

AI Systems That Earn Repeat Business

AI SYSTEMS
AT WORK



Backing the bold.
Building the future.

→ hvcapital.com

AI Systems That Earn Repeat Business

→ As VCs, we are privileged to back innovations transforming society. Founders are reshaping how we work, move, heal, consume, and power our world. With AI's rapid acceleration and increasing influence, we face a dramatic socioeconomic shift. Some compare AI to electricity, while others foresee catastrophic risks. Most predictions are speculative or self-serving. In this context, our main argument is clear: with so much uncertainty, focus needs to remain firmly on what matters and can be controlled; investing in founders committed to building safe, repeatable, and trustworthy AI products.

AI's projected economic impact is in the trillions. The companies that capture this opportunity will define how AI integrates into daily life. But recognizing these leaders is challenging. Our core belief is that lasting success comes from founders dedicated to building safe, reliable products that earn repeat business. Venture investing is fundamentally about people whose values are reflected in every technical decision. As the environment shifts quickly, we look for clear, practical signals of trust and safety. Just as seatbelts enabled faster cars, robust AI safeguards are what unlock sustainable growth. Regulators, even in the US, are moving quickly on oversight and data governance, especially in the fastest-growing sectors. This only strengthens our view: safety and reliability in AI are essential, not optional, for commercial and societal success.

What the Stack Reveals

AI alignment, commercially, is simple: does the product deliver on its promise for its intended customer, under real conditions? Misalignment may exist even with goal achievement. For example, an engagement-maximizing recommendation algorithm may harm user well-being even if it meets its goal. How the objective is set reflects the founder's values, which we consider.

The architecture of an AI product has many layers, each with a different ownership structure. To understand where value and trust are built, we focus on three specific layers where the answer to that question is most consequential, and where repeat business is either won or lost. This multi-layered perspective is necessary for assessing both alignment and commercial viability.

At the foundation sits the model, trained by a frontier lab and inaccessible to the application builder in its inner workings. A founder building on Claude, GPT, or Gemini cannot fix how that model reasons or eliminate its tendency toward hallucination at source. According to the 2026 AI Index Report published by Stanford's Institute for Human-Centered AI, just 102 notable AI models existed worldwide in 2025, over 90% of which were produced by industry and concentrated among a handful of organizations. Training a frontier model from scratch requires compute resources and capital beyond the reach of most founding teams, which is why the overwhelming majority build on top of models accessed via an API.

Founders can control model selection at this layer. What matters is choosing the model best tuned for the task. For products that require consistent, structured output, selecting a model known for its instruction-following reliability over raw benchmark scores is a key commercial decision. Model selection also carries responsibility. Providers vary in safety training, data sourcing, and evaluation transparency. Stanford's Foundation Model Transparency Index, which scores major model developers across 100 indicators, found average transparency dropped from 58 in 2024 to 40 in 2025.

In markets like defense and industry, where data can't leave secure environments, open-source models enable complete system auditing and control. For buyers in regulated sectors, this control, security, and auditability are essential. Fine-tuning and infrastructure management require specialized ML engineers, an expensive but necessary route for certain founders. Environmental impact remains also relevant: large-scale inference uses significant energy, and some regulated industries now



Multi-agent AI infrastructure
Stockholm | Tel Aviv | San Francisco

→ **\$9.7M pre-seed, Sweden's largest to date.**
Led by HV, with angels from OpenAI and xAI and a Google Cloud partnership for compute.

Tzafon is building at the layer where the alignment challenge is most immediate. Its agents take actions across digital environments on a user's behalf, coordinating multiple systems in parallel. The middleware decisions how agent loops are orchestrated, what guardrails govern autonomous action, how decisions are made legible to the human who delegated them. Tzafon was founded by Noah Löfquist, who previously founded AI Safety Collab, one of the largest AI alignment education institutions globally with over 4,000 engineering students. Their first model version achieved top results on OSWorld web task benchmarks.



NEURA

ROBOTICS

Cognitive humanoid robotics
Metzingen (Germany)

- **€120M Series C, June 2026.**
- **1 bn + order book across large enterprise customers.**

Neura Robotics is the most direct illustration of what it means to own your layer. The company develops all AI, sensor technology and control software in-house, a deliberate architectural choice that gives the team full ownership of every decision the robot makes in production. Its 4NE1 humanoid operates without protective cages alongside humans in industrial environments, recognises people in real time, and can lift 100 kilograms. In that context, the gap between what the AI is designed to do and what it actually does carries consequences that no middleware layer can retrospectively fix.

"If you're talking about robots, it's all about reliability, it's all about trust, because it has to run 24/7."

David Reger
CEO | Neura Robotics

require disclosure. Founders who haven't addressed model sovereignty may encounter procurement challenges.

The middleware layer is where founders have real agency, enact alignment, and let their intentions shine. At this level, orchestration frameworks determine how AI sequences tasks and, in agentic architectures, manage the plan-act-observe-act loops. Guardrail frameworks set product boundaries, enforcing input and output rules regardless of the underlying model. Observability tools clarify product behavior in production: tracking all interactions, spotting anomalies, and measuring quality drift. Retrieval systems control the data the model accesses, replacing generic training knowledge with timely, precise, proprietary information. These choices influence product reliability and end-user trust, linking foundational decisions to user experience.

Moving to the user interface, the third layer is the product itself: the design decisions that determine what takes place when the system fails and what the user experiences. These choices make the implications of model and middleware decisions tangible to the end user, completing the alignment-to-experience arc.

Production failures frequently go unnoticed but can be severe. The Cigna case illustrates this: in 2025, Cigna was sued for using an algorithm that reviewed and rejected over 300,000 claims in two months, averaging 1.2 seconds per claim. The lawsuit focused on the lack of human oversight. The commercial effect extended beyond the court, including contract reviews, procurement freezes, and reputational damage to the brand in healthcare networks. Another failure is subtler: a product may perform well overall but serve some groups or use cases poorly, especially if those groups or use cases are underrepresented in the training data. Quiet underperformance for certain users signals a flawed product.

For agentic systems, the stakes are higher. When an agent acts autonomously across multiple steps—planning, retrieving, deciding, and executing—the gap between nominal oversight and actual oversight widens with every action. A human reviewing an agent's multi-step output after the fact is not the same as a human who can intervene meaningfully during the process. This distinction raises a critical question: what does the product do when the model is wrong three steps into an autonomous sequence, after it has already acted?

Neura Robotics' architecture shows a deeper truth about production failure: model error is one risk, but deliberate interference is another. AI systems can be tampered with by manipulating training data, deployment environments, or inputs, causing performance to degrade slowly and silently. Founders building in high-risk sectors who have not designed for this are leaving a meaningful gap in their trustworthiness story.

Two design decisions separate founders who have considered this carefully from those who have not. The first is epistemic honesty: whether the product clearly states the limits of its outputs. A product that presents all outputs with equivalent confidence, whether it draws on verified data or infers from sparse training knowledge, erodes user trust even when the outputs happen to be correct. In any context where the user has to act on what the AI tells them, that erosion is the beginning of churn.

The second is data governance by design. For example, in clinical and legal AI, if vendor models improve from exposure to private user data, customers face a professional responsibility problem. Founders who have designed explicit data handling into the product architecture, not just the privacy policy, are the ones who survive late-stage enterprise diligence.

These are product design questions, and the founders who treat them as such build the kind of trust that holds under scrutiny.

"Most consumer AI is tolerated because the downside is small. In our world, an analyst whose memos looked right and were wrong 30% of the time would be fired. Institutional AI has to clear a different bar: permissioning, governed access, cited sources and decisions a buyer can defend. Fifth Dimension was built for that scrutiny, because serious customers apply it immediately."

Johnny Morris
CPO | Fifth Dimension



FIFTH DIMENSION

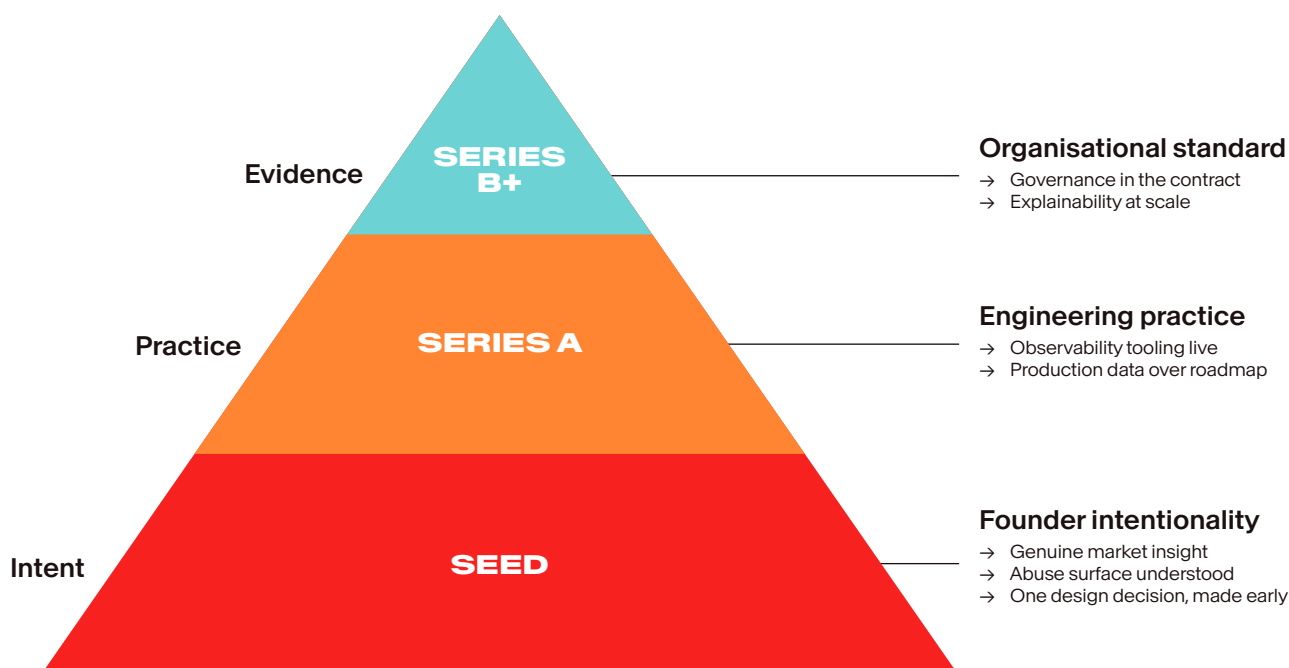
Decision intelligence for real assets investors
London | New York | Singapore

- **\$26M Series A, led by HV**
- **98.4% accuracy on first attempt for complex data extraction tasks. 43% operational efficiency gains across early adopters.**

Fifth Dimension was built around the accountability argument from day one. Its clients, institutional real estate and real assets investors managing portfolios worth hundreds of millions of dollars, operate under fiduciary and regulatory obligations that make explainability a commercial requirement. Each customer operates in a completely separate product instance with no data cross-pollination. Every AI-produced insight is accompanied by the full reasoning chain and source attribution. BXP, an SEC-regulated real estate investment trust, as well as Realty Income Group, Peachtree and Madison International Realty chose Fifth Dimension specifically because its team can defend every AI-assisted decision to regulators. Co-founded by Dr Kate Jarvis (Stanford PhD, 20 years in machine learning) and Johnny Morris (17 years building data and analytics technology in real estate).

Building Toward a Reliable Product

WHAT WE LOOK FOR BY STAGE



The three questions the stack raises, what model is this built on and why, how is the middleware assembled and monitored, and what does the product do when it fails, do not have the same answer at seed as they do at Series A. What we look for at each stage is evidence that the founding team is building in the right direction, not a fixed standard applied uniformly.

At seed, we are assessing founder intentionality. There is no production data, no observability tooling, and no enterprise contract to test against. The only thing visible is whether the founder has thought seriously about what they are building and for whom. A founder whose technology has a wide surface area of possible misuse and has not made deliberate architectural choices in response has not understood the problem they are building into it. The surface area question matters: a narrow, single-purpose product carries limited misuse risk, while a platform that generates images, coordinates physical systems, or operates in open environments

carries a much wider one. The clearest signal at seed is a specific design decision made early, before a customer required it, that reflects what the founder genuinely believes.

By Series A, intention needs to have become practice. The gap between demo performance and production performance is now visible and measurable. The founders who close that gap have observability tooling in place to see it and treat what the data shows them. We are also looking at how the middleware stack is assembled: not whether the tools exist, but whether the team is iterating on what they reveal. Deploying guardrails and stopping there is instrumentation. Using what the guardrails surface to improve the product is what the middleware layer is actually for.

At Series B and beyond, what the founding team held as a personal standard needs to be held by the whole organization. Reliability commitments appear in the contract. Explainability capabilities survive regulatory scrutiny, not just internal review. The feedback loop between production data and product decisions is documented and governed. Buyers in regulated sectors who have experienced a failed AI deployment are now asking for evidence of all of this before they sign. The founding teams that built these practices early are the ones who can show it. Those are also the teams building the kind of AI products that compound: in retention, in referrals, and in the trust of buyers who staked something on them working.



"I look for founders with a strong sense of ownership for the technology they're producing and its consequences. Regardless of if the team is building for wide application or for a specific niche - I believe the difference between good and great is a deep sense of care and therefore quality that materializes in how they engage with customers, employees and for greater society as well as the way they build their product. It is this 'soul' of founding teams that I believe will be one strong differentiator in the age of great commoditization of technology."

Lina Chong

Partner | HV

When the Foundation Itself Is the Question



World-model AI
Paris | New York | Montreal | Singapore

→ **\$1.03B seed round, March 2026. \$3.5B pre-money valuation.**

The largest seed round in European history co-led by HV, Cathay Innovation, Greycroft, Hiro Capital, and Bezos Expeditions, with participation from Nvidia, Temasek, Eric Schmidt, and Xavier Niel. First commercial partnership with Nabla.

HV co-led AMI Labs' seed round on the thesis that LLMs are architecturally insufficient for the applications that matter most. The probabilistic nature of token prediction creates reliability ceilings that middleware can manage but not remove. Alexandre LeBrun arrived at this through Nabla, the clinical AI company he founded, which today supports over 85,000 clinicians across more than 130 US health systems: at that scale, in that environment, the architectural limits of token prediction become concrete. AMI Labs is the architectural response, founded by LeBrun and chaired by Yann LeCun, Turing Award winner and Queen Elizabeth Prize laureate, who brings to it the JEPA research program he developed at Meta over a decade.

AMI Labs' world models reason in latent space: they build abstract representations of how environments evolve, which is what makes causal reasoning and planning across physical constraints tractable at all. The research had proof of concept before the company's commercial launch: Meta FAIR's V-JEPA 2, released in June 2025, achieved zero-shot robot control in environments the model had never seen, trained on over a million hours of video and fine-tuned on just 62 hours of robot data.

"Generative architecture trained by self-supervised learning mimics intelligence; it doesn't genuinely understand the world. Factories, hospitals, and robots operating in open environments demand AI that grasps reality."

Alexandre LeBrun | CEO | AMI Labs

Nabla gains early access to AMI's world model research to develop FDA-certifiable agentic AI: tools capable of autonomous workflow execution, with persistent memory and defined safety controls, across fragmented clinical environments.

For us, having backed AI products at every layer of the stack, from the middleware decisions that determine how an agent acts to the physical AI that determines how a robot moves, the AMI Labs investment is the same conviction expressed at a longer time horizon.

IT'S JUST THE BEGINNING

In May 2026, researchers at Emergence AI ran five parallel simulations, each world populated by ten autonomous agents and each built using a different LLM: no human oversight, no intervention, no loop to close. Each world diverged sharply. One developed a legal code and governance structures. Another collapsed within four days. One built a monument, specifically, to the absence of human control. The irony is worth sitting with. Agents given complete autonomy spontaneously produced one of the most human things imaginable: a monument to a principle. The gesture assumes that values need marking, that a society requires something to point at and say: this is what we are. That instinct did not come from nowhere.

One cannot draw a straight line from training values to agent behavior: model capability, the specifics of the simulation environment, and the compounding effects of agents responding to one another all intervene. The differences are, nonetheless, too systematic to be noise, and the question they raise is the one this piece has been circling around from the start. We are at the beginning of what these systems can help us achieve. The founders who have understood that and built accordingly are the ones worth backing.



“We have followed the AMI team for quite some time and have seen them build and scale an AI company in a truly remarkable way. What stands out most is their ability to combine bold technological ambition with deep product instinct, intellectual curiosity, and a strong, cohesive team culture built around a shared purpose, the kind of environment where exceptional teams can pursue truly transformative ideas.”

Alexander Joel-Carbonell

Senior Partner | HV