

Research Statement

Andrew Funck

PhD Candidate, Economics and Decision Sciences, HEC Paris

andrew.funck@hec.edu | andrewfunck.com

I am an economist studying entrepreneurship, education, and innovation, and how artificial intelligence is re-shaping each of these domains. Methodologically, I combine formal modeling with econometrics, machine learning, and natural language processing (NLP), including large language models (LLMs), to analyze rich micro-level observational data, and I organize, implement, and analyze large-scale randomized field experiments to obtain causally identified, policy-relevant evidence.

My current work can be organized into two interrelated themes:

1. Evaluation of startups

The first theme studies the evaluation of young ventures in the context of an accelerator, focusing on how artificial intelligence differs from humans in this task ([Funck, 2026](#)), how human judges aggregate their beliefs into their assessments ([Åstebro et al., 2026](#)), and how evaluation accuracy is shaped by expertise and complexity ([Åstebro et al., 2025](#)).

2. Entrepreneurship training and the formation of educational choices

The second theme turns to large-scale field experiments testing how information and training shape entrepreneurial and educational decisions: a randomized experiment with France Travail, the French public employment agency, testing whether online entrepreneurial training delivered with artificial intelligence tutors affects business creation among job seekers ([Åstebro et al., 2022](#)), and a field experiment testing whether role models addressing gender stereotypes shift enrollment in STEM education among Swedish teenagers ([Altmejd et al., 2024](#)).

I. Evaluation of Startups

I study venture evaluation in the context of the admissions process of a leading French business school's startup accelerator, with which I have an ongoing research partnership. In this process, professional judges score each applicant venture on a fixed battery of business criteria to determine admission.

In my job market paper, *Human vs. AI in Startup Evaluation: Evidence from an Early-Stage Accelerator* ([Funck, 2026](#)), I compare human and artificial intelligence assessments of the same ventures, building on recent evidence that generative AI can serve as a standalone evaluative tool in strategic decision-making ([Doshi et al., 2025](#)). I develop a theoretical framework in which an evaluator's grade responds to information in proportion to the ratio between the precision of its prior knowledge and the cost of processing that information. Prompting a general purpose large language model with the same applications, criteria, and instructions given to human judges, I show that AI grades are systematically more responsive than human grades to the information available in each application. Linking both sets of grades to hand-collected business performance of ventures (survival, employment, capital raised), I then show that AI scores are at least as strongly associated with venture performance as human scores. More recent AI models converge toward human evaluative patterns, with no gain in their association with performance.

In *Decomposing Venture Evaluation: Learning about Belief Formation* ([Åstebro et al., 2026](#)), with Thomas Åstebro and Francesco Giordano, recently accepted for publication in *Strategic Entrepreneurship Journal*, we turn to human evaluation alone and ask how much judges' criteria ratings inform their recommendation for admission, and how informative each is of venture quality, following the literature on evaluation in finance ([Gompers et al., 2020](#)). We show that judges' beliefs, as recorded in their ratings, explain their recommendations well but are only weakly tied to quality, and that several criteria most associated with quality receive

comparatively little weight in recommendations. This limited informativeness is stark: a model built purely from ventures' raw application attributes predicts quality about as well as judges' criteria ratings do.

A companion working paper, *Cognitive Complexity, Expertise and Prediction Accuracy in Venture Selection* (Åstebro et al., 2025), asks whether this weak accuracy comes from judges failing to use good information well, or from the information itself being too weak to support better decisions. Building on the definition of optimal evaluation in the Bayesian entrepreneurship tradition (Agrawal et al., 2025), we construct a theoretical benchmark for the best possible use of the same information, and find that it reproduces judges' decisions 85% of the time, and thus reaches the same limited accuracy in classifying venture quality ($\sim 60\%$). We then study what drives this accuracy: cognitive complexity, measured from the text of the application and from how much judges disagree with one another, lowers it; expertise raises it, but mainly for moderately complex ventures, and mostly because more able judges take on more evaluations rather than because judges improve with practice.

II. Entrepreneurship training and the formation of educational choices

The second block moves from the accelerator to the level of public policy, using large-scale randomized field experiments to study how training, information, and social influence shape entry into entrepreneurial and STEM careers, and how artificial intelligence can help deliver such interventions at scale.

With Thomas Åstebro, Marcos Balmaceda, Bruno Crépon, Mona Mensmann, Naja Pape and Mathis Schulte, I am running *Employment Creation through Psychological Entrepreneurship Training* (Åstebro et al., 2022), a large RCT with the French public employment agency, France Travail. Building on evidence that psychological training can outperform traditional business training (Campos et al., 2017), the trial evaluates whether online training courses in personal initiative and negotiation skills we have designed and produced affects entrepreneurial mindset, negotiation skills, and personal initiative, and in turn business creation, employment, and labor-market outcomes among job seekers. Beyond the training's treatment effect, we test two further margins: whether the narrative framing of invitation emails affects take-up and completion, and whether pairing the training with different AI chatbot tutor variants improves learning outcomes. The full experimental infrastructure, courses, platform, chatbot variants, IRB approval, and randomization protocol, is in place, and all legal agreements with France Travail covering job-seeker recruitment, data access, and data security are signed and in place for the current and coming years. The experiment is currently in its pilot phase, while the experiment at scale is planned for Spring 2027.

With Adam Altmejd, Thomas Åstebro, Mona Mensmann, Ali Mohammadi and Karl Wennberg, I study role models as a lever against the gender gap in STEM (Altmejd et al., 2024), in a pre-registered RCT covering roughly 12,000 Swedish students across about 300 schools. Motivated by social role theory and by evidence that exposure to relatable role models can reshape occupational stereotypes (Breda et al., 2021; Asanov et al., 2026), classes are randomly assigned to a presenter, female role model, male role model, or study counselor, who delivers a one-hour presentation on working in STEM addressing one of three gender-stereotypical beliefs, work goals, the "STEM worker" stereotype, or perceived skill requirements, or to a no-visit control. Students are surveyed before and after the visit to collect their beliefs and preferences about STEM careers, and are then followed in their subsequent educational decisions through administrative records. The study, supported by the Royal Swedish Academy of Sciences and the Kamprad Family Foundation, is in its fourth data-collection wave, rolling out across southern and central Sweden through winter 2026–2027.

Future work

In future work, I plan to expand my research agenda along two lines: the interaction between human and artificial intelligence evaluators, and the broader consequences of AI adoption inside organizations. Two projects speak to the first line. With the partner accelerator, I have designed and am piloting this fall a field experiment with MBA-student reviewers that varies, within reviewer and within venture, how much of the venture information file is shown, how many criteria are scored, and what form of AI support is offered, tracing how human evaluation responds to information quantity, evaluation workload, and AI assistance. I am also discussing with the accelerator a further experiment in which admission would depend randomly on AI or on human evaluations. That design would let me compare the accelerator's treatment effect between AI-evaluated and human-evaluated ventures, identified by a regression discontinuity exploiting the rank-based admission rule of the accelerator, and measure how AI adoption in screening changes the quality and composition of the ventures ultimately selected. Further questions I intend to pursue include what information artificial intelligence draws on when it assesses a venture, opening the black box of LLM evaluation by mapping assessments back to the material supplied to the model, and how the diffusion of AI reshapes human roles inside organizations, in particular the value of holding a person responsible for a decision the machine recommends.

References

- Agrawal, A., Camuffo, A., Gambardella, A., Gans, J., Scott, E. L., and Stern, S. (2025). *Bayesian Entrepreneurship*. MIT Press, Cambridge, MA.
- Altmejd, A., Åstebro, T., Funck, A., Mensmann, M., Mohammadi, A., and Wennberg, K. (2024). The impact of role models on gender stereotypical beliefs about educational choices. Pre-registered RCT, AEA RCT Registry AEARCTR-0010411.
- Asanov, I., Åstebro, T., Buenstorf, G., Crépon, B., Flores-Taïpe, F. P., McKenzie, D., Mensmann, M., and Schulte, M. (2026). Remote delivery of stem and entrepreneurship role models at scale changes college major choice in ecuador. *Nature Human Behaviour*, pages 1–11.
- Åstebro, T., Balmaceda, M., Crépon, B., Funck, A., Mensmann, M., Pape, N., and Schulte, M. (2022). Employment creation through psychological entrepreneurship training. RCT in progress with France Travail, HEC Paris and GENES.
- Åstebro, T., Funck, A., and Giordano, F. (2025). Cognitive complexity, expertise and prediction accuracy in venture selection. Working paper, HEC Paris.
- Åstebro, T., Funck, A., and Giordano, F. (2026). Decomposing venture evaluation: Learning about belief formation. *Strategic Entrepreneurship Journal*, pages 1–36.
- Breda, T., Grenet, J., Monnet, M., and Van Effenterre, C. (2021). Do female role models reduce the gender gap in science? evidence from french high schools. IZA Discussion Paper 14833, IZA.
- Campos, F., Frese, M., Goldstein, M., Iacovone, L., Johnson, H. C., McKenzie, D., and Mensmann, M. (2017). Teaching personal initiative beats traditional training in boosting small business in west africa. *Science*, 357(6357):1287–1290.
- Doshi, A. R., Bell, J. J., Mirzayev, E., and Vanneste, B. S. (2025). Generative artificial intelligence and evaluating strategic decisions. *Strategic Management Journal*, 46(3):583–610.
- Funck, A. (2026). Human vs. AI in startup evaluation: Evidence from an early-stage accelerator. Working paper, HEC Paris.
- Gompers, P. A., Gornall, W., Kaplan, S. N., and Strebulaev, I. A. (2020). How do venture capitalists make decisions? *Journal of Financial Economics*, 135(1):169–190.