

Hugo Vo

Senior Product Manager: Data, AI/ML & GenAI Products

I take data, machine learning and **GenAI/LLM products** from discovery to production in regulated financial services, and I stay accountable for adoption, trust and **measurable business impact**. My focus is AI that model-risk, compliance and business leaders can explain, audit and sign off.

[linkedin.com/in/hugovo](https://www.linkedin.com/in/hugovo) Houston area, TX

Measured impact Results from products taken from concept to production.

10× Defect detection rate (~2% → ~20%) ML decision-quality review system · patent pending US 2026/0187638 A1	100% Review coverage , up from a ~5% manual sample Same system, now the foundation of an enterprise quality product	+30% Detection yield Automated rule-tuning framework · patent pending US 2026/0260242 A1
-40% Capability cycle time Same framework, scaled across two product lines	-60% False-positive volume Statistical threshold tuning with detection effectiveness preserved	~\$500K Fraud losses cut in year one Anti-fraud engine taken from concept to multi-channel deployment

Selected work Production products, plus independent prototypes built on synthetic data.

Responsible-AI layer for model governance PROBLEM Governance reviewers and executives needed to understand complex model-tuning analysis quickly and consistently. APPROACH An LLM turns deterministic statistical results into review critiques and executive summaries. Output is schema-constrained and advisory only; a named human signs off; the AI can be switched off, with a non-AI fallback. RESULT Shipped to production with model-risk governance gates. The AI can never override a deterministic pass/fail decision.	Automated rule-tuning framework PROBLEM Tuning detection rules by hand was slow, inconsistent between cycles and hard to defend. APPROACH Co-invented an automated framework that evaluates candidate settings, scores them against a stated risk appetite and produces a governance-ready review package. RESULT +30% detection yield, -40% cycle time. Patent pending; scaled across two product lines and through a company acquisition.
--	--

ML decision-quality review

PRODUCTION

- PROBLEM** Quality control relied on reviewing a small random sample of decisions, so most defects went unseen.
- APPROACH** Hybrid NLP and gradient-boosted models with SHAP explanations score every decision and point reviewers at the likeliest defects, reconciled against real outcomes.
- RESULT** **100% coverage (from ~5%) and 10× defect detection. Patent pending.**

LLM investigation copilot

INDEPENDENT PROTOTYPE

- PROBLEM** Open-ended LLM answers can't be trusted in an investigation unless every claim can be checked.
- APPROACH** Designed the grounding, citation and system-prompt architecture for an LLM copilot inside a full-stack case management prototype (FastAPI, React/TypeScript, Docker).
- RESULT** **Every answer links back to the evidence it came from.**

PII protection for unstructured text

INDEPENDENT PROTOTYPE

- PROBLEM** Free-text narratives are valuable for ML but full of personal data.
- APPROACH** Evaluated 7 architectural options on accuracy and operating cost, then chose a layered, fail-closed design with a pass / review / block release gate.
- RESULT** **Benchmarked on a 1,000-record synthetic corpus, with a governed release path.**

Champion–challenger & red-team lab

INDEPENDENT PROTOTYPE

- PROBLEM** New models often reach production without proof that they beat the current one or hold up under attack.
- APPROACH** A shared lab where challengers (graph algorithms, NLP methods, ML models) compete against the incumbent on accuracy, latency and adversarial cases.
- RESULT** **Nothing is promoted without passing a gatekeeper, and every result is reproducible.**

How I build AI products

- 1. **Start from the decision, not the model.** Define who acts on the output, what success looks like and how adoption will be measured.
- 2. **Rules first, LLM where it adds value.** Deterministic logic handles routine cases; the LLM covers the rest, with cost controls.
- 3. **The LLM recommends; it never decides.** Human sign-off and deterministic checks stay the backstop.
- 4. **Set kill criteria before the build.** Pilots are measured against a stated bar, and a failed bar is a valid answer.
- 5. **Explainability, fairness, PII controls and auditability are designed in,** not retrofitted for model risk review.

GENAI & ML

LLM product design

Claude API

Prompt engineering

RAG & grounding

Agentic workflows

XGBoost

SHAP

NLP / transformers

PRODUCT

Discovery → production

Roadmaps & PRDs

User stories & acceptance criteria

Agile / Scrum

KPI definition

Change management

PLATFORM & DATA

Python

SQL

FastAPI

React / TypeScript

Docker

REST APIs

AWS (working knowledge)

GOVERNANCE

SR 11-7 / OCC model risk

PII & data classification

Access controls

Audit trail design

CAMS

Independent prototypes use synthetic data only. Detailed case studies are available on request.