

From Data to Strategy

Transforming Investment Processes with a Self-Correcting Multi-Agent Framework

Shuen Mei
VP, Director of Software & Data Engineering

HARBOR CAPITAL ADVISORS

PART ONE — THE PROBLEM

Why AI is hard in finance

The same model that sounds right can still be wrong, unsupported, or impossible to audit.

01

Who we are

A technology-enabled investment firm with a strong focus on research, client outcomes, and operational discipline.



Curated ETF Selection

We offer a select suite of active ETFs focused on pursuing compelling, risk-adjusted returns.



Driven by Curiosity

Our approach combines market obsession with commitment to advisor success.



Manager of Managers

We bring objectivity and speed to decision-making, free from internal product bias.



Advisor Partnership

We pair investment insight with a clear commitment to advisor and client outcomes.

Why AI is difficult in investment management

Fluent language is useful, but regulated finance also requires accuracy, traceability, and reviewable evidence.

What AI is good at

- Drafting quickly
- Summarizing noisy inputs
- Producing a first pass
- Translating language across formats

What finance demands

- Factual accuracy
- Supported claims
- Regulatory compliance
- Consistency over time
- Auditability and source traceability

The gap

A model can sound confident and still be wrong. In finance, confidence is not enough.

Fluency ≠ truth

That gap is where risk, rework, and compliance concerns show up.

THE CONCRETE USE CASE

Trade commentary is a concrete, high-value use case

What is it?

Timely, informed narrative explaining market views and target-weight changes after a rebalance.

Why automate?

Manual synthesis can consume 10+ hours per week and creates inconsistent review cycles.

What is the goal?

Scale drafting while preserving truthfulness, compliance review, and a complete audit trail.

THE WORKFLOW

- 1 **Collect source documents and market context**
- 2 **Retrieve strategy-specific target weights**
- 3 **Generate a structured first draft**
- 4 **Decompose the draft into atomic claims**
- 5 **Validate each claim against source evidence**
- 6 **Route approved output to the business workflow**
- 7 **Human reviewer approves, edits, or escalates**

Human approval, evidence links, and model identifiers are captured in the audit trail.

PART TWO — THE APPROACH

Architecture and Act–Verify–Refine

Build reliability through a repeatable, scalable control loop

02

A self-correcting multi-agent framework

The workflow generates commentary, verifies it against evidence, and refines it until it is production-ready.

WHY “GOOD ENOUGH AI” FAILS IN FINANCE

Fluent output is not the same as reliable output.

01 FLUENCY ≠ TRUTH

LLMs can generate convincing language without factual support.

02 RISK COMPOUNDS

Hallucinations and unverifiable claims create compliance and reputational exposure.

03 CONTROL IS REQUIRED

Outputs must be grounded, auditable, and correctable before release.

Treat narrative like code

Generate statements that can be tested, verified, and repaired until every claim is supported.

DESIGN PRINCIPLE

Treat narrative like code

Move from “generate and trust”
to “test, repair, and retest.”

IDEA

Produce
testable
claims, not
just fluent text.

EXECUTION

Build a loop:
Extract → verify
→ fix → retest.

OUTPUT

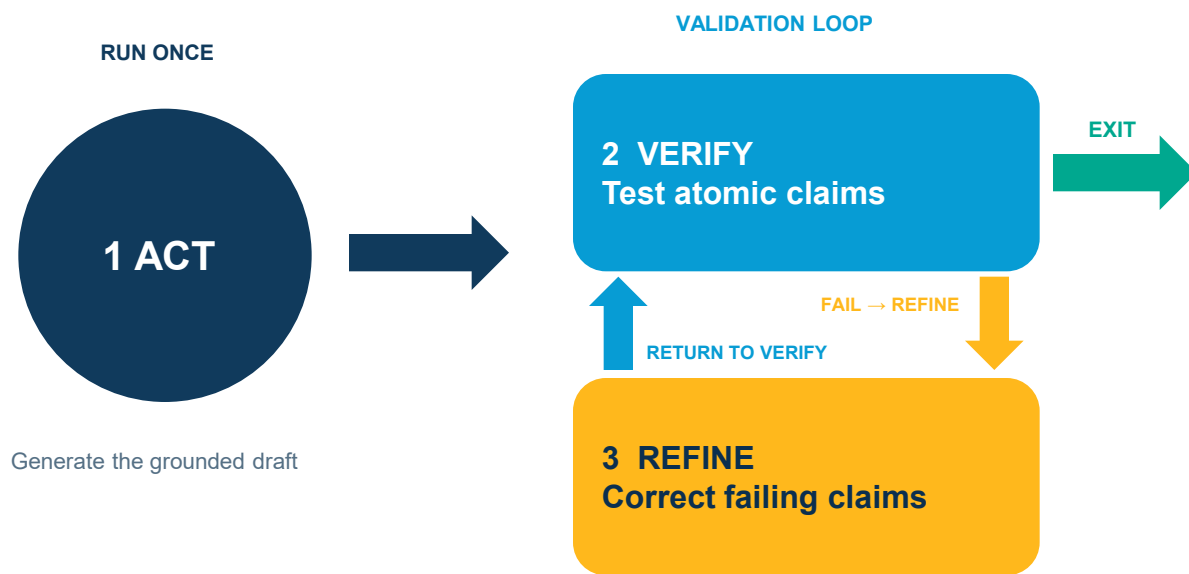
Reliable
automation
without
assuming the
model is
always right.

UNIT-TEST LOOP: EXTRACT → VERIFY → FIX → RETEST

ACT ONCE • VERIFY ⇌ REFINE

The loop makes reliability systematic

Every draft becomes a testable artifact rather than a final answer.



THE CONTROL LOOP

1

ACT

Generate the grounded draft once from governed inputs.

2

VERIFY

Extract atomic claims and test each one against ground truth.

3

REFINE

Correct failing claims, then return the revised draft to Verify.

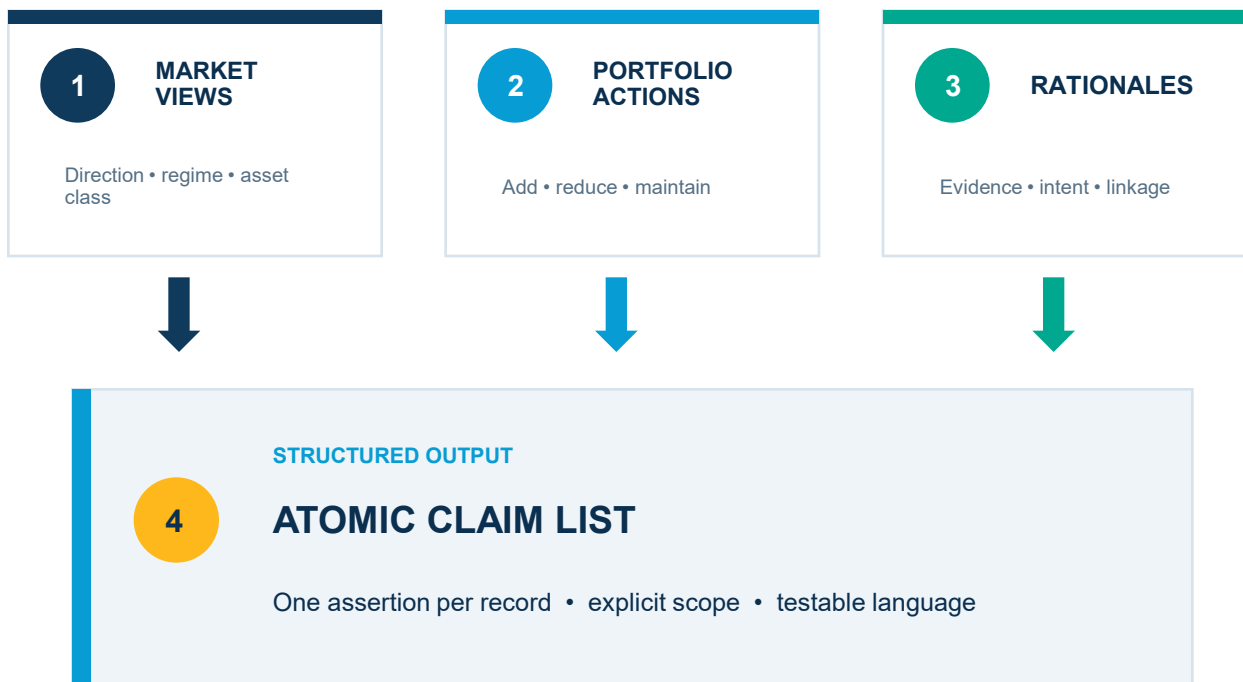
EXIT FROM VERIFY • All critical claims pass the reliability gate.

Reliability comes from repeated verification and refinement—not repeated generation.

ATOMIC CLAIM DECOMPOSITION

Atomic claims turn prose into testable assertions

Verification works only when statements are explicit, scoped, and independently testable.



VALIDATION READY

Structured claims flow directly into machine verification.

WHAT GETS EXTRACTED

- 1 MARKET VIEWS**
Macro environment, asset classes, and market direction.
- 2 PORTFOLIO ACTIONS**
Weight changes, additions, reductions, or maintained positions.
- 3 RATIONALES**
Evidence and strategy intent behind each portfolio action.
- 4 ATOMIC CLAIMS**
One explicit, scoped, independently testable assertion per record.

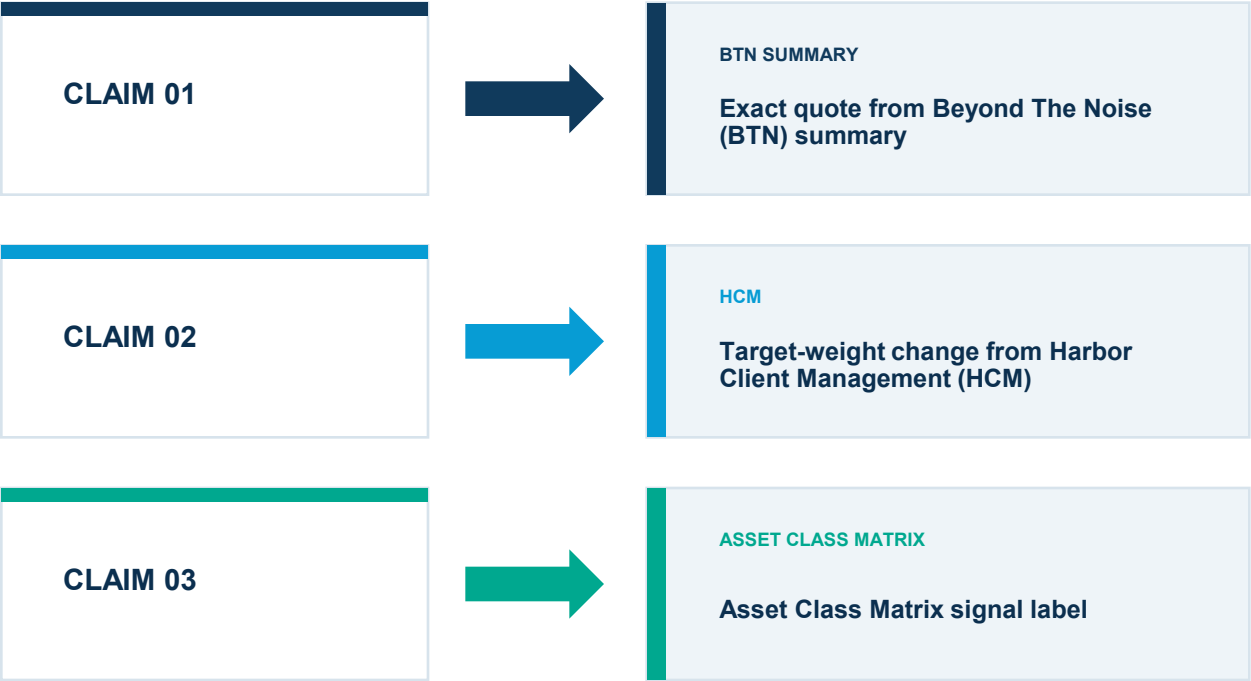
RESULT

A machine-ready list of explicit, testable claims.

TRACEABILITY

Every claim maps to exact supporting evidence

The validator records what supports the claim, where it came from, and how it was judged.



TRACEABILITY RECORD

Claim + status + evidence + source stay linked.

STRUCTURED VERIFICATION

- 1

CLAIM

The exact assertion extracted from the draft.
- 2

STATUS

Supported • Misleading/Incorrect • Missing Data
- 3

EVIDENCE

The precise quote or data point used for the decision.
- 4

SOURCE

Document, section, and lineage retained for audit.

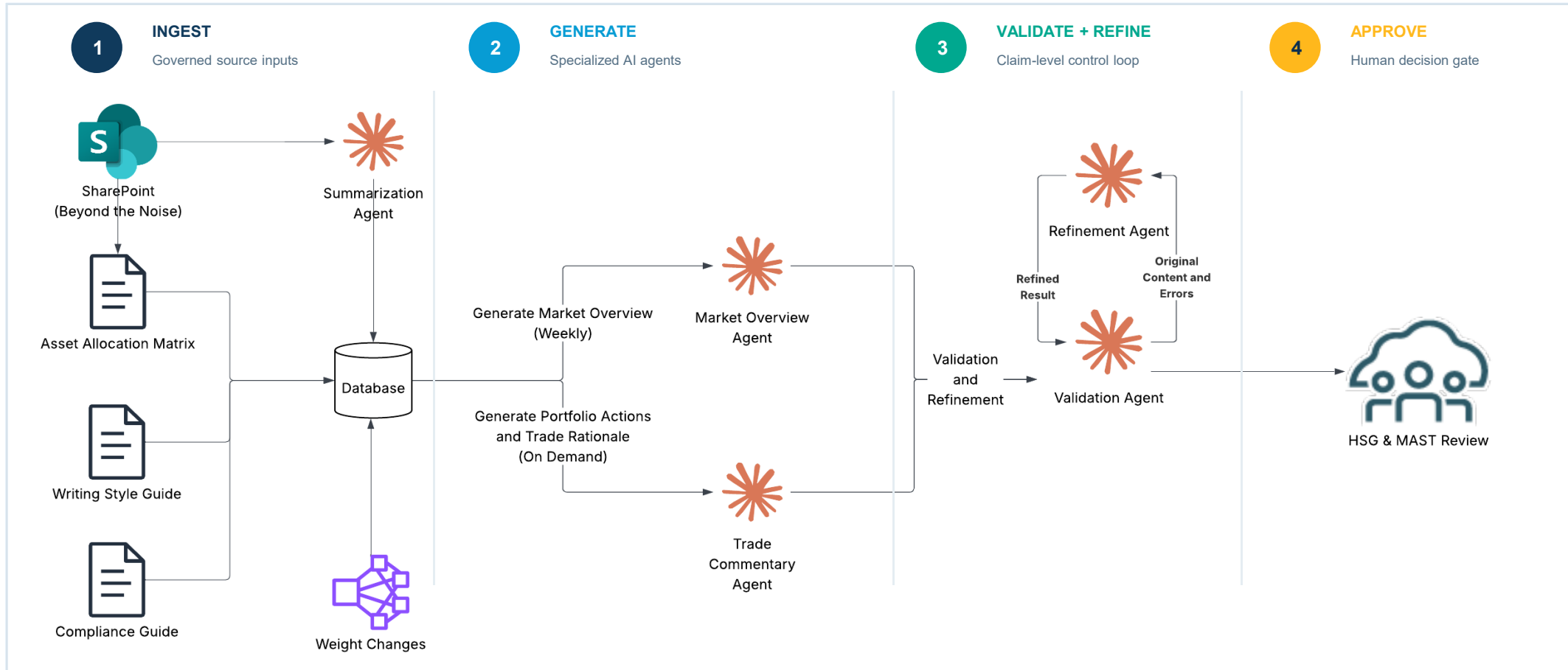
AUDIT READY

Document, section, and lineage retained for review.

SYSTEM ARCHITECTURE

A governed workflow connects data, generation, validation, and review

Source inputs flow through specialized agents, claim-level controls, and a final human approval gate.



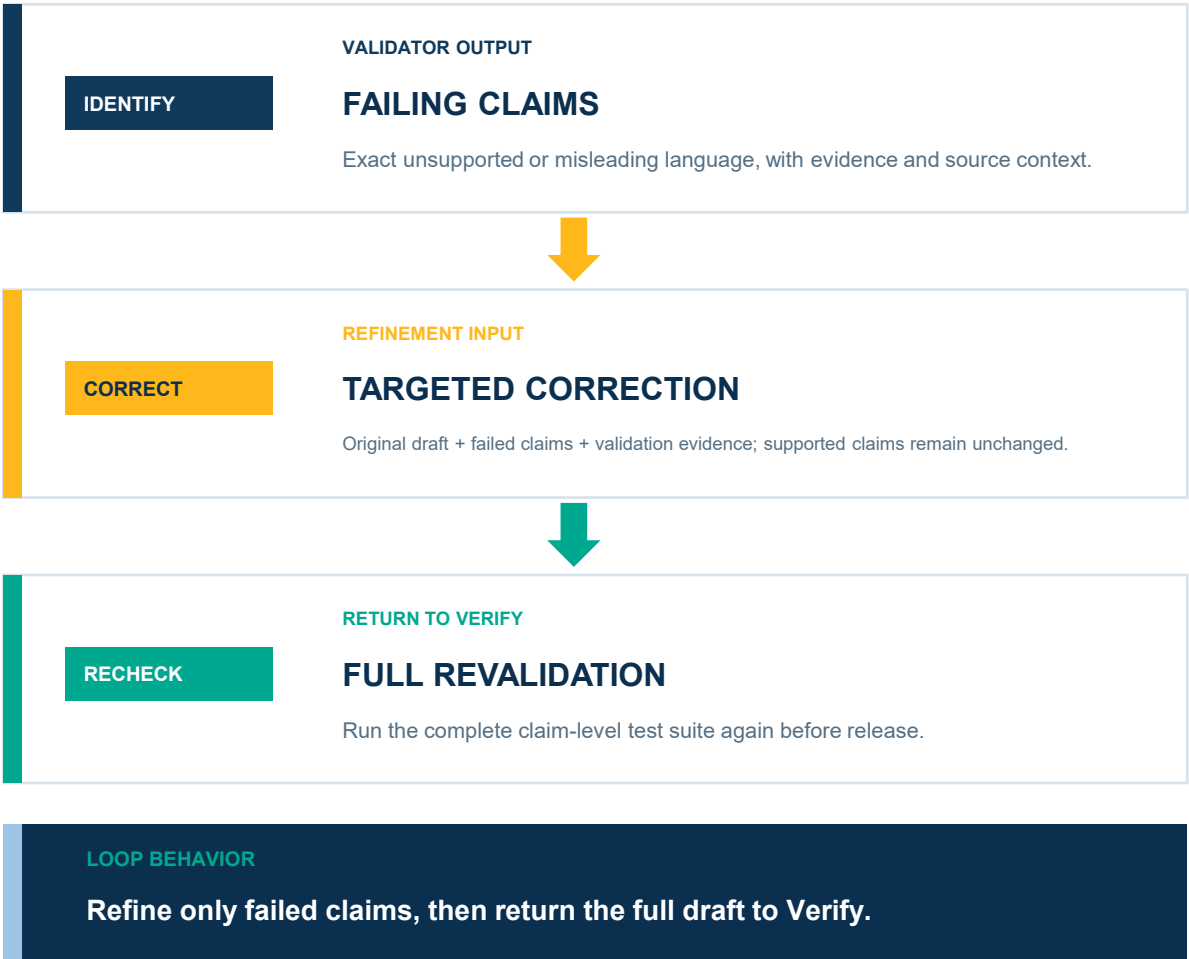
CONTROL PRINCIPLE

Automation produces the draft; evidence controls and human review govern release.

TARGETED REFINEMENT

Refinement fixes only what the evidence rejects

The model receives the original draft plus failing claims and their validation evidence.



REFINEMENT PROMPT

SUPPLY CONTEXT

Original draft, failed claims, evidence, and source context.

CORRECT SELECTIVELY

Correct only rejected claims; preserve content that already passed.

RETURN TO VERIFY

Return the revised draft to the complete claim-level test suite.

RELEASE ONLY WHEN CLEAR

Exit only when no misleading or unresolved missing-data claims remain.

RELEASE CONDITION

All critical claims pass; no missing-data issues remain.

Governance is embedded in the workflow

Controls are designed into the pipeline rather than added after generation.

EVIDENCE CONTROL

GROUNDING EVIDENCE

- Every factual statement is checked against approved source material.
- Exact evidence and source locations are retained.
- Unsupported statements cannot silently advance.
- Inputs and outputs remain strategy-specific.

DECISION CONTROL

HUMAN REVIEW

- Business users can review, tune, regenerate, and approve.
- Material exceptions are escalated rather than guessed.
- The workflow supports compliance review before distribution.

RECORD CONTROL

AUDIT TRAIL

- Validation output is stored as a structured record.
- Each claim retains status, evidence, and source.
- Review history can be reconstructed end to end.
- Reusable across future client-facing AI workflows.
- Control becomes shared infrastructure.

GOVERNANCE OUTCOME

Faster production with clearer accountability—not automation without oversight.

The framework defines explicit controls and reusable metrics

100%

Supported gate

Pass condition for every atomic claim

3

Claim statuses

Supported · misleading · missing data

12

Weeks of context

Rolling Beyond the Noise summaries

1

Audit trail

Evidence, validation, refinement, review

OPERATING CONTRACT

The release standard is explicit: supported claims, defined statuses, governed context, and retained evidence.

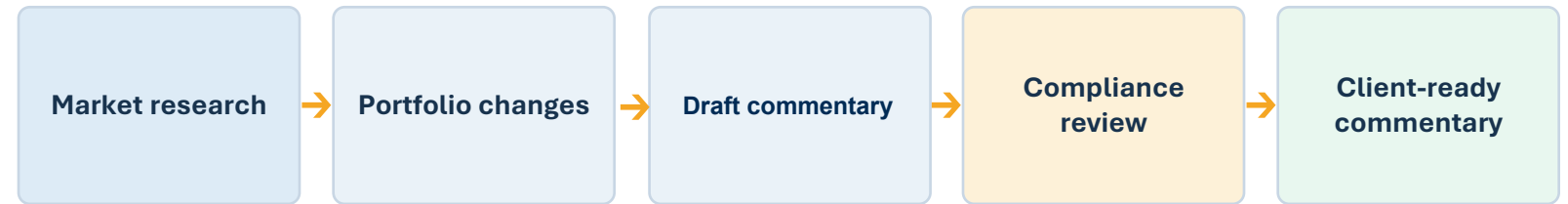
THE CONCRETE USE CASE RETURN

Automation saved 30–50 hours each quarter

10+
hours/week

spent synthesizing markets, portfolio changes, and commentary

Manual workflow



AI draft → human review

The Return:

Automate the generation of timely updates for clients while keeping the commentary grounded, consistent, and reviewable.

Estimated impact: roughly **30–50** hours per quarter saved by removing repetitive drafting and review cycles.

PART THREE — THE FUTURE

Lessons Learned and The Future

Turn one use case into reusable enterprise patterns

03

REUSABLE PATTERNS

Three design choices make reliable agentic AI repeatable

The framework applies anywhere high-stakes narrative must be grounded in enterprise evidence.

EVALUATE PATTERN

Atomic Claim Testing

Decompose statements and trace each one to an exact source before accepting the narrative.

PRECISION • AUDITABILITY

CORRECT PATTERN

Targeted Refinement Loop

Use structured validation failures to repair only unsupported content, then retest.

EFFICIENCY • CONTROL

SCALE PATTERN

Stable Reliability Contract

Separate the verification standard from any single model, provider, or use case.

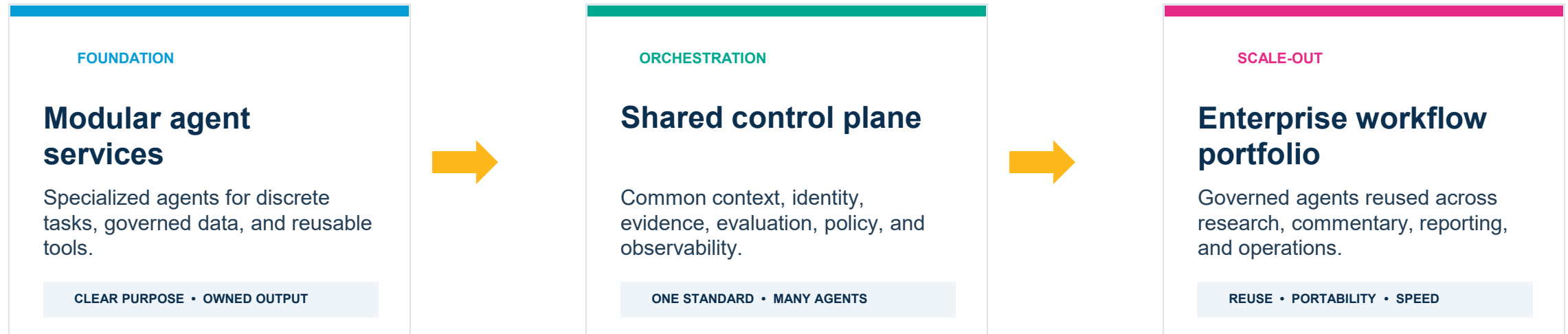
PORTABILITY • RESILIENCE

REUSABLE OPERATING MODEL

Evaluate claims, correct only failures, and keep the reliability standard independent of the model.

Scale comes from governed agent networks—not one-off copilots

Reusable controls become the shared platform for specialized workflows.



FUTURE STATE

Human judgment shifts from editing every draft to governing policies, exceptions, and outcomes across an agent portfolio.

QUESTIONS & CONVERSATION

Thank you

Reliable agentic AI scales when evidence, controls, and human judgment move together.

LET'S CONNECT

Shuen Mei

Technology • Agentic AI • Enterprise Controls

Continue the conversation on LinkedIn

SCAN TO CONNECT

Keep the conversation going

Contact me after the session.
Shuen.Mei@harborcpital.com

FINAL TAKEAWAY

Start with one governed workflow—then reuse the controls to scale the agent portfolio.

Disclaimer

Important Information

Utilizing Artificial Intelligence (AI) involves risks. AI systems, relying on complex algorithms and training data, may produce inaccurate results, introducing algorithmic uncertainty and potential biases. Lack of human oversight, security vulnerabilities, and the evolving regulatory landscape contribute to operational and legal risks. Ethical considerations, such as societal impact and fairness, further underscore the need for careful evaluation. Users should conduct thorough due diligence, monitor evolving risks, and seek professional advice to navigate the dynamic landscape of AI responsibly. Financial professionals should consult their firm's AI policy and compliance department for firm-specific requirements.

This material is for educational purposes only and is not investment advice or a recommendation by Harbor Capital Advisors.

This information has been provided by Harbor Capital Advisors in August 2026 for informational purposes only. It does not constitute or form part of any offer to issue or sell, or any solicitation of any offer to subscribe or to purchase, shares, units or other interests in investments that may be referred to herein and must not be construed as investment or financial product advice. Harbor nor AI in Finance has not considered any financial situation, objective or needs in providing the relevant information.

Investing entails risks and there can be no assurance that any investment will achieve profits or avoid incurring losses.

Harbor has taken all reasonable care to ensure that the information contained in this material is accurate at the time of its distribution, no representation or warranty, express or implied, is made as to the accuracy, reliability or completeness of such information.

The views expressed herein may not be reflective of current opinions, are subject to change without prior notice. This material does not constitute investment advice and should not be viewed as a current or past recommendation or a solicitation of an offer to buy or sell any securities or to adopt any investment strategy.

This material may contain forward-looking information that is not purely historical in nature. Such information may include, among other things, projections and forecasts. There is no guarantee that any of these views will come to pass.

Harbor Capital Advisors, Inc. is not affiliated with AI In Finance

Copyright © (2026) Harbor Capital Advisors, Inc. All Rights Reserved.

5804803