

THE AI-READY BOARD

by
Rene Schmidli

Rene Schmidli



Rene Schmidli, a financial expert with over 30 years of investment experience, covered the technology sector as a financial analyst as early as 1993 and accompanied a tech IPO in Switzerland in 2000. Today he works independently in various board and advisory roles. In December 2024, motivated by the rapid development in AI, he wrote his CAS thesis on Generative AI under Mateusz Dolata at the University of Zurich to provide a practice-oriented examination of AI strategy and governance issues in boards of directors.

Yariv Adan



Yariv Adan is a leading AI expert, builder, and investor with 17 years at Google as Senior Director of Product Management, where he co-founded Google Assistant and Google Lens and led Google Cloud Conversational AI and Applied Generative AI. He also served as site lead for Google Zurich. Today he is Founder and Managing Partner of Ellipsis Venture Capital, investing in early-stage AI deep tech companies across Europe and beyond.

A two People & AI Venture Capital Company - Doable?

*Turn AI into the “Operating System” to remove the human bottleneck
so the critical 20% of judgment delivers the full 100% impact!*

Rene Schmidli: Dear Yariv. Let's start with the provocation. You claim two people can cover what a 50-person venture capital firm does. Is that marketing, or is it literally true?

Yariv Adan: A 50-person firm isn't 50 people making decisions. It's maybe four or five people making decisions, and 45 doing the processing that leads to those decisions. Scanning emails, enriching company data, writing first-pass memos, scheduling calls, tracking action items. We automated all of that. So yes, two people can do the judgment work of a much larger team, because we stopped treating processing as if it requires human intelligence. It doesn't. What requires human intelligence is reading a room, sensing ambition, knowing when a founder is lying. We saved ourselves for that. And the infrastructure cost of all of this? Under 500 dollars a month.

Rene Schmidli: Walk me through what happens when a company name lands on your radar. What does the machine actually do before a human ever looks at it?

Yariv Adan: The pipeline runs in six stages without us touching it. First, web enrichment: the agent pulls the website, the founders' LinkedIn profiles, any press coverage. Second, a basic fit check: stage, geography, does it align with our thesis at all? Third, outreach. If it passes, the agent finds contact details and requests the pitch deck. Fourth, domain-specific assessment: is this a space we believe in, is the team credible at first read? Fifth, a confidence check. If the signal is strong, we run a full pre-meeting analysis, thesis, competitive

landscape, team scorecard, risk map. If it's weak, it gets flagged for human review or dropped. Only at step six does a meeting get scheduled. And even then, the founder receives a set of tailored questions before we sit down.

What we realized over time is that the system is not just a linear pipeline. It's a combination of a funnel and a loop. The funnel answers one core question: given any initial signal, a company name, URL, founder profile, or pitch deck: should we even spend time meeting this company?

But once that threshold is crossed, the process changes completely. We enter an iterative loop focused on building conviction and systematically closing information gaps. The system continuously identifies what we still do not know, then orchestrates the next-best action: follow-up questions, founder meetings, research, reference calls, additional documents, or data room analysis.

By the time I'm in the room, I've already read a structured brief on every meaningful dimension of that company. Every interaction across this entire process — emails, notes, calls, data room entries — compounds into a living document. The file grows smarter with each touchpoint.

And because the process is modeled as both a funnel and a loop, it becomes very natural to inject humans at every critical decision point. Human judgment can step in wherever conviction is weak, signals conflict, risks escalate, or strategic nuance matters more than pattern recognition. The AI handles orchestration and synthesis, while humans remain embedded in the highest-leverage moments of the process.

Rene Schmidli: The piece I find most striking is your simulated investment committee. AI models debating against each other. How does that actually improve the quality of decisions?

Yariv Adan: The problem with both humans and AI systems is that they tend to commit to a narrative early and then selectively reinforce it with supporting evidence. A strong pitch deck can create a halo effect that inflates every downstream assessment.

To counter this, we separate two mechanisms.

First, we run an internal investment committee simulation. A single model orchestrates multiple agent-like perspectives, each representing a different analytical mindset. One may focus on product and technology, another on market timing, another on execution risk or founder quality. They debate the opportunity from different angles before converging on a recommendation. It is less about theatrical "AI personalities" and more about stress-testing assumptions through structured internal disagreement.

Second, and more importantly from a reliability perspective, we use multiple independent models and skills redundantly across critical tasks. Market analysis, traction verification, technical defensibility, and competitive mapping are often evaluated separately by different systems that do not initially see each other's conclusions. This is designed to identify inconsistencies, weak evidence chains, and hallucinations before a human ever reviews the output.

The value is not that AI magically produces the correct answer. The value is that disagreement becomes visible and traceable. When multiple systems independently converge on the same conclusion, confidence increases. When they diverge, the divergence itself often reveals the real investment question.

Rene Schmidli: You use the word "conviction" a lot. How does an AI-native system build conviction if it's fundamentally a pattern-matching machine?

Yariv Adan: Conviction isn't a feeling, it's a state of reduced uncertainty. The machine helps us get there faster by doing what it does well: aggregating evidence, flagging inconsistencies, running the same checks hundreds of times without getting tired or biased by a founder's charisma. Venture itself is fundamentally pattern recognition: markets, founders, timing, distribution, product velocity, competitive dynamics. The

machine simply allows us to run that process with more structure, consistency, and memory than any human team could do manually.

We model this explicitly: every company gets a calibrated score, not a binary yes or no. Most seed-stage companies land between 2.5 and 3.5. A score of 4 demands verified, exceptional evidence. A 5 means best in class for that stage. And certain hard gates override the score entirely: no defensible moat, automatic pass. Team misses two of four outlier criteria, automatic pass. The numbers feel precise, but they're really just structured judgment.

Every interaction compounds. An email thread updates the scorecard. A reference call updates it again. Meeting notes, data room access, expert calls; all of it feeds back into a living document. What the machine cannot do is have the conversation that changes your mind about a person. That's where I come in. But I arrive better prepared, with every verified fact already in front of me, so I can spend the meeting asking the questions that matter, not the ones I should have researched in advance.

Rene Schmidli: This implies you'll miss some great companies, ones that score poorly but would have been exceptions. Doesn't that worry you?

Yariv Adan: It does, and I won't pretend otherwise. But here's the honest framing: every VC misses great companies. The boutique fund misses them because it can't process enough deal flow. The large fund misses them because it's too slow and the relationship becomes transactional. We miss some because our screening has hard gates that occasionally reject anomalies.

The question isn't whether you miss deals. You always do. The question is whether your process systematically improves your probability of seeing the right ones and having genuine depth when you do. The goal is not to find every great company. It's what quant funds figured out long ago: systematically better probabilities across many decisions beats chasing perfect hits on a few.

The important point is that the screening threshold itself is a strategic choice, not a system limitation. We can tune for more false positives or more false negatives depending on fund strategy, team capacity, and how much manual review we want in the process. Anything that doesn't fit the model cleanly goes to human review, not the bin. That's an important design choice.

Rene Schmidli: Let's talk about hallucinations. For a board or executive team considering AI in their decision-making, this is the central fear. How do you contain it?

Yariv Adan: You can't stop a model from being wrong about a fact. What you can do is stop it from sounding certain, stop it from contradicting itself, and stop it from reaching a conclusion its own evidence doesn't support. We built several layers. Every claim is tagged: VERIFIED, CLAIMED, or SUSPECT. The model is explicitly instructed that "not found" is a correct and valued answer. We trained it to flag gaps, not fill them with plausible-sounding guesses.

We also run plausibility benchmarks: if a seed-stage company reports an LTV/CAC ratio above 15x, that's a flag, not a feature. The system is trained to treat implausible precision as a red flag, not a green light. Every skill outputs structured JSON alongside its prose summary, so downstream systems consume structured data, not narrative. And we run cross-audits: independent skills analyze the same company, then we compare key facts. Discrepancies are mandatory log entries. Perfect agreement also gets flagged, because if every model always agrees on everything, something is wrong with the diversity of the analysis.

It's not a solved problem. But it's a managed one. And ultimately, the real question is not whether errors exist, but whether the system performs better, faster, and more consistently than the alternatives available to you. That's why we continuously run evaluations, inspect outputs, and measure failure modes instead of assuming reliability.

Rene Schmidli: You describe the human role as "the 20% that makes 100% of the difference." What specifically does that mean?

Yariv Adan: Four things. First: sourcing where there is no data. The best companies at the earliest stages don't exist in databases yet. They're at university dinners, at off-the-record conversations at conferences, in WhatsApp groups I'm part of. No agent monitors those. Second: reading people, not profiles. Six hours of conversation, including dinner and drinks, tells you things no LinkedIn page or pitch deck reveals. Does this person know what they don't know? Can they recruit? Will they be honest with you when things go wrong? Third: the left-field question. AI is extraordinary at structured analysis. It's still poor at the genuinely unexpected question that reframes the whole conversation. Fourth: winning the founder's trust. In a competitive deal, the best founder chooses their investor partly on conviction and partly on personal chemistry. That's irreducibly human. And there's a fifth dimension that often gets overlooked: we reinvest the time we reclaim. We code alongside founders. We go deep in ways that a processing-heavy team simply can't afford to. The machine doesn't just replace effort, it buys back attention for the things that actually matter.

Rene Schmidli: Do founders know that AI screened them before you met? And does it matter to them?

Yariv Adan: Some do, some don't, and increasingly it doesn't seem to matter, because the experience on their side is better, not worse. They get a faster response. If they're rejected, they get a thoughtful email that explains why, not silence. If they pass, they come into a meeting where I've already read everything and can ask real questions immediately. What founders hate is wasting time with investors who haven't done their homework. We've eliminated that. The AI does the homework. I do the conversation. From their perspective, that's an attentive, well-prepared investor. The fact that a machine helped prepare me is no different from having a great analyst on the team, except more consistent and available at any hour.

Rene Schmidli: Let's speak to the boardroom directly. What are the genuine risks that a governance body should understand before adopting this kind of model?

Yariv Adan: Five come to mind. First: building the wrong thing. Many organizations try to automate a legacy human-centric process instead of asking what an optimal AI-native process should actually look like. If the objective is poorly defined, the system can become highly efficient at producing the wrong outcome. Second: targeting the wrong benchmark. You're not trying to build a perfect system, you're trying to build one that performs materially better than the existing alternative across the metrics that matter. Third: systematic errors at scale. Most of the weaknesses people criticize in AI — hallucinations, bias, randomness, weak reasoning — already exist in human organizations too. The difference is that AI operates faster and at far greater scale, so small flaws compound rapidly. The system amplifies whatever the design gets wrong. This is why auditability isn't optional. Every decision must trace back to inputs, logs, confidence levels, and discrepancies. Fourth: losing control over quality and cost. AI systems actually create better audit trails than humans, but only if organizations actively review, evaluate, and act on them. Governance without continuous evaluation is just theater. And fifth: never launching anything. Many companies get trapped trying to boil the ocean. The better approach is to start small, measure aggressively, iterate quickly, and improve the system over time instead of waiting for perfection.

Rene Schmidli: Is this model transferable? Can a corporate strategy team, an M&A function, or a board-level decision process be redesigned the same way?

Yariv Adan: Yes, with one condition. You have to be honest about where humans are actually irreplaceable in your process, versus where they've simply always been involved. In most organizations, people do processing tasks; aggregating information, writing summaries, tracking action items, and call it strategic work because it leads to strategic decisions. That processing is automatable. The design question is: what are the three to five moments in your decision cycle where human judgment, trust, or relationship is genuinely the differentiating variable? Build the machine around protecting and amplifying those moments. Everything else is a candidate for automation. What we built for VC is transferable to M&A screening, hiring pipelines,

competitive intelligence, regulatory monitoring. Any domain where you're making repeated decisions under uncertainty with incomplete information. Which, frankly, is most of management.

Rene Schmidli: Last question. If this model works well, doesn't everyone copy it — and the edge disappears?

Yariv Adan: The tools will commoditize, yes. The system design won't. Because the real moat isn't the technology, it's the judgment that goes into the design. What questions do you ask? What do you classify as a hard gate and what's a soft signal? What does your investment thesis actually say, in sufficient precision that a machine can operationalize it? Most organizations don't know the answers to those questions. The process of building this system forced us to make our thinking explicit: our theses, our team scorecards, our investment committee memos, in a way that years of regular VC work hadn't required. We call it a brain dump, but it's really an externalization of judgment. That's the asset that compounds. And it's the one thing that can't be copied without doing the same hard thinking. And then there's the relationship layer — the network, the reputation, the founder who refers the next founder. You can copy my architecture. You can't copy that.