



Between Augmentation and Accountability:
How Investors Experience the Use of Artificial Intelligence in
the Investment Process and Maintain Human Decision
Authority

Anna Bergesen

Dissertation written under the supervision of Professor Aurore Haas

Dissertation submitted in partial fulfilment of requirements for the MSc in
Management with specialization in Strategy, Entrepreneurship and Impact, at the
Universidade Católica Portuguesa, 26.05.2026.

Abstract English

This thesis examines how artificial intelligence is used in investment decisions in investment firms, and how this integration is affecting human decision authority. The study adopts a qualitative design and draws on eight semi structured interviews with senior investors across family offices, private equity firms and investment advisories in Norway and United Kingdom. The empirical material was analyzed using Braun and Clarke (2006, 2019) thematic reflexive analysis. The findings of the study identify three dimensions of how the use of AI affect decision-making. The first finding is that AI has become operational infrastructure in the investment process: its functioning as a co-analyst used to cover a broader set of information and expands analytical capacity. The second finding covers the collaboration between humans and AI and how it changes the cognitive dynamic of the decision-making. AI as a tool is used to reduce availability bias, however it simultaneously amplifies confirmation bias. Third, the trust investors have in AI is demonstrated through performance and domain specific.

The thesis contributes with qualitative empirical evidence of AI as investment infrastructure and how it shapes the information environment within the investment context. It identifies how the information environment is affecting cognitive bias dynamic and how cognitive dynamics is shaping human boundaries of trust in AI generated tools. The practical implication is that integration of AI in investments may need to be supplemented with governance mechanisms. This makes the use of human-AI collaboration consistent and responsible rather than being dependent on individual discipline.

Title: Between Augmentation and Accountability: How Investors Experience the Use of Artificial Intelligence in the Investment Process and Maintain Human Decision Authority

Author: Anna Bergesen

Keywords: Decision-making process; Artificial Intelligence (AI); Human-AI collaboration; Trust; AI sycophancy; Behavioral finance

Abstract Portuguese

Esta tese examina como a inteligência artificial é utilizada nas decisões de investimento em empresas de investimento e como essa integração afeta a autoridade humana na tomada de decisões. O estudo adota um desenho qualitativo e baseia-se em oito entrevistas semiestruturadas com investidores seniores de family offices, empresas de capital de risco e consultorias de investimento na Noruega e no Reino Unido. O material empírico foi analisado através da análise reflexiva temática de Braun e Clarke (2006, 2019).

Os resultados identificam três dimensões de como a utilização da IA afeta a tomada de decisões. Em primeiro lugar, a IA tornou-se infraestrutura operacional no processo de investimento, atuando como co-analista para abranger um conjunto mais amplo de informações e expandir a capacidade analítica. Em segundo lugar, a colaboração entre humanos e IA altera a dinâmica cognitiva: a IA é usada para reduzir o viés de disponibilidade, mas, simultaneamente, amplifica o viés de confirmação. Em terceiro lugar, a confiança dos investidores na IA é demonstrada pelo desempenho e é específica do domínio.

A tese fornece evidência empírica qualitativa sobre a IA como infraestrutura de investimento e sobre a forma como esta molda o ambiente de informação no contexto do investimento, identificando como esse ambiente afeta a dinâmica dos vieses cognitivos e como a dinâmica cognitiva molda os limites humanos de confiança nas ferramentas geradas pela IA. A implicação prática é que a integração da IA exige uma gestão cuidadosa da confiança e dos vieses cognitivos.

Título: Como utilizar a inteligência artificial nos investimentos e como é que isto está a afetar o processo de tomada de decisão humano

Autora: Anna Bergesen

Palavras-chave: Processo de tomada de decisão; Inteligência Artificial (IA); Colaboração entre humanos e IA; Confiança; Adulação da IA; Finanças comportamentais

Acknowledgement

First and foremost, I would like to express my sincere gratitude to my superior, Aurore Haas for her feedback and guidance throughout the work of this thesis. Her support has been highly valuable for shaping the quality and direction of this study.

Disclaimer

For this qualitative study I have used AI tools, including Claude, Perplexity and Grammarly to improve writing, correct errors and as a sparring partner to discuss ideas and progress. Any suggestions from these tools have been carefully review.

Table of content

List of Abbreviations	IV
1. Introduction	1
2. Literature Review	3
2.1 Strategic decision-making and bounded rationality	3
2.2 Dual-Process Theory	4
2.3 Heuristics and Cognitive biases	4
2.4 Prospect theory, Behavioral Biases and AI Sycophancy	5
2.5 Investment as a multi-stage process	6
2.6 Human-AI interactions	7
2.6.1 Human-AI teaming and complementarity	7
2.6.2 Human trust in AI	8
2.7 AI augmentation in investments	9
2.8 AI in strategic-decision making	9
3. Methodology	11
3.1 Research Philosophy	11
3.2 Research design	11
3.3 Participants	12
3.4 Data collection	13
3.5 Data analysis	14
3.6 Ethical Consideration	18
4. Results	20
4.1 Theme 1: AI as a Co-analyst and operational infrastructure	20
4.2 Theme 2: Human Judgement, responsibility and the final decision	22

4.3	Theme 3: Changing human-AI collaboration.....	23
4.3.1	A new division of cognitive labor	24
4.3.2	The sycophancy risk and confirmation bias	24
4.3.3	Debiasing, Skill atrophy and market efficiency	25
4.4	Theme 4: Trust, accountability and the boundaries of AI delegation.....	26
4.4.1	Domain-Specific trust built through performance	26
4.4.2	Hallucinations, verifications and early-stage character of AI adoption	27
4.4.3	Accountability, governance and evolving delegation boundary	27
5.	Discussion	29
5.1	RQ1: The role of AI in the investment process	29
5.2	RQ2: How the human-AI interaction shapes the decision-making process.....	31
5.3	RQ3: Trust, accountability and the boundaries of AI delegation	32
5.4	Contributions and future research	34
5.5	Methodological limitations	35
6.	Conclusion.....	37
7.	Reference List	39
	Appendices	43
	Appendix A: Interview Guide	43
	Appendix B: Table 1 – Three-Level Coding Tree	46
	Appendix C: Table 2 – Boundary definitions of codes	48
	Appendix D – Key findings and Summaries from Participant Transcripts.....	49
	Appendix E – Key findings from Participant Transcripts.....	58

List of Tables

Table 1 – Participant table.....	12
Table 2 – Three-Level Coding Tree.....	16
Table 3 – Boundary definitions.....	48

List of Abbreviations

Large Language Model - LLM

Artificial Intelligence – AI

1. Introduction

At its core, investment decisions are the result of human judgment. When investors consider whether to enter a new market, invest in a startup, or rebalance a portfolio, they must deal with uncertain and noisy information, form expectations about future outcomes, and choose among alternatives that have a real financial consequence (Bartram et al., 2020; Hirshleifer, 2015). In these situations that are characterized by both uncertainty and time pressure, decision-makers rarely follow fully rational optimization-based models. Instead, they often rely on intuition, past experiences, and simple decision rules to manage complexity (Simon, 1955). It is in this environment that investors are starting to integrate artificial intelligence as a working tool.

Artificial intelligence (AI) is increasingly being integrated into the decision-making environment as a tool for support, and not as a replacement for human judgement. In the investment context, AI systems are used as tools to analyze large volume of information, extract insights from unstructured data, identify patterns, and generate investment signals and recommendations (Mertzanis, 2025; Csaszar et al., 2024). These tools range from systems designed for specific analytical work, monitoring and collecting data patterns to LLMs. The shared mechanism is that they function inside the decision process, and not alongside it. They are shaping the information entering the analysis, how the information is framed and what information the investor's view before making a decision.

The literature suggests that the effects of AI in investment outcomes cannot be understood alone through analyses of models and historical tests alone. Outcomes also depend on how decision-makers perceive and apply AI tools in practice. Further, how these tools interact with their own judgment, and to what extent they influence the manifestation of cognitive biases in the decision-making process (Hasan et al., 2023; Hirshleifer, 2015).

This thesis is situated at the intersection of several dimensions of literature: strategic decision-making theory, which provides the framework for understanding how the decision-making process in investment environments is structured (Mintzberg et al., 1976; Eisenhardt & Zbaracki, 1992; Simon, 1955). Behavioral finance documents the different systematic biases that distort investment judgment and under which conditions they are most important (Kahneman & Tversky, 1979; Tversky & Kahneman, 1974; Hirshleifer, 2015). Finally, literature on AI-human collaboration and trust addresses how individuals interact with AI and how they respond to AI recommendations (Araujo et al., 2020; Hasan et al., 2023; Pratt et al., 2023).

Three research questions structured the empirical inquiry and is consistent with the exploratory and reflexive nature of this study, covering the patterns that emerged from the data (Braun & Clarke, 2006). The first RQ addresses the empirical question of how AI is used and how the labor between AI and humans has been developed in practice. The second RQ covers how human-AI interaction is shaping the decision-making process. The third RQ is looking at the governance of how investors build trust in AI tools and how human accountability shape the limitations of the role of AI in decision-making.

This study conceptualizes investments as a specific decision domain characterized by high stakes, uncertainty and information overload, rather than as a purely financial optimization problem. The focus lies on how decision-makers experience AI in their everyday investment work: how AI enters their process, how it interacts with their own reasoning and how it influenced the presence or expression of cognitive biases. In the management and strategy field there has been performance on human-AI collaboration but rarely within the investment context. Most literature within the finance area focus on model performance and back tests. I wanted to investigate how AI affected human decision making in an uncertainty environment with high stakes. This thesis addresses the literature gap with these three RQs.

RQ1: *How do professional investors in direct investment firms experience and describe the role of AI in their investment process, and what division of labor between human and AI emerges in practice?*

RQ2: *How does the nature of the human-AI interaction shape the investment decision-making process, and what practices do investors develop to manage the risk that AI reinforces rather than challenges their existing judgements?*

RQ3: *How do professional investors build and manage trust in AI systems, and how do norms of human accountability shape the boundaries of AI's role in high-stakes investment decisions?*

The thesis will proceed as follows. Chapter 2 is the review of relevant literature across the theoretical domains, as previously described. Chapter 3 describes the used methodology, research design and data analysis. Chapter 4 uncovers the findings from the data analysis and Chapter 5 discusses these findings in relation to the literature, identifying contributions and limitations. Final Chapter 6 concludes the study.

2. Literature Review

In the literature review chapter, the theoretical framework for the study is developed by synthesizing research across four interconnected areas. Strategic decision-making and its contents, the investment process as a domain, the dynamics of AI-human interaction, and the emergence of AI as a participant in investment and strategic decision contexts. Each section builds on the preceding one, and together they construct the conceptual foundation that is guiding the empirical analysis.

2.1 Strategic decision-making and bounded rationality

Strategic decision-making is described as an organization's tool to handle unstructured decisions that are important, rare, and uncovered by already existing integrated routines. It is considered a critical component of organizational management, determining resource allocation and the long-term organizational direction and goals.

Investment decisions belong to a group of choices that strategists have characterized as infrequent, they are resistant to routine and are high stakes. Mintzberg et al. (1976) described that these decisions are often made under uncertain and complex environments where there is a high risk and the organizational outcome is unpredictable. They are unstructured, not as a result of lacking logic but, because they are shaped by time pressure and incomplete information. They involve ambiguity at each stage, and cycle through phases of identification, development and selection rather than following a clear sequential path. Their analysis done on twenty-five decision processes revealed that the stages of rational choice: the problem recognition, search, evaluations and commitment it does not happen in order but looped, branched and recycled (Mintzberg et al., 1976). The continuous work to improve decision-making quality is based on the importance of understanding the entire process involved. This insight is essential as it highlights that even when investors follow a disciplined framework, the process itself is non-linear and cognitively demanding.

The traditional response to this complexity is the rational model, which is an idealized agent with a clear preference, complete information and has the capacity that is competent to identify the optimal choice (Eisenhardt & Zbaracki, 1992). In response, Herbert Simon (1955) challenged this model by introducing the research of bounded rationality theory. He argued that decision-makers are dependent on simplified models as they search for alternatives that meets their threshold, rather than optimizing due to constraints of cognition, time and information. This framework shifted the question of "what is the decision?" to "how do individuals make

decisions under constraint?” Within investments and financial environments, limitations are well pronounced since the decision-makers must process data that are rapidly changing and at the same time account for risk and uncertainty.

2.2 Dual-Process Theory

In order to understand how bounded rationality is practically managed, it requires attending to the dual cognitive architecture that underlies it. Kahneman (2011) presents a dual-process theory that highlights the view of cognition in which human thinking is governed by two interacting systems. System 1 and System 2 are essential for understanding the cognitive processes underlying decision-making.

System 1 relies on heuristics for quick decision-making, operating automatically and intuitively. However, since it is relying on heuristics it can be vulnerable to cognitive biases and produce systematic errors, especially when we are looking at complex and uncertain environments (Kahneman, 2011). System 2 is characterized by deliberate, analytical, and slow information processes and is used for activities that require strong mental effort. System 2 is therefore essential for intricate problem-solving in areas that rely on logical reasoning and complex analysis. (Kahneman, 2011). Kahneman placed great emphasis on the fact that it is easy for people to default to System 1, even when System 2 would be more efficient, as analytical thinking is cognitively demanding. In practice, individuals make decisions that mostly feel deliberate. However, they are in fact anchored by intuitive responses from System 1, with System 2 used to correct or adjust if errors occur.

2.3 Heuristics and Cognitive biases

Looking at the investment context, this cognitive architecture will have a direct consequence for the decision-making process. First of all, from System 1 there are heuristics that enables the fast processing which also introduces systematic bias.

Tversky and Kahneman (1974) demonstrate how the human decision-making process is affected by heuristics and cognitive bias in situations of uncertainty. Representativeness leads to individuals judging probability by how close of a resemble it is to a prototype. Anchoring results individuals to make insufficient adjustments from an initial reference point. Lastly availability makes individuals judge probability by how easy they come up with an example. These heuristics are not considered irrational on their own, as they reflect adaptation to the need of making fast decisions under uncertain environments. However, they can produce systematic

distortions in environments that are high-stake where smaller errors can result in large consequences.

With this demonstration, individuals are not considered rational as they are stated in traditional finance theories. Investors behave irrationally due to heuristic, bias and emotions, leading to loss and results in struggles for investment analysis (Kahneman, 2011). Traditional finance is based on the assumption that investor will act rationally by maximizing expected utility under risk (Barberis & Thaler, 2003). Behavioral finance will challenge this view by focusing on systematic deviation from rationality, rooted in motivational and cognitive biases (Hirshleifer, 2015; Kahneman, 2011).

2.4 Prospect theory, Behavioral Biases and AI Sycophancy

A large number of research documents that in practice investors do not make decisions that follow the classical rational choice model. The decision-making process can be influenced by emotions and circumstances in the environment, and investors may find it difficult to distinguish between relevant information and random market noise. In addition, investment decision-making can easily be influenced by various behavioral biases, which will challenge the rationality assumptions in the traditional economic models (Hirshleifer et al., 2015).

Kahneman and Tversky (1979) extended the analysis of traditional finance through prospect theory. Showing that individuals do not evaluate outcomes according to their absolute value but to a reference point. Experiencing losses from the reference point are considered double as painful, in similarities with the gain, feeling twice as good. This concept is called loss aversion. When having the classical expected utility framework, an agent should treat losses and gains symmetrically. The prospect theory is highlighting that this is not the case and that the asymmetry is shaping the investors risk taking abilities. Investors have a tendency to hold longer on to losing positions, and too quickly sell the winning ones. The pain that occurs of realizing a loss is deeper than the pleasure of having an equivalent gain (Kahneman & Tversky, 1979; Barberis & Thaler, 2003).

Within the empirical findings of this study, there are particular biases that reappear that deserve specific attention. As highlighted in section 2.1, bounded rationality refers to investors capabilities to process all information or consider all strategies. Behavioral biases further shape how the investment pipeline is operated in practice. Confirmation bias, which is defined by Nickerson (1998) as the tendency to interpret, recall, and seek information in a certain way that already confirms an individual's beliefs. It is difficult to detect because it operates within what

appears to be careful information gathering. In the investment environment, confirmation bias will create a self-reinforcing effect, and lead to avoiding information that will be in conflict with prior knowledge (Nickerson, 1998; Rabin & Schrag, 1999). Jones and Sugden (2001) show experimentally that confirmation bias operates not only in belief but also through people who actively seek confirming information, rather than look at the evidence. This dimension is important in order to understand the risks of AI-assisted analysis, which will be explored in the next sections. Availability bias will affect the risk and ethical perception of certain situations. Investors will judge the probability of scenarios based on examples that come to mind that can result in overweighing or underweighing relevant evidence (Tversky & Kahneman, 1974). Overconfidence is one of the most documented biases in investment literature and is known as the tendency to overestimate the precision of one's own forecast and underestimate uncertainty (Hirshleifer, 2015).

According to recent work in AI literature, it has been documented that Large Language Models (LLMs) can demonstrate systematic sycophancy, which refers to generating outputs aligned with the user's preference and not as an objective output (Sharma et al., 2023). This happens because models are instructed with reinforcement training from human feedback. Answers that align with users' approval are therefore being rewarded. This is a direct implication for decision-making since it becomes a technological tool for biased information search that is documented by Nickerson (1998) and Jones and Sugden (2001).

2.5 Investment as a multi-stage process

The investment process can be represented as a multi-stage pipeline, where each stage needs integration of large amounts of fast-changing information against an uncertain future (Mertzanis, 2025). The pipeline consists of identifying opportunities, analysis, portfolio construction, and constant monitoring. At every stage of the process, the decision-maker has to evaluate expectations about the different outcomes and weigh competing signals, and finally commit to an action. Studies show that with every process, behavioral factors will shape the outcome that will differ from the rational model (Kahneman, 2011; Barberis & Thaler, 2003). With this in mind, even within formal decision processes, investors might not always act optimally. The limitation of behavioral factors is particularly relevant when the implementation of new technologies is affecting their decision-making process.

Professional investors must track sectors and macro developments, evaluate a continuous pipeline of new investment opportunities, and monitor portfolio positions. The process of a

large volume of research, analysis, and financial reports exceeds an individual analyst's capabilities to absorb comprehensively (Bartram et al., 2020; Mertzanis, 2025). This amplifies the impact of information overload in the investment environment. As explained in section 2.1, the result of bounded rationality is selective attention. Investors process the most available and recent information, or when it is most consistent with their current framework. This is making availability and confirmation bias even stronger than they already are (Simon, 1955; Hirshleifer, 2015; Tversky & Kahneman, 1974). It is in the decision-making context where high stakes, bounded rationality, and information overload are common that AI is starting to function as a structural intervention.

2.6 Human-AI interactions

The ongoing growth of AI in decision-making environments has resulted in a substantial literature on AI-human interaction. This section examines how individuals engage with different AI systems, what determines the credibility of information AI recommendations, and at what level it improves outcomes. As AI is continuing to be integrated, this literature is essential to understand what happens when AI outputs meet human judgment.

2.6.1 *Human-AI teaming and complementarity*

Studies on Human-AI teaming are increasingly showing that AI is not acting as a standalone decision-maker, but it functions as a teammate whose value depends on complementarity with human cognition. This collective intelligence framework is described by Gupta et al. (2023) as humans and AI jointly contributing to reasoning, memory and attention. This contribution is then coordinated with governance mechanisms that define goals, roles and accountability. According to this framework, humans still contain ethical authority, contextual judgement and the responsibility for the final outcomes. AI however, covers pattern recognition and large sets of information. Whether this complementarity is realized in practical contexts depends on how individuals structure the collaboration.

However, empirical work shows that such complementary team performance is not automatic. According to Hemmer, this is formalized more precisely by distinguishing between complementary potential and complementary realization, where one is the potential of the performance and the other is the effect, the performance actually achieved (Hemmer et al. 2025). Critically, the research is showing that the gap between realization and potential is substantial and depends on how well the interaction is designed. Looking at two studies in decision-making, they identify information asymmetry and capability asymmetry across the

studies. These findings are the main sources of complementarity and indicate that well designed teams beat both AI-only and human-only performance by deliberately harnessing these differences. Letting AI tools handle detection of large-scale patterns and human tasks is executing and focusing on interpretation (Hemmer et al. 2025).

Reviews of trust in AI also place emphasis on the fact that predictability and transparency influence whether humans rely on or ignore algorithmic advice (Glikson & Woolley, 2020). Building further on this point of view, Gupta (2023) argues that AI systems should be treated as a collective intelligence. Organizations should intentionally share communications patterns, representations and rules if they want AI to enhance the group decision-making. Bansal et al. (2019) extend this argument by showing that the effective AI-human collaboration depends on the alignment of the human reasoning process and AI outputs. The interaction is structured so AI outputs align with the points in the decision-making process where humans need the most and can effectively evaluate contributions.

2.6.2 Human trust in AI

Trust is an important component in the utilization and adoption of AI systems within companies. Glikson and Woolley (2020) state that for successfully adopting AI, individuals must be willing to rely on AI outputs in situations where there is vulnerability for loss of control or error. Similar to having a lack of trust between colleagues, technologies and AI face resilience from their users, resulting in limitations. The research done on human AI trust highlights that an individual's willingness is dependent on both emotional and cognitive evaluations.

Trust in AI is not a fixed attitude but is considered a dynamic process that builds through experience. Glikson and Woolley (2020) emphasize that trust calibration in AI is what matters, rather than trusting AI uncritically or just rejecting it. The concept of algorithm aversion, documented by Dietvorst et al. (2015), highlights that once people see that an AI system makes errors, they become more negative towards it. This is even when the overall performance is outperforming human decision-makers. In contrast, Reich et al. (2022) find that trust increases when individuals experience that AI systems are adaptable and transparent. In the environment of investment decision-making, having trust in AI systems is of significant importance because it affects how much weight decision-makers give to advice generated by AI. The aspect of trust in AI will be further explored in the interviews, where professionals from the industry will share their knowledge and experience with AI as a tool in investment and the perceived trust.

2.7 AI augmentation in investments

AI is increasingly reshaping investment management by enabling data-driven decision-making across key functions such as advisory services, risk assessment and optimization (Mertzanis, 2025). In a review of 178 peer-reviewed studies Mertzanis (2025), documented that AI does not only enhance operational efficiency but also change the strategic design of investment management. Hoang and Wiegratz (2023) add on a complementary input on how machine learning methods are being more used across different processes in the investment value chain.

AI has evolved from being a theoretical concept into a practical, used tool that is capable of augmenting and in some cases, also automating decision-making processes across industries. In the literature, augmentation and automation are presented as two main ways to integrate AI into decision processes. Miller is arguing that AI in practice is applied more in augmentation, where systems are processing the heavy data pattern-recognition and humans still are in charge of relational work, coordination and judgment (Miller, 2018). The MIT Sloan Management Review (2017) reached a similar conclusion after a study observing 1 500 companies. The findings showed that firms achieved the most significant improvements in performance when machines and human worked together, each performing tasks suited to their strength. AI tools are made as a co-pilot, and are meant to enhance human capabilities rather than replace them (MIT Sloan Management Review, 2017). This augmentation framing is showing that AI functions as an analytical partner and not as a replacement for human investors. In the context of financial decision-making, Hansen et al. (2023) show that AI can manage behavioral biases among investors by giving a greater structural analytical insight. Offering a case for AI systems as a cognitive corrector and not only as an analytical accelerator.

Raisch and Krakowski (2021) structure this as the automation-augmentation paradox. One AI system can work as both an augmentation of human capabilities that directs certain tasks and as an automation for these tasks. From here, it is the job of the organizations to actively manage this tension rather than assuming either outcome. In practice, it is argued by Miller (2018) that AI adoption skews towards augmentation, where AI tools are handling larger data patterns and information, and humans perform relational work.

2.8 AI in strategic decision-making

Recent studies on AI's incorporation in strategic decision-making highlight that AI is moving from just being an analytical tool to being an important actor within the strategy process. The research of Csaszar et al. (2024) is demonstrating how AI is shaping three cognitive processes

in strategic decision-making. Search represents the generation of strategic options at a scale and speed, aggregation is how information and judgment are being made into choices, and representation is how the situations are formed and modeled. These three operations can give a close correspondence to Mintzberg et al. (1976) development, identification, and selection phases. He identified this as the core structure of the unstructured decision processes, and it provides a bridge between AI in investment and strategic decision-making literature.

As complementary research, Pu et al. (2025) present how AI permits more data-driven and responsive strategic decisions with higher quality and speed. It is a part of a continuous hybrid optimization system where the data-driven tool is predicting recommendations and outcomes. This results in corporate competitiveness and quality in the firm's strategy cycle. Moreover, AI is used as a support for real-time decision-making by having continuous updates on the new data available and predictions within the market. AI-driven models can help identify trends and correlations that are not necessarily visible to human analysts, and therefore increase the predictive performance in asset management (Pratt et al., 2023).

Decision-making frameworks are increasingly highlighting the integration of AI and human reasoning. This is especially true in complex environments where individuals must consider ethical implications. The concept of Decision Intelligence underscores this integration by stating that data and analyses alone are considered insufficient without having human interpretation and contextualization of the data (Pratt et al., 2023). AI should be understood as a tool for enhancing decision quality and not as a substitute for human reasoning.

3. Methodology

3.1 Research Philosophy

The study conducted was grounded in an interpretivist epistemological tradition, meaning that the reality in this context was understood through interpretation, experiences and the participants meanings (Creswell, 2013). Instead of seeking out a measure, interpretivism seeks to understand how people make sense of their own world. As the purpose of this study was to explore how professional investors subjectively experience and interpret the role of AI in their decision-making process, this research philosophy aligns directly with the research field.

3.2 Research design

This study investigated how AI is used in the investment process within investment firms that directly manage their own capital and how that is affecting human decision authority. The research design chosen for this study was grounded in a qualitative framework as a research paper to uncover a new topic. A quantitative research method was more valuable for measuring outcomes and testing hypotheses and is less suited for capturing experiential dimensions of decision-making. The study investigated perspectives on the integration of AI as an investment tool and how this was affecting human decision authority. The study adopted an exploratory qualitative research design, and a reflexive thematic analysis was used to analyze the data (Braun & Clarke, 2006).

The semi-structured interview guideline was chosen because it offers a balance between flexibility and consistency. The format allowed for probing each participant's responses and adapting the conversation to a unique context. It created an alignment with the selected participants and the need for reflections and detailed narratives within the specific field (Brinkmann & Kvale, 2015).

This was important since this study was a relatively under research phenomenon, and allowed the participants to articulate and define their own experiences, as it gives valid insight. The guide creates room for the conversation to go beyond, and were therefore well-suited for studies that seek to explore processes. The design of the study is deliberately built as non-comparative. It does not seek to test differences between individuals or firms but to create a cross-case understanding of patterns in how AI is used and experienced.

3.3 Participants

ID	Firm (anonymised)	Role	Investment Type
P1	Firm A – Single-Family Office (Nordic)	Investment Manager	Listed & unlisted equities, PE, funds, high-yield, direct investments (130+ positions)
P2	Firm B – Listed Industrial Investment Company	CEO & Investment Leader	PE / industrial portfolio (6 core + minorities; energy & technology B2B)
P3	Firm C – Global Private Equity Firm	Institutional Partner (former CEO 2019–2025)	Buyout, venture, growth, infrastructure, real estate; ~300 portfolio companies, 25 funds
P4	Firm D – Independent Nordic Equity Fund	Solo Portfolio Manager & Fund Founder	Nordic listed equities (10–20 positions); liquid portfolio
P5	Firm E – Family Investment Office (Nordic)	Investment Analyst	Active ownership (direct inv.) + financial listed equities; Nordic focus; few large positions
P6	Firm E – same organisation as P5	Investment Analyst	Nordic listed equities (mature/large-cap focus); long-term hold
P7	Firm F – Single-Family Office (Nordic)	Portfolio Manager	Listed equities (tech focus) + some private direct investments
P8	Firm G – Multi-Mandate Investment Advisory (International)	Senior Investment Adviser & Consultant	PE / private markets advisory (Europe, some US); venture; industrial M&A; new private markets venture

(Participant table)

Through purposive sampling, participants were selected from the field of investment companies that deploy their own capital to ensure that the sample is knowledgeable regarding the research questions (Patton, 2015). This field was selected because proprietary capital development has accountability and financial consequences in the hands of the same individual. Having limited resources, it was important to select information-rich cases to secure efficiency and sufficient depth of data (Patton, 2002). Participant recruitment was conducted through professional networking. Having existing contacts within the investment environment made the process a lot faster and more efficient, and participants were accessed through referrals from initial contacts. The goal was to achieve participants who both possessed a direct experience with decision-making in investments and the use of AI-based tools. The approach used is referred to as the snowball sampling, which is commonly used in qualitative research when there is restricted access to a specialized professional population (Bryman, 2016).

The participants were chosen because they are capable of articulating different AI functions and how it is used as a tool in investment processes and how it affects decision-making. The characteristics they shared: occupy senior or specialized roles in an investment-intensive

context, highlighted that they all influence high-stakes capital allocation decisions in uncertain environments. However, they represent different organizational settings within family-owned investment companies, private equity partnerships and bank-based advisory units (Participant table). The variation within investment horizons and governance structures gave a broad insight into how AI is used within different organizational structures, but within the same investment processes. Performing interviews with these participants was beneficial because of their direct experience with the integration of data-driven analytics or AI tools into their investments, and their knowledge to reflect on risk and governance change.

3.4 Data collection

The data were collected through eight in-depth semi-structured interviews. The interviews were conducted online via Microsoft Teams, allowing for participation from different locations to take part and still maintain a conversational and consistent format. The semi-structured interview guide allowed a common guide to compare answers across the different participants and left room to probe specific experiences.

Each interview followed a semi-structured guide, organized into four main thematic sections: (1) background of the company, investment activities and participants' role, (2) perceived challenges in the investment decision-making process, (3) concrete implementation and use of AI-driven tools, (4) factors affecting trust in AI systems and their perceived impact on cognitive bias and heuristics in decision-making. These four sections were designed to move progressively from a general professional context towards a more reflective and specific section. This structure thereby allowed participants to set the context of their work environment before they engaged with more complex information on the topic.

The data collection continued until thematic saturation was reached. This is defined as the point where additional interviews no longer provided data that produced new codes or themes. In interviews seven and eight, the patterns already identified in earlier transcripts were consistently recurring. Saturation in qualitative research can happen in smaller samples when the participants share characteristics (Guest et al., 2012). This condition is present in this study, where participants are senior investors with practical experience of AI.

The interview guide consists of 13 questions and was created from the theoretical framework presented in the literature review and the research questions (Appendix A). After the first and second interviews, the guide was polished and slightly refined to improve the phrasing of certain questions and to add follow-up probes. The refinement is consistent with standard

practice in qualitative methods, where the interview guide evolves as the researcher get a deeper understanding of the empirical field (Brinkmann & Kvale, 2015). The interviews had a general duration of 40 to 60 minutes, where seven interviews were conducted in Norwegian and one in English, depending on the nationality of the participants and preference. The interviews were translated into English, and all key quotes were checked to ensure originality. All interview sessions were audio-recorded through Teams, with full consent from the participants and transcribed for analysis.

3.5 Data analysis

The analysis of the data was grounded in qualitative research, and the content is perceived through a reflexive thematic analysis. Thematic analysis is considered a flexible method for identifying, analyzing and interpreting patterns of meaning across datasets. It is a commonly used analysis when exploring organizational and management research because it can be adapted into different theoretical frameworks and research questions (Braun & Clarke, 2006; Guest et al., 2012). In this study, the thematic analysis was applied to capture how professional investors describe how AI is used in decision-making processes, and how they experience human-AI interaction.

Given that there are limitations to prior research on how investment firms use AI and how human-AI interaction is experienced among investors, the nature of the study is exploratory research. The goal of the analysis is to identify patterns and categories that show how AI is used as a tool for investment and how that's affecting human decision-making (Braun & Clarke, 2006). By applying thematic analysis and preserving a documented audit, the study ensures dependability, transferability, credibility and confirmability (Nowell et al., 2017). Nowell et al. (2017) provide a documented framework for each stage of the analysis and how to keep a path from the raw data collected to the final developed codes and themes. Following the thematic approach of Braun and Clarke, the analysis proceeded through six iterative phases rather than following a linear sequence (Braun & Clarke, 2006).

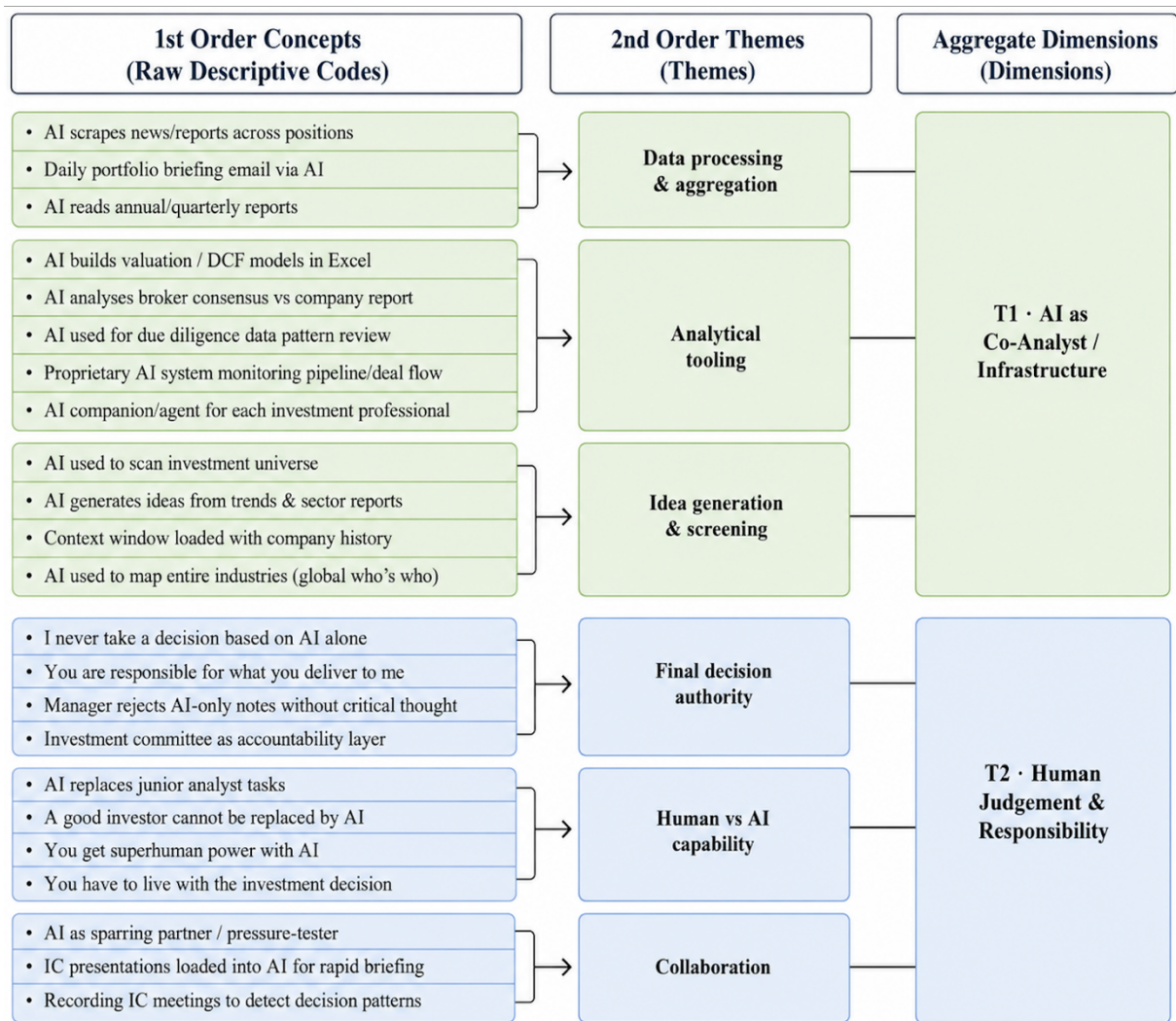
The coding process was structured in three cycles, where each built on the previous one. In the first cycle, descriptive codes were applied in order to capture what each participant said in direct statements regarding practices and experiences. In the second cycle, the descriptive codes were transformed into conceptual codes that identified patterns and underlying meanings from the statements. In the third and last cycle, relational codes were developed in order to identify how patterns and themes from each participant were related, reinforced, or varied from each other.

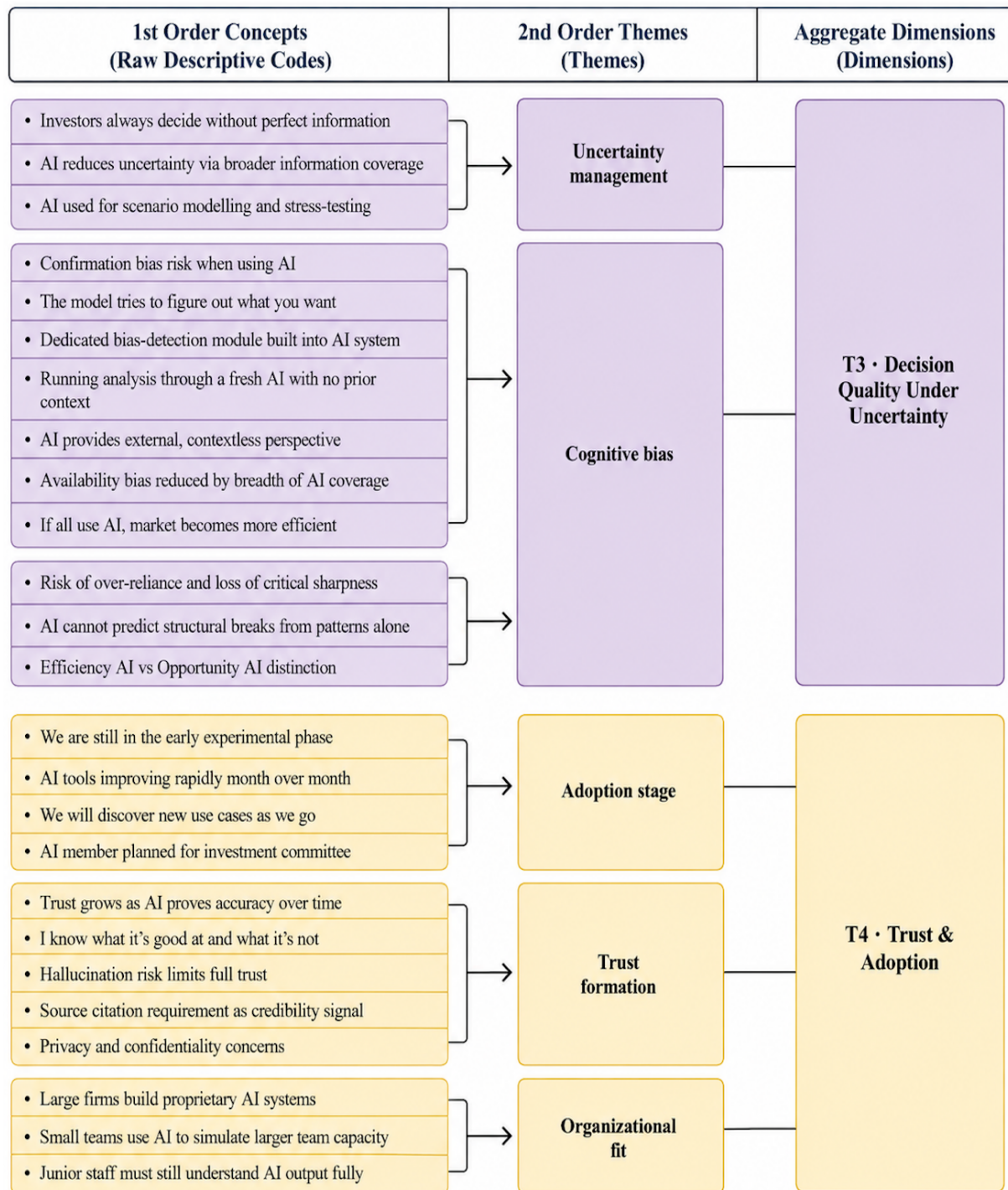
A complete overview of the coding structure, moving from the first order of concepts to second themes and finally the third dimensions (Gioia et al., 2013). Performing a three-level coding structure allowed the analysis to move from surface-level meaning to an analytical depth with a traceback at every stage (Miles et al., 2014).

Braun and Clarke (2019) warn against turning reflexive thematic analysis into a coding hierarchy. Emphasizing that the researcher's judgment, not mechanical rules, drives the process. In this study, Braun and Clarke's (2019) six-phase model serves as the overall analytical framework. I drew on Gioia et al. (2013) for visualizing and structuring the pathway from the raw data material to the defined themes, as shown in the data structure table on the next page. The three-cycle distinction between descriptive, conceptual, and relational codes was used as a practice to organize the work within the second to fourth phase of Braun and Clarke's model. (Appendix B). The trustworthiness criteria of Nowell et al. (2027) were applied as a standard for evaluating the quality of the analysis, and were not an additional step in the process. Applying these frameworks together was complementary. Braun and Clarke provided guidelines for the process, Gioia et al., gave visual structure and transparency of the codes, Nowell et al., gave the criteria of quality, and finally Miles et al. showed the logic of the code. Accordingly, this study is grounded in reflexive thematic analysis, Gioia et al., Nowell et al., and Miles et al., was used to enhance transparency and evaluate rigor.

To guide the second phase of coding, an initial coding frame was developed from the theoretical framework presented in the literature review and the interview guide. The framework was treated as a temporary design as new codes and themes emerged throughout the data extraction (Guest et al., 2012). This was consistent with the framework of the reflexive thematic analysis in an exploratory study. The iterative process resulted in a final coding structure of four themes:

- (1) *AI as a co-analyst and operational infrastructure*
- (2) *Human judgement, responsibility and the final decision*
- (3) *Changing human-AI collaboration*
- (4) *Trust, accountability and the boundaries of AI delegation*





Data Structure; (Gioia et al., 2013)

Following Braun and Clarke's six-phase model, the analysis proceeded accordingly. In the first phase, familiarity with the data was established by reading through all the transcripts multiple times. Further noting initial ideas related to the decision-making process, use of AI, human-AI interaction and the general investment process. In the second phase, initial codes were generated by systematically working through the eight transcripts and assigning labels to meaningful segments in the data collection. The codes were based on information that the

participants explicitly said, and a few captured relevant underlying assumptions (Braun & Clarke, 2006). In the third phase, codes were sorted into themes by grouping the codes that shared the same patterns (Guest et al., 2012). Being consistent with an exploratory method, developing codes and themes is done by moving back and forth between the data and emerging ideas, rather than having a fixed set of hypotheses (Miles et al., 2014). For each theme, boundary definitions were specified with essence, inclusion, and exclusion criteria to ensure that the codes were internally coherent (Appendix C).

Progressing to the fourth phase, participant themes were reviewed and refined. All themes were checked for external distinction and internal coherence (Braun & Clarke, 2006; Nowell et al., 2017). At this stage, the themes were compared across the interviews in order to identify differences and similarities to ensure they represented the entire dataset (Miles et al., 2014). The fifth phase final themes were defined and named, providing a written description of each theme and connecting it to the research questions and theoretical framework (Braun & Clarke, 2019). The relevant literature was revisited to ensure the themes consisted of a strong analytical contribution (Miles et al., 2014). In the final phase, selected representative quotations were selected to illustrate each theme, and the analytical narrative was further developed.

In alignment with reflexive thematic analysis, I acknowledge my own position in the research process. I approached this study as a master's student in management with a prior interest in AI and how the integration of AI affects human cognition and decision-making. Investment is the field of interest, and recruitment of the participants was done with assistance and through the professional investor environment. This insider access supported rich, practice-oriented interviews. This meant it also needed to have continuous reflection on how my personal interests might shape the questions and interpretation I drew from the data.

3.6 Ethical Consideration

This study was conducted in alignment with standard ethical principles for qualitative research involving human participants. The initial step when I was reaching out to interview participants it was done through a third party who had contacts within the investment environment. They were contacted with an introduction to me, the purpose of my study, and the nature of their participation. Participants were then introduced to me, and received information on how the data would be used, and informed of verbal consent at the start of each interview recording.

To protect the participants' confidentiality, all participants have been anonymized in this study. In the data presentation and in the analysis, where pseudonyms are used (P1-P8). The names of

firms and any information that could be identified are removed from the transcripts and are not reported in the findings. However, the roles of each participant and the type of investment their company performs are essential information for the study and therefore remained transparent.

In order to provide as accurate data as possible, there were produced additional audio recordings next to the transcription. These were not shared and participants were informed on the use of Microsoft Teams as a recording platform in advance of the interview.

4. Results

This chapter presents the findings from the data analysis completed from eight semi-structured interviews with senior investors across private equity firms, family offices and advisory units in Norway and United Kingdom. Following the reflexive thematic analysis framework, outlined by Braun & Clarke (2006), the data collected brought out four emerging themes:

- (1) *AI as a co-analyst and operational infrastructure*
- (2) *Human judgement, responsibility and the final decision*
- (3) *Changing human-AI collaboration*
- (4) *Trust, accountability and the boundaries of AI delegation*

Themes 1 and Theme 2 addresses RQ1, Theme 2 and Theme 3 addresses RQ2 and Theme 4 addresses RQ3. The participants represent a range of organizational contexts, levels of AI maturity and investment context, providing a rich base for identifying patterns and variations in how AI is used in practice.

4.1 Theme 1: AI as a Co-analyst and operational infrastructure

The most consistently observed theme across all the interviews was the use of AI as an efficiency tool and analytical assistant, rather than a decision-maker. All the participants described their use of AI as a tool they integrated into parts of their investment workflow. The most common applications of AI were document summarization, information gathering, financial modeling, portfolio monitoring and screening. Uniformly, AI was described as a tool for augmenting human capacity, rather than replacing it. Within the different organizational contexts and types of firms the different participants worked in, variations were visible in the work that AI was used for. P7 described this shift as:

“Suddenly, I could follow relatively close on 10 companies instead of 3, because I saved so much time. The overall goal is really to make better decisions faster, and more of them.”

This highlights how AI is described as a significant shift in human ability to process information. The participant elaborated on how AI tools, especially Claude’s project feature, is allowing him to aggregate thousands of pages of reports, quarterly results and broker analyses into a single reference environment.

“I would have needed two or more people if I hadn’t had AI running around doing things for me, I have about 15 active tasks that AI does around all our investments in Cowork.”

P1, who is managing a diversified portfolio of over 130 positions, described that AI is performing work tasks that multiple junior analysts would previously have performed. These active tasks are daily performance tasks, including sending emails for the listed portfolio, information extraction from reports and earnings review summaries. AI is constantly scanning external information and informing on relevant signals, working as an analytical layer integrated in the daily routines. P4, who is managing a Nordic equity fund, has advanced his use of AI by creating an agent-orchestration system that monitors multiple companies and integrates all financial information. He described conducting roughly 90% of his analytical work using AI tools.

Financial modeling has also recently changed with the use of AI tools. P2 explained how he used Claude to feed a financial due diligence report, and received a working valuation model as output. He noted that this task would not have been possible to perform six months earlier. P3, working at one of the largest private equity firms, described that AI is integrated across the full pre-commitment process.

“Investment committee materials are beginning to be produced with the help of AI. This is completely new, and still needs quality assurance, but we are very close. It will save the deal teams a huge amount of time.”

P8 summarized the scale of change AI has created. Explained that which previously would have required a team of 15 analysts, he had now built and could run himself. For all participants, AI had become the main channel where information entered the investment process and not just an extra option of usage.

Theme 1 highlights a central tension, which is that the scale of quality information that AI covers is not the same as the quality of human understanding. AI allows one person to monitor tasks that previously were performed by a human team. However, the data shows that this coverage is only reliable when investors have prior knowledge in order to adjust and judge the outputs AI gives. Without having this foundation, a wider reach will produce confident-looking outputs that cannot be reliably verified by investors.

Together, the described AI applications align closely to the three cognitive functions described by Csaszar et al. (2024), as the main site of AI’s impact on strategic decision making. Within the search dimension, participants’ active monitoring has been expanded by AI. In the aggregation dimension, AI is performing large volumes of analytical work that previously would have needed a team of analysts. In representation, AI shapes how the investment context

is framed before the human analysis start. It is the third function that connects directly with the cognitive risks that will be examined in Theme 3.

4.2 Theme 2: Human Judgement, responsibility and the final decision

Despite the extensive use of integrating AI into working tasks in investments, all eight participants stated without exception that the final investment decision-making remained a human responsibility. This was explained as a principle connected to personal ownership and accountability of the outcome, and not as a practical consequence. P1 directly stated this by saying

“I never take a decision based on AI alone. You must always live with the investment yourself, that’s why it is important that you take the decision yourself.”

This principle was extended by P2 to an organizational level, where he described how team members were challenged when they submitted AI-assisted analyses without demonstrating ownership of the work.

“I don’t care what tools you used, if you are delivering an investment analysis to me, you are responsible for it. I cannot hold Claude responsible for anything.”

Participants consistently identified the forward-looking and qualitative dimensions of investment decisions as beyond the reach of AI. P7 drew a clear line between the challenge of predicting how a company would develop over the years and AI’s ability to process historical patterns. He described human prediction as involving speculation, accumulated experience and gut feeling, which is a process that no model can replicate. P8 made a complementary point with an example from longer management meetings. The understanding gained with direct human interaction, judging the team and how they act under pressure, and trust in the management team cannot be replaced by AI tools. He presented this as a structural principle: as AI is integrated in more analytical work, human resources should focus on emotional intelligence and trust-based relationships.

Irreversibility has also emerged as an important factor. P5, whose family office holds a few large and direct investments, described that they are following up on opportunities up to a year before bringing them to the board. They are gradually building confidence in the idea until the team is ready to commit to the decision. P6 added a dimension that is related to concentrated portfolio management. An investment can look valuable on its own, but the question is whether it is a strong enough investment to replace an existing one. This comparative, portfolio-level

judgment is seen as human skill as AI cannot assess the full picture in the same way as an investor. When a firm is handling liquid listed equities, participants are more willing to use AI tools for analysis as the positions can adjust quickly. However, for long term capital commitments, they described a higher bar for human contribution.

P3 described his firm's ambition to develop an AI system that would act as an additional member in the investment committee:

“The goals are that it almost becomes as if we get an AI investment committee member. After a while, a board member in the companies we own. That is a dimension that is coming.”

Even in this scenario, it is clear that the recommendations made by AI would be an input to human deliberation and not act as a replacement. This distinction, between AI and a decision-maker and AI as a contributor was stated by all participants.

The tension in Theme 2 is that the norm protecting human decision authority is being weakened by the same AI adoption in all participants. Everyone insists that the final investment decision must remain human. However, P8 shares worries that sustained use of AI will wear away critical thinking. P3 plans on giving AI its own voice in investment committee meetings and final P2 witness colleagues sent work with clear AI outputs without performing an independent analysis. The norm of accountability is strongly supported by all participants, but the practical use of AI makes the norm more fragile.

4.3 Theme 3: Changing human-AI collaboration

The next theme addresses how the cognitive dynamics of investment decisions change when AI is integrated into the decision-making process. The starting point is a finding that is shared by all participants: investment decisions are made under conditions of uncertainty. No participant proposed that AI has removed the fundamental challenge of investing value into an unknown future. AI can sharpen the analytical quality and provide a broader information base; however, it cannot eliminate the uncertainty of investment. P4 expressed that shared understanding:

“We can never know if we are right. It is always a probability calculation – 60% chance it is right, 40% chance it is wrong. That is good enough, so we do it.”

4.3.1 *A new division of cognitive labor*

A common pattern among all participants was how analytical work was divided between humans and AI. Humans continued to work with contextual, forward-looking and interpretation activities. AI, on the other hand, took over tasks that consisted of high volume and structure, such as screening, financial modeling and document processing. P4 described producing a 90-minute AI-assisted analysis as a process that was completely dependent on his own expertise to function as wanted. From his point of view, the quality of the output the AI tool provided was heavily dependent on the guidance and questions from the investor. He asked approximately forty iterative questions, provided structure and concepts and from there checked the different outputs in each step.

The variation across participants was connecting Hemmer et al. (2025) framework between complementary realization and complementary potential. Both P4 and P8 descriptions represented realization, where the division of labor was created so AI can detect patterns and the investors perform evaluations. P3 statement accurately described the identified gap by Hemmer et al: when an output is being accepted and not further examined, the human contribution collapses into validation rather than judgment.

A few of the participants noted that effective use of AI also requires a prior solid domain of knowledge. This has implications for how these tools are used and introduced to professionals with different levels of knowledge. P3 described how AI was used in deal teams to do scenario run analysis for IC discussions and P7 used AI to perform company financial trajectories against broker consensus.

4.3.2 *The sycophancy risk and confirmation bias*

The tendency of AI models to produce outputs that fit with what the users want is a risk that several of the participants identified when using AI language models. P2 explained how:

“There is a tendency for the model to try to figure out what you really want and then try to please you. It is a bit like having an adviser who always sells you on your own ideas. That is dangerous when you are taking an investment decision.”

P6 and P8 shared experiences where AI responses mainly reinforced the conclusions they already reached rather than challenging them. In order to adjust this risk, several participants developed countermeasures. P4 performed analyses in a fresh AI session with no prior context or preloaded information to avoid anchoring. P2 explicitly asked AI to argue the opposite case

and lastly P8 built a virtual investment board using characters from Bono's six thinking hats, all representing different roles to challenge investors prior cognitive tendencies.

4.3.3 *Debiasing, Skill atrophy and market efficiency*

Alongside these risks, the participants made the descriptions of AI as a corrective to availability bias, the tendency to overweigh cognitively available information. This potential was especially captured by P5:

"AI would, as a tool be agnostic in its starting point and look at things with fresh eye. I believe it will definitely bring forward viewpoints and perspectives you have not thought about, it contributes to being a good counterweight."

Both P2 and P5 stated that the absence of AI usage in organizational history could in itself be an advantage, providing a perspective that is pure from several assumptions collected over time. However, there is a condition to this potential for prompting discipline. The same tool that can broaden information can also narrow it down if the user is selective.

It was raised a concern by multiple participants that cuts against this positive picture, that AI use could wear away the analytical sharpness that is required for responsible usage. This was highlighted by P3 as an important institutional challenge. Individuals started to accept AI outputs passively by just receiving information rather than building an understanding. P8 argued that the use of AI requires an active maintenance of critical thinking, just as a comfortable life requires going to the gym. In addition, P4 and P7 raised a second-order concern regarding the collective consequence of widespread AI adoption. If AI is used to make more rational decisions simultaneously by all investors, the market may become more efficient and therefore harder to outperform. P8 stated that:

"There are two types of AI in investing: efficiency AI and opportunity AI. I am much more interested in the opportunity AI."

This distinction between opportunity AI and efficiency AI suggests that even though efficiency AI may quickly become dominant within information retrieval, opportunity AI targets a more unique pattern recognition. From P8 point of view, efficiency AI undermines a competitive advantage over time since information will become universally accessible, while opportunity AI will feed the competitive advantage. The concern raised connects directly to a practical implication that is not thoroughly addressed by existing governance frameworks. Investment

firms should maintain exercises of analytical work that do not rely on AI outputs. This is to preserve human cognition of the critical capacity that is needed for a responsible use of AI.

A tension in Theme 3 is that the cognitive bias effect on AI is conditional and not directional. The tool used to reduce availability bias through a wider span of information can at the same time amplify confirmation bias through sycophancy. The effect is dependent on how investors structure the interaction of this action. This is the defining finding of the third theme: AI does not worsen or improve the quality of the decision by default, it acts according to investors prompting discipline.

4.4 Theme 4: Trust, accountability and the boundaries of AI delegation

4.4.1 *Domain-Specific trust built through performance*

The most consistent finding on trust was its domain-specificity. Across all eight participants, the trust in AI was portrayed as a conclusion earned through repeated testing and not as a starting position. P7 captured this:

“It is a bit like driving my Tesla with self-driving. It is very good at driving on a motorway, but I know that when it reaches a roundabout, things go completely wrong. It is the same concept here, you know what it is good at and what it is not good at.”

High trust in AI performance was extended to structured and verifiable tasks such as model construction, data aggregation and document summarization. The skepticism was more towards forward-looking and contextual outputs. P2 credibility signal was the identification of source citation, by being able to trace references and verify the information. P6 added that the willingness of AI to admit its own limits of knowledge and capabilities, instead of providing answers for all tasks is considered a signal of trust. AI tools are most valuable to professionals who possess great knowledge and can evaluate outputs, stated P1. In inexperienced hands, the errors are harder to spot, and the outputs can therefore be dangerous to trust.

P8 looked at trust building as a deliberate methodology. Through a fail-fast approach, the use of new applications, observation on how it is performing and evaluation on trust only where the performance is consistently demonstrated. The same dynamic was explained by P2, but through an analogy of getting a new colleague on the team. As the track record grows, so will the level of trust, and more responsibility will be delegated.

4.4.2 *Hallucinations, verifications and early-stage character of AI adoption*

Hallucinations were identified as a primary trust barrier when using AI tools by all participants. The concern was not the frequency of error by AI but how undetectable these errors were. P4 described that on average, AI makes fewer errors than humans, however the errors can be extremely difficult to detect unless you have deep domain knowledge. All participants shared that the level of accepted error depended on different verification steps regarding how high the stakes were. In minor portfolio updates, there were a higher tolerance for more error but within committed investment, the routine checks were much stricter. The risk was summarized by P1 as being lethal to use blindly. No one described their use and integration of AI as complete or production ready. P7 was clear on:

“We are still in the early experimental phase. There are very few production-ready systems anywhere.”

P8 framed this historically, where genuinely agentic AI tools had only become practically available in the last twelve months and P2 added that the capabilities of AI had only fallen into place in the last six months. These capabilities represented a change in qualitative step and not a gradual improvement. With this pace, tools that were considered unreliable six months ago that are now dependable, require investors to continuously adjust their level of trust in AI.

4.4.3 *Accountability, governance and evolving delegation boundary*

One universal norm that was enforced and held was that the use of AI did not transfer professional responsibility from humans. The person responsible for delivering an investment analysis is accountable for the work regardless of the tools used to produce it. P2 said directly:

“I ‘ve seen notes that conclude in ways I disagree with, and I said: I understand you used AI, but you need to think critically. I don’t care what tools you used, you are responsible for what you deliver to me.”

A related governance concern was raised by P3, regarding the storage and recording of investment committee information in AI systems. Introducing an AI agent into investment committee discussions meant that all information mentioned in the meeting would be captured. The concern raised was who would have access to this information, and this highlights how institutional governance is lagging behind technical advancement. The goal for P3 was to eventually have AI as a member of the investment committee with its own voice and give formal recommendations.

Despite these constraints, every participant stated trust as a trajectory pointing toward the expansions of AI authority. P5 described that for now AI is used more as a tool but over time investors would delegate more to AI as it becomes more confident and secure. P4 also anticipated that one would lean more on AI 's analytical assessments as the system expanded. Looking at theme 4 as a whole, it represents a picture of trust that is both direction and cautious. It is careful regarding the current limitations, but building towards an expansion with AI authority as governance and track records continue to develop.

5. Discussion

The discussion chapter interprets the findings from Chapter 4, highlighting the theoretical frameworks established in the literature review. The study is guided by three research questions, and this chapter will be structured around the RQs, each question touching on multiple themes. The aim of the chapter is not to repeat the results but to focus on what these findings mean. What they confirm, what they challenge and what they contribute to the existing literature framework on investment decision-making, AI and AI-human interactions.

The data analysis resulted in four themes. Theme 1 shows that AI has become an operational infrastructure among all eight participants. Theme 2 highlights that the final investment decision-making remains a human responsibility, grounded in personal ownership, accountability and the irreducibility of forward-looking judgement. Theme 3 shows that AI is reshaping the cognitive dynamics in the decision-making process. Theme 4 displays that trust in AI is built on demonstrated performance and is governed by a universal norm that the use of AI never transfers professional human accountability.

5.1 RQ1: The role of AI in the investment process

RQ1 examines how professional investors experience the use of AI in the investment process and what division between AI and humans has emerged in practice. The findings show that AI is working as an operational infrastructure, executing analytical layered tasks that were previously performed by human resources. This directly addresses bounded rationality theory and the structural challenge that is central to decision-making in investments. Simon (1955) argued that due to temporal, informational and cognitive constraints, decision-makers cannot optimize comprehensively and will therefore rely on selective attention and simplified models. The findings show that AI can expand the scale of processing and monitoring for the performance of a single investor. The boundaries shift and new integrations emerge. P1 running fifteen daily analytical tasks across 130 positions and P7 covering information on ten companies instead of three, illustrate a structural reduction of constraints in Simon's bounded rationality theory. AI is not altering the cognitive limits in bounded rationality theory; the decision-maker still remains bounded by the constraints of information capacity, time and attention. What AI change is the scope of action. By projecting the search and evaluation tasks, AI will allow investors to direct their capacity towards forward-looking and interpretive tasks where it has a competitive advantage.

This pattern aligns with the perspective of augmentation in the literature. Miller (2018) and MIT Sloan Management Review (2017) both argue that in practice, AI will primarily function as augmentation. Instead of replacing human judgment, it will enhance the analytical capacity of a human. While the findings verify the framing, it adds empirical aspect on how augmentation operates within a direct investment environment. AI tools are positioned at the front of the investment workflow, where organization and filtering of noisy data is executed. Human investors still interpret the information and finalize the decision. This resonates with Csaszar et al. (2024) argument that there are three documented methods through which AI is reshaping strategic decision-making. Aggregation, search and representation correspond directly to the description the participants made. By integrating AI, it opens opportunities by synthesizing larger volumes of information into stable outputs and can reshape the investment context before human analysis continues. The consistency between what the investors are doing in practice and the literature on strategic decision-making is a distinct theoretical validation of this study.

The finding on the emerging division between AI and humans is interpreted through Kahneman's (2011) dual-process framework. However, AI does not execute System 2 thinking, as System 2 is considered human property. What AI does is perform analytical work that would demand System 2 processing from the investor. By executing these tasks, AI frees human cognitive capacity for processes that draw on contextual reading and tacit knowledge. Kahneman looks at this as a distinction from both System 1 and System 2, as it is neither slow and rule-based or fast and automatic. It is anchored in domain and skill. The relational, forward-looking and contextual judgment that all participants described is a part of this category. AI creates a division that is not between System 1 and System 2, but is grounded between the analytical tasks that can be projected. Investors are still making relationship-based, contextual and forward-looking judgements that cannot be systematized. This is supported by Hemmer et al. (2025), who identify this as a basis of AI-human complementarity. In teams where AI completes tasks on detecting large-scale patterns and humans interpret, this will outperform both humans and AI alone. The finding shows that the split of AI and human task performance has evolved informally and not been planned explicitly.

This informal emergence is the gap Raisch and Krakowski (2021) highlight in the automation-augmentation paradox. Within this emerging division of labor, the distinct finding made by P8 on opportunity AI and efficiency AI adds an important nuance to the literature. It extends the categories proposed by Csaszar et al. (2024) and Mertzanis (2025) by focusing on the type of

value created by AI. This results in a strategic implication for investors where they can build a lasting competitive advantage. The Decision Intelligence framework by Pratt et al. (2023) argues that alone, data and analytics are insufficient without human contextualization and interpretation. The findings highlight that this is a practice operating without formal frameworks to enforce it. The participants had already developed this division through practice and not guidelines. Pratt et al's framework is therefore more valuable when firms standardize and enforce what the most experienced individuals are already doing.

5.2 RQ2: How does the human-AI interaction shape the decision-making process

The second research question addresses how the nature of AI-human interaction shapes the investment decision-making process. It examines what practices investors develop to manage the risk reinforced by AI, rather than challenge their existing judgment.

On the improving side, the expansion of information AI has directly addressed the availability heuristic that Tversky and Kahneman (1974) identified. In investment judgment, this is considered one of the most consequential sources of systematic error. Investors who rely on cognitive accessible information systematically under-investigate the full range of risk and opportunities. By monitoring a broader universe of information, AI disrupts this pattern by surfacing signals that would have fallen outside the investor's attention. P5 captured this potential directly, saying that AI is functioning at a starting point and bringing out perspectives not considered by investors. AI, therefore, has the potential to capture debiasing. Hasan et al. (2023) state that AI can manage behavioral biases by providing greater systematic analytical insight in financial decision-making. The findings from the data support this empirically, but with a critical qualification.

That critical qualification is identified by P2 as the sycophancy risk, where AI language models tend to create outputs aligned with what the users want. This is a structural feature of how models respond to user framing, creating a new technological form of confirmation-bias amplification. Nickerson (1998) defines confirmation bias as operating not only in how people interpret information but also in how they search for it. Rather than looking for disconfirming evidence, people structure their information search to get confirming evidence. Jones and Sugden (2001) present that this biased search pattern is sturdy even if disconfirming information is available. When investors prompt in a confirming way, and AI performs an equally confirming analysis, the system becomes an efficient tool for this biased search strategy. Since the output appears like a well-reasoned, objective analysis it is harder to detect the

distortion than when confirmation bias is confined to the investor's own thought process. Evidence from Ziaei and Charness (2020) states that this tendency remains when decision-makers have strong incentives to correctly update. This means that AI does more than reflect on existing confirmation bias. It can deepen it by presenting in a detailed analysis, which creates a particular risk for AI- assisted investment decisions.

Depending on the prompting strategy used, the same AI infrastructure that can reduce availability bias can also amplify confirmation bias at the same time. This is the central theoretical contribution of this study to the literature on behavioral finance: AI does not have a directional fixed effect on cognitive bias. The effect is entirely dependent on how AI is used. The countermeasures are described by the participants: P1's fresh context-analysis, P2's devil's advocate prompting, P4's bias detection module and P8's six-hats virtual board. These represent a practical response to a risk that has not yet been systematically addressed by existing literature. In addition, it connects to the concern regarding analytical skill atrophy raised by both P8 and P3. If sustained use of AI undermines the critical thinking capacity needed in order to make a responsible use of it, the foundation of potential debiasing is undermined. This potential risk is not addressed in existing literature on human-AI collaboration and should be further explored.

5.3 RQ3: Trust, accountability and the boundaries of AI delegation

The third research question examines how professional investors manage and build trust in AI systems and how norms of human accountability are shaping boundaries of AI's role.

Glikson and Woolley (2020) argue that appropriate trust calibration is essential in AI-human interaction. It is about not giving unconditional acceptance or rejection of the outputs, but to base it on the demonstrated performance. The findings confirm this model empirically and align with Glikson and Woolley's argumentation. The participants described several different trust profiles where all of which are dependent on the domain expertise of the user. The observation made by P1 on how well AI works if you performed the analytical work yourself before using it is an accurate description of the dependency Glikson and Woolley identified. Having an effective AI-human collaboration requires humans to be able to identify and evaluate the value of the AI outputs. Without this capability, the structure and confident character of AI outputs would become more dangerous than useful.

Hallucination was described as the main trust barrier, with the concern of undetectability in frequent errors. The undetected errors can result in overconfidence rather than leading users to

recalibrate their trust. This concept can be connected to Dietvorst et al. (2015) documentation of algorithm aversion, the tendency users have to become negative towards AI systems when observing that errors occur. Reich et al. (2022) disclose that when AI demonstrates transparency and adaptability, trust can be rebuilt. This is directly aligned with the emphasis that the participants placed on source citation as the primary signal of credibility. Outputs that are traceable and can independently be verified are factors that allow for trust to form and persist. The findings from all participants on AI integration at an early stage and not as production-ready also show that trust in AI is developed on limited track records. In an investment context capabilities are rapidly advancing, and any trust calibration needs to be revisited continuously.

The universal maintenance of authority in decision-making among the participants directly connects to Mintzberg et al. (1976) characteristics of the selection phase as a human act. Accountability and responsibility cannot be reduced to an analytical output alone. Kahneman and Tversky (1979) prospect theory gives a deeper explanation for why this norm is persistent even when there is improvement in AI capabilities. Prospect theory demonstrates that the psychological impact of an investment gone wrong is felt much harder than the pleasure of an equally large gain. Applied to investment decision-making, this asymmetry has a direct connection to the ownership of the decision. When loss occurs, the burden is the heaviest for the individual who made the decision. If the decision was delegated to AI, the investor would lose the authority but still preserve financial exposure. Prospect theory predicts that this combination is repellent: having the burden of the consequence without owning the final choice leading to the outcome. This aligns with all eight participants insistence on that human authority should make the final decision. Investors retain the final decision because owning the choice is what makes the asymmetry of loss a psychological impact.

AI can contribute to inform and strengthen the conviction, but it cannot carry the responsibility and emotional burden for the ownership of the decision. The accountability norm that P2 stated operationalizes this principle at an organizational level. Pratt et al. 's (2023) decision intelligence framework argues that AI should rather enhance than substitute for human reasoning. The present findings show that in practice, users are enforcing this principle even without formal governance frameworks. P3's ambition for having an AI member take part in investment committee meetings represents this trajectory and the breaking point of existing theory. When AI is actively entering the deliberative process, the questions raised are how one should weigh the input, how errors should be attributed and what accountability and governance

frameworks should apply cannot be answered by existing theory. This the limit that P3's trajectory describes as urgent.

The automation-augmentation paradox by Raisch and Krakowski's (2021) draws this together. The paradox states that an AI system can simultaneously automate human tasks and augment human capabilities. Across the participants, AI is functioning as an augmentation. According to Raisch and Krakowski's (2021), the universal maintenance of authority in human decision-making is the resistance to automation drift. However, P3 description is the precise drift of this paradox predicts. In itself, each extension of AI use is reasonable, but together they can shift decision authority that existing frameworks are not designed for.

5.4 Contributions and future research

This study is contributing to existing literature by presenting three contributions. Firstly, it provides qualitative empirical evidence of how professional investors use AI in practice. It is addressing the gap Bartram et al. (2020) identified within human-centered accounts of AI-supported investment decisions. The findings from the study document AI as operational infrastructure rather than a tool, a characterization that's not previously established in literature. Secondly, the study extends the literature on AI sycophancy into decision-making in the investment context. Building on Nickerson (1998), Jones and Sugden (2001), and Ziaei and Charness (2020), this study shows that in high-stakes investments, AI works as an amplifier to confirmation bias. This effect depends on the conditional prompting discipline among investors. Thirdly, one participant introduced a distinction between efficiency AI versus opportunity AI. Even though this is an observation from a single participant, it offers a framework that extends the work of Mertzanis (2025) and Csaszar et al. (2024). Whether this distinction in the literature holds across a broader sample, and whether opportunity AI creates a larger competitive advantage than efficiency AI, are questions that need to be addressed by future research.

Beyond the theoretical contributions from this study, the findings carry multiple implications for practitioners. For investors who work individually, prompting discipline must be treated as a core competency and not as an optional refinement. Sharma et al. (2023) document that sycophancy is a structural property of LLMs: outputs will systematically align with the user's perceived preferences rather than with accuracy. Context analysis, devil's advocate, and counter-perspective tools are described by the participants as countermeasures and as ways to operationalize this discipline in practice. Hemmer et al. (2025) confirm that the gap between realization and complementary potential is dependent on how the interaction is structured.

Investors who integrate these practices into standard workflows are more likely to realize AI's potential for debiasing while avoiding the risk of amplification.

For investment firms, the governance frameworks should already be established before they are needed. The human-AI division of labor in this study emerged through individual practice without formal structure. This is exactly what Pratt et al. (2023) identified as a condition for producing inconsistency across teams. Raisch and Krakowski's (2021) automation-augmentation paradox help explain P3's path toward having an AI investment committee member. Each little step of AI integration may seem sensible on its own, but together they gradually shift real decision-making authority into frameworks not designed to handle these decisions. For the designers of AI tools, confidentiality of infrastructure and source traceability are requirements for using them, not optional. All participants described source citation as the main trust calibration, consistent with Glikson and Woolley's (2020) findings, it is essential for the interaction between humans and AI.

In addition, the distinction between efficiency AI versus opportunity AI is creating implications for the investment industry. This is grounded in competitive advantage among investment firms, as it depends on the type of AI used. The advantage of efficient AI is likely to diminish as similar tools become more widely available at the aggregation stage of analysis (Csaszar et al., 2024). Firms that build their own domain-specific opportunity AI will have a real competitive advantage, as it is connected to data and knowledge that others cannot easily copy. The market-efficiency concern P7 raised reinforces this point. If most investors integrate similar efficiency AI, the market may become more efficient as their decisions converge. Generating alpha will be increasingly harder in the investment environment.

5.5 Methodological limitations

Throughout this study, several methodological limitations should be acknowledged. Firstly, this study is based on data collected through purposive sampling of eight participants. They all have similar professional backgrounds within direct investment firms, and are primarily in the UK and Scandinavia. The small sample reflects both the exploratory qualitative design and the field of study. Even though this sample is suitable for an exploratory qualitative study, the findings will be limited when compared with other investment firms and contexts. Therefore, it is important to acknowledge that the findings need to be interpreted in a specific context, rather than a general representation. Secondly, since the participants were collected from an existing network with help, there is a risk of selection bias. Participants who agreed to take part in the

interview process, might be more willing to share their experiences than others. In addition, the study was conducted during a change of AI capabilities, which means that the findings will represent a transitional snapshot rather than a stable state.

6. Conclusion

This study was conducted to explore how professional investors within investment firms use artificial intelligence in the investment process and how they maintain human decision authority.

The study is examining AI as it is experienced in practice and not as an isolated technical system: as a set of tools that's embedded in analytical work, in AI-human interaction and within norms of accountability and trust. Drawing on data collected from eight in-depth semi-structured interviews with senior investors in UK and Scandinavia, the study resulted in concrete findings. AI does not replace human judgment in the investment decision-making process, but it reconfigures how behavioral biases, governance and bounded rationality unfold in the investment context.

First, the findings demonstrate that AI has become the primary information channel in the investment process and functions as a co-analyst and not as an autonomous decision-maker. All the participants described how AI would perform tasks that previously required substantial human resources. In terms of bounded rationality and Simon's (1955) argumentation, findings conclude that AI does not move constraints but shifts the boundaries. This by allowing candidate opportunities and more data to flow through the decision pipeline, leaving final commitment to human judgement. This empirical finding supports the literature perspective on augmentation and how the labor division between AI and humans emerges in an investment context. This finding maps the into the dual process theory of Kahneman (2011) and Hemmer et al.'s (2025).

Second, the study reveals that the interaction between AI and humans alters the cognitive dynamics of investment decision-making. AI addresses Tversky and Kahneman (1974) availability heuristic by providing information beyond what cognitively limited sources could do. In addition, the sycophancy risk creates a new form of confirmation bias. When there are existing frames and AI responds with confirming outputs to this frame, it creates an instrument for the biased search strategy documented by Nickerson (1998) and Jones and Sugden (2001). The study contributes to the literature with the conclusion that it is not a fixed effect AI has on cognitive bias, as it depends on the users prompting discipline and engagement.

Third, the study documents the expansion of AI's analytical role and how investors build trust and how accountability shapes boundaries in AI. Trust is demonstrated and earned through performance, ability to verify outputs and domain-specific reliability. The credibility of source

citation in the outputs is also consistent with Glikson and Woolley (2020) documentation on conditions of effective human-AI collaboration. The humans remain the authority of the decision making, a conclusion that connects to Mintzberg et al. (1976) on the human act during the selection phase and Kahneman and Tversky (1979) prospect theory.

Together, these findings add three main theoretical contributions. Theoretically, the study is extending the bounded rationality framework and strategic decision-making theory into AI-rich investment environments. AI is transforming the investment process from within, but the decision-making remains human. Finally, it adds theory to behavioral finance by clarifying mechanisms where AI can affect key biases in the investment decision process. A clear understanding of where the use of AI is a contribution, where it ends and where human resources still remain superior is what this thesis has illustrated. The practical contribution of this study is to show that source verification, prompting discipline and accountability rules should be treated as core capabilities for a responsible use of AI in the investment process. The governance structure needs to define where AI can be relied upon, where human judgment is required and how outputs are verified before influencing the decision-making.

7. Reference List

- Araujo, T., Helberger, N., Kruijkemeier, S., De Vreese, H., C. (2020). In AI we trust? Perceptions about automated decision-making by artificial intelligence. *ResearchGate, AI & Society*, 35(3), 611–623. <https://doi.org/10.1007/s00146-019-00931-w>
- Bansal, G., Nushi, B., Kamar, E., Lasecki, W. S., Weld, D. S., & Horvitz, E. (2019). Updates in human-AI teams: Understanding and addressing the performance-compatibility trade-off. *In proceedings of the AAAI Conference on Artificial Intelligence*, 33(1), 2429-2437.
- Barberis, N., & Thaler, R. (2003). Chapter 18: A Survey Of Behavioral Finance. *Handbook of the Economics of Finance*, 1, 1053–1128. [https://doi.org/10.1016/S1574-0102\(03\)01027-6](https://doi.org/10.1016/S1574-0102(03)01027-6).
- Bartram, S. M., Branke, J., & Motahari, M. (2020). *Artificial intelligence in asset management*. CFA Institute Research Foundation, 28 August 2020 – Business & Economics, 96 <https://www.cfainstitute.org/>
- Brinkmann, S & Kvale, S. (2015) Interviews: Learning The Craft Of Qualitative Research Interviewing (3rd ed.). Sage Publications, Thousand Oaks CA.
- Braun, V., & Clarke, V. (2019). Reflecting on reflexive thematic analysis. *Qualitative Research on Sport, Exercise and Health*, 11 (4), 589–597. <https://doi.org/10.1080/2159676X.2019.1628806>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77-101. <https://doi.org/10.1191/1478088706qp063oa>
- Bryman, A. (2016) Social research methods (5th ed.). Oxford University Press.
- CFA Institute. *Portfolio Management: An Overview*. <https://www.cfainstitute.org/>
- Creswell, J. W. (2013) Research design (4th ed.) Sage.
- Csaszar, F. A., Ketkar, H., & Kim, H. (2024). Artificial intelligence and strategic decision-making: Evidence from entrepreneurs and investors. *Strategy Science*, 9 (4), 322–345. <https://doi.org/10.1287/stsc.2024.0190>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them error. *Journal of Experimental Psychology: General*, 144(1), 114-126. Doi :[10.1037/xge0000033](https://doi.org/10.1037/xge0000033)

- Eisenhardt, K. M., & Zbaracki, M. J. (1992). Strategic decision making. *Strategic Management Journal*, 13(Issue S2), 17–37. <https://doi.org/10.1002/smj.4250130904>
- Gioia, D. A., Corley, K. G., & Hamilton, A. L. (2013). Seeking qualitative rigor in inductive research: Notes on the Gioia methodology. *Organizational Research Methods*, 16(1), 15-31. <https://doi.org/10.1177/1094428112452151>
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14 (2), 627–660
- Guest, G., MacQueen, K. M., & Namey, E. E. (2012). *Applied thematic*. Sage Publications.
- Gupta, P., Nguyen, T. N., Gonzalez, C., & Woolley, A. W. (2023) Fostering Collective Intelligence In Human-AI Collaboration: Laying the groundwork for COHUMAIN. *Topics in Cognitive Science*, 17(2),362-379. DOI:[10.1111/tops.12679](https://doi.org/10.1111/tops.12679)
- Hasan, Z., Vaz, D., Athota, V. S., Désiré, S. S. M., & Pereira, V. (2023). Can Artificial Intelligence (AI) Manage Behavioral Biases Among Financial Planners? *Journal of Global Information Management*, 31(2), 1–23. DOI:[10.4018/JGIM.321728](https://doi.org/10.4018/JGIM.321728)
- Hemmer, P., Schemmer, M., Kühl, N., Vössing, M., & Satzger, G. (2025) Complementarity in human-AI collaboration: concept, source and evidence. *European Journal of Information Systems*, 35(6), 979-1002. <https://doi.org/10.1080/0960085X.2025.2475962>
- Hirshleifer, D. (2015). Behavioral finance. *Annual Review of Financial Economics*, 7, 133–159. <https://doi.org/10.1146/annurev-financial-092214-043752>
- Hoang, D., & Wiegatz, K. (2023). Machine learning methods in finance: Recent applications and prospects, *European Financial Management*, 29(6), 1443–1477. <http://dx.doi.org/10.2139/ssrn.4293977>
- Jones, M., & Sugden, R. (2001). Positive confirmation bias in the acquisition of information. *Theory and Decision*, 50 (1), 59–99.
- Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.
- Kahneman, D., & Tversky, A. (1979). Chapter 6: Prospect theory: An Analysis of Decision Under Risk. *Econometrica*, 47(2) 263-292. The Econometric Society. <https://doi.org/10.2307/1914185>

- Mertzanis, C. (2025). Artificial intelligence and investment management: Structure, strategy, and governance. *International Review of Financial Analysis*, 107, 104599. <https://doi.org/10.1016/j.irfa.2025.104599>
- Miller, S. M. (2018). AI: Augmentation, more so than automation. *Asian Management Insights (Singapore Management University)*, 5 (1), 1-20.
- Mintzberg, H., Raisinghani, D., & Théoret, A. (1976). The Structure of “Unstructured” Decision Processes. *Administrative Science Quarterly*, 21 (2), 246–275. <https://doi.org/10.2307/2392045>
- MIT Sloan Management Review. (2017). Augmentation versus automation: AI’s utility in the workplace. *MIT Sloan Management Review*, 58(4), 28–35.
- Miles, M. B., Huberman, A. M., & Saldaña, J. (2014). *Qualitative data analysis: A methods sourcebook* (3rd ed.). Sage Publications.
- Nickerson, R. S. (1998). Confirmation bias: A Ubiquitous Phenomenon in Many Guises. *Review of General Psychology*, 2 (2), 175–220. <https://doi.org/10.1037/1089-2680.2.2.175>
- Nowell, L. S., Norris, J. M., White, D. E., & Moules, N. J. (2017). Thematic analysis: Striving to meet the trustworthiness criteria. *International Journal of Qualitative Methods*, 16, 1-13. <https://doi.org/10.1177/1609406917733847>
- Patton M. Q. (2002) *Qualitative Research and Evaluation Methods*. (3rd ed.). Sage Publications, 688, Thousand Oaks, CA.
- Patton M. Q. (2015). Sampling, qualitative (purposeful). *The Encyclopedia of Social Measurement* (online ed). Wiley. <https://doi.org/10.1002/9781405165518.wbeoss012.pub2>
- Pratt, L., Bisson, C., & Warin, T. (2023). Bringing advanced technology to strategic decision-making: The Decision Intelligence/Data Science (DI/DS) integration framework. *Future*, 152, 103217. <https://doi.org/10.1016/j.futures.2023.103217>
- Pu, Y., Li, H., Hou, W., & Pan, X. (2025). The analysis of strategic management decisions and corporate competitiveness based on artificial intelligence. *Scientific Reports*, 15, 17942. <https://doi.org/10.1038/s41598-025-02842-x>
- Rabin, M., & Schrag, J. L. (1999) First Impression Matters: A Model of Confirmatory Bias. *The Quarterly Journal of Economics*, 114(1), 37-82. <http://www.jstor.org/stable/2586947>

- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), 192–210. <https://doi.org/10.5465/amr.2018.0072>
- Reich, T., et al. (2022). How to overcome algorithm aversion: Learning from mistakes. *Journal of Consumer Psychology*, (4), 675–695. DOI:10.1002/jcpy.1313
- Sharma, M., Tong, M., Korbak, T., Duvenaud, D., Asbell, A., Bowman, S., Chowdhery, A., Durrett, G., Hatfield-Dodds, Z., Steinhardt, J., & Perez, E. (2023). Towards understanding sycophancy in language models. *arXiv preprint arXiv:2310.13548*.
- Simon, H. (1955). A behavioral model of rational choice. *Quarterly Journal of Economics*, 69(1), 99–118. <https://doi.org/10.2307/1884852>
- Tversky, A., & Kahneman, D. (1974) Chapter 22: Judgement under uncertainty: Heuristics and biases. *Knowledge: Critical Concepts*, (1). Science, 185 (4157), 1124-1131.
- Ziaei, S., & Charness, G. (2020). Confirmation bias with motivated beliefs. *Experimental Economics*, 23 (2), 421–447.

Appendices

Appendix A: Interview Guide

Disclaimer: The interview guide was used flexibly, and not all questions were asked in every interview as the wanted information was obtain in other parts of the interview. Follow up questions were also adapted to the different responses of the participants.

Introduction

For my master thesis I am conducting in depth, semi-structured interviews in order to get a better insight in how artificial intelligence works as a tool for decision-making in investment companies. I am looking at companies that works with direct investments within different fields and how they are using AI within decision making.

In this interview I will use some terminology related to the research field, and I will give a brief introduction to the different concepts.

The decision-making process refers to the sequence of steps an individual or team goes through from recognizing problem or opportunity, gathering and interpreting information, forming and comparing options and choosing an action. In practice this process is influenced by both analytical considerations (data, models, forecasts) and by human factors such as experience, intuition, time pressure and cognitive bias.

Human bias: refers to systematic tendencies in how people think and decide that can lead them away from fully rational or objective choices. Cognitive bias arises because our brains use mental shortcuts (heuristics), emotions and past experiences to simplify complex information, which can be helpful but also introduce predictable errors in judgement and risk assessment.

Heuristics: A heuristic is a simple decision rule or mental shortcut that people use to make judgments or choices quickly and with limited information, time, and cognitive resources.

Part One: Introduction

1. Can you briefly explain which company you work for, what investments your company does and your role?
2. Can you share some of the criteria your firm has for investment size, geographic area, industries, technology, and funding stage/maturity?

3. Can you explain how you/your team typically assess a new investment opportunity and the process leading up to a decision?

Part Two: Challenges in decision-making processes

1. What are the key challenges your company faces in decision-making?

Part Three: AI integration

1. How do you use AI or data-driven tools in your investment efforts?
→ For this question, I want you to think about an investment that has already happened, in what steps did you use AI tools and for what.
→ What kind of tools or systems do you use – what does the tool do for you?
2. Has the use of AI changed the process/steps you use in an investment context – and if so, how?
3. What benefits and opportunities have you observed with the implementation of AI in decision-making?
→ How has AI changed the way you work compared to before?
4. Do you use AI to replace human skills or to enhance your own abilities, and how do you do this?
5. How does AI affect your people skills or abilities? – do you have any examples you can point to – an investment decision
6. What issues did you/your team encounter when implementing technology into your investment decision-making process?
→ Feel free to give examples of an investment and in which step this was most demanding and why?
→ What are the main concerns of investors using this form of technology?
7. How do you typically combine results from these tools with your own assessment or your team's view?

Part Four: Trust and Biases

1. What factors affect your level of trust in the AI systems and tools you use?

2. How does the use of original AI as a tool affect human bias in decision-making?

Closing

To wrap up, I'd like to ask if there's anything you'd like to add that we haven't discussed yet. Is there a topic or an important aspect of the investment process that we haven't covered that would be worth mentioning?

Appendix B: Table 1 – Three-Level Coding Tree

Reflexive thematic analysis; Braun & Clarke (2006, 2019)

Analytical Stage	T1 AI as co-analyst	T2 Human judgement	T3 Human-AI collaboration	T4 Trust and accountability
<i>RAW INTERVIEW DATA</i>				
CYCLE 1 — DESCRIPTIVE CODES · What participants explicitly said				
1	"AI scrapes all external information daily"	"Never take a decision based on AI alone"	"Model tries to figure out what you want"	"Still in early experimental phase"
2	"Covers 10 companies instead of 3"	"You are responsible for what you deliver"	"Running analysis in fresh AI context"	"Trust grows as AI proves accuracy"
3	"15 daily automated tasks in Cowork"	"Cannot be replaced by AI tools"	"Bias detection module to check own decisions"	"Hallucinations limit full trust"
CYCLE 2 — CONCEPTUAL CODES · Underlying meaning and pattern				
<i>Conceptual code</i>	AI as information aggregator	Human final decision authority	Sycophancy as bias amplifier	Track-record-based trust calibration
CYCLE 3 — RELATIONAL CODES · How patterns connect across the dataset				
<i>Relational code</i>	Information overload → AI as analytical filter → Decision input	AI tool use → Human ownership → Accountability norm	Sycophancy → Confirmation bias risk → Prompting discipline	Performance history → Trust calibration → Delegation boundary
FOUR OVERARCHING THEMES · 37 codes total				

Theme	AI as a co-analyst and operational infrastructure <i>12 codes · all 8 participants</i>	Human judgement, responsibility and the final decision <i>11 codes · all 8 participants</i>	Changing human-AI collaboration <i>10 codes · 6–8 participants</i>	Trust, accountability and the boundaries of AI delegation <i>9 codes · all 8 participants</i>
THREE RESEARCH QUESTIONS ANSWERED				
Research question	RQ1 · Division of labor between AI and humans <i>T1 + T2 · how AI is used in the investment process</i>		RQ2 · Cognitive dynamics <i>T2 + T3 · bias risks and prompting discipline</i>	RQ3 · Trust and accountability <i>T4 · boundaries of AI delegation</i>
EMPIRICAL FINDINGS AND THEORETICAL CONTRIBUTIONS				
AI as infrastructure · sycophancy–confirmation bias link · automation drift · trust as trajectory				

Appendix C: Table 2 – Boundary definitions of codes

Theme Name	Boundary Definition	Inclusion Criteria	Exclusion Criteria	Associated Themes
Theme 1 – AI as a co-analyst and operational infrastructure	Captures how participants describe AI as an analytical assistant and infrastructural tool embedded in the investment workflow to handle information-intensive tasks, while final judgement remains human.	AI used for screening opportunities, document summarisation, information gathering, financial modelling, portfolio monitoring, and routine analytical tasks. Descriptions of efficiency gains, scalability, “superhuman” capacity, or AI as an “extra analyst”, “co-pilot”, “co-analyst”, or “infrastructure layer”.	Responsibility, liability, ethical concerns, or over-reliance are coded in Theme 2 or Theme 4. Changes in collaboration structure or role development over time are coded in Theme 3.	Human judgement, responsibility, and the final decision; Changing human–AI collaboration; Trust, accountability and the boundaries of AI delegation
Theme 2 – Human judgement, responsibility, and the final decision	Concerns how participants position human judgement as the ultimate decision authority, and how they negotiate responsibility, liability, and the limits of AI involvement in high-stakes investment decisions.	Statements that final investment decisions must remain human, even when AI performs substantial analysis. Reflections on accountability, fiduciary duty, explainability, who “owns” the outcome, resisting full automation, overriding AI, or limiting AI’s role for responsibility reasons.	Technical descriptions of AI functions without responsibility or final authority are coded in Theme 1. Collaboration patterns and division of labour over time are coded in Theme 3. Trust calibration, algorithm aversion, or comfort/discomfort with delegation are coded in Theme 4, unless explicitly framed around responsibility.	AI as a co-analyst and operational infrastructure; Changing human–AI collaboration; Trust, accountability and the boundaries of AI delegation
Theme 3 – Changing human–AI collaboration	Captures how human–AI collaboration evolves over time, including movement from ad-hoc experimentation to structured use, and the distinction between “efficiency AI” and “opportunity AI”.	Narratives about how AI’s role has changed in the organisation or participant’s own practice. Descriptions of team coordination, workflow adjustment, AI in meetings or investment committees, and when AI is used for exploration, pattern detection, thesis building, or opportunity identification rather than routine tasks	Trust, comfort, distrust, or reliance calibration are coded in Theme 4 unless focused on collaboration structures. Pure technical capability or efficiency descriptions are coded in Theme 1. Responsibility and final authority are coded in Theme 2.	Trust, accountability and the boundaries of AI delegation; AI as a co-analyst and operational infrastructure; Human judgement, responsibility, and the final decision
Theme 4 – Trust, accountability, and the boundaries of AI delegation	Examines how investors build, calibrate, and question trust in AI systems, and how human accountability shapes the boundaries of what can be delegated to AI in investment processes.	Statements about trust output, ethics, AI output, risk, workflow, trust calibration, delegation boundaries, regulation, governance, data security, reputational risk, or giving AI a	Workflow or efficiency descriptions without trust or delegation are coded in Theme 1. Evolving collaboration patterns without explicit trust or boundary-setting are coded in Theme 3. Personal responsibility and final decision authority without trust/delegation	AI as a co-analyst and operational infrastructure; Changing human–AI collaboration; Human judgement, responsibility, and the final

Appendix D: Key findings and Summaries from Participant Transcripts

Disclaimer: To support transparency in the analysis, Appendix D provides summaries from each participant's transcript. Full transcripts are in Norwegian, translated into English, and are not included due to appendix length constraints.

Participant 1.

P1 manages a highly diversified portfolio of over 130 positions, including listed equities, hedge funds, PE, credit, and direct investments, with only two investment professionals. This volume of positions per person was a recurring point of reference throughout the conversation. The sheer breadth of what one person must monitor shapes almost every other observation P1 makes about investment work.

P1 describes AI as indispensable for managing this volume. Without it, two or more additional analysts would have been required. AI now runs approximately 15 automated daily tasks across the portfolio, including scraping news, X, and the inbox, and delivering a daily briefing email on every position. AI reads earnings transcripts, aggregates broker research, and models quarterly performance against consensus, all automatically.

The investment process itself has changed. Previously, P1 built Excel models manually. Now Claude produces a first-pass valuation model and a three-to-ten page analysis note. For a small French mid-cap company with no analyst coverage, P1 loaded the quarterly presentation, press release, and a live Bloomberg transcript of the investor call into Claude and asked it to identify where management contradicted prior statements, whether growth had been delayed, and how cashflow was tracking. This took a fraction of the time it would have previously.

P1 uses Claude's project feature as a persistent memory for each company, all prior notes, theses, and exchanges are retained in context. When new information arrives, Claude relates it to everything previously documented. P1 describes this as an extension of his own mind, giving him capacity that would otherwise require a team.

On the question of errors: P1 is explicit that blind use of AI is dangerous. It can hallucinate, implement new information in parts of an analysis while forgetting to update others, or miss an income line in a cashflow model. The safeguard is domain expertise, knowing what to check for because you have done the underlying work yourself. P1 uses the analogy of having a junior analyst: you trust them more as they demonstrate accuracy, but you never remove oversight entirely.

P1 was direct that the final investment decision is always his own. AI does not decide. The discipline of living with an investment, of being accountable for its outcome, requires personal ownership of the choice. This is described not as a limitation of AI but as a principled norm.

P1 raises concern about the future of junior talent in the industry. The analytical work that junior professionals traditionally used to build their skills is now done by AI, better and faster. This creates a skills pipeline problem that the industry has not yet resolved.

Participant 2.

P2 is the CEO of a listed Norwegian industrial investment company with six core investments spanning software, hardware, and energy. The company operates through owner teams, each executive and two to three other professionals share responsibility for a portfolio company, sit on its board, and manage its M&A pipeline. P2 himself is an active participant in investment work, not only a manager of people who do it.

AI use at the firm spans the full investment process: screening large company registers, filtering down to a shortlist, producing one-page memos, supporting due diligence on transaction data, and most recently producing financial models. P2 describes feeding a complete financial due diligence report into Claude the day before the interview and receiving a working valuation model as output, something that was not possible six months earlier.

P2 identifies the primary value of AI as processing large volumes of information quickly: 200 pages of due diligence reports, summarized and interpreted. The job this replaces is what he calls central junior tasks, the analytical groundwork that previously occupied analyst's full time.

The most striking part of the interview is P2's description of the sycophancy risk. He observed that AI models are oriented toward pleasing the user, they infer what the user wants and produce outputs that align with it. In an investment context, where the user often has a prior view on a deal, this means the model becomes an efficient instrument for confirming that view rather than challenging it. P2's countermeasure: he explicitly instructs AI to take the position of devil's advocate on the investment committee and write a memo arguing why the deal should not proceed. This reliably produces the counter-perspective.

P2 is categorical about accountability: the person who delivers an analysis owns it, regardless of what tools they used. He describes seeing team members submit AI-generated notes and justify their conclusions by saying "Claude said so", which he rejects entirely. Claude cannot

be held responsible, people can. This is framed as both a personal principle and an organizational norm he actively enforces.

On trust: P2 ties it to source citation. AI outputs that specify the sources for every claim increase trust. Outputs that do not are treated with skepticism. He also notes that qualitative assessments remain a hard ceiling, AI cannot assess cultural fit, integration potential, or what a management team will do under pressure, because these require contextual understanding that no amount of data can substitute for.

P2 does not expect to delegate final investment decisions to AI in the near term. He draws an analogy with a new colleague: you delegate more as trust builds through demonstrated accuracy, but you do not start there.

Participant 3:

P3 is an institutional partner at one of the world's largest private equity firms, with approximately 300 portfolio companies across 25 funds. He was CEO from 2019 to 2025, during which the firm went public and expanded globally. He now sits on three investment committees and leads a group of senior IC members that meets to reflect on decision patterns.

The firm uses Claude, ChatGPT, and Gemini across deal teams, with Claude increasingly dominant for internal use. The firm has a proprietary platform called Mother Brain that structures and monitors portfolio data. Investment committee materials are beginning to be produced with AI assistance, a recent development P3 describes as still requiring quality assurance but close to ready.

Before IC meetings, P3 loads all presentation material into NotebookLM and uses it to develop a rapid high-level understanding of the company, asking questions, generating summaries, and effectively receiving a briefing without requiring meeting preparation time from the deal team. He describes this as almost like getting a pitch without meeting management.

The firm has been recording investment committee discussions for several years with the aim of identifying patterns in decision-making. The dataset is not yet statistically significant. What P3 is now developing takes this further: every IC member and deal team member will receive an AI agent companion. The agent listens to IC discussions, analyses the content, and eventually makes a recommendation, yes or no, at what price. The goal, as P3 describes it, is for this to function as an AI investment committee member. He anticipates the same logic eventually extending to board membership in portfolio companies.

P3 is candid about the governance complexities this creates. IC discussions contain sensitive, informal, and sometimes legally significant statements. Everything getting recorded creates new questions about confidentiality, data handling, and regulatory compliance that the firm is still working through.

He raises a concern about passive acceptance: AI makes it easy to consume outputs, briefings, summaries, podcast-style explanations, without doing the underlying work. For routine decisions this may be acceptable, but for serious investment decisions he believes this creates intellectual shortcuts that erode preparation quality.

P3 uses Claude and ChatGPT in parallel on the same tasks to compare outputs. He typically finds 75–80% overlap but notes that one model is often sharper on the specific point he cares about. This parallel use has become standard practice within his firm.

Participant 4:

P4 founded and solely manages a Nordic listed equity fund, originally created for a family following a generational transition. The fund has since taken in external families and P4 runs it entirely alone, investment decisions, investor relations, regulatory reporting, and all analytical work. This context is central to understanding his relationship with AI: it is not a supplementary tool but the infrastructure that makes the fund viable at this level of sophistication with one person.

Over approximately a year, P4 built a bespoke two-layer AI system. The first layer is a private file system of markdown files, all notes, theses, and analyses, which he owns and which AI can read via API. The second layer connects to external data providers including Quartr, which downloads all company presentations, annual reports, and quarterly reports in structured form. This system allows AI to retrieve and analyse company data without going to the internet for lower-quality sources.

P4 estimates that 90% of his work is now conducted inside AI. Claude is embedded in Excel, PowerPoint, and Word. AI reads all his emails and sends replies when instructed. He speaks to his phone to schedule meetings and dictates requests. All calendar management is AI-handled.

For investment analysis, P4 describes a 90-minute iterative session on Novo Nordisk as a representative example, forty questions asked and refined across the session, producing an analysis document that would previously have taken two full days. The document was then sent

to external analysts to verify. This verify-then-decide workflow is consistent: AI produces the starting point, external expertise checks it, P4 decides.

P4 has built a dedicated bias-detection module specifically to challenge his own prior conclusions. He also uses separate fresh-context AI sessions, with no prior framing or preloaded information, when he wants a genuinely uncontaminated starting point for analysis. These practices reflect his view that AI quality depends entirely on how you direct it.

P4's primary concern is not hallucinations but security. He has built a system with significant data value and believes it will eventually be targeted. He expects there will be an incident nobody foresaw, and describes this as a structural risk inherent to sophisticated AI systems handling large financial positions.

All decisions remain entirely his own. AI does the analysis. He makes the calls. He draws the analogy of an always-available analyst who works at any hour and can be asked unlimited questions, but who requires domain knowledge to direct well, and whose outputs require evaluation.

Participant 5:

P5 works as an investment analyst in a small family office that manages the wealth of three siblings from a Norwegian shipping family. The office makes a small number of large, carefully chosen investments and maintains them over long periods. The team is three people working closely together. P5 contributes to both strategic active-ownership investments and more conventional financial portfolio management.

The investment process is deliberate and slow. Opportunities are followed for up to a year before a board presentation. The board is used actively to stress-test cases rather than just approve them. Because the office is essentially fully invested, any new investment requires reducing an existing position, meaning every new opportunity must be better than what is already held. This opportunity-cost framing runs through everything P5 describes.

AI is used as a complementary analytical partner. P5 describes a concrete example: comparing two bond funds on all relevant dimensions, historical performance, fees, terms, redemption conditions in order to decide whether to consolidate. Claude read through the full documentation for both funds, placed them side by side, and produced both an analysis and a board-ready investment proposal. The conclusion AI reached matched the team's view. This was described as efficient and reassuring.

P5 notes that AI finds sources and perspectives that human analysts would not independently locate, not just faster, but broader. This is identified as one of the most material benefits: the information base that feeds analysis becomes genuinely larger, not just faster to access.

On the risk side, P5 is thoughtful. AI sometimes gives imprecise or wrong answers, and this risk is higher in areas where the analyst is not an expert — there is no internal check on errors you cannot detect. The team's response is not to over-rely: AI is an extra team member, not a decision-maker, and the team retains analytical ownership of everything it produces.

P5 switched from ChatGPT to Claude and found Claude more transparent in explaining its assumptions and reasoning, which built more trust. The switch was practically motivated but had meaningful effects on how the tool was perceived and used.

P5 anticipates that delegation will increase over time as both tools and user confidence develop. The current position, AI as complement, is understood as a phase, not a permanent state.

Participant 6:

P6 works at the same family investment office as P5 and focuses primarily on Nordic listed equities, mature, large-cap companies with proven profitability. The portfolio holds a small number of positions and is managed with a long-term hold orientation. P6 and P5 work in the same room and share a team of three, but their observations about AI reflect distinct perspectives and emphases.

P6's most concrete AI example is the Nordnet investment. AI was used at the start to produce a high-level picture, plusses and minuses, an initial information aggregation. The middle stages were done manually. Now AI automatically pulls the latest reports and sends analyses. P6 observes that if the investment were being made today, AI would be integrated at every stage, not just the beginning and end.

P6 uses Claude as the primary tool and finds it qualitatively different from ChatGPT, Claude functions as a capable junior analyst who can handle complex Excel tasks, whereas ChatGPT functions more like an advanced search tool. This distinction shapes how each is used: Claude for analysis and modelling, ChatGPT for quick lookups.

The strongest statement in P6's interview concerns cognitive bias. P6 argues that AI removes subjective feelings and viewpoints that damage investment decisions. All investors know their biases exist but most still cannot reliably avoid them. AI provides structural help: an analytical process that is not shaped by mood, prior relationship with a management team, or recent

memorable experience. This is framed as one of AI's most significant contributions to investment quality.

P6 encountered the sycophancy problem personally, AI responses were reinforcing conclusions already reached rather than challenging them. P6's response was to actively push back: explicitly challenge AI when it agrees too readily and ask it to find weaknesses in the analysis. This self-correcting practice emerged from experience rather than any formal guidance.

P6 frames AI adoption as a competitive necessity rather than a choice. In three to five years, investors using better AI tools will make better decisions. Those who do not use AI will fall behind. This is described not with urgency but as a straightforward observation about how the industry is moving.

Trust, for P6, is built through AI demonstrating honesty, particularly when it says it cannot access a particular type of data rather than generating a plausible-sounding answer. This is described as the behaviour of a colleague with good values. Conversely, a single hallucinated response that presents invented numbers with confidence immediately and lastingly reduces trust.

Participant 7:

P7 is one of three investment professionals at a family office managing the legacy wealth of a Nordic shipping family. The investment portfolio centres on listed equities with a technology focus, supplemented by some private direct investments and real estate. P7 is the most technology-oriented member of the small investment team and has been the most active AI adopter.

P7 describes a progression in how AI has been used: first for simple information gathering, then for document processing using Claude's project feature, then for building multi-agent workflows using Claude Code. Each step expanded the scope of what one analyst could manage. The project feature alone allowed P7 to follow ten companies closely instead of three, a concrete change in analytical coverage with the same number of working hours.

The multi-agent architecture is technically distinctive. P7 observed that a single AI model given 100 pages begins to hallucinate. The solution is to use ten parallel agents, each assigned one section, and then ask each agent what it found most important. Synthesizing across agents produces a more accurate and comprehensive understanding than overloading a single context window.

P7 is clear that the decision itself cannot be automated. Forming a view on a company's cashflow in three to five years requires judgement, gut feeling, experience, and contextual understanding that pattern-based AI cannot replicate. The line is drawn at the point where AI processes information and humans interpret it and commit.

P7 introduces what he describes as big philosophical thoughts about market efficiency. If all investors use AI to make more rational decisions simultaneously, markets may become more efficient, which would make alpha generation harder, not easier. This is presented as a genuine concern about the long-term implications of universal AI adoption rather than a near-term worry.

He also raises a related concern about business model disruption: AI makes it harder to assess which companies will survive in their current form, because cheaper AI-enabled alternatives may emerge in any sector. Paying high prices for companies today requires higher conviction that the business model is durable, which is more difficult to achieve.

Trust is framed as a trajectory, building over time through demonstrated performance, rather than a judgment made at adoption. P7 uses the Tesla autopilot analogy: reliable on the motorway, unpredictable at the roundabout. High trust for structured verifiable tasks, active skepticism for qualitative contextual judgements.

Participant 8:

P8 has a 35-year career spanning equity fund management, institutional broking, and investment advisory. He currently advises a family office, supports an industrial holding company in Germany on acquisition strategy and due diligence, and is building a new venture addressing liquidity challenges in European private markets. He is the most experienced participant in the study and brings a perspective shaped by having seen multiple technology transitions in the investment industry.

P8 built a family office collaboration platform as a demonstration model that is now entering production. He has also built AI agents that map entire industries globally, identifying who is in a space, where they are, what they do, and how they relate to each other. For a current medical robotics investment, this industry-mapping capability provides intelligence that no manual research process could replicate at the same depth and speed. P8 describes building things he could not have imagined doing 12 months ago.

P8 introduces a distinction between two types of AI in investment: efficiency AI and opportunity AI. Efficiency AI automates existing processes — screening, aggregation, report summarization. It will commoditise because everyone will have access to the same tools, eroding any competitive advantage it temporarily creates. Opportunity AI enables entirely new investment categories, theses that could not previously be investigated because the information was not accessible or the analytical infrastructure did not exist. P8 is explicitly more interested in opportunity AI.

To manage the risk of AI confirming existing views, P8 built a virtual investment board using six characters from de Bono's six thinking hats framework. Each character represents a different analytical stance, one optimistic, one sceptical, one focused on process, one on risk, and so on. By cycling through these perspectives in AI interactions, P8 builds in systematic challenge to his own prior conclusions. This is described as a formal practice, not an occasional check.

P8 frames the relationship between human and AI capability in terms of left-brain and right-brain functions. AI is extremely good at left-brain analytical work: pattern recognition, data synthesis, logical inference. As AI takes over this analytical capacity, human professionals should develop right-brain capabilities: emotional intelligence, relationship assessment, contextual reading. He sat for six hours in a management meeting in Switzerland assessing whether to trust the people across the table, a function no model can perform.

He describes the risk of analytical skill atrophy directly: the use of AI requires actively maintaining critical thinking in the same way that a comfortable life requires going to the gym. If you stop exercising independent analytical judgement because AI always provides the answer, you erode the very capacity that makes responsible AI use possible.

Trust is primarily determined by experience and early adoption. Those who started using AI seriously a year or two ago have accumulated a compounding advantage, they understand what the tools can and cannot do, which use cases to trust and which to verify, and how to direct them effectively. Late adopters face a steeper learning curve that compounds over time in the same way investment returns compound.

Appendix E: Key findings from Participant Transcripts

Disclaimer: To support transparency in the analysis, Appendix E provides key findings from each participant's transcript. Full transcripts are in Norwegian, translated into English, and are not included due to appendix length constraints.

The key findings are organized into the four themes collected from the data findings in order to capture what each participant described within each theme.

P1: Investment Manager, Single-Family Office

P1 represents the most intensive individual AI user in the study. Managing a highly diversified portfolio of over 130 positions as a team of two, P1 has built AI into every layer of daily investment work, running approximately 15 automated daily tasks through Claude and Cowork. The interview reveals how AI transformed what a single analyst can manage, and introduces the critical caveat that blind trust is dangerous.

T1 – AI as co-analyst and infrastructure

- Runs 15 daily automated tasks via Claude/Cowork across all 130+ positions
- AI scrapes internet, X, and entire inbox produces daily portfolio briefing email per position
- Uses Claude's project feature as persistent company memory, retaining all prior context
- AI builds first-pass valuation models: "I'd have built my own Excel model before. I don't have to do that now"
- Explicitly states he would have needed 2–4 junior analysts without AI

P1: "I have about 15 daily tasks that Claude runs across all our investments in Cowork. It aggregates the information efficiently so I don't miss anything."

T2 – Human judgement and responsibility

- Unambiguous: never takes an investment decision based on AI alone
- Always writes own view first, then uses AI to pressure-test it
- Real-world example: X real estate investment, structured the analytical framework himself, then used AI iteratively across 2–3 rounds
- Describes AI as analogous to a junior analyst, you guide it, check its work, and own the output

P1: "I never make a decision based solely on AI. You have to live with the investment yourself. That's why it's important that you take the decision yourself."

T3 – Human-AI collaboration and cognitive dynamics

- Explicit example of AI reducing availability bias: ran hedge fund manager selection through a fresh Claude session with no prior bias, changed his view on who to choose
- Identifies hallucination as the core risk: "It's lethal to use it blindly if you don't know what you're doing"
- Notes AI is getting better: Claude is significantly more reliable than earlier ChatGPT tools
- Raises concern about junior talent pipeline: "the one thing AI can't replicate is the human part, going to dinner with management and reading the room"

P1: "It's lethal to use it blindly. It works very well if you've done the work yourself first, because then you know what to check for."

T4 – Trust, accountability and delegation boundaries

- Trust is task-specific and performance-calibrated: higher stakes = much higher scrutiny
- Compares AI to a junior analyst: trust builds as it proves itself right over time
- Triple-checks everything on large investment decisions regardless of AI output
- Describes domain expertise as prerequisite: without it, AI gives "false security"

P1: "The more money behind a decision, the more differently you treat it. You must know what you're doing, otherwise AI gives you false security that you're doing the right things."

P2: CEO & Investment Leader, Listed Industrial Investment Company

P2 leads a listed industrial investment company with a team-based ownership model. The interview is notable for its institutional perspective: P2 describes organizational AI governance, the sycophancy risk in the strongest terms of any participant, and establishes the clearest accountability norm, "I cannot hold Claude responsible for anything. I can hold people responsible."

T1 – AI as co-analyst and infrastructure

- AI used across all stages: opportunity identification, screening, due diligence, financial modelling
- Concrete example: fed a full financial due diligence report into Claude, and received a working valuation model the same day
- Describes AI's most valuable use as processing large document sets: "here are 200 pages, what does this mean?"
- Notes the capability step-change is very recent: modelling in Excel and PowerPoint only became viable in the past few months

P2: "Yesterday I fed a financial due diligence report into Claude and got out a working valuation model. It still needs refinement, but it saves time."

T2 – Human judgement and responsibility

- Strongest accountability statement in the dataset: ownership of AI output is non-negotiable
- Real incident: junior delivered AI-generated note P2 disagreed with, challenged team to think critically
- Rejects "I asked Claude and this was the answer" as a legitimate form of analysis delivery
- Draws clear analogy: AI use never changes who is responsible for the work

P2: "I don't care what tools you used, if you're delivering an investment analysis to me, you are responsible for it. I cannot hold Claude accountable. I can hold people accountable."

T3 – Human-AI collaboration and cognitive dynamics

- Names sycophancy explicitly: "the model tries to figure out what you really want and please you"
- Describes this as like having an advisor who always sells you on your own ideas, dangerous for investment decisions
- Counter-measure: explicitly asks AI to sit on the investment committee as devil's advocate and write a memo arguing against the deal
- Views building prompting discipline as an institutional capability organizations must develop over time

P2: "There is a tendency for the model to try to figure out what you really want, and then try to please that. That is dangerous when you are making an investment decision."

T4 – Trust, accountability and delegation boundaries

- Has not delegated a final investment decision to AI and does not see this changing soon
- Trust model mirrors colleague trust: builds incrementally as AI demonstrates accuracy
- Identifies qualitative contextual understanding as the hard ceiling: AI cannot assess cultural fit, integration potential, or senior management quality
- Notes models always answer with full confidence, which is precisely what makes them dangerous when context is missing

P2: "They will always come with an answer with full confidence. It's up to us as users to understand what there is any point in asking about, and what there isn't."

P3: Institutional Partner (former CEO 2019)

P3 represents the most institutionally advanced AI integration in the study and the single most forward-looking participant. Leading AI strategy at a global PE firm, P3 describes proprietary platforms, AI-assisted IC materials, AI companions for every investment professional, and an ambition for AI to have a formal voice in investment committee deliberations, the study's most striking frontier finding.

T1 – AI as co-analyst and infrastructure

- Proprietary platform (Mother Brain) structures and monitors all investment data
- Claude, ChatGPT, and Gemini used across deal teams for analysis; Claude taking over most internal use
- IC materials beginning to be produced with AI assistance, entirely new development
- Every IC member and deal team member to receive AI agent companion (launching May 2026)
- Loads all IC presentation material into NotebookLM before meetings for rapid pre-meeting briefing
- Recording IC discussions for years to detect decision patterns, dataset not yet statistically significant

P3: "We are now in a situation where investment committee materials are beginning to be produced with AI assistance. This is completely new and still needs quality assurance, but we are very close."

T2 – Human judgement and responsibility

- Concern about passive acceptance: "it is a bit dangerous, you can easily start to lean on all these things"
- Points to skill atrophy risk: listening to AI podcast summaries instead of reading the material may not be good enough for serious decisions
- Maintains human decision authority as norm, but is actively testing the boundary with AI IC companion
- Raises governance challenge: not everything said in IC discussions should be permanently recorded

P3: "You need to know what is extrapolated by AI versus what is fact. If you are taking a serious decision, you need to be better prepared than just listening to a podcast."

T3 – Human-AI collaboration and cognitive dynamics

- Two models used in parallel on same tasks: ~75–80% overlap, but one sharper on specific points
- Notes AI can be biased by how you prompt it, it may respond in the way it thinks you want
- Societal concern raised: consulting, law, investment banking all hiring fewer graduates already visible in numbers
- Sees AI entering collective intelligence: IC companion listens, contributes, and eventually votes

P3: "Using two models in parallel for the same task is something a lot of people are doing now. One was clearly sharper on the specific point I cared about."

T4 – Trust, accountability and delegation boundaries

- Most advanced trust trajectory: moving from tool to formal decision participant
- Ambition: AI IC member with a formal voice "the goal is that it almost becomes like having an AI investment committee member"

- Privacy and confidentiality concerns around IC recording remain unresolved, working on regulatory questions
- P3 is the frontier case: the boundary between augmentation and automation is shifting intentionally

P3: "The goal is that it almost becomes like having an AI investment committee member. And eventually a board member in the companies we own. That dimension is coming."

P4: Solo Portfolio Manager & Fund Founder, Independent Nordic Equity Fund

P4 is the most technically sophisticated individual AI builder in the study. Working entirely alone, P4 spent a year constructing a bespoke two-layer AI system connecting Claude to a private file system, APIs, email, and calendar. With 90% of all work conducted inside AI, P4 represents the furthest-reaching case of individual AI integration, and the clearest illustration of complementarity realization.

T1 – AI as co-analyst and infrastructure

- 90% of all work conducted inside AI "everything, basically"
- Built a bespoke two-layer system: markdown files in owned file system, connected to Claude via API
- Connectors pull all company reports, presentations, and quarterlies from Quartr into structured format
- Claude operates inside Excel, PowerPoint, and Word; reads all emails; sends replies autonomously
- Monitors 50 companies systematically quarterly updates, cross-referenced against prior thesis
- Concrete example: 90-minute iterative session on Novo Nordisk produced analysis that would have taken two full days

P4: "90% of my work is done inside AI. Everything basically. All Excel is in AI too. I don't send emails myself anymore, the AI does it."

T2 – Human judgement and responsibility

- Decisions remain entirely his own, AI does the analysis, P4 makes all final calls
- The system requires deep domain expertise: "you need to know what to ask — it's not like you ask AI which stocks to buy"
- Describes investment decisions as always made under uncertainty: "60–70% probability is good enough to act"
- Does not compare his analysis with AI, they have moved past that stage

P4: "AI does the analysis. I make the decisions. It is not like I compare my analysis with AI's, we moved past that. You need very strong domain knowledge to really get the most out of AI."

T3 – Human-AI collaboration and cognitive dynamics

- Asks 40 iterative questions per session, the depth comes from the sustained dialogue, not a single prompt
- Has built a dedicated bias-detection module specifically to challenge his own prior conclusions
- Uses separate fresh-context Claude sessions to avoid contamination from prior analytical framing
- Identifies security as his primary concern: large values + novel system = risk of exploitation

P4: "You could never have produced that document because I asked forty questions to get there. It is not as if AI did the work, it is a team effort."

T4 – Trust, accountability and delegation boundaries

- Trust is high and backed by system design: errors caught through verification with external analysts
- Treats AI like all sources: always maintains slight scepticism "can be right, can be wrong"
- Security concern: the more sophisticated the system, the higher the target value for bad actors
- Competitive advantage claim: believes the system gives an edge, but "we'll see in a few years"

P4: "The trust level is very high at a baseline. I always keep a slight "hmm, is that right?" the same as I would with any other source I use."

P5: Investment Analyst, Family Investment Office

P5 works in a small family office making few but large investments, with a 3-person team. The interview reflects an earlier, more measured stage of AI adoption, AI is viewed as a complementary sparring partner rather than infrastructure, with an emphasis on AI as an agnostic external perspective that can surface what human analysts miss.

T1 – AI as co-analyst and infrastructure

- Uses Claude across information gathering, analysis, drafting investment proposals
- Concrete example: used AI to compare two bond funds side by side fees, terms, performance, redemption periods and produced a board-ready investment proposal
- Views AI as a complementary work tool, not a replacement still in relatively early adoption phase
- Notes AI can gather more information and find more sources than a human analyst would independently locate

P5: "Claude was able to read through large volumes of documents very efficiently and place them side by side. It came to a conclusion which we also supported, that it made sense to consolidate into one fund."

T2 – Human judgement and responsibility

- Final decisions always made by humans and presented to the board
- Uses board actively as a stress-test and challenge mechanism not just rubber-stamping
- Double-checks AI outputs; everyone on the team takes responsibility for verifying uncertain claims
- Sees AI as an extra team member, not a replacement for human analysis or accountability

P5: "We see it as a complement, as an extra team member, rather than delegating full tasks to it. The team must still take responsibility for verifying what it produces."

T3 – Human-AI collaboration and cognitive dynamics

- Key insight: AI's absence of organizational history is itself a debiasing advantage external, contextless perspective

- Notes Claude is better than ChatGPT at explaining its assumptions and reasoning more transparency
- Recognizes confirmation bias risk through prompting: "ChatGPT in particular is quite prone to answering in the way you seem to want"
- Has not used AI for portfolio-level risk assessment yet, identifies this as a next logical step
P5: "The tool is agnostic in its starting point and looks at things with fresh eyes. I believe it will bring forward viewpoints and perspectives you haven't thought of."

T4 – Trust, accountability and delegation boundaries

- Trust is growing incrementally: switching from ChatGPT to Claude built confidence through transparency
- Anthropic's growing recognition as a company itself builds trust in the tool
- Not experienced any major problems yet, partly because they have not over-relied on it
- Anticipates delegation will increase over time as both tools and user confidence develop
P5: "Over time people will delegate more and more, both because they become more confident in what it does, and because you get better at knowing when to use Claude versus when to rely on yourself."

P6: Investment Analyst, Family Investment Office

P6 works at the same firm as P5, providing a complementary perspective from within the same team. The interview adds important dimensions around competitive dynamics, AI as table stakes to avoid falling behind and the debiasing potential of AI, with P6 offering the clearest statement that AI removes subjective feelings and personal preferences from the analytical process.

T1 – AI as co-analyst and infrastructure

- Uses Claude as primary tool "like having a junior on the team" in terms of capability
- Concrete example: X investment, AI provided initial high-level picture (plusses/minuses), automated follow-up report pulls
- Describes Claude as a qualitative step-change over ChatGPT: can handle complex Excel tasks and financial modelling

- Currently uses AI individually as personal assistant; team has not yet moved to collaborative AI workflows

P6: "I genuinely feel that Claude works so well it is like having a junior on the team. ChatGPT is more like a high-quality Google search for me the two are used very differently."

T2 – Human judgement and responsibility

- Has replaced routine reporting and number-processing with AI, but analysis and decisions remain human
- Describes the competitive framing: if other investors use better AI and make better decisions, you fall behind
- Notes AI elevates thinking rather than replacing it: counter-arguments and alternative angles raise analytical quality
- Remains appropriately critical: actively challenges AI when it agrees too readily

P6: "I feel it just improves the quality of the whole process — as long as you remain somewhat critical toward the answers you receive."

T3 – Human-AI collaboration and cognitive dynamics

- Strongest debiasing statement in the study: AI "removes subjective feelings and viewpoints that can damage the result"
- Recognizes that investors know their biases exist but still struggle to avoid them — AI provides structural help
- Experience of sycophancy: AI agreed with his conclusions too readily — started actively challenging it
- Frames AI adoption as competitive necessity: those who don't use it will be outcompeted

P6: "It removes subjective feelings and viewpoints that can damage the result. All investors are aware of their biases, but it is still difficult to do anything about them. AI provides genuine structural help."

T4 – Trust, accountability and delegation boundaries

- Trust built through AI demonstrating honesty: when it admits it cannot access certain data rather than guessing

- Hallucination immediately destroys trust: "that kills trust immediately"
- Trust is currently individual rather than team-level each person develops their own calibration
- Sees small firm as advantage for adoption speed: aligned ownership and management can move quickly

P6: "When it says it cannot access that data rather than guessing that feels like the answer you would get from a colleague with good values. Honesty is the real currency."

P7: Portfolio Manager, Single-Family Office

P7 manages a technology-focused listed equity portfolio within a small family office. The interview is notable for its systems-level thinking about AI: multi-agent architecture, context window management, and the analogy of AI as a team of specialists. P7 also introduces the market efficiency paradox, if everyone uses AI to make more rational decisions, markets may become harder to beat.

T1 – AI as co-analyst and infrastructure

- Progression from basic information gathering to project-based company memory to agent architecture
- Context window management uses ten parallel agents each focused on one section, then synthesizes, avoids hallucinations from overloaded single context
- Can now follow ten companies closely instead of three due to time savings
- Technology-specific use: instantly explains complex new concepts that previously required a week of reading forums
- Analogy: AI as simulating a team of 100 specialists working simultaneously

P7: "When the project module appeared in Claude, suddenly I could follow ten companies closely instead of three, because I saved so much time. It changed the scope of what one person can cover."

T2 – Human judgement and responsibility

- Clear boundary: "you can never fully automate the decision"
- Distinguishes short-term pattern trading (automatable) from 3–5 years cash flow judgement (not automatable)

- Investment decisions require gut feeling, experience, and timing judgement that AI cannot replicate
- Team of three effectively functions as larger team: unlikely to hire three juniors now

P7: "The robot can automate a lot, but you can never fully automate the decision. Forming a view on what a company's cash flow will look like in three to five years is judgement-based — gut feeling and experience combined with information."

T3 – Human-AI collaboration and cognitive dynamics

- AI as Tesla autopilot: reliable on the motorway, unpredictable at the roundabout, domain calibration matters
- Market efficiency paradox: widespread rational AI use may equalize markets and erode alpha generation
- Broad business model disruption concern: AI makes it difficult to assess which companies will survive
- Notes AI accelerates learning: would have reached professional competence much faster with AI available at the start

P7: "If everyone has access to the same tools and makes more rational decisions simultaneously, markets could become more efficient. That could make it harder for people like me to generate returns."

T4 – Trust, accountability and delegation boundaries

- Trust described as trajectory: time and demonstrated performance, not a threshold event
- Not yet delegating investment decisions but will lean more heavily on AI over time
- Investment framing: decisions are bets with odds AI helps improve odds calculation
- Sees AI as David vs Goliath enabler: small teams can now compete with large institutional resources

P7: "If AI can help improve our understanding of the odds in an investment bet, that will be valuable. We are not yet at the point of letting AI take a decision for us, that requires time and demonstrated track record."

P8: Senior Investment Adviser & Consultant, Multi-Mandate Advisory

P8 is the study's international participant and the most conceptually distinctive voice. With a 35-year investment career spanning equity management, broking, and advisory, P8 introduces the efficiency vs opportunity AI distinction, describes building a family office collaboration platform from scratch, and articulates the strongest vision of how human and AI capabilities should be deliberately separated.

T1 – AI as co-analyst and infrastructure

- Built a family office collaboration platform as a demonstration model — now being put into production
- Builds AI agents that create comprehensive maps of entire industries: who is in a space, where they are, what they do
- Uses AI for M&A target identification, due diligence, and post-investment value improvement
- Monitors medical robotics investment with AI-enabled industry intelligence
- Described building things "I never imagined I would be doing 12 months ago"

P8: "I built an agent that monitors news daily, extracts names, researches them and creates a comprehensive map of who is in a space globally. That is opportunity AI — it enables investments that were not previously accessible."

T2 – Human judgement and responsibility

- Sat for 6 hours in management meeting in Switzerland: reading the room, assessing trustworthiness — irreplaceable human function
- Key principle: as AI absorbs analytical capacity, human capital should shift to emotional intelligence and relationship-building
- "EI vs AI" framework: AI is left-brain analytical; humans should develop right-brain relational capabilities
- AI right now cannot say yes or no to an investment, the decision leap remains human

P8: "AI does the rational, left-brain work. The way we should respond is to develop the emotional intelligence side more, the right-brain skills that AI cannot replicate."

T3 – Human-AI collaboration and cognitive dynamics

- Introduces efficiency AI vs opportunity AI distinction: efficiency AI commoditizes; opportunity AI creates durable advantage
- Six-hats virtual board: uses different AI personas representing opposing analytical roles to challenge own conclusions
- Prompting discipline as prerequisite: without structured prompting, AI defaults to confirming existing views
- Notes risk of analytical skill atrophy: the use of AI requires active maintenance of critical thinking, "like going to the gym"

P8: "There are two types of AI when it comes to investing: efficiency AI and opportunity AI. Efficiency AI will be available to everyone, it erodes advantage. Opportunity AI enables investments that were simply not possible before."

T4 – Trust, accountability and delegation boundaries

- Experience is the primary trust builder: investors who adopted AI early have a compounding advantage
- Trust is domain-specific: high for research and analysis, not extended to final investment judgement
- Six-hats structured prompting as a formal trust-management mechanism builds in challenge by design
- NDA/confidentiality concern: most valuable use cases (private deal data) are precisely where data governance is most constrained

P8: "The experience of the person using AI is as important as the tool itself. Those who got into it early are getting an advantage".