

Man versus machine: comparing human analysts and agentic AI in early-stage venture evaluation

1. The question

Venture capital firms screen thousands of startups a year to find the handful worth investing in, a process that is slow, expensive, and heavily dependent on individual analyst judgment. Conducted as a master's thesis at Rotterdam School of Management, this study tests a practical question: how does Kuanta compare with human analyst judgment on early-stage ventures? It is examined along two lines: how closely the two channels agree, and how well each predicts real-world outcomes.

2. The approach

The study draws on the participants of the Philips Innovation Awards (PHIA), the Netherlands' largest entrepreneurship competition, across three cohorts (2022–2024). A representative sample of 150 startups was taken from those already assessed by an independent jury of human analysts, who scored each business plan on four dimensions: Innovation, Team, Market, and Financials. The same documents were evaluated independently by Kuanta, whose specialised agents score each venture across seven dimensions and a composite, with no access to the analysts' scores. Funding and survival outcomes were traced through a scraping exercise of public sources and verified by hand. A set of machine learning models were then trained on each channel's scores across 18 configurations per outcome measure, cross-validated on ventures neither channel had seen, and compared on predictive power.

3. What the results show

Kuanta's scores align moderately with the analyst panel, most strongly on team quality and least on harder-to-verify dimensions such as market opportunity and financial

projections. On overall predictive accuracy the two channels are comparable. However, closer examination of the confusion matrix reveals that the difference appears where it matters most, in which ventures actually go on to succeed.

Kuanta is better at selecting winners. Across every model tested, Kuanta identifies more of the startups that go on to raise funding than the analyst panel does, consistently achieving higher recall. Translated through an expected value framework that accounts for VC's power-law payoff structure, where missing a future winner is far costlier than one extra review, this recall advantage converts into higher expected value per startup screened. Two caveats sit alongside these results: several comparisons rest on small subgroups, and some outcome data may carry look-ahead bias from Kuanta's enrichment process. The clearest structure sits in where the two channels agree and disagree, shown in Figure 1.

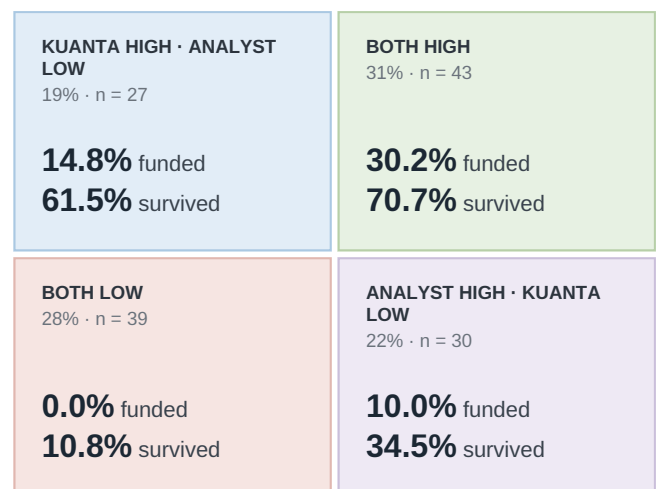


Figure 1. Consensus is the strongest signal. Funding and survival rates by agreement quadrant; rows show Kuanta's rating (high on top), columns the analyst panel's (high on the right). Where the channels disagree, ventures favoured by Kuanta show higher funding and survival rates than those favoured by analysts, though on modest subgroup sizes.

Agreement carries the strongest signal of all. When Kuanta and the analyst panel both rate a startup highly, the funding rate rises to 30.2% and survival to 70.7%, well above either channel on its own. When both rate it weakly, funding falls to 0.0% and survival to 10.8%. The consensus of AI and human judgment predicts far better than either score alone.

The disagreements are equally revealing. The two channels split on roughly four in ten startups, and in those cases the ventures Kuanta favours outperform the ones the panel favours on both outcomes: they survive at 61.5% against 34.5%, and raise funding at 14.8% against 10.0%, a gap best read with some caution given the modest subgroup sizes. A more limited edge appears at the dimension level: Kuanta's models show a marginally significant advantage in predicting survival overall ($p = .046$), though this result should also be interpreted with the same caution.

Each channel also captures winners the other misses. Kuanta enriches the pitch with external data on product and market; the analysts add tacit judgment on people and narrative. Because the two read partly different signals, they carry independent, complementary information rather than one reproducing the other, which is why their agreement is so much stronger than either alone.

4. What it means in practice

The takeaway points toward a dual-screening framework rather than a contest between human and machine. Run both channels in parallel, fast-track the ventures they both rate highly, deprioritise those they both reject, and give the disagreements closer review, since that is where the highest-value opportunities tend to sit. Used this way, Kuanta widens the funnel and sharpens the odds of catching the winners a single reviewer would miss.

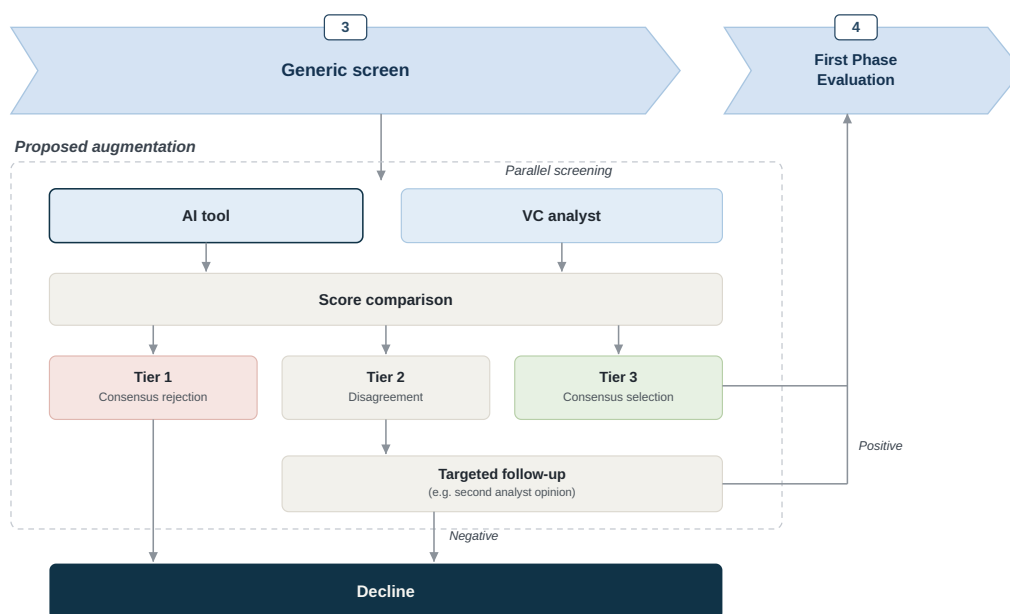


Figure 2. A dual-screening framework. Applied to the venture screening funnel described by Fried and Hisrich (1994). The AI tool and the VC analyst screen in parallel; consensus decisions are handled at once and only disagreements draw a targeted follow-up before a venture advances to first-phase evaluation.

The author has no financial or employment relationship with Kuanta and received no compensation for the research; the design, analysis, and interpretation were conducted independently. Startup names were used with the data provider's permission; no personal data on founders or analysts were processed. As a study of 150 startups from a single competition, the findings read as early evidence for AI-augmented evaluation rather than a final word. Several limitations temper this evidence: subgroup sizes in the disagreement analysis are modest, some outcome data carries a risk of look-ahead bias from Kuanta's enrichment process, and the sample, while representative of PHIA, may not generalise to other competition formats or investment contexts.

Source: L. Meijer (2026). *Man versus machine: comparing human analysts and agentic AI in early-stage venture evaluation*. Master's thesis, Rotterdam School of Management, Erasmus University.