



**INVESTPUPPY
UNVARNISHED**



InvestPuppy

The AI revolution — automating bad practice at scale

*For anyone else who's been in **that** room.*

Serious about outcomes. Not serious about ourselves.

About InvestPuppy

InvestPuppy builds Vektor, a systematic portfolio management platform for wealth management professionals.

What this series is

Every experienced practitioner thinks it. Almost nobody publishes it: *there has to be a better way.*

We thought it long enough that we went and built one.

Before that, we spent decades in the rooms.

The same project was failing for the third time. The project plan had been signed off before anyone had understood the first requirement.

The relationship was too warm for anyone to say so.

It isn't incompetence.

It's the incentive structure doing its job perfectly.

Nobody planned it this way. Nobody needed to.

We got frustrated enough to write it down. Then frustrated enough to build something. These papers are the writing-it-down part.

The something we built is a separate conversation.

AI doesn't fix broken processes. It executes them faster, at greater scale, with fewer opportunities for the human hesitation that might have caught the problem.

The question isn't whether your AI implementation will work. The question is what, exactly, it will be working very hard to do.

1. The New Excitement

The pattern is familiar. A powerful new technology arrives. The capabilities are real – genuinely transformational for the problems it is designed to solve. 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 have watched this happen before. Each time the technology was real. Each time a significant proportion of the implementations were not.

Agile.

A set of principles for software teams working on problems too complex to specify in advance. It worked for that. Then organisations began applying it to regulatory compliance programmes and financial software implementations – work that was not complex in the Agile sense, it was simply uncertain. The certifications multiplied. The outcomes did not.

Blockchain.

Between 2017 and 2020, a significant portion of financial services strategy documents contained the word. Most no longer do. The methodology phase has, mercifully, passed. What remains are the specific applications where the underlying technology actually solved a real problem – and the silence about everything else that was attempted.

RPA.

Robotic process automation was deployed to accelerate processes that had never been examined. The robots

executed flawlessly. They executed the wrong thing flawlessly, at speed, continuously. The process reviews that should have happened before deployment happened after the first audit finding.

AI is different in scale. It is not different in kind.

2. What AI Actually Does

AI excels at pattern recognition, systematic optimisation across large parameter spaces, anomaly detection, and the consistent execution of defined rules at volume. For tasks where the problem is well-defined, the data is clean and representative, and the output can be evaluated objectively, AI performs at a level that human analysts cannot match consistently over time.

This description contains three conditions. Well-defined problem. Clean, representative data. Objectively evaluable output. In a surprising number of AI implementations in financial services, none of these conditions has been verified before deployment begins.

The boundary between what AI should do and what it should not is not drawn by the technology. It is drawn by the people deploying it — or it is not drawn at all.

AI does not ask whether the process it is automating is correct.

It asks how to execute it more efficiently.

AI will fill the space. It will not ask permission.

3. The Knowledge Base That Isn't

The presentation was confident. The vendor had done this before. Slides were clean. The capability was real. And then, offered as evidence of thoroughness, came the proudest claim: all of the client's existing data, processes, target operating models, and documentation had been ingested into a central repository. This repository would form the knowledge base for the new AI tools. The organisation's institutional knowledge, captured and ready.

Nobody in the room asked the obvious question.

The obvious question is: knowledge of what, exactly?

What was recorded is not what was true.

Every organisation accumulates documentation. Process maps written for a regulatory submission. Target operating models produced by a consulting firm, approved by a steering committee, and filed without ever being operationalised. Policy documents that describe how decisions should be made, written by people who were not the people making them. Workflow diagrams that reflect the process as it was designed, not the process as it was executed — which diverged in week three and was never updated.

This documentation is not institutional knowledge. It is the institution's paper record of its intentions. The two things are related. They are not the same thing. An AI system cannot tell them apart.

The 'etc.' is where the problem lives.

The vendor's repository contained data, processes, TOMs, and — in the word that deserves more attention than it received — et cetera. The scope of the ingestion was not defined. It was not examined. It was not qualified. Volume was the metric presented to the client. Completeness was the selling point. The more that was ingested, the more comprehensive the knowledge base. The more comprehensive the knowledge base, the more capable the AI.

This logic is not wrong. It is precisely backwards.

The repository contained five years of process documentation. It contained the TOM from 2019 that was replaced in 2021. It contained the compliance procedures that were superseded when the regulation changed. It contained the workflow diagrams from the system that was decommissioned. It contained the decisions that were made under a market regime that no longer existed. It contained everything that had been recorded, without distinction between what was current and what was obsolete, between what was accurate and what was aspirational, between what reflected how the organisation worked and what reflected how someone had once hoped it might.

The AI treated all of it with equal confidence.

Completeness as a substitute for accuracy.

The vendor called it capturing institutional knowledge. What they captured was institutional amnesia — everything the organisation had recorded, including everything it had recorded incorrectly, presented as ground truth to a system with no mechanism for doubt.

The absence of a quality check was not a gap in the implementation. It was the implementation. The value

proposition was the ingestion itself — the comprehensiveness, the speed, the fact that it had all been done. Examining the data before ingesting it would have been slower. It would have required decisions about what to exclude. It would have required the organisation to confront, in detail, the quality of its own documentation. This is uncomfortable, slow, and not in scope. The vendor was not paid to do it.

The vendor called it institutional knowledge.
The organisation called it a foundation.
It was years of undifferentiated decisions —
some of them wrong, none of them examined,
all of them now running at scale.

4. How Bad Practice Gets Industrialised

A financial services firm runs a manual client onboarding process. The process is slow, inconsistent, and prone to error. The decision is made to automate it with AI. The AI is trained on historical onboarding data. The automation goes live. Throughput triples. Error rates fall.

Three things have happened. The speed of onboarding has increased. The inconsistency has decreased. And every bias, every undocumented workaround, every regulatory shortcut that was present in the historical data has been encoded into the model, given the weight of systematic validation, and scaled to three times the previous volume.

The process was not examined before it was automated. It was accelerated.

The three mechanisms of industrialisation.

The first is speed. A human executing a flawed process has natural friction — pauses, second thoughts, the occasional question that surfaces an inconsistency. AI removes this friction. The flaw executes without hesitation, at volume, continuously.

The second is authority. A decision made by an AI system carries institutional weight that a decision made by a junior analyst does not. The model's output is treated as objective — derived from data, validated by results, free from the biases of individual judgment. The flawed assumption embedded in the training data is invisible. The model's confidence is not. Challenging an AI output requires a different kind of courage than challenging a colleague's recommendation — it requires arguing

against the appearance of rigour itself.

The third is embedding. A manual process can be changed by changing the instructions. An AI model trained on years of historical data requires retraining, re-validation, and redeployment to change. The bad practice, once encoded, is no longer a habit. It is infrastructure.

We have seen this in compliance screening that encoded historic review thresholds as permanent parameters. In portfolio rebalancing logic that systematised the manual workarounds of a single analyst. In client risk profiling that automated the assumptions of a process nobody had documented. In credit decisioning models trained on approval histories that reflected the lending preferences of a desk that no longer existed — faithfully reproduced, at scale, for a different market.

In each case the AI was working correctly. The process it was executing was the problem.

5. The Process Audit Nobody Does

Before any AI implementation, there is a question that should be mandatory and is almost never asked: is the process we are proposing to automate the process we should be running?

This question is not asked for three reasons. The first is that it is uncomfortable — it implies the existing process may be wrong, which implies the people who designed and run it may have been doing something wrong. The second is that it is slow — answering it requires genuine process examination, which the implementation timeline does not allow for. The third is that it is not in scope — the AI vendor was engaged to automate the process, not to audit it.

The incentive structure, again, doing its job perfectly.

The vendor was not paid to ask if the process was right.

The client did not want to hear that it wasn't.

The AI automated it anyway.

6. Thoughtful Integration

The organisations that benefit materially from AI are not the ones that move fastest. They are the ones that ask the

process question before the procurement question. This is not a counsel of slowness. It is a description of what separates implementations that compound in value from implementations that compound in debt.

Thoughtful AI integration has a sequence. It begins with understanding the process in its current state — not the documented version, but the actual one: the workarounds, the undocumented exceptions, the decisions that are made by feel rather than by rule. This requires talking to the people who run the process, not the people who designed it. They are not always the same people.

The second step is determining whether the process is correct. Not optimised — correct. Whether it is doing what the organisation intends, producing the outcomes the organisation needs, and carrying the risks the organisation has chosen to carry. If the answer is no, the process should be redesigned before it is automated. An AI implementation is not a process redesign. It cannot substitute for one.

The third step is identifying which elements of the process are genuine pattern-recognition and optimisation problems — the tasks where AI adds material value. The fourth is designing the AI layer around those elements specifically, with explicit boundaries at the points where human judgment, accountability, and contextual interpretation are required.

The boundary between what AI should do and what it should not must be drawn before the model is built. It cannot be retrofitted. Once a model is in production, the boundary is wherever the model's outputs end — which is almost always further than anyone intended.

We built Vektor with this sequence explicitly. Systematic portfolio construction — signal selection, deterministic max-Sharpe optimisation, lot sizing — are pattern-recognition and calculation problems. We used systematic methods. Mandate interpretation, investment judgment, client relationship management — these are not. We did not automate them. The boundary was drawn before the first line of code was written, not after the first implementation complaint was received.

7. The Test

Before any AI or ML implementation, three questions. Not as a framework. As a minimum.

Is the process we are proposing to automate the process we should be running?

If the answer requires more than a day to verify, the process has not been examined. Stop. Examine it. The implementation can wait. The organisation's processes cannot afford not to be examined before they are encoded at scale.

Which elements of this process are pattern-recognition problems, and which require human judgment?

If the boundary has not been drawn explicitly, it has not been drawn. AI will fill the space. It will not ask permission.

If this AI implementation encodes the current process perfectly and executes it at ten times the volume, is that a good outcome?

If the answer is not an unqualified yes, the process is not ready to be automated. It is ready to be examined.

Adding an AI layer to a broken process

does not fix the process.

It fixes it in place.

The Unvarnished Series

Thirteen field notes on why financial software implementations fail.

Published under our own name. All freely available, no registration required.

Why Implementations Fail (UNV-01)

The same implementation failures repeat across the industry, unrecorded, because no party benefits from writing them down. This paper names the pattern openly, once, rather than let it repeat quietly again.

Nobody Owns It (UNV-02)

Ask a vendor who owns an implementation and they say the client. Ask the client and they say the vendor -- both are right, and that shared, undefined ownership is the actual failure mode, not any single decision either side made.

Complexity as Cover (UNV-03)

A consultant's real product is often the document justifying the advice, not the advice itself -- built to survive the project's failure, not prevent it. Complexity becomes the cover story for that gap.

Plan Last (UNV-04)

Project plans are written before anyone understands the workflow they claim to schedule, in columns spanning weeks one through thirty-six. Planning first and discovering later gets the sequence backwards.

The Change Request Trap (UNV-05)

A change request is a formal admission that the original specification was wrong, and also a commercial mechanism for charging separately for what should have been included from the start. Neither party is surprised when it happens.

Invisible Stakeholders (UNV-06)

Every implementation has a second set of stakeholders who never attended a workshop or reviewed a specification, yet are the ones who have to live with what was built. Their absence from the room doesn't make them absent from the consequences.

The Hammer, the Tool, and the Methodology (UNV-07)

A tool gets picked up, used for the job it suits, and put down. A methodology requires certification, governance, and a steering committee -- and the industry keeps reaching for the second when it only ever needed the first.

Life After Live (UNV-08)

Go-live is the final milestone on the Gantt chart, but it's the beginning of the harder problem: an organisation built around a legacy system now has to actually live inside the new one. The project plan ends; the real work doesn't.

The U in UAT (UNV-09)

UAT sign-off documents carry the right signatures -- from consultants testing scripts written by other consultants, not from anyone who actually uses the system day to day. The U in UAT is the part everyone skips.

Nine Women (UNV-10)

Fred Brooks named it in 1975 and the same boardroom conversation still happens fifty years later: a project is late, so the proposal is more people. The logic feels intuitive and has always been wrong.

The AI Revolution — Automating Bad Practice at Scale (UNV-11) (you are here)

AI doesn't fix a broken process -- it executes it faster, at greater scale, with fewer chances for the human hesitation that might have caught the problem. The real question isn't whether an AI implementation will work, but what it will end up working very hard to get wrong.

Model Bank — the perfect solution for banks that don't exist (UNV-12)

A 'model bank' configuration is a composite of institutions that don't exist, built from someone else's defaults before your institution was ever part of the conversation. Every subsequent customisation is just correcting an assumption that was never really about you.

The parameter that broke the camel's back (UNV-13)

A platform demo answers every question with yes, right up until one specific parameter turns out to be the one nobody actually tested. This paper is the account of exactly where that line sits, and what breaks on the other side of it.

Our Better Way

This is what we built instead.



"So, what did you do about it?" — a fair question after thirteen papers on what goes wrong. InvestPuppy builds Vektor, a systematic portfolio management platform for wealth management professionals — the direct answer, series by series.

These four papers are the how to the thirteen whys above — not a sales pitch, but a technical account of the specific choices Vektor made where the industry's default pattern would have produced the same failures just described.

- Technical Debt (BW-01)
- Vektor House Implementation (BW-02)
- The Configurability Gap (BW-03)
- AI as Instrument (BW-04)

More at investpuppy.com