

AI in Research

Part Two: From Collaboration to Autonomy

June 2026



Anton Holst,
Quantitative Researcher

The Natural Next Step

In Part One, we described how our researchers collaborate with AI agents throughout the research workflow. The natural next step is to ask whether these agents can move beyond assistance and conduct the research process end-to-end. This requires the agent to decide on what investment idea to develop, turn it into testable model code, evaluate it and either accept or reject it on good grounds. That requires exploration, implementation, evaluation, and judgment to work together without blurring into one another. This is the challenge our evolving agentic research framework is designed to address.

The objective is not to replace the researcher, but to produce documented candidate strategies that can be inspected, challenged, refined, or discarded by human researchers. The researcher's role therefore evolves from hands-on collaborator to strategic director: setting priorities, supervising the research process and reviewing the backtest results, and determining which results warrant further attention.

This remains a rapidly evolving field. What we describe here reflects where we are now. It is not a fixed blueprint. It is a living system, shaped by daily use, continued experimentation, and rapid advances in the underlying AI models. Our understanding of where AI creates genuine value in systematic investing continues to evolve alongside the technology itself.

Why One Agent Is Not Enough

The simplest version of an agentic research system would rely on a single AI agent handling the entire research process. In practice, however, this quickly proves inadequate. A single agent tasked with exploring ideas, writing code, debugging data issues, running backtests, interpreting results, and deciding whether to continue must simultaneously balance competing objectives: creativity versus precision, speed versus rigor, exploration versus skepticism.

The path from an initial investment hypothesis to a validated strategy involves fundamentally different types of work. Idea generation requires broad exploration, synthesis across information sources, creativity, and an understanding of market mechanics. Implementation requires precision, code quality, and familiarity with internal research frameworks. Evaluation requires statistical discipline, balanced skepticism, and judgement. These are distinct skills that even among human researchers rarely reside in the same person. When too much context and too many objectives are assigned to a single agent, reasoning quality deteriorates and outputs become less reliable.

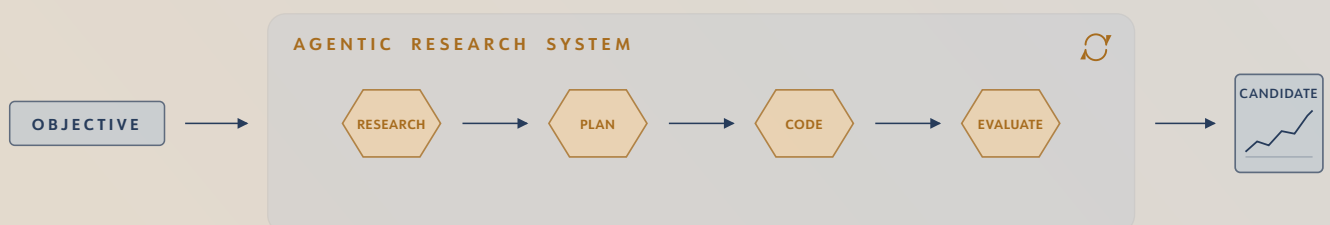
Our current work is focused on building a research architecture around specialized agents with clearly defined responsibilities. This modular design allows each agent to operate with narrower context, clearer objectives, and greater focus while enabling deterministic validation of the output between the different stages of the research pipeline.

The Structure of The Pipeline

Structuring the system as a pipeline of modules with clearly defined responsibilities allows us to optimize each stage for the specific task at hand.

The Fully Agentic Pipeline

An orchestrated team of specialised AI agents takes an investment idea from first sketch to a tested candidate strategy. The researcher's role shifts from hands-on collaborator to strategic director.



The system starts off by launching several research agents with access to our internal knowledge base, historical market data and prior research to identify candidate ideas. An agent might identify a concept discussed in a podcast, trace its theoretical foundation in academic literature, and connect it to recent methodological developments. This ability to follow a line of thinking across different types of material, much as a human researcher would, makes these agents powerful. The results from these parallel threads are then synthesized into a candidate hypothesis.

Once an idea is deemed sufficiently promising, it is passed to a planning agent that formulates a prior, a set of explicit expectations for how a strategy should behave and why. This step matters. It forces the system to articulate its reasoning before any code is written, much as a human researcher would.

The next stage is implementation. A specialized coding agent translates the research plan into model code within our proprietary research framework, using the same infrastructure, data structures, and conventions as our human researchers. Restricting the agent's scope in this way improves reliability and reduces implementation errors. Implementation is an iterative process, the agent writes code, runs tests, identifies issues, and refines the model through repeated feedback cycles. Like human researchers, agentic systems benefit from rapid feedback rather than perfect first attempts. The resulting feedback can trigger another research cycle, allowing the strategy to evolve through successive iterations. Alternatively, the idea may be rejected if the evidence does not support further development.

Once the model behaves as intended, an evaluation agent analyzes the backtest results and critiques the strategy against the original investment case. It assesses robustness, identifies weaknesses, and suggests areas for further investigation. The resulting feedback can trigger another research cycle, allowing the strategy to evolve through successive rounds. Alternatively, the idea may be rejected if the evidence does not support further development.

From The Viewpoint of The Researcher

The result of a full pipeline run is a candidate strategy accompanied by a complete research story: the prior, the implementation decisions, the alternatives explored, the backtest results, and the key diagnostics. Fully agentic research can increase the speed and scale at which ideas are generated and tested, but only if the output remains structured, transparent, and manageable. The researcher can review not only what was built, but why it was built, how it behaved, which alternatives were considered, and where the system itself was uncertain. In this setup, the researcher's role shifts from guiding each decision to setting objectives, evaluating evidence, and deciding which results deserve further attention.

Expanding the Capabilities of the System

As the system is used, successful runs produce examples that can guide future work. Each completed model adds to a growing library of reference implementations, showing how investment ideas can be translated into models within our research framework.

The system's capabilities are also expanded through ongoing development. Progress in agentic research is driven not only by advances in the underlying AI models, but also by the skills, templates, conventions, and validation mechanisms built around them.

As more research tasks become automated in this new paradigm, the role of researchers shifts toward higher-value activities. Engineering the environment in which the agents operate and embedding more of our research process and institutional knowledge into it, creates leverage across all agentic research runs, amplifying the productivity of both researchers and agents.

Important Information

This document is provided for informational purposes only and describes certain internal technologies, workflows, and research practices related to the firm's use of artificial intelligence. Such practices may evolve over time, and no assurance is given as to the accuracy, completeness, or future applicability of the information. Nothing in this document should be construed as investment advice, an offer to sell, or a solicitation to purchase any security, financial instrument or service. The views expressed are those of the author and do not necessarily represent the views of Lynx Asset Management AB. There is no guarantee that any approach described will result in the achievement of any investment edge or alpha, and systematic investment strategies involve inherent risks.