



VEKTOR

BY
INVESTPUPPY

OUR BETTER WAY · VEKTOR BY INVESTPUPPY

**We Don't Say AI.
Let Us Explain Why.**

*What we used instead, where we used it, and why the boundary was
drawn before the first line of code.*

UNV-11 described what we watched others do with AI. This paper describes what we decided when it was our turn.

Written for DPMs, compliance officers, risk committees, and institutional partners evaluating Vektor's approach to machine learning governance.

OUR BETTER WAY · THE SERIES

The Unvarnished series documented what we witnessed in the rooms where financial software implementations fail — thirteen papers, published under the InvestPuppy name, naming the failure modes clearly because naming them seemed more useful than another optimistic vendor promise.

Our Better Way is the other half of that account. Each paper describes what we built in response to what we witnessed: the architectural decisions, the engineering discipline, and the operational choices that define how InvestPuppy approaches the problems the Unvarnished series named. Not theory. A record of how we applied what we learned.

This is Paper 04. Paper 01 (BW-01) covers the cost of technical debt and the architectural discipline that prevents it. Paper 02 (BW-02) covers the Vektor : House institutional implementation pathway. Paper 03 (BW-03) covers the configurability architecture that separates data-driven systems from code-driven ones. This paper covers the question that preceded all of them: where does machine learning belong in a systematic investment management platform — and where does it not?

1. The Moment Every Builder Faces

At some point in building Vektor, we faced the question that every financial technology team faces now.

Not whether to use AI. That question is not asked anymore. The pressure — from investors, from marketing, from the general direction of the industry — is not whether but how quickly. AI is the answer. The question of what problem it is the answer to has become secondary.

We asked a different question. We asked where.

Where does machine learning outperform human analysis on a well-defined, measurable task with clean inputs and an objectively evaluable output? We used it there.

Where does the decision carry accountability, require contextual judgment, and need to be explainable to a client, a regulator, or a compliance reviewer by a named human being? We did not use it there.

The boundary between those two categories was drawn before we wrote the first line of code. It has not moved.

This is not a conservative position on AI. It is what understanding AI looks like.

2. What UNV-11 Named

UNV-11 documented the pattern: a powerful technology arrives. The capabilities are real. Organisations begin asking not whether it fits their problem but how quickly they can deploy it. The question of fit is replaced by the urgency of adoption.

We watched this happen with Agile. With blockchain. With RPA. In each case the technology was genuine. In each case a significant proportion of implementations were not — because the implementations began before anyone had examined whether the process being automated was the process that should be running.

AI is not different in kind. It is different in scale. A flawed process automated with AI executes the flaw faster, at greater volume, with fewer opportunities for the human hesitation that might have caught the problem. The speed is the problem.

UNV-11 named three mechanisms: speed removes friction, authority makes AI outputs hard to challenge, and embedding converts a bad practice from a habit into infrastructure. Once encoded at scale, the process is no longer a habit that can be broken by a decision. It is a system that requires a project to change.

We documented these failure modes in other organisations. Then we had to decide what to do with them when it was our turn.

3. The Governing Principle

One sentence. Every decision that follows derives from it.

AI in Vektor is an instrument of computation. It does not make investment decisions, client decisions, or compliance decisions. Humans make those decisions — every time, at every gate, with full information.

This principle was not adopted in response to regulatory pressure. It was the starting position. The question at the outset was not "what do regulators require?" It was "where does machine learning outperform human analysis on a clearly defined, measurable task — and where does human judgment remain essential?" The boundary between those two categories is structural. It cannot be toggled in configuration. It cannot be overridden by a developer. It is enforced by the workflow architecture.

A system in which a human can be removed from a decision gate by configuration is not a system with human oversight. It is a system with optional human oversight. That is a different thing. We built the former.

WHY WE SAY MACHINE LEARNING

Every platform is AI-powered now. The term has been applied so broadly that it has stopped carrying information — it describes everything from a spreadsheet formula to a large language model generating autonomous decisions. We know, because UNV-11 documented the pattern.

We say machine learning because it is the accurate term for what we built. Machine learning names specific techniques: statistical analysis and classification applied to well-defined tasks with defined inputs, defined outputs, and measurable performance. That is what Vektor uses — in two places, for two specific purposes.

Calling it AI would describe a broader category than we built. We do not make broader claims than we can demonstrate.

In this industry, precision is not modesty. It is the only form of credibility that compounds.

4. Where We Used It

Two applications. Both bounded. Both with defined inputs, defined outputs, and a human review step before any capital-affecting consequence.

APPLICATION 1 — UNIVERSE SCREENING

The SageMaker statistical analysis environment applies liquidity, sector, and dividend-continuity screens to an exchange universe of 500 or more instruments and returns a ranked, filtered investable set.

What it does: ranks instruments against explicit, stated criteria. The reason any instrument is included or excluded can be stated precisely.

What it does not do: it does not select the final portfolio. The screened universe is the input to optimisation. A human analyst reviews the screening output before the optimisation runs.

The output is a list. The decision is human.

APPLICATION 2 — XGBOOST SIGNAL VALIDATION

The XGBoost classification model validates live trading signals against three years of daily price history and historical signal outcomes before those signals are presented to the portfolio manager for approval.

What it does: outputs a confidence score per signal. Signals below the configured threshold are flagged and filtered before PM review.

What it does not do: it does not approve strategies, assign clients, execute trades, or make any capital-affecting decision. Its output is one input to the PM review — visible, labelled, and overridable.

The output is a confidence score. The decision is human.

5. Where We Did Not Use It

The exclusions are design decisions. They will not change as the platform develops.

The portfolio manager approves every strategy before client assignment. That approval is an explicit human action — not a confirmation of an AI recommendation, not a threshold check, not an exception-based override of a default. The PM approves, or the strategy does not proceed.

The operations team records every cash funding event. No ML model determines whether capital is ready to deploy. A named human records the fact.

The portfolio manager confirms the allocation before any orders are submitted. The allocation engine calculates. The PM confirms. These are separate steps and the separation is structural.

No position is established without three explicit human actions in sequence: strategy approval, cash funding, allocation confirmation. None can be automated. None can be skipped. The system will not proceed without them.

This is not a workflow policy that can be changed in settings. It is an architectural property of the platform.

6. Why the Boundary Has to Be Drawn First

UNV-11 named it: once a model is in production, the boundary is wherever the outputs end — which is almost always further than anyone intended.

Boundaries drawn before implementation are structural. Boundaries drawn after implementation are aspiration.

A compliance framework that says "human oversight is required" but does not enforce that requirement at the architecture level produces human oversight when people remember to apply it. That is not the same thing. A system in which the audit trail is a feature — added by a developer who remembered to log — will have gaps where a developer forgot. A system in which the audit trail is enforced at the infrastructure level, below the service code, cannot have those gaps.

The same logic applies to the AI boundary. We drew it before we wrote the code because a boundary drawn after the code is written has to fight the code to hold its position. A boundary drawn before is the architecture.

7. The Regulatory Convergence

Three financial regulators — MAS in Singapore, the SFC and HKMA in Hong Kong, and the FCA and PRA in the United Kingdom — have independently published AI governance frameworks for investment management. All three converge on the same four requirements.

Explainability: the reasoning behind an AI output must be statable in plain terms.

Human oversight: a human must be present at investment decision points.

Accountability: there must be a named human accountable for AI-assisted decisions.

Auditability: AI behaviour must be reviewable after the fact.

Vektor's architecture satisfies all four. Universe screening criteria are explicit and statable. Signal validation confidence scores are visible and labelled. The PM approval gate is a human action, not a threshold. The audit trail captures every action, timestamped, attributed, append-only.

We built this framework in 2024 and 2025, before these guidance documents were finalised. Not because we anticipated the regulation. Because we understood the technology well enough to reach the same conclusions the regulators did.

We didn't follow the regulation. We got there first. Because we understood what we were building.

8. What This Means in Practice

For a DPM or family office principal evaluating Vektor:

The investment decisions remain yours. The allocation methodology is systematic and documented. The signals are validated and confidence-scored. The portfolio construction is reproducible and auditable. At every point where a decision is made that affects your clients' capital, a named human makes it.

Vektor makes that decision better-informed and faster to execute. It does not make the decision.

For a compliance officer or risk committee:

The AI governance documentation — the two bounded ML applications, the human approval gates, the append-only audit trail — is available for review. WP-08 covers the full technical and governance detail. WP-06 covers the audit and compliance architecture. The system produces the compliance record as a structural output of normal operation. It does not require a separate compliance process run after the fact.

For an institutional partner or regulator:

The framework is aligned with MAS FEAT, SFC/HKMA, and FCA/PRA AI governance principles across all four requirements. The alignment is architectural, not claimed in a policy document. The policy follows from

the architecture. Not the other way around.

'AI-powered' has become a marketing claim.

We made it an engineering decision.

ABOUT INVESTPUPPY

InvestPuppy Pte Ltd is a Singapore-based financial technology company building Vektor — a systematic listed equity portfolio management platform for boutique discretionary portfolio managers, independent wealth managers, and family offices. The platform gives wealth management professionals institutional-grade systematic portfolio construction at a price point accessible to boutique practices.

This is Paper 04 of the Our Better Way series — InvestPuppy's account of what was built in response to the failure modes documented in the Unvarnished series. The companion paper for this subject is UNV-11: The AI Revolution — Automating Bad Practice at Scale. All thirteen Unvarnished papers are available ungated at investpuppy.com/unvarnished.