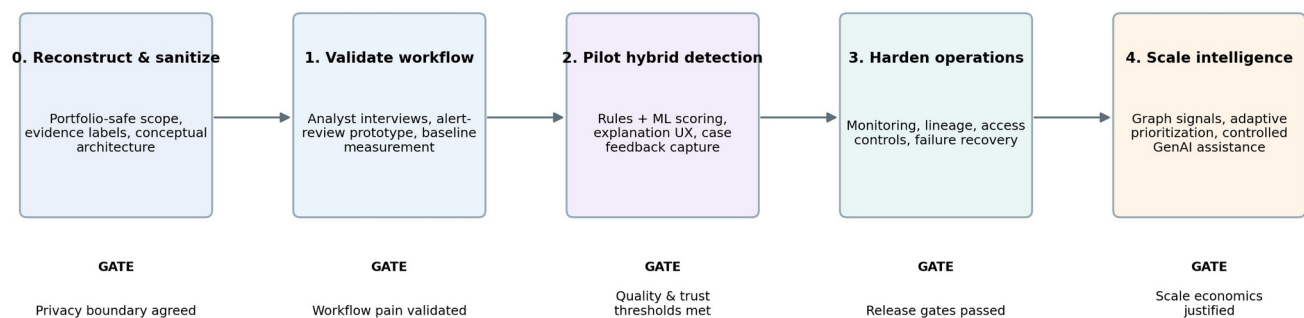
Product Vision and Principles

The long-term vision is a trusted risk-intelligence layer that connects transaction behavior, entities, evidence, analyst judgment, and governance.

Principle	Design implication
Human authority	Automation prioritizes and assists; authorized humans approve high-impact decisions.
Evidence before narrative	Every explanation should point to concrete signals, events, or policy reasons.
Risk-calibrated workflow	Thresholds and queues are designed around operational capacity and business impact.
Privacy by design	Access, retention, minimization, and auditability are product requirements, not afterthoughts.
Feedback with governance	Analyst outcomes improve evaluation and policy only after label-quality controls.
Fail safely	Provider, pipeline, or model failure produces explicit states and recoverable analyst workflows.

Vision roadmap

Evidence-Gated Product Roadmap



Roadmap sequence is a portfolio recommendation and does not reveal the confidential company delivery plan.

Figure AFM-01.2 - Evidence-gated product evolution, shown as a portfolio recommendation.

Value Proposition

The product helps risk teams focus on the cases that matter, understand the evidence, and complete defensible investigations more efficiently.

Customer pain	Product capability	Expected value
Too many low-value alerts	Risk-ranked queue and configurable policy	More analyst attention directed to meaningful cases.
Scattered evidence	Unified transaction, entity, relationship, and history context	Faster investigation and more consistent decisions.
Unclear model reasoning	Reason codes, evidence links, confidence, and model/version metadata	Greater trust and reviewability.
Repetitive case writing	Controlled GenAI draft with mandatory review	Less manual reporting effort without removing accountability.
Slow learning from outcomes	Structured disposition and feedback loop	Better evaluation, tuning, and rule maintenance over time.

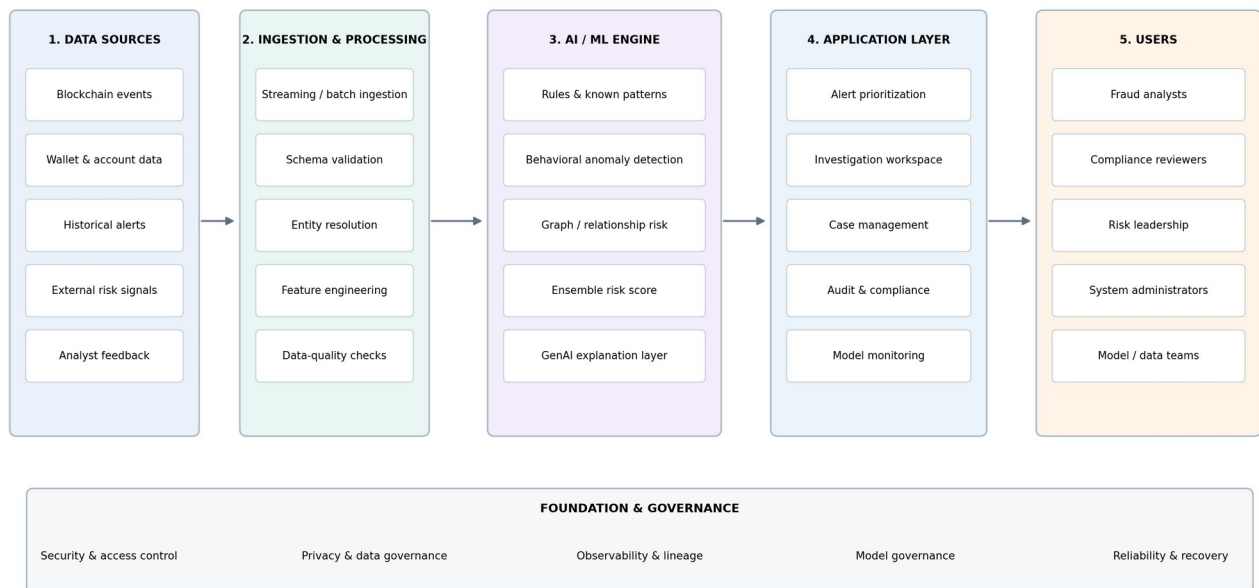
07

Solution Concept

The product is a layered system that separates data processing, risk detection, explanation, workflow, and governance responsibilities.

Sanitized Conceptual Architecture

AI-Assisted Blockchain Transaction Monitoring



Illustrative reconstruction for portfolio use - not the production architecture

Figure AFM-01.3 - Sanitized conceptual system architecture.

Core product modules

- Data ingestion, validation, entity resolution, and feature generation.
- Rules engine for known and policy-defined risk patterns.
- Machine-learning services for behavioral, anomaly, and relationship signals.
- Ensemble risk scoring and alert policy.
- Evidence-linked explanation and optional Generative AI assistance.
- Alert queue, investigation workspace, case management, and audit trail.
- Model, data, product, and operational monitoring.

08

AI Opportunity Assessment

AI is justified where the system must recognize patterns across time, entities, and relationships, but it should not replace explicit policy or accountable judgment.

Decision area	Best-fit approach	Why
Known prohibited or policy-defined pattern	Rules	Deterministic, explainable, and easy to audit.
Unusual behavioral change	Anomaly / behavioral model	Can identify deviations not captured by static rules.
Wallet or entity relationship risk	Graph / relationship analytics	Captures network structure and fund-flow relationships.
Alert prioritization	Calibrated ensemble score	Combines complementary signals and business policy.
Alert explanation	Template + evidence-linked generation	Improves usability while preserving source evidence.
Final case disposition	Authorized human decision	High-impact decision requires accountability and contextual judgment.

09

Scope and Non-Goals

A portfolio-safe MVP focuses on the end-to-end alert-review loop rather than attempting fully autonomous fraud prevention.

In scope	Out of scope / non-goal
Ingest a defined set of blockchain and account events.	Expose or reproduce confidential data sources and vendor integrations.
Generate risk signals and a combined alert score.	Claim autonomous legal, regulatory, or enforcement decisions.
Prioritize alerts and present evidence to analysts.	Guarantee detection of every fraud pattern.
Support investigation notes, decisions, and audit history.	Publish exact thresholds, controls, or production architecture.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-01
In scope	Out of scope / non-goal	
Capture reviewed outcomes for evaluation.	Use unreviewed Generative AI output as the final case record.	

10

Success Criteria

Success must balance customer value, AI quality, operational health, and governance. The values below are targets for pilot design, not achieved results.

Dimension	Target indicator	Decision use
Customer value	Reduced median investigation time for comparable alert types	Continue only if efficiency improves without lower decision quality.
AI quality	Higher precision in the top-priority queue than the baseline policy	Validate that prioritization adds operational value.
Trust	Analysts rate explanations useful and can identify supporting evidence	Improve or restrict AI assistance when trust is weak.
Operational	Alert scoring and data freshness meet agreed service levels	Block expansion until reliability is adequate.
Governance	Required access, audit, review, and incident controls pass release gates	No release when mandatory controls fail.

11

Assumptions and Risks

The largest uncertainties are operational and organizational as much as technical.

ID	Assumption / risk	Impact	Validation / mitigation
A1	Analysts have meaningful pain from current alert volume and fragmented context.	High	Interviews, shadowing, baseline timing, and queue analysis.
A2	Available labels are sufficiently reliable for supervised evaluation.	High	Label audit, adjudication guide, and inter-rater agreement.
A3	Hybrid scoring improves top-of-queue quality.	High	Offline backtest and controlled pilot against baseline.
R1	False negatives create material exposure.	Critical	Segmented evaluation, conservative release gates, post-event review.
R2	Explanations overstate certainty or omit evidence.	High	Evidence-linked templates, confidence display, human approval.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION			AFM-01
ID	Assumption / risk	Impact	Validation / mitigation
R3	Sensitive data is exposed through access, logs, or GenAI prompts.	Critical	Minimization, access control, redaction, provider boundary, audit.
R4	Fraud behavior or data distribution changes.	High	Drift monitoring, champion/challenger evaluation, rollback policy.

12

Validation Plan

Progress through explicit evidence gates rather than building the full vision at once.

1. Validate analyst workflow pain, baseline investigation time, and alert triage behavior.
2. Create a sanitized interactive alert-review prototype and test information hierarchy and explanation usefulness.
3. Build a versioned benchmark with reviewed labels and segment definitions.
4. Compare baseline rules, candidate models, and hybrid scoring using offline evaluation.
5. Run a limited pilot with audit logging, human approval, and clearly defined rollback criteria.
6. Decide whether to scale, revise, or stop based on customer value, AI quality, operational readiness, and governance evidence.

Product Manager Reflection

The product is the controlled decision workflow, not only the model

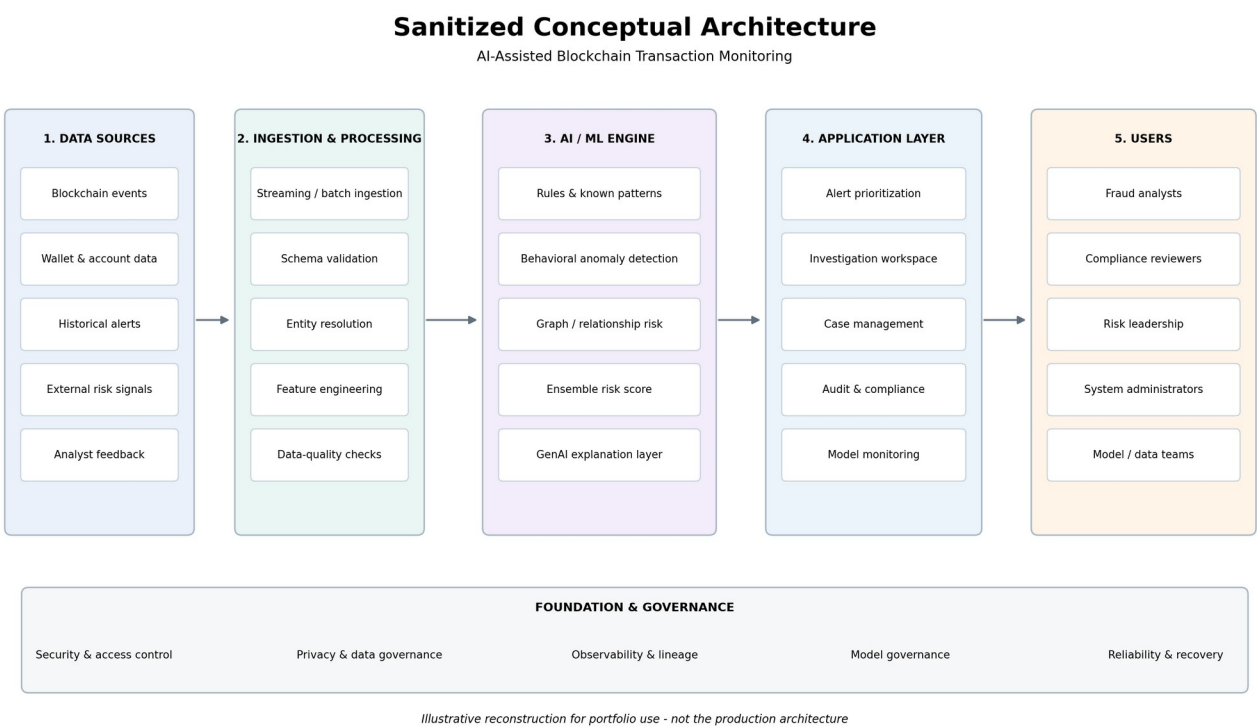
The architecture makes the AI models visible, but the customer outcome depends on the full operating loop: reliable data, calibrated prioritization, explainable evidence, analyst judgment, feedback capture, monitoring, and governance. A technically strong model can still fail as a product when alerts are unactionable, explanations are generic, workflows are fragmented, or outcomes are not captured for learning.

What this case study demonstrates

- Converting a high-risk business problem into measurable user and system requirements.
- Balancing fraud detection, false-positive cost, explainability, privacy, reliability, and analyst capacity.
- Designing hybrid AI where rules, machine learning, relationship analysis, and Generative AI have distinct responsibilities.
- Using evidence labels and phase gates to avoid presenting assumptions or targets as achieved outcomes.
- Working across product management and solution architecture while preserving confidentiality.

AI-Assisted Blockchain Transaction Monitoring System

Product Requirements Document - sanitized confidential-project portfolio edition



Sanitized conceptual architecture reconstructed from memory; not the exact production design.

DOCUMENT STATUS Final Portfolio Edition	VERSION 1.0	AUTHOR Min Thu Kyaw	DATE July 2026
ROLE AI Product Manager & Solution Architect	COMPANY Confidential Blockchain Company	DOCUMENT CODE AFM-02	PUBLICATION Portfolio-safe reconstruction
PORTFOLIO PURPOSE Translate the product concept into testable functional, AI, data, security, analytics, and operational requirements.	EVIDENCE BASIS User-provided high-level architecture and user journey	CONFIDENTIALITY No code, data, PRD, metrics, vendors, thresholds, or internal screenshots	USE Hiring and portfolio review only

Document Control

Purpose	Define a reviewable MVP specification and acceptance model for the analyst alert-to-disposition workflow.
Primary audience	Product, AI/ML, data, engineering, security, compliance, design, QA, and operations.
Evidence classification	Reconstructed fact / Illustrative design / Target / Hypothesis / Recommendation
Source limitation	Only high-level architecture and a representative user journey were supplied. No confidential artifacts are reproduced.
Related documents	AFM-01, AFM-03, AFM-04, AFM-05
Change control	Update only when the sanitized scope, evidence basis, or portfolio narrative changes.

How to Read the Evidence Labels

Label	Meaning
RECONSTRUCTED	High-level information recalled and generalized by the author; exact implementation details are intentionally excluded.
ILLUSTRATIVE	A portfolio artifact created to demonstrate product-management method; not an internal company artifact.
TARGET	A proposed threshold to validate; never presented as an achieved result.
HYPOTHESIS	A belief that requires a defined test and decision rule.
RECOMMENDATION	A portfolio recommendation derived from the reconstructed problem and architecture.

Contents

- 01 PRD Summary
- 02 Goals and Non-Goals
- 03 **Personas and Primary Workflow**
- 04 MVP Scope
- 05 Functional Requirements
- 06 AI Requirements
- 07 Data Requirements
- 08 Security, Privacy and Compliance Requirements
- 09 Non-Functional Requirements
- 10 Analytics and Instrumentation
- 11 Acceptance Criteria
- 12 Traceability Matrix
- 13 Release Gates and Change Control

01

PRD Summary

The MVP validates whether a hybrid risk-scoring and investigation workflow can improve alert prioritization and analyst efficiency without weakening trust, governance, or operational reliability.

Field	Definition
Product goal	Help authorized analysts identify, understand, investigate, and disposition suspicious blockchain activity using prioritized, evidence-linked alerts.
Primary user	Senior Fraud Analyst.
North Star target	Verified high-risk alerts resolved per analyst hour.
MVP hypothesis	A hybrid risk score plus unified evidence and controlled AI assistance improves top-of-queue quality and reduces investigation friction.
Release approach	Limited, evidence-gated pilot with mandatory human review and explicit rollback criteria.
Confidentiality boundary	This PRD contains generalized portfolio requirements only and is not the employer's internal PRD.

02

Goals and Non-Goals

The PRD is intentionally narrow enough to test customer value and the riskiest AI assumptions.

Goals	Non-goals
Prioritize suspicious activity using multiple risk signals.	Fully autonomous fraud or compliance decisions.
Present evidence and reason codes in one reviewable alert.	Reproduce exact confidential rules, thresholds, or model details.
Support investigation, disposition, escalation, and audit history.	Cover every blockchain, product, jurisdiction, or fraud typology.
Capture reviewed outcome labels for evaluation and improvement.	Automatically train or deploy models from raw analyst feedback.
Measure product, model, operational, and governance quality.	Publish actual production metrics or company performance.

03

Personas and Primary Workflow

The MVP centers on the analyst journey and includes supporting controls for compliance review, operations, and model monitoring.

End-to-End Detection and Investigation Workflow

Hybrid detection + human review + feedback loop

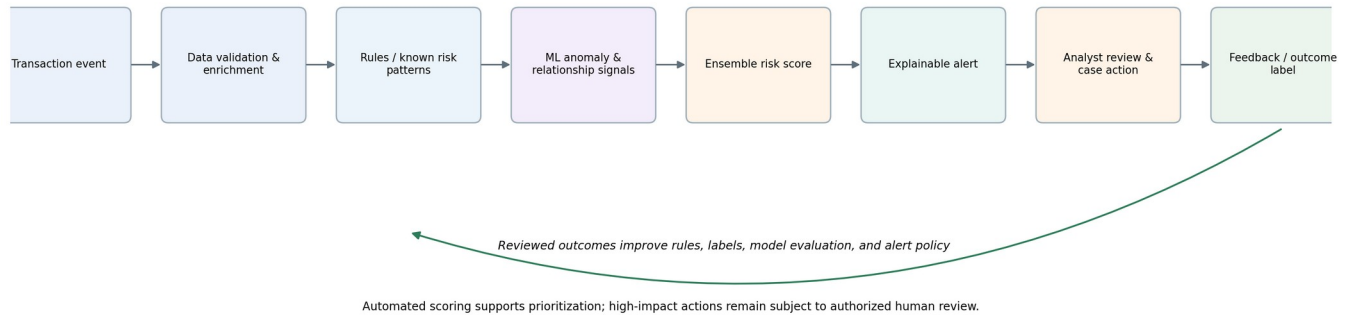


Figure AFM-02.1 - Primary alert-to-feedback workflow.

Primary use case

1. A transaction or behavior pattern enters the monitoring pipeline.
2. Data is validated, enriched, and converted into risk features.
3. Rules and machine-learning services produce versioned risk signals.
4. The alert policy creates or updates an alert with score, severity, reason codes, and evidence links.
5. The analyst reviews the alert, investigates related activity, records a disposition, and optionally escalates the case.
6. The system preserves the decision, evidence, audit history, and reviewed outcome label for authorized evaluation.

04

MVP Scope

The MVP is an end-to-end workflow slice, not a complete enterprise platform.

Capability	MVP	Deferred
Data ingestion	Defined blockchain events, account/wallet metadata, historical outcomes, and approved external risk signals.	Broad provider marketplace and every chain.
Detection	Rules, at least one behavioral/anomaly model, combined risk score, and alert policy.	Autonomous model selection and self-modifying rules.
Explanation	Reason codes, evidence links, confidence, and controlled draft summary.	Unrestricted conversational agent over all sensitive data.
Investigation	Queue, alert detail, timeline, related entities, notes, disposition, escalation, and audit history.	Full enterprise case orchestration across every external system.
Monitoring	Data freshness, scoring latency, alert volume, model/version quality, and failure states.	Complete cross-enterprise observability platform.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-02
Capability	MVP	Deferred

Administration	Role-based access, configuration, model/version visibility, and policy change logging.	Self-service tenant customization of confidential detection logic.
----------------	--	--

05

Functional Requirements

Requirements are written to be testable while avoiding confidential implementation details.

ID	Requirement	Priority
FR-101	The system shall ingest approved transaction events and preserve source timestamp, source identifier, and processing status.	Must
FR-102	The system shall validate required fields and route invalid or incomplete events to an observable exception state.	Must
FR-103	The system shall resolve approved wallet/account/entity relationships using versioned logic.	Should
FR-201	The system shall evaluate applicable deterministic risk rules and store rule identifier, version, outcome, and evidence reference.	Must
FR-202	The system shall produce versioned machine-learning risk signals for configured event or entity segments.	Must
FR-203	The system shall combine approved signals into a calibrated risk score and severity band.	Must
FR-204	The system shall create or update an alert according to a versioned alert policy.	Must
FR-301	The analyst queue shall support sorting and filtering by severity, age, segment, alert type, status, and assignment.	Must
FR-302	The alert detail shall present risk score, reason codes, supporting evidence, event timeline, and related-entity context.	Must
FR-303	The analyst shall be able to record notes, evidence, disposition, rationale, escalation, and follow-up action.	Must
FR-304	The system shall preserve immutable audit events for material alert and case changes.	Must
FR-305	The system shall support reviewer approval for configured high-impact case states.	Must
FR-401	The system may generate an evidence-linked	Should

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-02
ID	Requirement	Priority
	draft summary that is clearly marked as AI-assisted and editable.	
FR-402	The system shall not submit, close, or escalate a case solely because of Generative AI output.	Must
FR-501	Authorized users shall be able to view model, rule, policy, and explanation versions associated with an alert.	Must
FR-502	The system shall capture the reviewed outcome label and label provenance for approved evaluation use.	Must

06

AI Requirements

AI requirements define product behavior and governance, not confidential model implementation.

ID	AI product requirement
AIR-001	Every AI signal shall include model/service version, timestamp, eligible segment, confidence or calibrated score, and trace identifier.
AIR-002	The alert score shall be evaluated against a documented baseline and reported by meaningful segment, not only as an aggregate.
AIR-003	The product shall expose human-readable reason codes and supporting evidence appropriate to the signal type.
AIR-004	Model or policy deployment shall require versioning, approval, rollback criteria, and monitoring configuration.
AIR-005	The evaluation set shall separate training, tuning, and held-out assessment data and document label provenance.
AIR-006	Generative AI output shall be grounded in authorized case evidence and must distinguish facts, inferences, and recommended next checks.
AIR-007	The system shall protect against prompt injection, cross-case leakage, unsupported claims, and sensitive-data disclosure.
AIR-008	When an AI dependency fails, the interface shall show a clear failure state and preserve a recoverable manual workflow.
AIR-009	Analyst feedback shall not automatically change a production model or rule without data-quality review and release governance.
AIR-010	Quality monitoring shall support drift, calibration, alert-yield, override, and segment-level analysis.

07

Data Requirements

Data design must support lineage, reproducibility, privacy, and evaluation.

ID	Requirement
DR-001	Maintain a data dictionary for approved event, entity, feature, alert, case, and outcome fields.
DR-002	Preserve source lineage and transformation version for features used in a material risk signal.
DR-003	Enforce data minimization and purpose limitation for analyst views and AI prompts.
DR-004	Separate production operations, analytics, evaluation, and development access according to policy.
DR-005	Define retention, deletion, correction, and legal-hold behavior for each data class.
DR-006	Detect data freshness, schema, completeness, duplication, and distribution anomalies.
DR-007	Record label source, reviewer, review status, timestamp, and adjudication state.
DR-008	Prevent direct use of unreviewed or ambiguous outcomes as ground truth for performance claims.

08

Security, Privacy and Compliance Requirements

Controls are expressed at product-requirement level to avoid exposing the company's exact security design.

ID	Requirement
SPC-001	Enforce role-based and least-privilege access to alerts, cases, configuration, and model information.
SPC-002	Log authentication, access, export, configuration, disposition, escalation, and approval events.
SPC-003	Encrypt sensitive data in transit and at rest using organization-approved controls.
SPC-004	Define approved data boundaries for external AI providers and prevent unauthorized data transmission.
SPC-005	Support configurable redaction or masking of sensitive identifiers in portfolio, training, and non-production contexts.
SPC-006	Maintain an incident path for data exposure, model misuse, incorrect

ID	Requirement
	high-impact action, or material service failure.
SPC-007	Support evidence retention and export according to authorized compliance and audit needs.
SPC-008	Require human approval for configured high-impact case actions and policy exceptions.

09

Non-Functional Requirements

The product must remain usable and safe when data, models, providers, or infrastructure are imperfect.

Category	Requirement / target
Reliability	The alert-review workflow remains available or degrades explicitly when optional AI assistance is unavailable.
Performance	Scoring and alert freshness targets are defined by operational risk tier and validated under representative load.
Scalability	Throughput and storage capacity are measured against a documented volume model before expansion.
Observability	Pipeline freshness, errors, latency, alert volume, model/version, and queue health are observable.
Explainability	Analysts can identify the evidence and version behind each material risk reason.
Accessibility	Core queue, alert, and case workflows support keyboard navigation, readable hierarchy, and meaningful status communication.
Localization	Time, number, terminology, and language behavior are defined for supported analyst markets.
Recovery	Failure, retry, replay, rollback, and data-reconciliation procedures are tested before pilot expansion.

10

Analytics and Instrumentation

Measurement must connect events to decisions without creating a second uncontrolled sensitive-data store.

Event family	Representative events	Decision enabled
Queue	queue_viewed, alert_opened, filter_applied, assignment_changed	Understand workload and prioritization behavior.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-02
Event family	Representative events	Decision enabled
Investigation	evidence_opened, relationship_viewed, note_added, next_check_selected	Identify friction and useful evidence.
Decision	disposition_saved, escalation_requested, approval_completed, case_closed	Measure workflow completion and outcome quality.
AI assistance	explanation_viewed, draft_generated, draft_edited, suggestion_accepted/rejected	Measure usefulness without equating use with correctness.
System	ingestion_delay, scoring_latency, provider_failure, retry, rollback	Control operational risk and release readiness.
Governance	permission_denied, export_requested, policy_changed, model_version_changed	Support audit and control monitoring.

11

Acceptance Criteria

The release must prove an end-to-end controlled workflow, not just visible screens.

Area	Acceptance criterion
End-to-end	A representative event can be ingested, scored, converted to an alert, reviewed, dispositioned, and audited.
Evidence	Every material risk reason links to approved evidence or a clearly identified derived signal.
Human control	Configured high-impact states cannot be completed without the required authorized review.
AI quality	Offline evaluation and pilot monitoring meet agreed target thresholds by critical segment.
Privacy and security	Mandatory access, logging, provider-boundary, and incident controls pass review.
Failure behavior	Data, model, or provider failure is visible, logged, recoverable, and does not silently create unsupported decisions.
Measurement	Required product, model, operational, and governance events are captured and validated.
Documentation	Known limitations, versions, rollback steps, and owner responsibilities are current.

12

Traceability Matrix

Each product goal connects to a user capability, requirement set, evidence, and metric.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION				AFM-02
Goal	Capability	Requirements	Validation evidence	Metric
Prioritize meaningful risk	Hybrid scoring and alert policy	FR-201-204, AIR-001-005	Offline comparison and pilot queue analysis	Top-of-queue precision, alert yield
Reduce investigation friction	Unified evidence and workflow	FR-301-305	Usability study and pilot timing	Median investigation time, completion rate
Maintain trust	Reason codes, provenance, human review	AIR-003, AIR-006-009, SPC-008	Explanation review and control tests	Usefulness, edit/override rate, audit completeness
Operate reliably	Observability and recoverable failure	FR-501, NFRs	Load, failure, retry, rollback tests	Latency, freshness, availability, recovery
Improve responsibly	Outcome capture and governed feedback	FR-502, DR-007-008	Label audit and evaluation pipeline	Label quality, drift, segment performance

13

Release Gates and Change Control

Scope or model changes require explicit impact review across customer value, AI quality, operations, security, and governance.

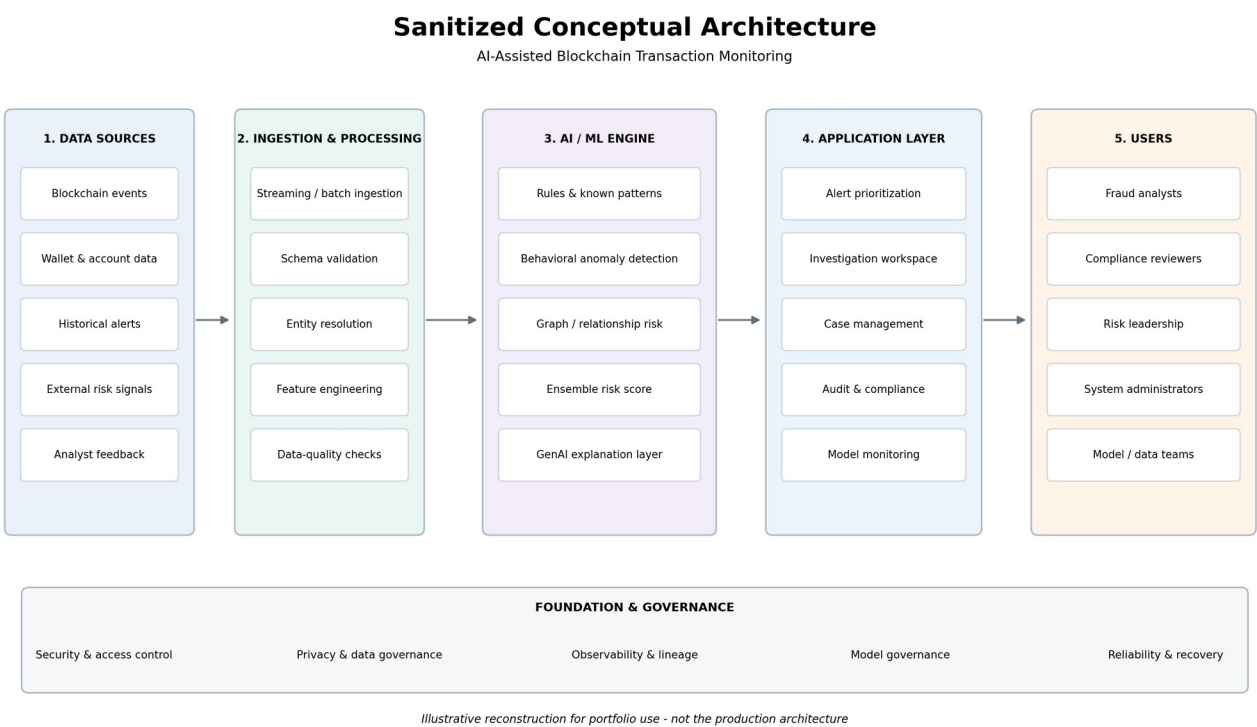
Gate	Minimum evidence
G0 - Problem / workflow	Validated user pain, baseline, target workflow, privacy boundary.
G1 - Prototype	Representative tasks completed; explanation and evidence hierarchy understood.
G2 - Offline AI	Versioned data, label review, segment metrics, calibration, limitations, rollback candidate.
G3 - Pilot readiness	End-to-end workflow, access/audit controls, monitoring, failure recovery, support ownership.
G4 - Expansion	Customer value, AI quality, operational stability, governance, and economics meet decision rules.

Change-control triggers

- A new data source, external provider, blockchain network, or jurisdiction is introduced.
- A new model family or material feature category affects scoring or explanations.
- The alert threshold, severity policy, reviewer requirement, or case action changes.
- Data retention, access, export, or provider boundary changes.
- A deferred feature is moved into the pilot without updated evidence and acceptance criteria.

AI-Assisted Blockchain Transaction Monitoring System

AI Strategy & Evaluation Plan - sanitized confidential-project portfolio edition



Sanitized conceptual architecture reconstructed from memory; not the exact production design.

DOCUMENT STATUS Final Portfolio Edition	VERSION 1.0	AUTHOR Min Thu Kyaw	DATE July 2026
ROLE AI Product Manager & Solution Architect	COMPANY Confidential Blockchain Company	DOCUMENT CODE AFM-03	PUBLICATION Portfolio-safe reconstruction
PORTFOLIO PURPOSE Explain the role of rules, machine learning, relationship analytics, Generative AI, human review, and model governance.	EVIDENCE BASIS User-provided high-level architecture and user journey	CONFIDENTIALITY No code, data, PRD, metrics, vendors, thresholds, or internal screenshots	USE Hiring and portfolio review only

Document Control

Purpose	Define how AI creates product value and how quality, trust, safety, and operational readiness will be evaluated.
Primary audience	AI Product, Data Science, ML Engineering, Risk, Compliance, Security, QA, and Model Governance.
Evidence classification	Reconstructed fact / Illustrative design / Target / Hypothesis / Recommendation
Source limitation	Only high-level architecture and a representative user journey were supplied. No confidential artifacts are reproduced.
Related documents	AFM-01, AFM-02, AFM-04, AFM-05
Change control	Update only when the sanitized scope, evidence basis, or portfolio narrative changes.

How to Read the Evidence Labels

Label	Meaning
RECONSTRUCTED	High-level information recalled and generalized by the author; exact implementation details are intentionally excluded.
ILLUSTRATIVE	A portfolio artifact created to demonstrate product-management method; not an internal company artifact.
TARGET	A proposed threshold to validate; never presented as an achieved result.
HYPOTHESIS	A belief that requires a defined test and decision rule.
RECOMMENDATION	A portfolio recommendation derived from the reconstructed problem and architecture.

Contents

- 01 AI Product Strategy**
- 02 Hybrid Detection Design**
- 03 Feature and Signal Taxonomy**
- 04 Model and Policy Responsibilities
- 05 Data and Label Strategy
- 06 Offline Evaluation
- 07 Online and Pilot Evaluation
- 08 Explainability and Human-in-the-Loop
- 09 Generative AI Assistance
- 10 Monitoring, Drift and Incident Management
- 11 Model Governance
- 12 Evaluation Scorecard
- 13 AI Risk Register

01

AI Product Strategy

The AI strategy assigns each technique a narrow responsibility and measures whether it improves the analyst decision workflow.

AI is successful only when it improves risk prioritization and investigation quality inside a controlled, explainable, human-reviewed workflow.

Layer	Responsibility	Not responsible for
Rules	Known patterns, explicit policy, deterministic controls	Discovering every novel behavior.
Behavioral / anomaly ML	Identify unusual changes or patterns across time	Making a final case decision.
Graph / relationship analytics	Assess entity connections and fund-flow structure	Inferring legal or regulatory conclusions.
Ensemble scoring	Combine signals and calibrate priority	Replacing operational alert policy.
GenAI assistance	Summarize authorized evidence and draft reviewable narrative	Creating unsupported facts or executing high-impact actions.
Human analyst	Interpret context, investigate, disposition, escalate, and approve	Manually reconstructing every raw signal from separate systems.

02

Hybrid Detection Design

A hybrid design reduces dependence on any single signal type and makes policy, model quality, and analyst capacity explicit.

End-to-End Detection and Investigation Workflow

Hybrid detection + human review + feedback loop

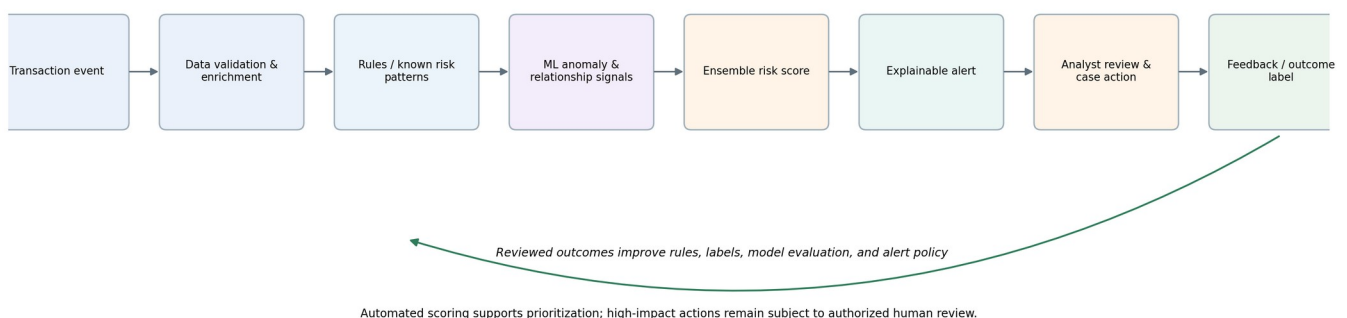


Figure AFM-03.1 - Hybrid detection and feedback loop.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-03
Signal type	Illustrative examples	Evaluation focus
Rules	Velocity, prohibited sequence, policy threshold, known relationship	Coverage, precision, stability, explainability.
Behavioral	Change from historical frequency, amount, timing, counterparties, device or account behavior	Segmented discrimination, calibration, drift.
Anomaly	Rare combination or deviation from learned normal pattern	Alert yield, review usefulness, stability under data shift.
Relationship / graph	Exposure to risky entities, unusual network structure, rapid fund movement	Evidence traceability and false-positive concentration.
External risk	Approved sanctions, reputation, or intelligence signals	Freshness, provenance, conflict resolution.
Analyst feedback	Reviewed disposition, rationale, escalation, confirmed outcome	Label quality, bias, delay, and adjudication.

03

Feature and Signal Taxonomy

The taxonomy communicates what the system reasons about without exposing the exact proprietary feature set.

Category	Representative feature families	Sensitivity / governance note
Transaction	Amount, frequency, direction, timing, asset, sequence	Potentially sensitive; minimize and version transformations.
Behavioral	Deviation from entity history, changes in routine, burst patterns	Requires stable entity history and segment-aware baselines.
Counterparty	New counterparties, concentration, recurrence, risk exposure	Relationship evidence must be reviewable.
Graph / flow	Path length, clustering, fan-in/fan-out, flow velocity	Avoid opaque scores without supporting paths or entities.
Account / wallet	Age, activity history, linked identifiers, prior alerts	Access and retention must follow approved policy.
External intelligence	Approved watchlists or risk providers	Freshness, provenance, licensing, and conflict handling.
Operational context	Channel, product, geography, system or event quality	Use only when legitimate, necessary, and governed.

04

Model and Policy Responsibilities

Models estimate risk signals; product policy determines how those signals become an alert and how humans must respond.

Component	Input	Output	Key control
Rule service	Validated event/entity data	Rule hits and evidence	Versioned rule and policy approval.
Behavioral model	Time-windowed features	Risk score / anomaly score	Segment definition and drift monitoring.
Relationship model	Entity graph / flow context	Relationship risk and evidence path	Traceable supporting nodes/edges.
Ensemble / calibration	Approved model and rule signals	Calibrated priority score	Champion/challenger comparison and rollback.
Alert policy	Score, severity, capacity, rule conditions	Create, suppress, merge, escalate, or queue	Policy owner, version, and simulation.
Explanation service	Reason codes, evidence, authorized context	Analyst-facing explanation	Grounding, uncertainty, and human review.

05

Data and Label Strategy

Fraud labels are often delayed, incomplete, or influenced by analyst workflow. The evaluation plan must treat label quality as a first-class risk.

Area	Plan
Label definition	Define alert disposition, case escalation, confirmed outcome, uncertain, duplicate, insufficient evidence, and policy-only categories.
Adjudication	Use reviewed guidelines and secondary review for ambiguous or high-impact examples.
Splitting	Use time-aware and entity-aware splits to reduce leakage; preserve a held-out evaluation set.
Class imbalance	Report precision-recall metrics, top-K / threshold yield, and operational workload rather than accuracy alone.
Sampling	Include representative negatives, difficult cases, edge segments, and changed behavior.
Provenance	Store source, reviewer, review status, timestamp, model/policy exposure, and later outcome.
Bias / concentration	Check quality across legitimate product segments without using prohibited or unjustified sensitive attributes.

Offline Evaluation

Offline metrics determine whether a candidate is worth piloting; they do not prove live customer value.

Metric	Why it matters	Required breakdown
Precision / positive predictive value	Controls wasted review effort at the chosen threshold.	Alert type, segment, severity, model/rule source.
Recall / sensitivity	Estimates missed known-positive events.	Critical-risk classes and time periods.
PR-AUC	Useful under strong class imbalance.	Candidate versus baseline with uncertainty.
Precision at K / top percentile	Measures quality of the limited analyst queue.	Capacity scenarios and alert policy.
Calibration / Brier score	Supports meaningful risk bands and policy decisions.	Segment and model version.
Alert yield	Connects alerts to reviewed outcomes.	By rule/model mix and operational cohort.
Stability / drift sensitivity	Shows behavior across time and distribution change.	Time window, chain/product segment, data-quality state.
Explanation faithfulness	Checks whether reason codes match underlying signals.	Signal type and error category.

Offline decision rule

Advance a candidate only when it demonstrates a meaningful improvement over the approved baseline in the target operating region, does not create unacceptable degradation in critical segments, has a defined alert-volume impact, and can be explained and monitored.

Online and Pilot Evaluation

The pilot measures whether AI quality translates into better analyst decisions and sustainable operations.

Dimension	Pilot metric	Interpretation
Customer value	Verified high-risk alerts resolved per analyst hour	Primary productivity and value indicator.
Queue quality	Precision / yield in the reviewed priority bands	Does prioritization improve the analyst queue?
Investigation	Median time to first evidence and disposition	Does the workflow reduce friction?
Trust	Explanation usefulness, edit rate, override	Are outputs useful without being blindly

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-03
Dimension	Pilot metric	Interpretation
	rate	accepted?
Operational	Freshness, scoring latency, failure/retry, backlog	Can the product operate reliably at pilot volume?
Governance	Audit completeness, access exceptions, policy/model changes	Are mandatory controls working?
Risk	Critical missed-event review and adverse incident count	Are unacceptable harms or blind spots appearing?

08

Explainability and Human-in-the-Loop

The explanation design should help an analyst inspect evidence and decide, not merely provide a fluent story.

Explanation element	Requirement
Reason code	Concise category describing the material risk signal.
Supporting evidence	Specific events, entities, relationships, or policy references.
Score / severity	Calibrated risk band with clear interpretation and version.
Confidence / uncertainty	Where appropriate, communicate uncertainty and data limitations.
Source provenance	Model, rule, data, and explanation versions associated with the alert.
Human action	Review, investigate, request more evidence, escalate, close, or override with rationale.
Review policy	Mandatory approval for configured high-impact states or exceptions.

09

Generative AI Assistance

Generative AI is a controlled productivity layer for summarization and investigation support, not the core risk detector.

Allowed use	Control
Draft an alert summary from authorized evidence	Ground only in retrieved evidence; cite or link the source context.
Summarize an activity timeline	Preserve chronology and distinguish facts from inferred patterns.
Suggest next investigation checks	Present as recommendations, not required or authoritative actions.
Retrieve similar reviewed cases	Enforce access control and avoid copying unsupported conclusions.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-03
Allowed use	Control	
Draft a case narrative	Clearly mark AI-assisted text; require analyst edit and approval.	
Natural-language query over case evidence	Use constrained retrieval, prompt-injection defenses, and output validation.	

Prohibited behavior

- Invent evidence, wallet relationships, risk reasons, or policy requirements.
- Reveal data from another case, customer, user, or access scope.
- Close, escalate, report, or block a case without required human authorization.
- Treat a fluent explanation as proof that the underlying risk score is correct.
- Send sensitive context to an unapproved provider or retain it outside the approved boundary.

10

Monitoring, Drift and Incident Management

Monitoring connects model behavior to data quality, product usage, analyst outcomes, and operational reliability.

Monitor	Signal	Response
Data	Freshness, schema, missingness, duplication, distribution shift	Pause or flag affected scoring; investigate upstream source.
Model	Score distribution, calibration, segment quality, drift	Compare baseline/challenger; adjust policy or rollback.
Alert policy	Volume, suppression, merge rate, age, backlog	Simulate policy change; adjust capacity or threshold.
Explanation	Missing evidence, unsupported statement, low usefulness, high edit rate	Restrict template/GenAI behavior and review prompts/retrieval.
User workflow	Abandonment, overrides, escalations, review time	Investigate UX, model quality, or policy mismatch.
Security / privacy	Unauthorized access, leakage, provider boundary, suspicious query	Contain, revoke, investigate, notify, and document.
Incident	Material incorrect action or missed critical event	Preserve evidence, assess impact, rollback, and complete review.

11

Model Governance

Every material model, rule, policy, and Generative AI configuration change requires an accountable owner and reproducible evidence.

Artifact	Minimum contents
Model / service card	Purpose, owner, version, data scope, segments, metrics, limitations, prohibited uses, monitoring, rollback.
Evaluation report	Dataset/version, label process, baseline, metrics, uncertainty, segment analysis, errors, decision.
Release record	Approved version, dependencies, change summary, testing, owner, rollout, rollback, monitoring.
Decision log	Why a model, threshold, provider, feature family, or policy changed and what evidence supported it.
Incident review	Impact, detection, containment, cause, corrective actions, revalidation, and communication.

12

Evaluation Scorecard

The scorecard below is a target framework and does not report production performance.

Area	Target gate	Status in public portfolio
Baseline comparison	Candidate materially improves the target operating region.	Target - not measured.
Critical segment recall	No unacceptable degradation in defined high-impact segments.	Target - not measured.
Top-queue precision	Meets threshold at realistic analyst capacity.	Target - not measured.
Calibration	Risk bands support consistent alert policy.	Target - not measured.
Explanation quality	Evidence-linked and useful to analysts.	Target - not measured.
Operational readiness	Freshness, latency, failure recovery, and monitoring pass.	Target - not measured.
Security / privacy	Mandatory controls and provider boundary pass review.	Target - not measured.
Human control	Required review and override paths are effective.	Target - not measured.

13

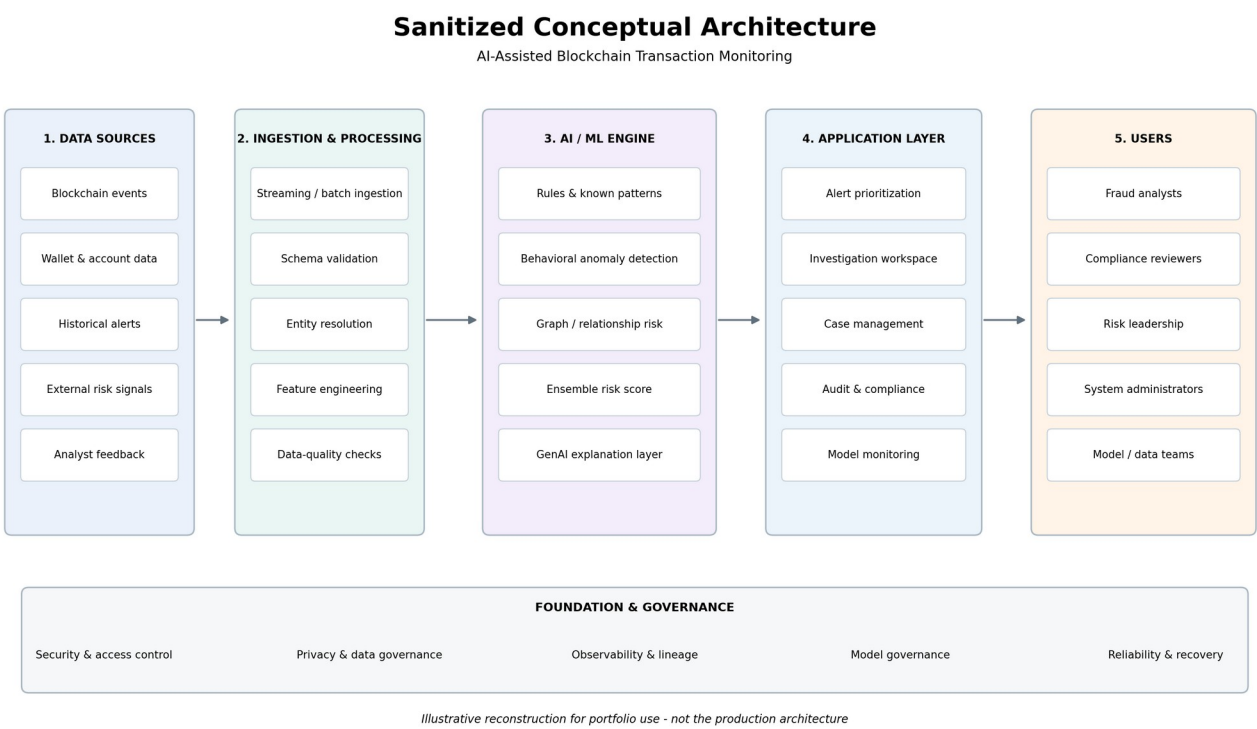
AI Risk Register

The system must manage both detection errors and product/system failure modes.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-03
Risk	Severity	Control
False negative / missed critical pattern	Critical	Segment evaluation, conservative policy, retrospective review, monitoring.
False positive overload	High	Precision-at-capacity, suppression/merge policy, workload monitoring.
Opaque or misleading explanation	High	Evidence-linked reasons, faithfulness review, analyst control.
Label bias or contamination	High	Guidelines, adjudication, provenance, time/entity-aware splits.
Data or concept drift	High	Drift monitoring, challenger evaluation, rollback.
Sensitive-data leakage through GenAI	Critical	Approved boundary, minimization, access control, redaction, testing.
Prompt injection or unsafe retrieval	High	Constrained retrieval, sanitization, instruction hierarchy, output validation.
Automation bias	High	Uncertainty, evidence access, mandatory review, override capture.
Provider / service failure	Medium to High	Explicit degraded mode, retry, manual workflow, recovery testing.

AI-Assisted Blockchain Transaction Monitoring System

As-Built Architecture & Traceability Report - sanitized portfolio reconstruction



Sanitized conceptual architecture reconstructed from memory; not the exact production design.

DOCUMENT STATUS Final Portfolio Edition	VERSION 1.0	AUTHOR Min Thu Kyaw	DATE July 2026
ROLE AI Product Manager & Solution Architect	COMPANY Confidential Blockchain Company	DOCUMENT CODE AFM-04	PUBLICATION Portfolio-safe reconstruction
PORTFOLIO PURPOSE Document the reconstructed high-level architecture, component responsibilities, product decisions, traceability, limitations, and governance boundaries.	EVIDENCE BASIS User-provided high-level architecture and user journey	CONFIDENTIALITY No code, data, PRD, metrics, vendors, thresholds, or internal screenshots	USE Hiring and portfolio review only

Document Control

Purpose	Connect the sanitized product requirements to a high-level system design without claiming access to internal code or exact production architecture.
Primary audience	Hiring managers, product and technical leaders, architecture reviewers, risk and security stakeholders.
Evidence classification	Reconstructed fact / Illustrative design / Target / Hypothesis / Recommendation
Source limitation	Only high-level architecture and a representative user journey were supplied. No confidential artifacts are reproduced.
Related documents	AFM-01, AFM-02, AFM-03, AFM-05
Change control	Update only when the sanitized scope, evidence basis, or portfolio narrative changes.

How to Read the Evidence Labels

Label	Meaning
RECONSTRUCTED	High-level information recalled and generalized by the author; exact implementation details are intentionally excluded.
ILLUSTRATIVE	A portfolio artifact created to demonstrate product-management method; not an internal company artifact.
TARGET	A proposed threshold to validate; never presented as an achieved result.
HYPOTHESIS	A belief that requires a defined test and decision rule.
RECOMMENDATION	A portfolio recommendation derived from the reconstructed problem and architecture.

Contents

- 01 Evidence Basis and Limitations**
- 02 Architecture Overview**
- 03 Layer Responsibilities**
- 04 End-to-End Data Flow
- 05 Application and User Experience
- 06 Security and Governance Foundation
- 07 Requirement-to-Component Traceability
- 08 Key Product and Architecture Decisions
- 09 Known Limitations and Portfolio Boundaries
- 10 Operational Readiness Checklist
- 11 Product Manager Reflection

01

Evidence Basis and Limitations

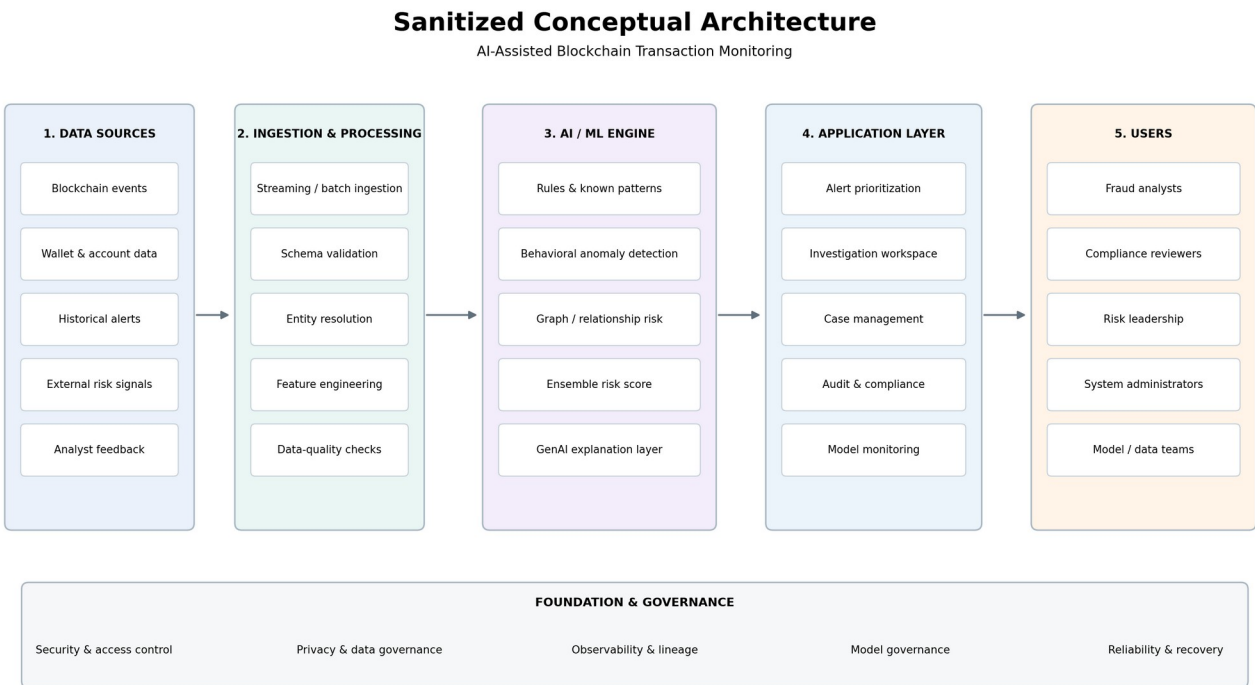
This report is reconstructed from the author's high-level memory and the supplied conceptual diagrams. It is not an internal implementation report.

Classification	Evidence
Reconstructed	Layered architecture: data sources, ingestion/processing, storage, AI/analytics, application, user interfaces, external services, and governance.
Reconstructed	User journey: monitor, review, investigate, consult AI assistance, and preserve a decision.
Illustrative	Component boundaries, requirement mapping, monitoring model, and failure states developed for portfolio demonstration.
Excluded	Exact cloud services, databases, vendors, schemas, API contracts, code, thresholds, models, feature definitions, security controls, and production metrics.

02

Architecture Overview

The architecture separates transaction data processing, detection, explanation, analyst workflow, and governance to support independent control and evolution.



Illustrative reconstruction for portfolio use - not the production architecture

Figure AFM-04.1 - Sanitized conceptual architecture.

03

Layer Responsibilities

Each layer has a clear product responsibility and evidence boundary.

Layer	Responsibility	Representative outputs
Data sources	Provide approved blockchain, wallet/account, historical, external-risk, and analyst-feedback data.	Versioned source events and source-quality status.
Ingestion & processing	Validate, normalize, enrich, resolve entities, and create features.	Clean events, relationships, feature records, exceptions.
Storage	Preserve raw/curated events, features, model artifacts, alerts, cases, and audit history according to policy.	Lineage, reproducibility, retention, and reviewable evidence.
AI & analytics	Evaluate rules, behavioral signals, relationship risk, ensemble score, and explanation context.	Versioned risk signals, score, severity, reasons, monitoring telemetry.
Application	Prioritize alerts, support investigation, manage case state, and expose monitoring/admin controls.	Queue, alert detail, investigation record, audit events.
User interfaces	Serve analysts, reviewers, risk leadership, administrators, and model/data teams.	Role-appropriate decisions and controls.
Foundation	Security, privacy, observability, model governance, reliability, and recovery.	Control evidence and operational readiness.

04

End-to-End Data Flow

A material alert must be reproducible from source event through processing, risk signals, policy, analyst action, and outcome label.

1. Receive an approved source event and assign a traceable processing identifier.
2. Validate schema, required fields, time ordering, duplication, and source freshness.
3. Enrich with approved entity, relationship, and historical context.
4. Create versioned features and evaluate applicable deterministic rules.
5. Invoke eligible model services and record model/version outputs.
6. Combine signals through a versioned calibration and alert policy.
7. Create or update an alert with score, severity, reason codes, evidence, and versions.
8. Present the alert to an authorized analyst and preserve investigation actions and decisions.
9. Store reviewed outcomes separately from raw model predictions for governed evaluation.

End-to-End Detection and Investigation Workflow

Hybrid detection + human review + feedback loop

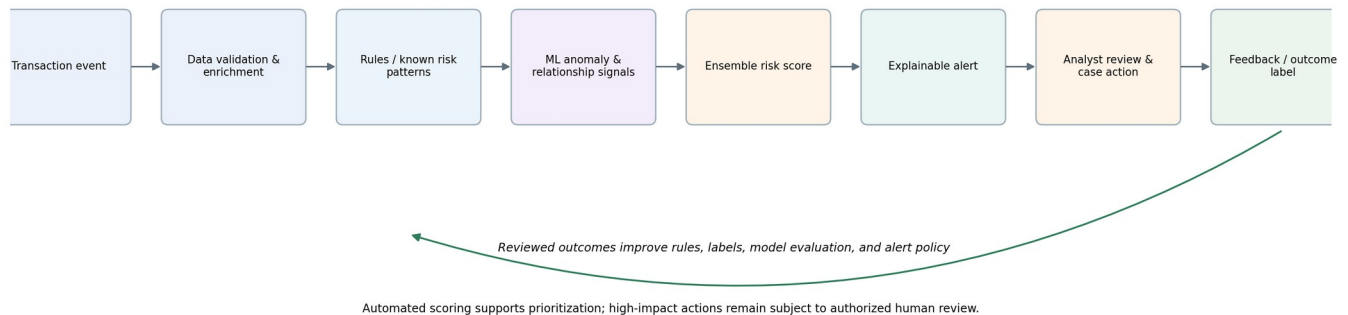


Figure AFM-04.2 - Reconstructed end-to-end flow.

05

Application and User Experience

The application layer converts technical signals into a manageable operational workflow.

Surface	Primary user need	Required information / control
Alert queue	Know what to review next	Severity, age, reason, assignment, status, segment, SLA, sorting/filtering.
Alert detail	Understand why the alert exists	Evidence, timeline, entities, relationships, score, reasons, versions, uncertainty.
Investigation workspace	Form and test a case hypothesis	Notes, evidence collection, related activity, prior outcomes, next checks.
Case decision	Record a defensible outcome	Disposition, rationale, escalation, approval, policy reference, timestamps.
Risk operations	Manage quality and capacity	Volume, backlog, yield, aging, analyst load, failure states.
Model / data monitoring	Understand system behavior	Data quality, drift, calibration, model version, segment metrics, incidents.
Administration	Control configuration and access	Roles, policies, model/rule versions, change history, export controls.

06

Security and Governance Foundation

Security and governance span every layer and are not represented as one downstream component.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-04
Control domain	Architecture implication	
Identity and access	Role-based authorization at data, application, configuration, export, and AI-assistance boundaries.	
Data privacy	Purpose limitation, minimization, retention, masking, deletion, and provider-boundary enforcement.	
Auditability	Material access, configuration, model, policy, alert, and case actions produce immutable audit events.	
Model governance	Versioned artifacts, approval, evaluation, monitoring, rollback, and incident review.	
Observability	Data freshness, processing errors, scoring latency, alert health, provider failure, and user-impact telemetry.	
Reliability	Retry, replay, idempotency, reconciliation, degraded mode, recovery, and tested rollback.	
Separation of duties	Different authority for model development, release approval, case action, compliance review, and access administration.	

07

Requirement-to-Component Traceability

The matrix demonstrates how product requirements drive architecture decisions.

Requirement group	Primary components	Evidence produced
FR-101-103 Data ingestion	Connectors, validation, entity resolution, feature pipeline	Source lineage, processing status, exception log, data-quality metrics.
FR-201-204 Detection	Rules, model services, calibration, alert policy	Versioned signals, score, severity, alert reason, simulation result.
FR-301-305 Investigation	Queue, alert detail, investigation, case workflow, approvals	Usability evidence, audit events, disposition records.
FR-401-402 GenAI assistance	Retrieval, prompt/orchestration, evidence templates, output validation	Grounding test, leakage test, edit/reject telemetry.
AIR-001-010 AI governance	Model registry, evaluation pipeline, monitoring, release control	Model card, evaluation report, drift alert, rollback record.
DR-001-008 Data governance	Catalog, lineage, retention, label store, access control	Dictionary, lineage, label provenance, access evidence.
SPC-001-008 Security / compliance	IAM, audit, encryption, provider boundary, incident workflow	Control review, access test, incident drill, export log.

08

Key Product and Architecture Decisions

The decisions below are public portfolio reasoning, not a disclosure of internal company decisions.

Decision	Rationale	Trade-off
Use hybrid rules + ML	Combines policy transparency with pattern detection.	More components, monitoring, and calibration complexity.
Separate risk score from alert policy	Allows product/operations to simulate workload and business impact.	Requires disciplined versioning and policy governance.
Make evidence first-class	Supports analyst trust, auditability, and investigation speed.	Additional data lineage and UI complexity.
Use GenAI only as an assistance layer	Improves summarization and reporting without delegating authority.	Requires grounding, security controls, and human review.
Capture outcome labels with provenance	Supports responsible evaluation and improvement.	Labels remain delayed, noisy, and operationally expensive.
Design for degraded mode	Preserves safe operation when optional AI or providers fail.	More states and test scenarios.
Version models, rules, policy, and explanations	Enables reproducibility and incident analysis.	Increases platform and governance overhead.

09

Known Limitations and Portfolio Boundaries

Honest limitations increase credibility and prevent the public artifact from implying access or evidence that does not exist.

- No source repository, internal PRD, architecture specification, dataset, model card, test report, or production telemetry was used.
- The diagrams intentionally omit exact vendors, protocols, security controls, model families, feature logic, and infrastructure topology.
- No actual performance, cost, latency, alert-volume, false-positive, fraud-loss, or compliance outcome is reported.
- The architecture is suitable for demonstrating reasoning, but it cannot be treated as an approved build specification.
- All requirements, metrics, and roadmap items must be revalidated against the real organization, jurisdiction, operating model, and data.

10

Operational Readiness Checklist

A pilot should not begin until end-to-end value and mandatory controls are testable.

Area	Readiness evidence
Data	Source agreement, schema, lineage, freshness monitoring, exception handling, replay.

Area	Readiness evidence
AI	Baseline, versioned evaluation, segment quality, calibration, limitations, rollback.
Workflow	Representative tasks, permissions, approval, escalation, audit, support ownership.
Security / privacy	Access review, provider boundary, encryption, logging, retention, incident path.
Reliability	Load profile, latency, failure injection, degraded mode, reconciliation, recovery.
Measurement	Metric definitions, instrumentation QA, dashboard ownership, decision rules.
Change control	Release record, approvers, rollout cohort, rollback trigger, communication.

Product Manager Reflection

The product is the controlled decision workflow, not only the model

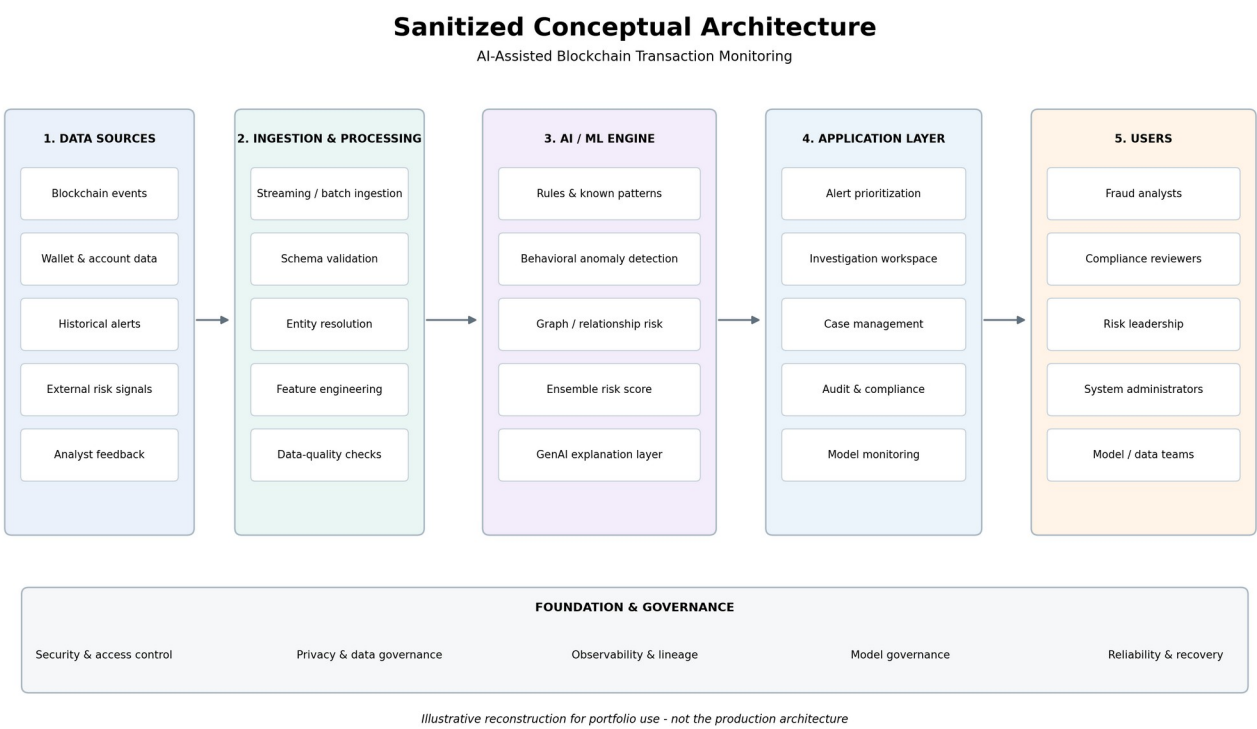
The architecture makes the AI models visible, but the customer outcome depends on the full operating loop: reliable data, calibrated prioritization, explainable evidence, analyst judgment, feedback capture, monitoring, and governance. A technically strong model can still fail as a product when alerts are unactionable, explanations are generic, workflows are fragmented, or outcomes are not captured for learning.

What this case study demonstrates

- Converting a high-risk business problem into measurable user and system requirements.
- Balancing fraud detection, false-positive cost, explainability, privacy, reliability, and analyst capacity.
- Designing hybrid AI where rules, machine learning, relationship analysis, and Generative AI have distinct responsibilities.
- Using evidence labels and phase gates to avoid presenting assumptions or targets as achieved outcomes.
- Working across product management and solution architecture while preserving confidentiality.

AI-Assisted Blockchain Transaction Monitoring System

Pilot Measurement & Roadmap Plan - sanitized confidential-project portfolio edition



Sanitized conceptual architecture reconstructed from memory; not the exact production design.

DOCUMENT STATUS Final Portfolio Edition	VERSION 1.0	AUTHOR Min Thu Kyaw	DATE July 2026
ROLE AI Product Manager & Solution Architect	COMPANY Confidential Blockchain Company	DOCUMENT CODE AFM-05	PUBLICATION Portfolio-safe reconstruction
PORTFOLIO PURPOSE Define how the product hypothesis, AI quality, analyst workflow, operational readiness, and governance would be measured before scale.	EVIDENCE BASIS User-provided high-level architecture and user journey	CONFIDENTIALITY No code, data, PRD, metrics, vendors, thresholds, or internal screenshots	USE Hiring and portfolio review only

Document Control

Purpose	Provide an evidence-gated pilot, metric tree, instrumentation plan, roadmap, backlog, and decision rules.
Primary audience	Product, risk operations, AI/ML, engineering, analytics, security, compliance, and leadership.
Evidence classification	Reconstructed fact / Illustrative design / Target / Hypothesis / Recommendation
Source limitation	Only high-level architecture and a representative user journey were supplied. No confidential artifacts are reproduced.
Related documents	AFM-01, AFM-02, AFM-03, AFM-04
Change control	Update only when the sanitized scope, evidence basis, or portfolio narrative changes.

How to Read the Evidence Labels

Label	Meaning
RECONSTRUCTED	High-level information recalled and generalized by the author; exact implementation details are intentionally excluded.
ILLUSTRATIVE	A portfolio artifact created to demonstrate product-management method; not an internal company artifact.
TARGET	A proposed threshold to validate; never presented as an achieved result.
HYPOTHESIS	A belief that requires a defined test and decision rule.
RECOMMENDATION	A portfolio recommendation derived from the reconstructed problem and architecture.

Contents

- 01 Pilot Purpose and Hypotheses**
- 02 Pilot Cohort and Scenarios**
- 03 Metric System**
- 04 North Star and KPI Tree
- 05 Instrumentation Plan
- 06 Pilot Method
- 07 Exit and Stop Criteria
- 08 Roadmap
- 09 Prioritized Backlog
- 10 Risk and Decision Register
- 11 Operating Cadence
- 12 Product Manager Reflection

01

Pilot Purpose and Hypotheses

The pilot should test customer value and risk controls together. A technically promising model is not enough.

Hypothesis	Evidence	Decision rule
H1 - Priority quality	The hybrid queue contains a higher proportion of useful high-risk alerts than the baseline workflow.	Continue only if improvement is material at realistic analyst capacity.
H2 - Investigation efficiency	Unified evidence and explanations reduce time to meaningful evidence and disposition.	Continue only if speed improves without lower review quality.
H3 - Analyst trust	Analysts can understand, challenge, edit, and appropriately rely on risk explanations.	Revise or restrict assistance when usefulness is low or automation bias is observed.
H4 - Operational readiness	Data, scoring, alert policy, UI, and monitoring operate reliably during the limited cohort.	Block expansion if mandatory freshness, failure, or support controls fail.
H5 - Governance	Access, audit, review, provider boundary, and incident controls work in real workflow conditions.	No expansion when a critical governance gate is unmet.

02

Pilot Cohort and Scenarios

Use a narrow cohort and representative alert types to reduce risk and produce interpretable evidence.

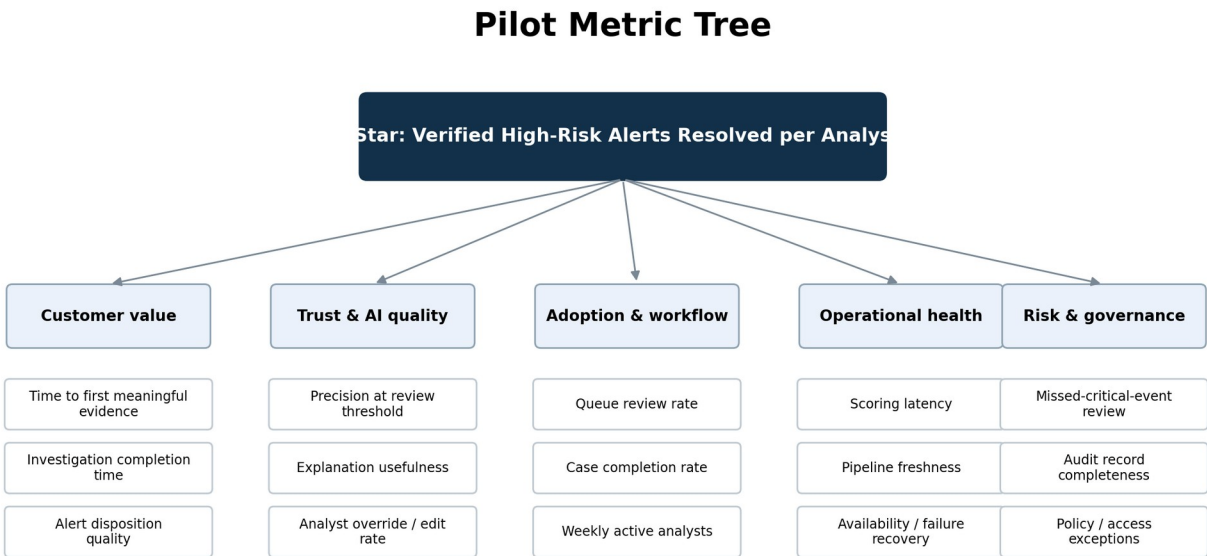
Element	Portfolio recommendation
Participants	A small group of experienced fraud analysts, one operations lead, one compliance reviewer, and named product/engineering/model owners.
Duration	Long enough to observe repeated workflow use, representative alert variation, and at least one operational issue; exact duration set by data volume.
Scope	Limited data sources, alert types, user roles, and case actions with mandatory human review.
Baseline	Current approved workflow or policy measured using the same definitions and comparable cohorts.
Scenarios	High-priority true risk, difficult false positive, insufficient evidence, duplicate/related alert, data delay, model/provider failure, escalation and approval.
Safeguards	No autonomous high-impact action; clear rollback; support owner;

daily review of critical incidents and quality signals.

03

Metric System

The metric system links business value to analyst behavior, AI quality, system performance, and governance.



All thresholds in this portfolio document are targets for validation, not reported production achievements.

Figure AFM-05.1 - Pilot metric tree. All thresholds are targets.

04

North Star and KPI Tree

The North Star should reward verified outcomes, not raw alert creation or model score volume.

Level	Metric	Definition / caution
North Star	Verified high-risk alerts resolved per analyst hour	Count only reviewed outcomes meeting the defined evidence and quality rule.
Customer value	Median time to first meaningful evidence	Do not equate a fast click with a useful investigation.
Customer value	Median time to disposition for comparable alert types	Segment by complexity and avoid mixing unrelated case classes.
AI quality	Precision / yield in priority bands	Evaluate at realistic queue capacity and by segment.

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-05
Level	Metric	Definition / caution
AI quality	Critical-segment recall proxy	Use reviewed historical outcomes and explicit uncertainty.
Trust	Explanation usefulness and evidence findability	Collect structured rating plus edit/reject reasons.
Workflow	Alert review, case completion, escalation, and approval rates	Interpret abandonment with qualitative evidence.
Operational	Data freshness, scoring latency, queue age, error/retry rate	Tie to risk tier and user impact.
Governance	Audit completeness, access exceptions, policy/model changes	Zero tolerance may apply to defined critical failures.
Economics	Cost per evaluated event / reviewed alert / completed case	Do not optimize cost before validating quality and value.

05

Instrumentation Plan

Instrumentation must be minimal, privacy-aware, versioned, and tied to a decision.

Event / measure	Properties	Owner / use
alert_created	alert type, severity, model/rule/policy versions, segment, trace ID	AI/Product - quality and volume.
alert_opened	analyst, assignment, queue position, age, severity	Product/Ops - prioritization behavior.
evidence_viewed	evidence type, source, time from open	Product - evidence usefulness.
ai_assistance_used	feature, retrieved sources, output version, edit/reject state	AI/Product - controlled usefulness.
disposition_saved	outcome, rationale completeness, escalation, reviewer status	Risk/Analytics - workflow and labels.
case_closed	time, approvals, final outcome, follow-up	Product/Ops - completion and value.
system_failure	component, severity, duration, affected alerts, recovery	Engineering/Risk - reliability and incident response.
access_or_policy_event	actor, object, action, result, version	Security/Compliance - control evidence.

06

Pilot Method

Use mixed quantitative and qualitative evidence to understand why the metrics moved.

1. Confirm definitions, baseline period, representative segments, privacy boundary, and pilot owners.

2. Run task-based usability sessions on the alert queue, evidence view, investigation, disposition, and AI-assistance controls.
3. Validate offline AI quality and simulate alert-volume impact before exposing analysts to the candidate policy.
4. Release to a limited cohort with daily critical-signal review and controlled configuration changes.
5. Collect product telemetry, model quality, operational incidents, analyst feedback, and reviewed case samples.
6. Perform weekly error analysis across false positives, missed positives, explanation issues, workflow friction, and data problems.
7. Complete a phase-gate decision: expand, revise, pause, or stop with evidence and owners.

07

Exit and Stop Criteria

Pilot success requires all critical dimensions to pass; strong performance in one area cannot compensate for a failed safety or governance gate.

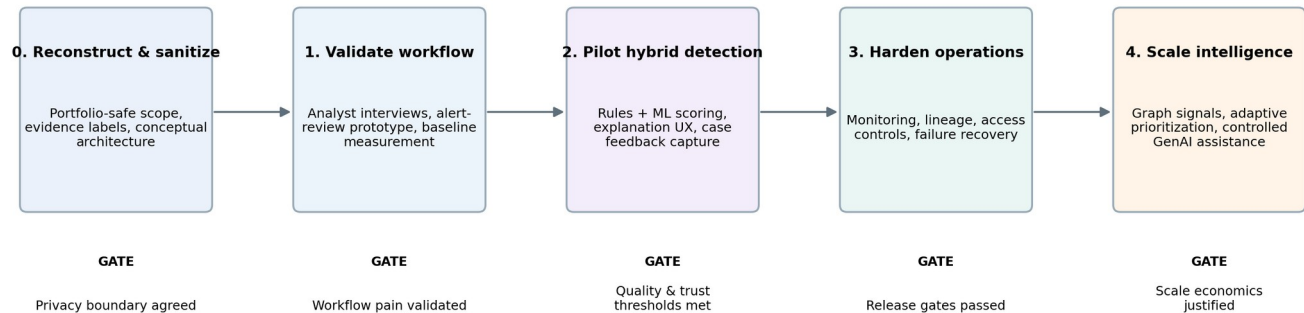
Decision	Criteria
Expand	Customer value improves; priority quality meets target; analysts understand evidence; mandatory controls pass; operations are stable; economics are acceptable for next stage.
Revise and repeat	Value signal exists but workflow, explanation, model segment, alert policy, data quality, or monitoring requires a bounded change.
Pause	Evidence is insufficient, baseline is unreliable, labels are not trustworthy, or an operational dependency prevents valid evaluation.
Stop	Critical security/privacy failure, unacceptable missed-risk exposure, no meaningful workflow value, or cost/complexity is disproportionate to validated benefit.

08

Roadmap

The roadmap sequences learning and control before feature expansion.

Evidence-Gated Product Roadmap



Roadmap sequence is a portfolio recommendation and does not reveal the confidential company delivery plan.

Figure AFM-05.2 - Evidence-gated roadmap recommendation.

Horizon	Objective	Key evidence
Phase 0 - Sanitize and align	Public portfolio boundary, stakeholder map, problem statement, architecture, evidence labels.	Confidentiality review and artifact approval.
Phase 1 - Workflow validation	Baseline, user interviews, journey, prototype, usability findings.	Validated pain, workflow, and information hierarchy.
Phase 2 - Offline AI validation	Versioned data, labels, baseline, candidates, calibration, error analysis.	Candidate decision with limitations and alert-volume simulation.
Phase 3 - Limited pilot	End-to-end workflow, controls, monitoring, support, measurement.	Customer value, AI quality, reliability, governance evidence.
Phase 4 - Hardening	Scale testing, broader segments, incident drills, cost and operations.	Release readiness and sustainable operating model.
Phase 5 - Controlled expansion	Graph intelligence, richer assistance, selected integrations, broader roles.	Incremental value and risk approval for each expansion.

09

Prioritized Backlog

The backlog focuses on evidence and release risk before optional intelligence features.

Priority	Backlog item	Outcome
P0	Traceable ingestion, validation, exception state, and replay	End-to-end data reliability
P0	Versioned rule/model signals, calibration, and alert policy	Reproducible risk decision

CONFIDENTIAL AI PRODUCT PUBLIC PORTFOLIO EDITION		AFM-05
Priority	Backlog item	Outcome
P0	Analyst queue, alert evidence, disposition, escalation, and audit	Core customer workflow
P0	Access control, provider boundary, logging, retention, incident path	Mandatory governance
P0	Benchmark, label guide, segment metrics, rollback candidate	AI release evidence
P0	Instrumentation QA and pilot dashboards	Decision-quality evidence
P1	Evidence-linked explanation templates	Trust and efficiency
P1	Controlled Generative AI summary and case draft	Productivity with review
P1	Relationship / graph evidence view	Complex investigation support
P1	Champion/challenger and policy simulation	Safer model and threshold change
P2	Similar-case retrieval and recommended next checks	Advanced analyst assistance
P2	Broader integrations and administration	Scale and enterprise fit

10

Risk and Decision Register

The pilot requires explicit ownership of both product uncertainty and operational risk.

ID	Risk / decision	Trigger	Owner action
D-01	Set the alert policy separately from the model score.	Model candidate chosen.	Simulate volume and approve policy version.
D-02	Require human approval for high-impact states.	Workflow design.	Configure role and audit requirements.
R-01	Labels are inconsistent or delayed.	Low agreement or ambiguous outcomes.	Revise guide, adjudicate, restrict supervised claims.
R-02	Pilot queue overloads analysts.	Backlog age or volume exceeds capacity.	Pause/adjust policy; do not hide the issue through silent suppression.
R-03	GenAI produces unsupported or sensitive output.	Grounding/leakage test or incident.	Disable or restrict feature, investigate, revalidate.
R-04	Critical segment quality degrades.	Threshold breach or adverse review.	Rollback, segment policy, or stop pilot.
R-05	Data freshness or lineage fails.	Monitoring breach.	Flag affected alerts, reconcile, and assess impact.

11

Operating Cadence

The operating cadence keeps product, AI, risk, and engineering decisions synchronized.

Cadence	Review
Daily during pilot	Critical incidents, data freshness, scoring failures, queue overload, privacy/security events.
Twice weekly	Alert samples, false positives, missed-risk review, explanation issues, workflow friction.
Weekly	Metric dashboard, segment quality, product feedback, backlog, configuration changes, decision log.
At every release	Evaluation report, model/rule/policy versions, control evidence, rollout and rollback plan.
Phase gate	Customer value, AI quality, operations, governance, risk, economics, and next decision.

Product Manager Reflection

The product is the controlled decision workflow, not only the model

The architecture makes the AI models visible, but the customer outcome depends on the full operating loop: reliable data, calibrated prioritization, explainable evidence, analyst judgment, feedback capture, monitoring, and governance. A technically strong model can still fail as a product when alerts are unactionable, explanations are generic, workflows are fragmented, or outcomes are not captured for learning.

What this case study demonstrates

- Converting a high-risk business problem into measurable user and system requirements.
- Balancing fraud detection, false-positive cost, explainability, privacy, reliability, and analyst capacity.
- Designing hybrid AI where rules, machine learning, relationship analysis, and Generative AI have distinct responsibilities.
- Using evidence labels and phase gates to avoid presenting assumptions or targets as achieved outcomes.
- Working across product management and solution architecture while preserving confidentiality.