



ANCHORING AND LOSS AVERSION IN HUMAN-LLM INTERACTIONS: A PROSPECT THEORY APPROACH TO AI-ASSISTED DECISION-MAKING

Biplab Paul¹, Gourab Majumder², Aditi Barai³, Priyangshu Goswami⁴,
Kushal Chanda⁵

¹Research Scholar, Kalyani University

²Research Scholar, University of Calcutta

³Research Scholar, University of Calcutta

⁴Research Scholar, University of Calcutta

⁵Research Scholar, University of Calcutta

Article DOI: <https://doi.org/10.36713/epra31560>

DOI No: 10.36713/epra31560

ABSTRACT

Large Language Models (LLMs) are increasingly used to conduct human decision-making, raising concerns regarding their role in potentially reinforcing behavioural biases. The research compares anchoring and loss aversion in Human-LLM interaction supported by the theoretical foundation of Prospect Theory.

A quantitative, between-subjects experimental design was employed with 251 participant responses. Numerical anchors, gain/loss framing, AI trust and AI influence were examined as influences on decision-making outcomes. For statistical analysis, descriptive statistics, reliability analysis, independent t-test, chi-square test, paired t-test, correlation and regression test, ANOVA were done.

Results showed that AI-generated recommendations played a significant role on participants' decisions, whereas gain/loss framing heavily affected risk preferences. Moreover, participants exhibited greater sensitivity to losses compared to gain equivalent to Prospect Theory. Moreover, AI trust but significantly predicted AI influence and reliance.

Overall, the study's findings have extended findings of the existing literature on behavioural decision-making and human-AI interaction and provided evidence that LLMs may influence judgements through anchoring, framing, and trust-related mechanisms.

KEYWORDS: Large Language Models, Anchoring Bias, Loss Aversion, Prospect Theory, AI-Assisted Decision-Making, Human-AI Interaction, AI Trust

1. INTRODUCTION

1.1 Background of the Study

Large Language Models (LLMs) have revolutionised artificial intelligence, evolving rapidly from computational assistance to interactive decision-making tools. Intelligent tools such as ChatGPT and other generative AI systems are becoming popular among individuals and organisations for tasks such as information retrieval, financial planning and decision-making, professional advice, and problem-solving. The current research findings reveal that generative AI is becoming a common application, with approximately 16.3% of the world's population reporting that they have used generative AI tools, AI is increasingly moving into the mainstream of digital assistance [1]. In addition, international survey analysis showed that AI is consistently being adopted in various sectors and has growing importance for both the workplace and consumer usage [2]. Nevertheless, as LLMs grow more intrusive in the human decision-making process, there are significant behavioural implications. Compared to traditional information sources, LLMs can generate tailored recommendations that can shape user perceptions, values, and decisions [3]. While AI systems can enhance efficiency and lessen cognitive load, users might become overly reliant on the AI-provided suggestions, believing them to be more accurate,

knowledgeable, and authoritative than the input from their own thoughts and experiences.

Behavioural economics also shows that cognitive biases can influence human decisions, such as the anchoring effect and loss aversion [4]. Loss aversion refers to the tendency to weigh losses more heavily than equivalent gains, whereas anchoring refers to when individuals heavily rely on initial information. The effect of such loss-sensitive recommendations on whether they act as behavioural anchors and affect behaviour has become relevant. This study investigates the effect of AI-created information on human judgment and decision-making in AI assistance.

1.2 Problem Statement

Although AI systems rely on LLMs to provide guidance and suggestions for decision-making, there is a significant research problem as to whether and how AI recommendations subtly influence human judgement due to behavioural biases. While a huge amount of research has been conducted on decision-making biases, including anchoring and loss aversion in the human context, the introduction of LLMs creates a different context where the initial source of information is produced by an artificial agent rather than another person or classic information media. Users may interpret the numerical



estimates, recommendations, rankings, and evaluations derived from LLMs as authoritative guidance [5]. For instance, an AI-given investment recommendation may establish an initial reference point for further economic decisions even when they have other information that could lead to a different choice. Likewise, how AI systems present the wins and losses can shape the risk preferences and tolerance shown by their users when faced with uncertainty. This issue also becomes significant with consistently increasing trust in AI. A study conducted by AI Trust and conducted with over 48,000 people in 47 countries confirmed that a growing number of people have recognised AI-oriented advantages, yet concerns are also present that impact trustworthiness, reflecting the complicated mix between trusting and relying on AI [2]. Hence, despite the growing use of AI, there is minimal empirical evidence related to the relationship of AI-generated recommendations with other well-known behavioural mechanisms such as anchoring and loss aversion. Hence, this research addresses the issue of whether LLMs merely provide information or shape human decisions by reinforcing behavioural biases.

1.3 Research Gap

Although extensive behavioural economics literature has focused on exploring anchoring, loss aversion, and the use of heuristics in decision-making processes in behavioural economics, most of the current research is centred around human-generated information or conventional decision settings or algorithms with minimal interactive possibilities [6]. The development of LLMs introduces a completely different research field, where artificial intelligence systems actively create recommendations, explanations, and even numerical anchors that might potentially affect human judgment. Similarly, little experimental research has been done to find out whether LLM-generated anchors cause any noticeable shifts in decision-making results or whether increased trust in AI systems strengthens these shifts [7]. Therefore, the linkage between prospect theory-based loss aversion and AI-generated framing has also been limitedly explored. This study addresses this gap by experimentally examining anchoring and loss aversion in human-LLM interactions.

1.4 Research Aim and Objectives

This research aims to investigate the impact of Large Language Model (LLM)-generated recommendations on human decision-making behaviour through the mechanisms of anchoring and loss aversion. It will consider Prospect Theory as the main theoretical framework for AI-assisted decision-making.

The research objectives will be:

- To investigate whether the numeric recommendations generated by LLMs create anchoring effects and impact the human decision-making outcome.
- To assess the effect of gain-framed versus loss-framed information on users' risk preference and decision-making behaviour in an AI-assisted scenario.
- To explore how trust in LLMs affects the use of AI recommendations in people's decision-making process.
- To determine the effectiveness of Prospect Theory principles in explaining the behavioural patterns in

human-LLM interactions through the lens of anchoring and loss aversion mechanisms.

1.5 Research Questions

- To what extent LLM-generated recommendations influence human decision-making through anchoring effects?
- How does the outcome framing (gain versus loss) impact risk-taking behaviour and decision-making preference in AI-assisted decision-making scenarios?
- Is there an association between trust in LLMs and reliance on the AI recommendations in decision-making?
- How effectively does the Prospect Theory explain human behavioural responses, including anchoring and loss aversion, in human-LLM interactions?

2. LITERATURE REVIEW

2.1 Large Language Models and AI-Assisted Decision-Making

Large Language Models (LLMs) can be defined as a significant advancement in AI, which allows the production of human-like responses, comprehending intricate information, and assisting in decision-making in various fields. Unlike traditional AI systems focusing on predictable tasks, LLMs are adaptive, use contextual information, and give recommendations [8]. As they become more prevalent in financial planning, healthcare advice, education, and consumers' decisions, AI has shifted from being an analytical tool to an active decision-support system. A recent study by Eigne and Händler (2024), suggested that LLM-assisted decision-making is influenced by technological factors, user characteristics, task complexity, and psychological factors such as trust and reliance [9]. However, the strength of LLM-generated content can also impact human judgment, propagate cognitive biases, and impact unbiased decision-making.

2.2 Human-AI Interaction and Trust in LLMs

Trust is a pivotal element influencing their interaction with AI systems when users depend on LLMs for recommendations and decision support. The existing research on HAI interactions has found that people tend to judge the credibility of AI based on competency, accuracy, transparency, and communication [10]. LLMs can assist in building user confidence by providing information with elaborate explanations and fluent responses. Recent research illustrates instances of users expecting more from LLMs and their growing trust in AI-generated content, emphasising the need for realistic expectations and confidence [11]. Additionally, the level of trust and acceptance of users for AI recommendations can be greatly affected by AI systems' explanations [12]. Therefore, trust acts as a crucial mechanism that could reinforce the reliance on AI and enhance the behavioural influence generated by LLM recommendations.

2.3 Prospect Theory

The principal theoretical framework for understanding decision-making in risk and uncertainty is deeply rooted in Prospect Theory. The theory, developed by Kahneman and Tversky (1979), suggested that individuals compare outcomes based on a reference point rather than focusing exclusively on

their absolute value [13]. Its value function further suggests that losses have more psychological impact than equivalent gains,

while individuals tend to exhibit risk aversion for gains and greater risk-seeking tendencies when confronting losses.

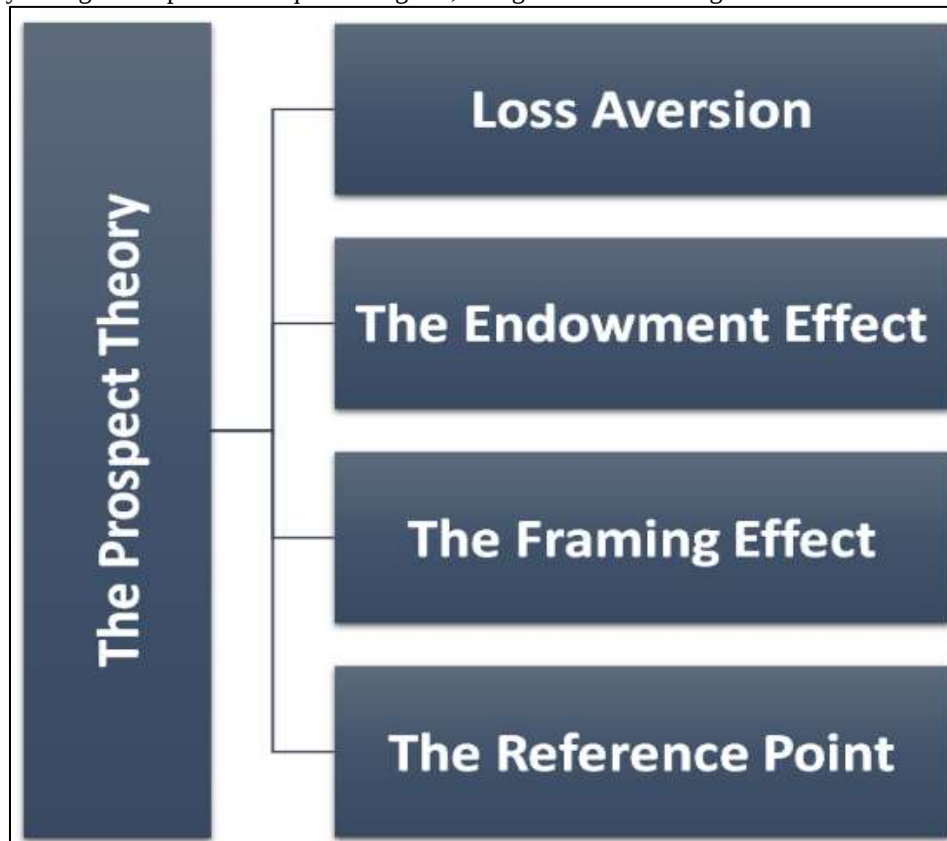


Figure 1: Contexts of Prospect Theory [38]

These principles are also particularly relevant to AI-assisted decisions where LLM-generated recommendations can establish or modify the reference point for comparison when evaluating alternatives. In this regard, evidence by Suri et al. (2024) stated that human-like framing and decision effects are further related to LLMs, showcasing the importance of behavioural decision theories in AI contexts [14]. The present framework of Prospect Theory is used as a vital framework to investigate the change in behaviour of risk preference for human-LLM interaction that occurs with gain and loss framing.

2.4 Anchoring Bias

Anchoring bias describes the tendency for information and an initially presented value to have disproportionate effects on subsequent judgement, whereas adjustments made from the reference point are rational. In human-LLM interaction, the numerical recommendations produced by an LLM can be referred to as anchors for making independent decisions for users. This has become more relevant with growing empirical evidence. Nguyen (2024) identified strong anchoring effects with various LLMs, whereas Carter and Liu (2025), who conducted two controlled experiments with 775 managers,

showed that AI-generated recommendations influenced the subsequent efficacy assessment [15,16]. Therefore, anchoring can function on two levels, where outputs from LLM models may be anchor-sensitive, while on the other hand, their recommendations can then anchor human users. This study aims to investigate whether exposure to an LLM-generated financial recommendation systematically shifts participants' decision-making processes toward the value suggested by the AI, thus questioning the generalisability of this effect.

2.5 Loss Aversion in Decision-Making

A key aspect of prospect theory is loss aversion, which suggests that people tend to feel more strongly about a loss than positive experiences achieved from equivalent gains (Kahneman & Tversky, 1979) [17]. Here, individuals will tend to steer clear of potential losses, even where the expected value of outcomes might be similar, which can impact financial decisions, consumer choices, and risk evaluations. In the context of human-LLM interactions, loss aversion becomes especially relevant as AI systems often convey information in varying ways, like emphasising potential gains or highlighting potential losses [18].

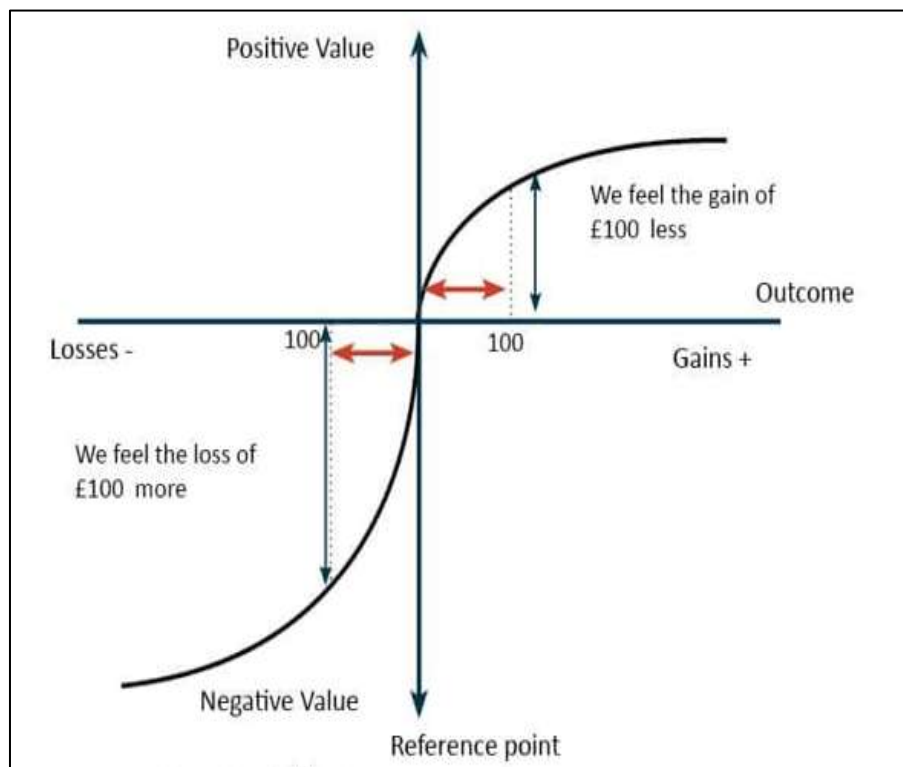


Figure 2: Loss Aversion [39]

Overall, an LLM-generated recommendation framed around avoiding losses is more likely to elicit a behavioural change than an equivalent gain-oriented recommendation. Therefore, it is significant to comprehend this mechanism because AI-supported interventions can either enhance or diminish natural loss-sensitive behaviour based on how information is presented and understood by the users.

2.6 Heuristic Decision-Making and Cognitive Biases

Heuristics, or cognitive shortcuts, are often important in human decision-making to simplify making complex decisions in uncertain conditions with limited information. Heuristics can facilitate quick decision-making, resulting in systematic errors such as anchoring, availability bias, and confirmation bias [19]. LLMs can also contribute to heuristic processing in AI-assisted settings, such as decreasing the cognitive effort to assess options. For instance, individuals may accept AI recommendations as convenient shortcuts, particularly when the system provides confident explanations and structured outputs [20]. This phenomenon can be related to automation bias, when users tend to trust the automated recommendation more than the information analysed by themselves. In this regard, exploring heuristic decision-making is vital in understanding whether suggestions from LLMs facilitate rational decision support or maintain existing cognitive bias.

2.7 Empirical Studies on AI and Behavioural Biases

In recent years, several empirical studies have examined the impact of AI systems on subjects' judgments and behavioural outcomes. Existing research has shown that users' ratings, decisions, and dependence on the algorithmic recommendations depend on the algorithmic system's competence and trustworthiness [21]. There is also emerging research indicating that LLMs themselves can display

behaviours that are similar to cognitive biases, such as anchoring, underscoring a broader issue of behaviours and biases in AI systems. Moreover, research on generative AI adoption shows that human trust and advice acceptance have profound impacts on their reliance on AI recommendations [22]. However, there have been limited empirical studies that have looked at the interaction between AI-generated anchors and existing human biases in a motivating experimental context [23]. This gap highlights the need for research investigating behavioural mechanisms in direct human-LLM decision interactions.

2.8 Research Gap

Despite significant literature that has focused on investigating behavioural biases such as anchoring, loss aversion, and heuristics in decision-making, there is not extensive empirical research concerning the impact of these factors when artificial intelligence systems make decisions. Although many behavioural studies have highlighted human-made decision-making processes, existing literature does not consider situations when the decision-making process is done by LLMs and AI-generated recommendations [24]. In addition, despite many studies concerning trust and reliance on human-AI interaction, there is not sufficient knowledge about the effect of trust in LLMs on the psychological impact of AI-generated information [25]. Furthermore, previous studies have explored trust and reliance in human-AI interaction, insufficient attention has been given to whether trust in LLMs strengthens the psychological influence of AI-generated information. This study tries to investigate this issue by experimentally examining the impact of LLM-generated anchors and gain/loss framing on human behavior and decisions.



3. THEORETICAL FRAMEWORK AND HYPOTHESES DEVELOPMENT

3.1 Prospect Theory as Theoretical Foundation

Prospect Theory serves as the main theoretical basis for this research since it addresses the problem of decision making under conditions of risk, uncertainty and gains or losses. Unlike the expected utility theory, which assumes rational evaluation of alternatives with regard to their expected outcomes, Prospect Theory suggests that decision making depends on psychological factors such as reference points, loss aversion, diminishing sensitivity and outcome framing [31]. Current research continues to prove the relevance of Prospect Theory for the decision-making process.

The first principle, according to which decision makers tend to be reference dependent, is the basic component of Prospect Theory. Decisions are made not with regard to the outcome itself but rather in terms of whether this outcome represents a gain or a loss relative to some reference point. In the current research, the AI-generated numerical investment recommendation can serve as an external reference point for evaluating one's decision about investments. It is especially relevant for human-LLM interaction due to the fact that recent research proves that numerical information provided by LLMs can affect people's judgements. Nguyen (2024) [32], for example, found strong anchoring effects of GPT-4, Claude 2, Gemini Pro and GPT-3.5 with respect to previously presented numerical values.

The second principle – loss aversion – means that losses are psychologically more powerful than equivalent gains. Recent empirical data support the continuing relevance of this idea but the degree of loss aversion depends on experimental conditions. According to Brown et al. (2021) [33] who carried out a meta-analysis of 607 loss-aversion estimates obtained in 150 experiments, the loss-aversion coefficient varied. Also, Bleichrodt and L'Haridon (2023) [34] demonstrated that loss aversion can be robust across different stake sizes while admitting that there is still discussion about its generalisability.

Another principle of Prospect Theory is that of framing effects that imply different decisions resulting from the same economic outcome but depending on whether the latter was framed as gain or loss. Recent research confirms that individuals tend to respond differently to the same objectively equivalent information depending on its framing [35]. Further research that is based on Prospect Theory reveals that risk preferences can be quite different in gain and loss domains [36].

Thus, Prospect Theory provides an adequate framework for investigation of AI-induced anchoring, gain-loss framing and loss aversion in the context of AI-assisted investment decisions. This theory allows investigating whether the numerical recommendation of the LLM impacts the reference point of an individual, whether the framing of gains/losses influences risk preferences and whether losses are psychologically weighted more heavily than gains.

3.2 Conceptual Framework

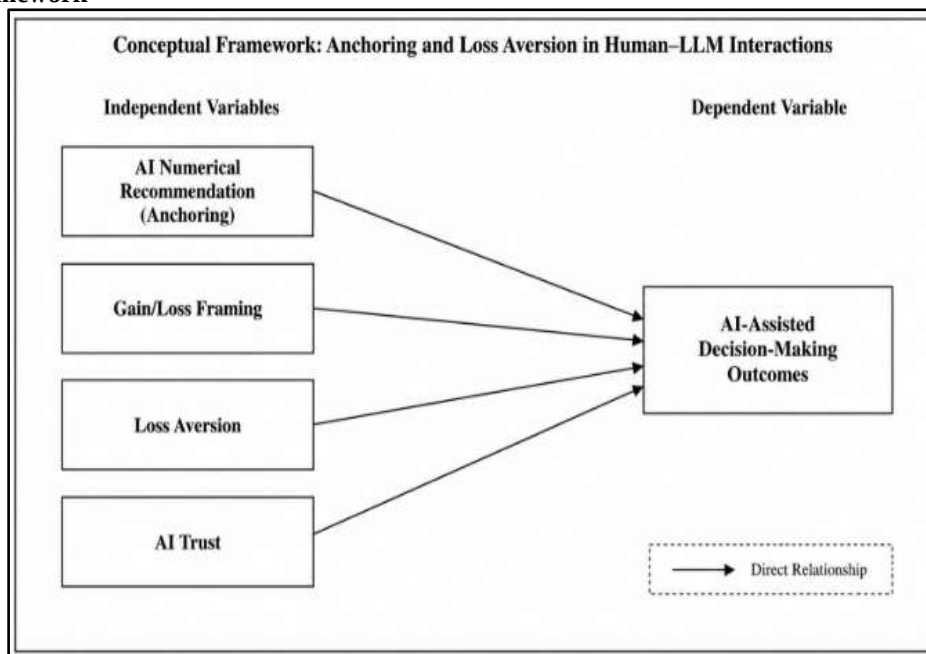


Figure 3: Conceptual Framework

The conceptual framework illustrates the relationships among the variables of interest in the current research. AI numerical recommendations (anchoring), gain/loss framing, loss aversion, and AI trust are the core independent variables. These variables will affect decision-making in an assisted way, especially decisions on investment, risk preferences, and usage of AI recommendations. Prospect Theory is the theory used as a foundation for the current study, as it explains the process of decision-making based on reference points, gains, losses, and

various cognitive biases. Moreover, AI trust will affect acceptance and integration of LLM recommendations in decision-making.

3.3 Hypotheses Development

The hypotheses are derived from the research gap and Prospect Theory. While much research has been done into anchoring, loss aversion and framing as behavioral biases in normal human decision making, there is relatively little empirical analysis explaining these biases in terms of LLM-



generated recommendations. In particular, not much work has been done regarding whether LLM-generated numerical recommendations affect human decision making as an anchor point and whether trust in AI leads to more reliance on recommendations generated by LLMs. Thus, H1 states that being exposed to an LLM-generated numerical recommendation would significantly impact one's decision on investing as an anchoring effect. H2 states that gain/loss framing significantly impacts risk-taking behavior. H3 states that humans are significantly more loss-averse than gain-averse. H4 states that high AI trust leads to reliance on LLM recommendations.

4. RESEARCH METHODOLOGY

4.1 Research Philosophy

The research philosophy used in this study is positivism which assumes that behavioural phenomena can be objectively measured, analysed and explained using empirical evidence [26]. As the study objectives aimed to numerically investigate the impact of LLM-generated anchors and loss framing on human decision-making behaviour, positivism is the most suitable approach. The present study aims to systematically collect the measurable aspects of AI-interaction factors through experimental design and uses statistical analysis to explore the association between AI interaction factors and the measurable aspects of cognitive biases, trust, and decision outcomes.

4.2 Research Approach

This research has followed the deductive research where the existing theories are used to establish hypotheses which will be tested with empirical data [27]. As per theoretical prospect of Prospect Theory, it can be observed how anchoring and loss aversion influence individuals' decisions in uncertain situations. Hypotheses related to the effects of LLM recommendations, outcome framing and trust in AI on decision behaviour are built on this theory. These proposed relationships are then explored by applying statistical methods on the data acquired from the experimental study.

4.3 Research Design

The quantitative, between-subjects experimental research design is chosen to investigate causal links between factors of LLM interaction and outcomes of human decision-making. The participants will be randomly allocated to different experimental conditions to receive an AI-generated anchor or no anchor, and gain or loss framed information. It allows the study to compare between groups and enables them to assess whether LLM recommendations affect the decision choices, risk preferences, and reliance behaviour, while controlling for the experimental conditions.

4.4 Population and Sampling

The target population in this study are adults with experience in using digital technologies and artificial intelligence-based tools for information search, recommendations or decision support. The method of sampling will be convenience sampling since it is exploratory experimental research and these participants can be accessed in the online platform [28]. The proposed sample size is approximately 200-300 participants, sufficient for the

experimental and testing research with statistical analysis especially for hypothesis testing in preliminary research. A formal power analysis will be performed to verify an adequate sample size and to provide statistical reliability.

4.5 Experimental Design

This study is conducted based on a 2×2 factorial experimental design, which investigates the effect of LLM-generated anchoring and outcome framing on human decision-making. The first factor is called "LLM anchoring", which consists of two conditions including "LLM generated numerical recommendation" and "no initial AI recommendation". Outcome framing is a second factor which has gain framed info and loss framed information. Each participant will be randomly sorted into one of four experimental groups and work on financial decision-making scenarios. Individual and interaction effects of the experiment factors could be analysed by this design.

4.6 Data Collection Instrument

Primary quantitative data will be gathered by an online experimental questionnaire that includes demographic questions, scenarios based on AI interaction, decision making tasks, and psychological measurement scales [29]. Participants will be shown controlled financial decision scenarios, varying LLM suggestions and gain/loss framing. They will be numerically scored for their responses, as well as by using Likert scale measures, on their investment choices, risk preference, level of trust in LLMs, and their willingness to follow LLMs recommendations. The online format allows for efficient data collection, while ensuring consistency between experimental conditions.

4.7 Measurement of Variables

The variables that will be measured in this study are associated with LLM influence, behavioural biases and decisions. The measurement of LLM anchoring will be based on the participants' final adjustments with the initial AI-generated recommendations to demonstrate the extend of adjustment towards the provided anchor. Loss aversion will be evaluated using participants' reaction in equivalent gain and loss situations, testing for differences in willingness to take risks. Risk preference will be measured using preference for certain and uncertain outcomes. Trust in LLMs and reliance on AI will be evaluated using Likert scale items measuring the reliability, trustworthiness, and confidence in recommendations made by AI.

4.8 Data Analysis Techniques

The obtained data will be subjected to statistical analysis to test the hypotheses formulated for which SPSS will be optimised. In order to summarise the characteristics of the participants and the overall decisions to be made, descriptive statistics will be used. The internal consistency of a multi-item scale of trust in LLMs and AI reliance will be tested using Cronbach's alpha. The decision difference between the anchored and non-anchored groups will be compared using independent-samples t-test and the main and interaction effect of LLM anchoring and gain/loss framing will be examined



using ANOVA. Regression analysis models will determine whether trust in LLMs predicts AI reliance.

4.9 Reliability, Validity and Ethics

The existing measurement scales will be used and internal consistency of the measurements will be examined using Cronbach's alpha as a measure of reliability [30]. Construct validity will be achieved by theoretical content congruence with Prospect Theory and previous studies in behavioural decision-making. The experimental design will include manipulation checks to ensure participants' correct perception of AI anchoring and outcome framing conditions. Ethical issues will involve informed consent, voluntary participation, anonymity of the participants and protection of data. Participants will be informed about the purpose of the study and their right to withdraw at any stage without

consequences. Appropriate institutional ethical guidelines for human participant research will be followed.

5. RESULTS AND ANALYSIS

5.1 Participant Demographics

The final dataset comprised 251 respondents with no missing data for the demographics. The gender distribution was relatively balanced though the males comprised the largest segment. Males were 141 (56.2%), females were 106 (42.2%), and non-binary individuals were 4 (1.6%). In the case of education level, the largest proportion, which comprised 138 (55.0%), were postgraduates, undergraduates were 93 (37.1%), while doctoral/professional qualification holders comprised 20 (8.0%). As for the AI experience level, there were 123 (49.0%) moderately experienced, 69 (27.5%) highly experienced, and 59 (23.5%) with low experience.

Table 5.1: Participant Demographic Profile

What is your gender?					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Female	106	42.2	42.2	42.2
	Male	141	56.2	56.2	98.4
	Non-binary/Other	4	1.6	1.6	100.0
	Total	251	100.0	100.0	

Table 5.2: Age Descriptive Statistics

Descriptive Statistics					
	N	Minimum	Maximum	Mean	Std. Deviation
What is your age? (Recorded as age in years)	251	18	55	37.11	10.796
Valid N (listwise)	251				

5.2 Descriptive Statistics

From the descriptive statistics above, the AI recommendations were viewed positively. From the AI Trust construct, the mean is 3.6882 (SD = 0.5511). Observed scores ranged from 2.25 to 5.00. This means that the respondents showed relatively high levels of trust towards the AI recommendation. For the AI Influence construct, the mean is 3.4369 (SD = 0.5798), with scores ranging from 1.83 to 4.83. The highest mean among the AI Influence construct items is perceived usefulness (M = 3.69, SD = 0.637). The lowest score among the items under the AI Influence construct is the

consideration of the AI recommendation when determining the investment amount (M = 3.36, SD = 0.759). Experimental results from the investment choices also reveal some differences between the AI-anchor group and the no-anchor group. Participants who saw the AI recommendation in numbers recorded a mean deviation of -₹629.37 from ₹10,000 while those that did not see any recommendation recorded a mean deviation of -₹1,069.60 from ₹10,000. For risk measures, respondents were more ready to accept gain gamble (M = 3.46, SD = 0.705) than loss gamble (M = 2.53, SD = 0.689).

Table 5.3: Descriptive Statistics of Main Variables

Variable	N	Minimum	Maximum	Mean	Std. Deviation
Age	251	18	55	37.11	10.796
AI-recommended investment amount (₹)	126	8,000	12,000	10,253.97	1,599.683
Final investment amount (₹)	251	5,000	13,600	9,151.39	1,496.993
Gain-gamble acceptance	251	2	5	3.46	0.705
Loss-gamble acceptance	251	1	4	2.53	0.689
AI Trust	251	2.25	5.00	3.6882	0.55106
AI Influence	251	1.83	4.83	3.4369	0.57976

5.3 Reliability Analysis

Reliability analysis was done through Cronbach's alpha. For the four-item AI Trust Scale, the Cronbach's alpha was .851, showing a relatively high internal consistency. Item-total correlations for each of the items varied between .686

and .703, showing that all items consistently added to the latent variable. For the six-item AI Influence scale, the Cronbach's alpha value was .874, again showing relatively high internal consistency. Five of the items were found to have relatively high corrected item-total correlations, varying between .740



and .785, whereas the perceived usefulness item had a value of .237.

Table 5.4: Reliability Analysis

Reliability Statistics		
Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.851	.851	4

Reliability Statistics		
Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.874	.867	6

5.4 Experimental Manipulation Checks

Two main manipulations were employed in the experiment exposure to a recommendation from an AI algorithm and gain and loss frames. There were 125 participants who formed part of the no-AI recommendation group and 126

participants who formed part of the AI recommendation group in the experimental sample. However, the provided SPSS output does not show statistical results of manipulation check questions on whether participants knew about the numerical recommendation and gain/loss framing.

Table 5.5: Experimental Manipulation Checks

Item	Strongly Disagree	Disagree	Neutral	Agree	Strongly Agree	N
I considered the AI recommendation when deciding how much to invest	1 (0.4%)	23 (9.2%)	128 (51.0%)	82 (32.7%)	17 (6.8%)	251
My decision was close to the recommendation provided by the AI	1 (0.4%)	22 (8.8%)	122 (48.6%)	89 (35.5%)	17 (6.8%)	251
I would follow the AI recommendation when making a similar decision	0 (0.0%)	24 (9.6%)	122 (48.6%)	90 (35.9%)	15 (6.0%)	251
I relied on the AI information when making my decision	1 (0.4%)	28 (11.2%)	109 (43.4%)	100 (39.8%)	13 (5.2%)	251

5.5 Hypothesis Testing

H1: AI-generated numerical recommendations create an anchoring effect

An independent-samples t-test compared the deviation from the ₹10,000 reference between respondents who received an AI numerical recommendation and those who did not. The no-anchor group had a mean deviation of -₹1,069.60, whereas the AI-anchor group had a less negative mean deviation of -₹629.37. The difference between groups was statistically significant. Because Levene's test was significant ($F = 5.419$, $p = .021$), the unequal-variance result was considered. The test produced $t(237.071) = -2.348$, $p = .020$, with a 95% confidence interval from -₹809.525 to -₹70.944.

H2: Gain/loss framing affects risk-taking behaviour

A Pearson chi-square analysis was conducted to find out whether there was a relationship between framing gain/loss and risky versus certain decision-making. In the gain frame condition, 63 respondents (50.4%) made risky decisions, while 62 (49.6%) went for a certain decision. However, 31 participants (24.6%) made a risky decision in the loss frame situation, while 95 (75.4%) chose the certain one. Pearson chi-square analysis yielded significant results, $\chi^2(1) = 17.826$, $p < .001$. There were no cells with an expected count less than 5, which means that the assumptions for chi-square were met. Thus, framing was significantly related to risk taking behavior.

H3: Individuals demonstrate loss aversion

The paired samples t-test was done for willingness to accept the equivalent gain and loss gambles. The average willingness to accept the gain gamble was 3.46, while that of the loss gamble was 2.53. The average paired difference was 0.932, and the confidence interval was 0.813 to 1.051. The difference was statistically significant, $t(250) = 15.404$, $p < .001$. It shows a much lower willingness to accept the equivalent loss gamble; thus, there is an indication of loss aversion. H3 is supported.

H4: AI Trust positively influences AI reliance/influence

There was a statistically significant positive relationship between AI Trust and AI Influence, $r = .348$, $p < .001$, as established by Pearson correlation test. Further, the regression test provided support to this relationship. AI Trust was a statistically significant predictor of AI Influence ($B = .366$, $SE = .063$, $\beta = .348$, $t = 5.852$, $p < .001$). The obtained results show that $R^2 = .121$, indicating that 12.1% of variance in AI Influence was explained by AI Trust. Therefore, more AI Trust means more AI Influence. H4 is supported.



6. DISCUSSION

6.1 Interpretation of Findings

These results overall provide some support for the proposed behaviour relations. H1 was found to be true, as participants receiving an AI-generated numerical recommendation were characterized by statistically significantly different investment deviation from ₹10,000 compared to participants without an AI recommendation. The participants exposed to AI had a mean deviation of -₹629.37, whereas the participants without an AI recommendation had a mean deviation of -₹1,069.60. The difference is statistically significant, $t(237.071) = -2.348$, $p = .020$. This result is consistent with the assumption about the possibility of an LLM-generated numerical value to affect the subsequent financial judgement. This finding is consistent with Nguyen (2024), who managed to demonstrate significant anchoring effect in several LLMs and emphasized the importance of numerical information in AI-based financial decisions.

H2 was also true, as there was a statistically significant association between gain/loss framing and participants' risk choices, $\chi^2(1) = 17.826$, $p < .001$. However, the relation must

be interpreted with caution. In the gain condition, 50.4% of participants made a choice of a risky option, while only 24.6% did it in the loss condition. Therefore, framing had a definite effect, although the relation was not completely consistent with the prospect theory prediction.

H3 got considerable support. The participants stated being more willing to take the gain gamble ($M = 3.46$) rather than the equivalent loss gamble ($M = 2.53$). The mean difference is 0.932, $t(250) = 15.404$, $p < .001$. This indicates a strong loss sensitivity and is consistent with the recent studies proving the reliability of loss aversion under certain conditions [34].

Finally, H4 is confirmed. There was a positive correlation between AI Trust and AI Influence, $r = .348$, $p < .001$. Thus, the relationship is moderate and positive. The regression analysis indicated that AI Trust positively predicted AI Influence, $B = .366$, $\beta = .348$, $t = 5.852$, $p < .001$, $R^2 = .121$. Hence, trust accounted for 12.1% of variance of AI influence. This finding is consistent with Klingbeil, Grützner and Schreck (2024) [37], who found a positive association between AI trust and reliance on AI advice.

Table 6.1: Summary of Hypothesis Findings

Hypothesis	Statistical result	Decision
H1: LLM recommendation $\rightarrow$ anchoring	$t = -2.348$, $p = .020$	Supported
H2: Gain/loss framing $\rightarrow$ risk preference	$\chi^2(1) = 17.826$, $p < .001$	Supported
H3: Loss aversion	$t(250) = 15.404$, $p < .001$	Supported
H4: AI Trust $\rightarrow$ AI Influence	$\beta = .348$, $p < .001$; $R^2 = .121$	Supported

6.2 Theoretical Implications

The results are consistent with the use of Prospect Theory in human-LLM decision making, expanding the application of Prospect Theory beyond human decision contexts. The difference between gain and loss acceptance confirms that loss aversion evaluation continues to be a reality even when the financial decisions are made within an AI-enhanced context. Likewise, the framing effect result shows that the way outcomes are framed could change behaviour, despite the direction being contrary to the traditional prediction of Prospect Theory.

Anchoring effect result contributes yet another theoretical contribution, as it has shown that the anchor does not need to be generated by a person or traditional information source but could originate from the LLM as well. This is crucial as anchoring within LLM-generated information has been already demonstrated in the literature [32], however, this research focuses on the effect of the information on human users.

6.3 Practical Implications

There are significant implications of the results for the use of AI in financial decision-making. First of all, organisations that design LLM-based financial decision support systems should understand that the number recommendations might serve as a behavioural anchor for users despite the presence of other data. The same issue is noted by Nguyen (2024) [32] who stresses the significance of conscious prompting and awareness of the concept of anchoring in relation to AI-based financial forecasting.

Second, the user interface should distinguish informational recommendations from the actual decision

recommendations. Users should be stimulated to think about alternative scenarios instead of taking for granted an AI-recommended value. Finally, there is a very significant correlation between the level of trust and influence of AI which demonstrates the need for correct calibration of the level of trust.

6.4 Limitations

There are several limitations in this study. Firstly, the experiment involved 251 participants who were recruited using convenience sampling, reducing the generalizability of the results. Secondly, the experimental situation was a simple form of decision making in terms of financial choices and does not fully reflect real-world financial choices with tangible financial implications. Thirdly, even though the anchoring result indicates the existence of a considerable difference in AI and non-AI conditions, the deviation from ₹10,000 was measured, and thus this cannot be considered as the conclusive proof of the movement to the particular AI anchor. Fourthly, the H2 result indicates the existence of a significant framing effect, however, this is contrary to conventional expectations of the Prospect Theory. Finally, the regression between AI Trust and AI Influence is associative and cannot be interpreted as causation of influence by trust.

7. CONCLUSION

This study has significantly explored the impact of LLM generated recommendations on human decision-making based on Prospect Theory with emphasis on anchoring effects, loss aversion, and AI dependency. The obtained empirical evidence indicated that AI recommendations in terms of numerical values affect subsequent decision-making significantly as



participants who were exposed to AI-generated anchors behaved as per recommended reference value. Moreover, the association between LLMs' recommendation influence and the level of people's trust in LLMs was found to be positive. Thus, psychological acceptance of the recommendations plays an important role for shaping users' behaviour. Overall, this research highlighted that LLMs can be considered both as information sources and as behavioural influences capable of shaping human judgements. Therefore, this finding revolutionarily contributed to the field by demonstrating the importance of designing AI systems that promote informed decision-making while reducing unintended behavioural biases.

8. FUTURE SCOPE

Future research can conduct further investigation by exploring human-LLM decision-making in various domains, such as healthcare, education, recruitment, and finance. It could help in obtaining generalised conclusions about individuals' behaviour under the influence of LLMs. Furthermore, a larger sample size and a longitudinal design could be used in order to discover whether increased exposure to LLMs leads to increased behavioural dependence on AI. It would be also essential to consider additional psychological aspects of decision-making, such as perceived AI expertise, explainability of AI, emotional connection, and cognitive reflection.

REFERENCE LIST

1. Microsoft Corporate Responsibility. (2026). Global AI Adoption in 2025 – AI Economy Institute | Microsoft. [online] Available at: <https://www.microsoft.com/en-us/corporate-responsibility/topics/ai-economy-institute/reports/global-ai-adoption-2025>
2. Reuters (2025). Emerging economies lead the way in AI trust, survey shows. Reuters. [online] 28 Apr. Available at: <https://www.reuters.com/business/emerging-economies-lead-way-ai-trust-survey-shows-2025-04-28/>
3. Deconcini, B.M., Coucourde, G., Console, L., Anouti, M., Gaudio, G. and Visciola, M., 2025. Recommender systems for renewable energy communities: tailoring LLM-powered recommendations to user personal values and literacy. In Proceedings of the 13th International Workshop on Behavior Change Support Systems (BCSS) 2025 (pp. 85-96). CEUR. https://iris.unito.it/bitstream/2318/2121315/1/BCSS25_Paper7.pdf
4. Taylor, O., Rossi, M., Lee, J. and Hernandez, M., 2024. Behavioral economics in consumer decision-making: analyzing the impact of cognitive biases. *International Journal of Management, Business, and Economics*, 1(1). <https://doi.journals.net/ijmbe>
5. He, J. and Liu, J., 2025. Investigating the impact of LLM personality on cognitive bias manifestation in automated decision-making tasks. *arXiv preprint arXiv:2502.14219*. <https://arxiv.org/pdf/2502.14219>
6. Korteling, J.E., Paradies, G.L. and Sassen-van Meer, J.P., 2023. Cognitive bias and how to improve sustainable decision making. *Frontiers in psychology*, 14, p.1129835. <https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2023.1129835/full#cite>
7. Li, Z., Zhu, H., Lu, Z., Xiao, Z. and Yin, M., 2025, April. From text to trust: Empowering ai-assisted decision making with adaptive llm-powered analysis. In Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (pp. 1-18). <https://doi.org/10.1145/3706598.3713133>
8. Zhao, H., Liu, Z., Wu, Z., Li, Y., Yang, T., Shu, P., Xu, S., Dai, H., Zhao, L., Jiang, H. and Pan, Y., 2024. Revolutionizing finance with llms: An overview of applications and insights. *arXiv preprint arXiv:2401.11641*. <https://arxiv.org/pdf/2401.11641>
9. Gimmelberg, D. and Ludviga, I., 2025. Strategic complexity and behavioral distortion: Retail investing under large language model augmentation. *International Journal of Financial Studies*, 13(4), p.210. <https://doi.org/10.3390/ijfs13040210>
10. Eigner, E. and Händler, T., 2024. Determinants of llm-assisted decision-making. *arXiv preprint arXiv:2402.17385*. <https://arxiv.org/pdf/2402.17385>
11. Steyvers, M., Tejada, H., Kumar, A., Belem, C., Karny, S., Hu, X., Mayer, L.W. and Smyth, P., 2025. What large language models know and what people think they know. *Nature Machine Intelligence*, 7(2), pp.221-231. <https://doi.org/10.1038/s42256-024-00976-7>
12. Sharma, M., Siu, H.C., Paleja, R. and Peña, J.D., 2024. Why would you suggest that? human trust in language model responses. *arXiv preprint arXiv:2406.02018*. <https://arxiv.org/pdf/2406.02018>
13. Regenwetter, M., Robinson, M.M. and Wang, C., 2022. Four internal inconsistencies in Tversky and Kahneman's (1992) cumulative prospect theory article: A case study in ambiguous theoretical scope and ambiguous parsimony. *Advances in Methods and Practices in Psychological Science*, 5(1), p.25152459221074653. <https://doi.org/10.1177/25152459221074653>
14. Suri, G., Slater, L.R., Ziaee, A. and Nguyen, M., 2024. Do large language models show decision heuristics similar to humans? A case study using GPT-3.5. *Journal of Experimental Psychology: General*, 153(4), p.1066. <https://arxiv.org/pdf/2305.04400>
15. Nguyen, J.K., 2024. Human bias in AI models? Anchoring effects and mitigation strategies in large language models. *Journal of Behavioral and Experimental Finance*, 43, p.100971. <https://doi.org/10.1016/j.jbef.2024.100971>
16. Carter, L. and Liu, D., 2025. How was my performance? Exploring the role of anchoring bias in AI-assisted decision making. *International Journal of Information Management*, 82, p.102875. <https://doi.org/10.1016/j.ijinfomgt.2025.102875>
17. Brown, A.L., Imai, T., Vieider, F.M. and Camerer, C.F., 2024. Meta-analysis of empirical estimates of loss aversion. *Journal of Economic Literature*, 62(2), pp.485-516. <https://www.aeaweb.org/content/file?id=20862>
18. Chen, Y., Kirshner, S., Ovchinnikov, A., Andiappan, M. and Jenkin, T., 2025. Decision Biases of Humans, of AI, and of Humans with AI. Available at SSRN 4943377. <https://papers.ssrn.com/sol3/Delivery.cfm?abstractid=4943377>
19. Rastogi, C., Zhang, Y., Wei, D., Varshney, K.R., Dhurandhar, A. and Tomsett, R., 2022. Deciding fast and slow: The role of cognitive biases in ai-assisted decision-making. *Proceedings of the ACM on Human-computer Interaction*, 6(CSCW1), pp.1-22. <https://doi.org/10.1145/3512930>
20. Spatola, N., 2024. The efficiency-accountability tradeoff in AI integration: Effects on human performance and over-reliance. *Computers in Human Behavior: Artificial Humans*, 2(2), p.100099. <https://doi.org/10.1016/j.chbah.2024.100099>
21. Lou, J. and Sun, Y., 2026. Anchoring bias in large language models: An experimental study. *Journal of Computational Social Science*, 9(1), p.11. <https://doi.org/10.1007/s42001-025-00435-2>



22. Fui-Hoon Nah, F., Zheng, R., Cai, J., Siau, K. and Chen, L., 2023. Generative AI and ChatGPT: Applications, challenges, and AI-human collaboration. *Journal of Information Technology Case and Application Research*, 25(3), pp.277-304.
<https://doi.org/10.1080/15228053.2023.2233814>
23. Campbell, H., Goldman, S. and Markey, P.M., 2025. Artificial intelligence and human decision making: Exploring similarities in cognitive bias. *Computers in Human Behavior: Artificial Humans*, 4, p.100138.
<https://doi.org/10.1016/j.chbah.2025.100138>
24. Lima, G., Grgić-Hlača, N. and Cha, M., 2021, May. Human perceptions on moral responsibility of AI: A case study in AI-assisted bail decision-making. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (pp. 1-17).
<https://doi.org/10.1145/3411764.3445260>
25. Mehrotra, S., Degachi, C., Vereschak, O., Jonker, C.M. and Tielman, M.L., 2024. A systematic review on fostering appropriate trust in Human-AI interaction: Trends, opportunities and challenges. *ACM Journal on Responsible Computing*, 1(4), pp.1-45. <https://doi.org/10.1145/3696449>
26. Ali, I.M., 2024. A guide for positivist research paradigm: From philosophy to methodology. *Ideology Journal*, 9(2).
<https://doi.org/10.24191/ideology.v9i2.596>
27. Kumar, S. and Ujire, D.K., 2024. Inductive and deductive approaches to qualitative research. *International Journal of Multidisciplinary Educational Research*, 13(1), pp.58-63.
<http://ijmer.in.doi./2024/13.1.69>
28. Stratton, S.J., 2021. Population research: convenience sampling strategies. *Prehospital and disaster Medicine*, 36(4), pp.373-374.
<https://doi.org/10.1017/S1049023X21000649>
29. Rahman, A. and Muktadir, M.G., 2021. SPSS: An imperative quantitative data analysis tool for social science research. *International Journal of Research and Innovation in Social Science*, 5(10), pp.300-302.
<https://www.academia.edu/download/79116055/300-302.pdf>
30. Bleichrodt, H. and L'haridon, O., 2023. Prospect theory's loss aversion is robust to stake size. *Judgment and Decision Making*, 18, p.e14.
doi:10.1017/jdm.2023.2https://www.cambridge.org/core/service/s/aop-cambridge-core/content/view/E36163758B43A3AF57F91DD60E46BF33/S1930297523000025a.pdf/prospect_theorys_loss_aversion_is_robust_to_stake_size.pdf
31. Brown, A.L., Imai, T., Vieider, F.M. and Camerer, C.F., 2024. Meta-analysis of empirical estimates of loss aversion. *Journal of Economic Literature*, 62(2), pp.485-516.
<https://www.aeaweb.org/content/file?id=20862>
32. Cicerale, A., Blanzieri, E. and Sacco, K., 2022. How does decision-making change during challenging times?. *PLoS One*, 17(7), p.e0270117.
<https://doi.org/10.1371/journal.pone.0270117>
33. Fisher, S.A., 2022. Framing Effects and Fuzzy Traces: 'Some' Observations: Fisher SA. *Review of Philosophy and Psychology*, 13(3), pp.719-733.
<https://doi.org/10.1007/s13164-021-00556-3>
34. Klingbeil, A., Grützner, C. and Schreck, P., 2024. Trust and reliance on AI – An experimental study on the extent and costs of overreliance on AI. *Computers in Human Behavior*, 160, p.108352. <https://doi.org/10.1016/j.chb.2024.108352>
35. Nguyen, J.K., 2024. Human bias in AI models? Anchoring effects and mitigation strategies in large language models. *Journal of Behavioral and Experimental Finance*, 43, p.100971.
<https://doi.org/10.1016/j.jbef.2024.100971>
36. Yang, Y.P., Li, X. and Stuphorn, V., 2022. Primate anterior insular cortex represents economic decision variables proposed by prospect theory. *Nature communications*, 13(1), p.717.
<https://doi.org/10.1038/s41467-022-28278-9>
37. Khanal, B. and Chhetri, D.B., 2024. A pilot study approach to assessing the reliability and validity of relevancy and efficacy survey scale. *Janabhavana Research Journal*, 3(1), pp.35-49.
<https://doi.org/10.3126/jrj.v3i1.68384>
38. Ansari, S. (2024). *Prospect Theory*. [online] Economics Online. Available at:
<https://www.economicsonline.co.uk/definitions/prospect-theory.html/>