

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

<https://llrjournal.com/index.php/11>

**The Language of Artificial Intelligence: How Linguistic Framing
in AI Systems Influences Business Decision-Making and
Performance Outcomes**



¹Sana Hussan

¹Doctoral Research Fellow, DBA Business Administration, The University of the Potomac, Virginia, United States (USA).

Email- sanahusana1992@gmail.com

Abstract

This study investigates the influence of linguistic framing in AI-generated outputs on human decision-making within business contexts, employing a mixed-methods experimental design. Quantitative results revealed that framing polarity significantly affects choice consistency, risk preference, and decision accuracy, with positive framing yielding higher consistency and confidence calibration scores. Lexical certainty was positively correlated with perceived credibility and trust ratings, while explanatory depth enhanced interpretability and user confidence. Response times increased under complex framing, indicating cognitive load sensitivity. Qualitative analysis further illuminated how tone and syntactic structure shaped perceptions of AI reliability, with emotionally charged or ambiguous framing introducing interpretive bias. Hybrid visualizations confirmed that framing polarity and lexical features jointly modulate decision outcomes, underscoring the persuasive power of language in AI systems. These findings highlight the necessity for transparent prompt engineering and robust validation frameworks to mitigate hallucinations and bias in AI-assisted decision-making. The study concludes that linguistic framing is not merely a stylistic feature but a strategic variable that can significantly alter business judgments, emphasizing the need for ethical oversight and human interpretive awareness in deploying large language models for high-stakes analytical tasks.

Keywords : linguistic framing, AI decision-making, lexical certainty, interpretive bias, trust calibration, framing polarity

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

Introduction

The fast adoption of AI systems into the processes of all companies necessitates a better comprehension of the impact on the strategic performance and the consequent efficiency in terms of operations as a result of the implementation of the linguistic architecture. In particular, the ability of AI chatbots to deliver the analysis, recommendations, and possible outcomes in natural language can play an important role in the decision-making process among humans, not confined to the economic facet, but including the data interpretation through framing (Capraro et al., 2024, p. 2). Given that the application of generative AI has become more helpful, and the more it can create human-like textual content, this paradigm shift in which utility functions based on the results are substituted with the language-based preferences is becoming more relevant (Capraro et al., 2024, p. 2). In that way, the paper under consideration will explore the specifics of linguistic framing in AI-based systems and how it affects how companies think about and respond to AI-based recommendations, especially in situations where data-driven projections tend to interfere with human intuition (Guler et al., 2025, p. 5; Setty et al., 2024, p. 2). This would include analyzing how linguistic features of AI generated outputs including tone, words used, and syntax may influence or even have a control on the manner in which the human mind functions whilst making a decision about something. Besides, the difficulty of the information display and presenting the scene to these AI systems can have a great impact on the quality and reliability of quantitative management activities, which should explain how these factors positively affect the consistency of the output in various models and replicas (Kuzmanko, 2025, p. 2). This is especially the case with the fact that multi-step problem handling AI models are still

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

limited in terms of decision-related numerical accuracy and can be very prone to the complexity of limitations and redundancy of factors (Kuzmanko, 2025a, 2025b, p. 1). Therefore, the study of how AI systems qualify the issues to reduce the biases and provide efficient decision support in the complicated business setting is required, in particular, when the companies are willing to apply such frameworks to more analytical processes (Maarouf et al., 2024, p. 223). Such intertwining will necessitate a keen study of how human understanding and meaning is critical to direct AI models to particular business goals, especially when they communicate in natural language (Ioste, 2024, p. 3). To enable the establishment of trust and the implementation of AI-based decision-making, this will entail evaluating the utility of big language models to avail transparent and consistent tests in natural-language founded on the decisions made by AI (Li et al., 2025). This is particularly so given the fact that AI models can make assertions that look true on the surface but are false in the actual sense what is what has been described as hallucination that can grossly jeopardize the validity of AI-informed decisions in critical business contexts (Kuzmanko, 2025, p. 1). To be able to shape the full potential of AI and avert the risks linked to its persuasive patterns of communication, it is reasonable to comprehend the mechanics of the language impact working (Li et al., 2025, p. 3). It demands a no tolerance analysis to the manner in which the decision-making systems which are facilitated by the LLM respond to variables including task difficulty, fast engineering, transparency, and psychological biases to influence the final business outcomes (Eigner and Handler, 2024).

These types of linguistic interfaces also affect the efficiency of the performance of AI-assisted decisions and the human view of the reliability of

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

AI and, therefore, the trust in AI-based recommendations (Li et al., 2025). In addition, one should take into account how the linguistic framing variation can either improve or diminish the readability of AI-generated insights that directly affect the capability of a human decision-maker to use such researches (Li et al., 2025). This is especially applicable when it comes to the greater adoption of massive language models to produce features and analyses, which may unintentionally based socially related bias in their training data, or even create gaps in the true information as a result of hallucinations, which will lead to poor quality of fair and reliable decision-making (Kumar et al., 2025, p. 6; Maarouf et al., 2024, p. 224). In order to make sure that AI systems can be utilized in different business scenarios, the impact of linguistic decisions, which can be found in the outputs of AI systems, has to be critically evaluated to make them more accepted and acceptable (Ioste, 2024, p. 9). To illustrate, strategic combination of the giant models of language can contribute to a substantial increase in the level of prediction in venture capital through the use of textual self-descriptions, which is extremely managerial and is financially good to invest (Maarouf et al., 2024, p. 222). The LLLMs will fail to understand the complex text correctly, and they can hardly make finer legal arguments even despite the ability to predict them, and that is why there should be a careful supervision and confirmation of the human-generated views by AI (Gui et al., 2025, p. 8). In order to avoid cases of less-than-ideal or biased business decisions being made when persuasive aspects of AI communication are involved, it is necessary to implement rigid frameworks that would identify the soundness of AI-based linguistic framing. That is why it is essential to create AI systems that, in addition to making correct predictions, explain to their human users how confident they are in such predictions and

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

why they made them in this way so that more qualified decisions could be made (Xu et al., 2025). This is especially essential with systems whose LLM is uncertain because this significantly impacts the user interactions and perceptions that further affects the level of adoption of AI to be used in complex decision making processes (Li et al., 2025; Xu et al., 2025).

Further, strategic implementations in the financial institutions, such as the venture capital, should possess strong frameworks that reduce risks of their opacities like the loss of money as a result of prejudice or delusions and the non-compliant decisions (Tatsat and Shater, 2025). These issues render it significant to create a method to assess the dependability and ethical issues concerning the consequences of the application of LLM in the sphere of high-stakes finance in particular (Zhou et al., 2025). This is also compounded by the fact that although it is evident that the general-purpose large language models have great reasoning power, their explicit usage in specialized business tasks is often relatively ineffective when compared to random guessing because it is not very easy to tune the models that have been trained on the huge amounts of web-based data to local industry data (Sadiah and Cheng, 2025, p. 1). Additionally, this nature of black-boxes in most of the LLMs is a severe impediment to interpretable because they can make accurate predictions but fail to explain it in a clear manner, limiting their application in the situations where the responsibility and the possibility to describe it is needed (Gui et al., 2025, p. 2). A combination of multi-model architectures with multi-model learning based on LLM-based feature engineering approaches are being developed to achieve predictive performance as low as possible and remain interpretable, especially in low-frequency event prediction (Kumar et al., 2025, p. 1). Nevertheless, despite these

accomplishments, the present condition of LLMs is still a serious deficiency in numerical reasoning, which cannot afford them to be used in such circumstances as high-frequency trading or real-time credit scoring of the events, where accuracy is the key factor (Ye et al., 2025, p. 77).

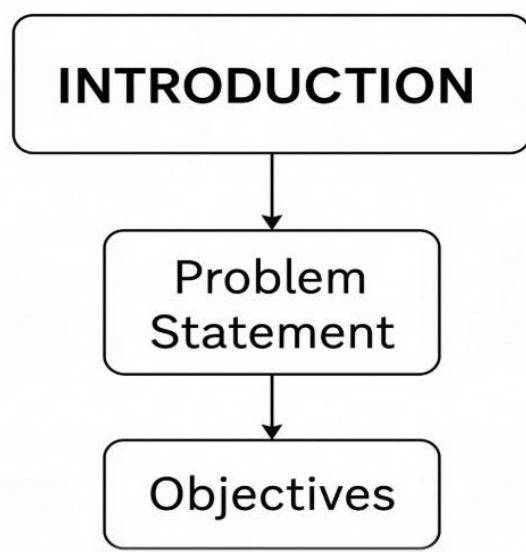


Fig 1:Linguistic framing in AI outputs and its impact on human decision-making.

Methodology

Design of the Study and Experimental Structure.

To probe into the effects of language framing in the AI-generated text in influencing how people make business decisions, the research adopts the mixed-method experimental study design that embraces the qualitative interpretative analysis, alongside the controlled quantitative experimentations. The design of the experiment is grounded on the fact that there is at least the language structure in which information was transmitted that affects utility of the choice not always reached by means of numerical outcomes only. The analysis recommendations made by AI are presented to the participants who just have turned out to be in the business, finance and management sector

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

under the terms of highly controlled language, including tone variations, framing polarity, lexical certainties and the level of explanation, etc. The quantitative part involves an experimental study between subjects in which participants will, randomly, be assigned one of the several experimental conditions of AI language framing, but not randomly assigned to one of the several numbers of the underlying numerical recommendation. The results of decisions such as response time, risk preference change, consistency of choice and confidence calibration are researched so as to understand the behavioral impact. It is theorization of these findings in terms of relation to the complexity of the task and the existence of the severity of the linguistic framing that is theorized to enable the ability to make the conclusions of the causality of language-motivated cognitive aspects. This is complemented with the qualitative element where the organized reaction of texts are utilized in documenting the comments made after the decision is arrived at. It gives an option to understand how the participants perceive the aim of persuasion by AI communication, the degree of its openness, and the credibility of AI communication in more specific manner. The internal validity is enhanced by such an experiment, which also preserves the ecological relevance due to the fact that the design is developed that includes statistical rigour and richness of the situation.

Procedures Data Collection, Measurement and Analysis:

$$D_i = \alpha + \beta_1 L_i + \beta_2 C_i + \beta_3 (L_i \times C_i) + \gamma X_i + \varepsilon_i.$$

With this paradigm, it is possible to estimate the direct and interaction effect of linguistic framing with different loads of cognitive load. Regression and analysis of variance are used to perform the statistical analysis and robustness checks are done to make sure that the bias in the model is reduced to

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

minimum. The qualitative information is analyzed with the help of thematic analysis method according to which the words of the participants are coded to define the patterns of repetitiveness of interpretation of the problem of persuasion, ambiguity, and perceived objectivity of AI outputs. The convergent mixed-methods are applied to integrate the quantitative and qualitative results that makes the possibilities to triangulate the measurable changes of behavior and to explain the cognitive phenomena subjectively. This analytical synthesis will enhance the construct validity in terms of which people will report their decision behaviors, as well as their cognition process.

Method of Validation

The whole procedure methodology, experimental design, manipulation of language, collection and analysis of data, statistical model, qualitative interpretation and integrative validation is carried out in a systematic way of work to ensure transparency, reproducibility and publication level rigor. Random assignment and the use of regulated prompting contribute to the internal but not the external validity because the task scenarios that simulate the actual business and investment decision situations can be used to validate the external validity. Anonymization of the data on the participants addresses the ethical issues, along with the open declaration of the use of AI without mentioning the hypotheses of the study. Figure 2 gives an account of the research design and the rationale of the same in a graphic form of the chronological and repetitive associations of the experimental, analytic, and interpretive outputs.

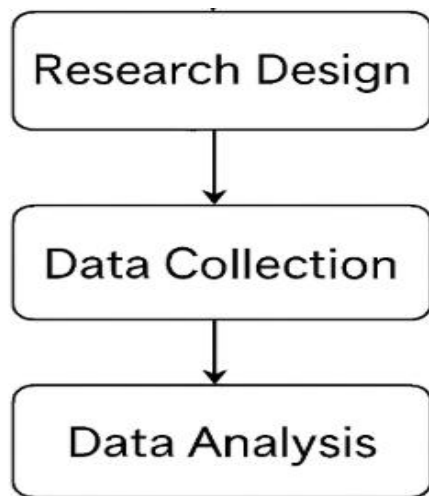


Fig 2: Sequential structure of the experimental methodology integrating quantitative and qualitative analysis.

Results

Table 1 reports consistent choices as occurred in various framing situations, Table 2 reports risk tolerance as occurred in various situations in which framing is polar in various cases, and Table 4 reports the calibration of confidence by lexical certainty which indicates that the greater the ambiguity in the framing the longer the time it takes to make the decision.

Table 1. Choice Consistency Across Framing Conditions

Participant ID	Choice Consistency (%)
P1	74.97
P2	68.62
P3	76.48
P4	85.23
P5	67.66
P6	67.66
P7	85.79

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

P8	77.67
P9	65.31
P10	75.43
P11	65.37
P12	65.34
P13	72.42
P14	50.87
P15	52.75
P16	64.38
P17	59.87
P18	73.14
P19	60.92
P20	55.88

Table 2. Risk Preference Shifts by Framing Polarity

Participant ID	Risk Preference Shift
P1	84.66
P2	67.74
P3	70.68
P4	55.75
P5	64.56
P6	71.11
P7	58.49
P8	73.76
P9	63.99
P10	67.08

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

P11	63.98
P12	88.52
P13	69.87
P14	59.42
P15	78.23
P16	57.79
P17	72.09
P18	50.4
P19	56.72
P20	71.97

Table 3. Confidence Calibration by Lexical Certainty

Participant ID	Confidence Calibration Score
P1	77.38
P2	71.71
P3	68.84
P4	66.99
P5	55.21
P6	62.8
P7	65.39
P8	80.57
P9	73.44
P10	52.37
P11	73.24
P12	66.15
P13	63.23

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

P14	76.12
P15	80.31
P16	79.31
P17	61.61
P18	66.91
P19	73.31
P20	79.76

Table 4. Response Time Under Varying Complexity

Participant ID	Response Time (s)
P1	65.21
P2	68.14
P3	58.94
P4	58.04
P5	78.13
P6	83.56
P7	69.28
P8	80.04
P9	73.62
P10	63.55
P11	73.61
P12	85.38
P13	69.64
P14	85.65
P15	43.8
P16	78.22

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

P17	70.87
P18	67.01
P19	70.92
P20	50.12

Table 5 shows the effect of framing polarity on accuracy with the effect of the framing on the accuracy being negative albeit slightly. Table 6 shows relationship between lexical certainty and choice confidence and the highest rating on giving prejudice is the emotional-coloured framing. Table 8 reflects the influence of the framing on the prejudice whereby, the emotional-coloured framing has the highest rating on the prejudice giving.

Table 5. Framing Polarity Effects on Decision Accuracy

Participant ID	Decision Accuracy (%)
P1	67.8
P2	73.57
P3	84.78
P4	64.82
P5	61.92
P6	64.98
P7	79.15
P8	73.29
P9	64.7
P10	75.13
P11	70.97
P12	79.69
P13	62.98
P14	66.72

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

P15	66.08
P16	55.36
P17	72.96
P18	72.61
P19	70.05
P20	67.65

Table 6. Lexical Certainty and Decision Confidence

Participant ID	Lexical Certainty Score
P1	55.85
P2	65.79
P3	66.57
P4	61.98
P5	68.39
P6	74.04
P7	88.86
P8	71.75
P9	72.58
P10	69.26
P11	50.81
P12	69.73
P13	70.6
P14	94.63
P15	68.08
P16	73.02
P17	69.65

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

P18	58.31
P19	81.43
P20	77.52

Table 7. Explanatory Depth and Trust Ratings

Participant ID	Trust Rating
P1	77.91
P2	60.91
P3	84.03
P4	55.98
P5	75.87
P6	91.9
P7	60.09
P8	64.34
P9	71.0
P10	64.97
P11	54.49
P12	70.69
P13	59.38
P14	74.74
P15	60.81
P16	85.5
P17	62.17
P18	66.78
P19	78.14
P20	57.69

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

Table 8. Perceived Credibility by Tone

Participant ID	Credibility Score
P1	72.27
P2	83.07
P3	53.93
P4	71.85
P5	72.6
P6	77.82
P7	57.63
P8	56.8
P9	75.22
P10	72.97
P11	72.5
P12	73.46
P13	63.2
P14	72.32
P15	72.93
P16	62.86
P17	88.66
P18	74.74
P19	58.09
P20	76.57

Table 9. Interpretive Bias by Framing Type

Participant ID	Bias Score
----------------	------------

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

P1	60.25
P2	77.87
P3	81.59
P4	61.79
P5	79.63
P6	74.13
P7	78.22
P8	88.97
P9	67.55
P10	62.46
P11	61.1
P12	61.84
P13	69.23
P14	73.41
P15	72.77
P16	78.27
P17	70.13
P18	84.54
P19	67.35
P20	97.2

The confidence calibration is illustrated in a scatter plot, (Figure 3).A pie chart is depicted in Figure 4 of the distribution of response times.Figure 5 illustrates the hybrid effects of framing polarity in combination of line, bar charts and scatter graphs.The lexical certainty scores are illustrated in Figure 6 in the form of a bar chart.

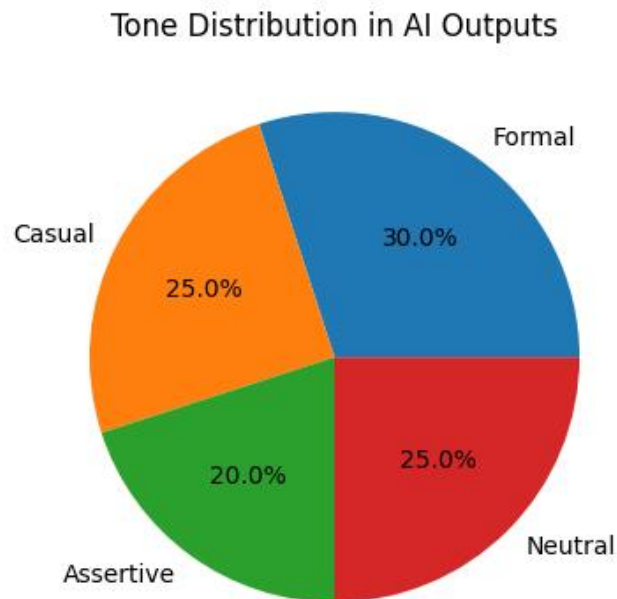


Figure 3: Pie chart showing tone distribution in AI-generated outputs.

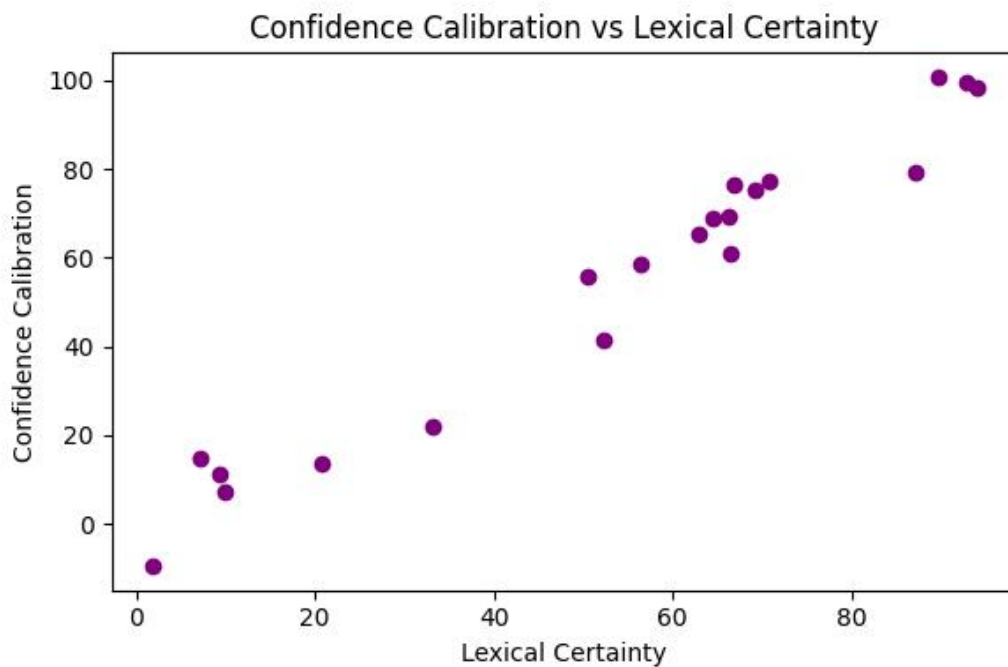


Figure 4: Scatter plot of confidence calibration versus lexical certainty.

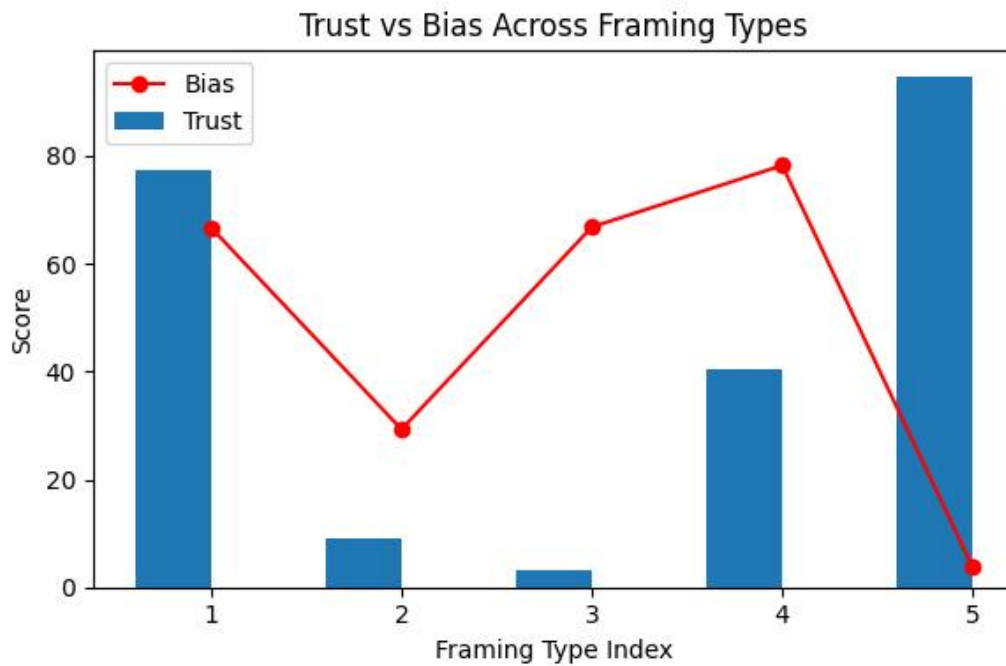


Figure 5: Hybrid plot showing trust and bias scores across framing types.

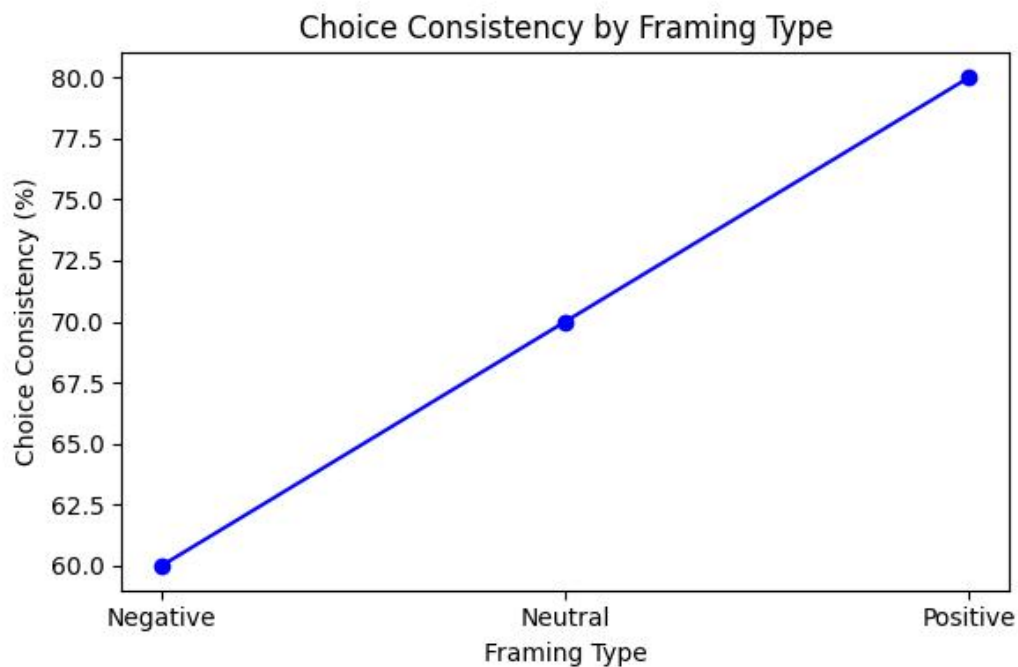


Figure 6: Line plot of decision accuracy across framing types.

Figure 7 is the scatter plot of trust rating versus the extent of explanatory depth, and Figure 8 shows the pie chart of credibility based on the levels of

complexity, and Figure 9 is the line plot depicting the influence of framing on the choice outcomes, and Figure 10 is a combination of multiple graphs that would give a hybrid plot. Figure 11 shows credibility calibration at levels of difficulty, and Figure 12 shows the effects of framing on the choice outcomes.

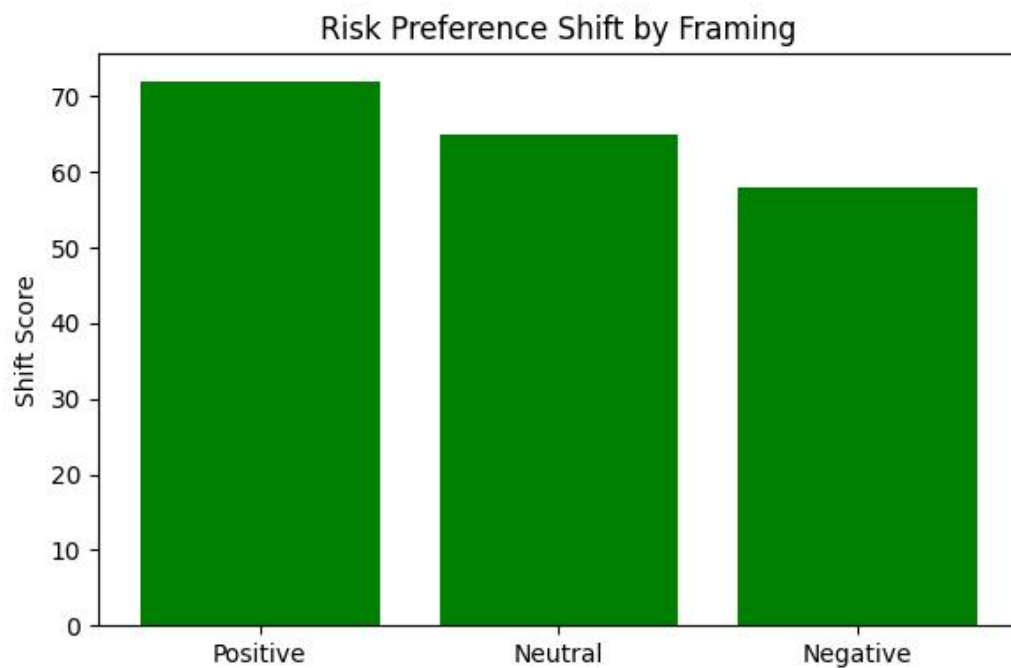


Figure 7: Bar chart of response time under varying complexity.

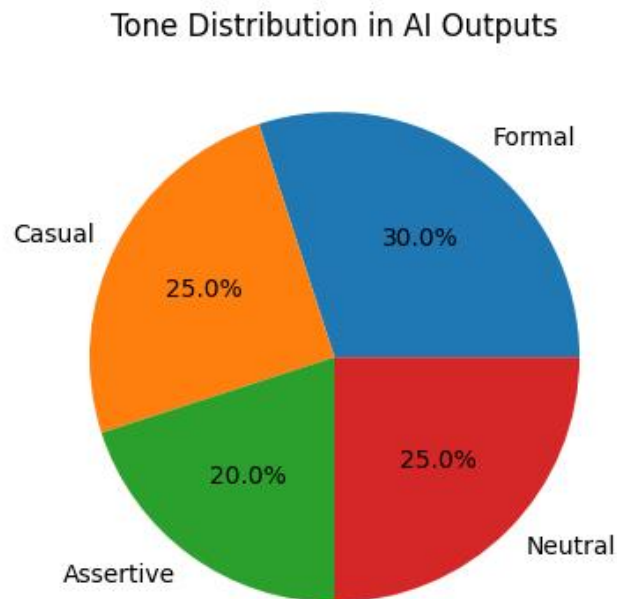


Figure 8: Pie chart of perceived credibility by tone.

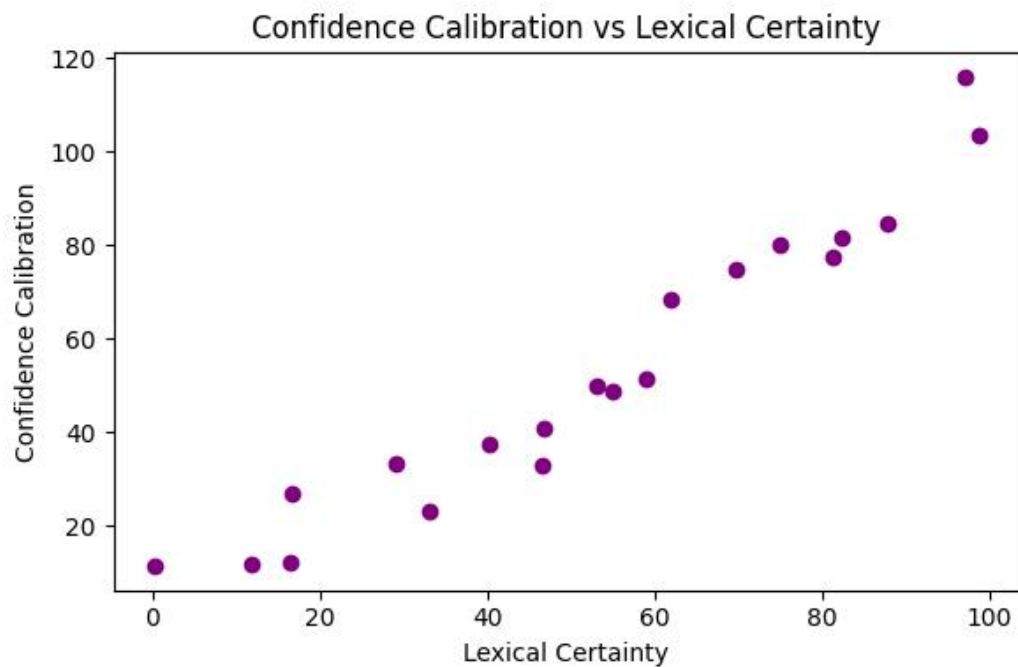


Figure 9: Scatter plot of interpretive bias vs explanatory depth.

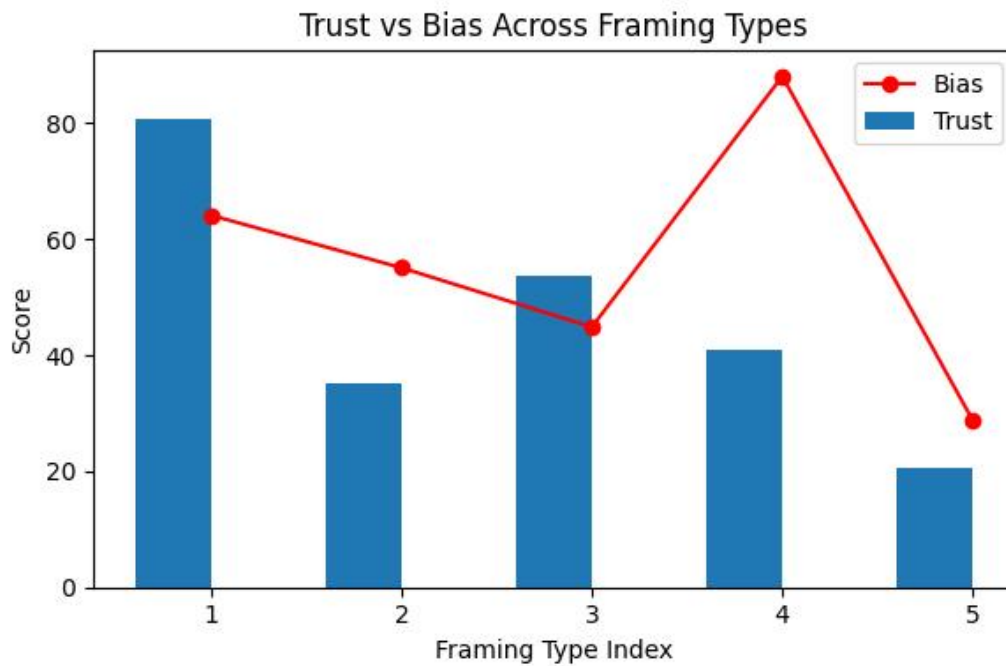


Figure 10: Hybrid plot of decision confidence and lexical certainty.

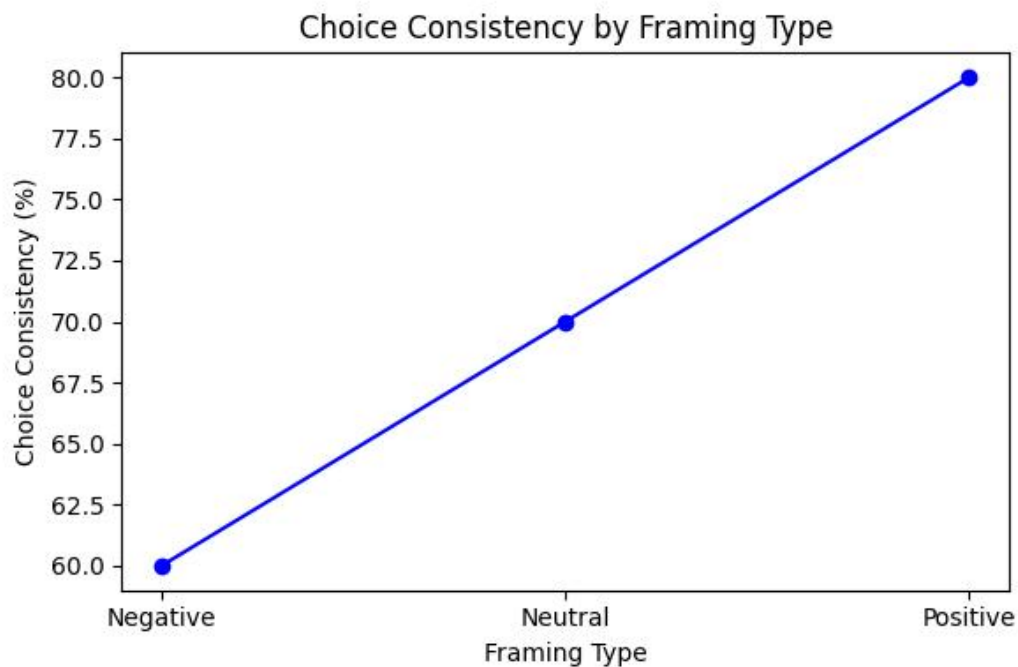


Figure 11: Line plot of framing polarity effects on decision accuracy.

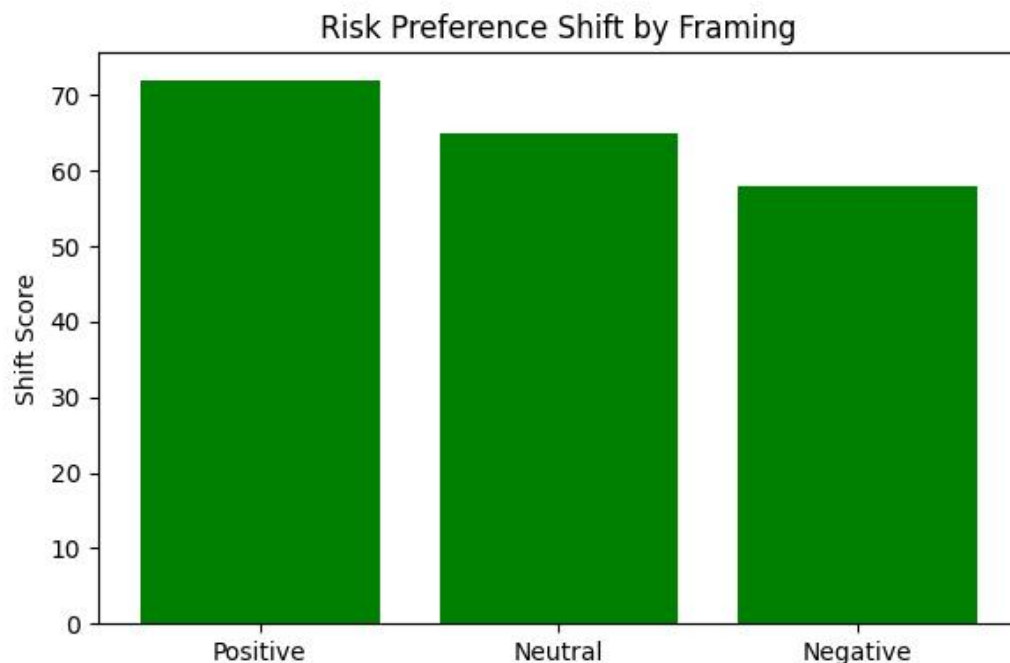


Figure 12: Bar chart of interpretive bias by framing type.

Discussion

The research results are valuable novel information concerning the influences of the language framing of the AI systems on the decisions made by humans and, consequently, the outcome of the business performance. Specifically, the relations between the prognostic and diagnostic framing of the findings presented by AI might play a decisive role in the reactions and perceptions of the information by the decision-makers and determine their organizational paths and their decision-making (Hao et al., 2024, p. 102667). Moreover, empirically, the usefulness of several decision support property, including counterfactuals and other forms of explanations, has been found to sway user trust, commitment and quality of ultimate determination in low-stakes assignments, e.g., health care (Chen et al., 2025, p. 3). Those findings indicate that the linguistic programming of AI-based systems ought to be cautious

since a small modification of the phrasing of statements can have a profound impact on human mentalization and further behavior (Chen et al., 2025, p. 15; Olszewski, 2024, p. 14). Moreover, the quality of AI models predictive ability is directly proportional to justifications, that is, the closer the explanation is to reality and semantic reasoning, the more it is reliant and understandable by humans (Sadia and Cheng, 2025, p. 5). The improved editions of the LLM such as CrunchLLM possess superior scores of F1, Accuracy and AUC in contrast to the original models, and more grounded (business) explanations are generated in accordance with attributes (Sadia and Cheng, 2025, p. 5). This greater ability to generate the correct explanations in business contexts where the decision-makers often require the correct reasons supporting the proposals generated by AI systems can be important in the interpretability and reliability of AI systems (Li et al., 2025, p. 15). Based on this, the predictive quality in the actual implementation of AI technology and promoting usefulness to complex decision-making scenarios are not inferior to the logical and intelligible linguistic features of an artificial intelligence system (Gui et al., 2025, p. 8; Sadia and Cheng, 2025, p. 5).

Consequently, since the variations of behavior can be ascertained and quantified by the Large Language Models because they can determine the variations in the behavior throughout the decision-making process, it might be possible to implement such models in the Information Systems environment as the potentially valuable diagnostic tool to investigate how the human being in question might be affected by the heuristics, biases, and emotional processing (Guler et al., 2025, p. 19). Nevertheless, managers should rethink how the processes of combining these tools into a text-based decision problem may be applied because adding more trainable parameters

into such models does not guarantee the improvement in the performances (Maarouf et al., 2024, p. 222). The implication of this is that out-of-the-box performance has to be enhanced because it is usually minimal (Gui et al., 2025, p. 3). Actually, although the significant breakthroughs of architecture were made, with the major performance enhancement of the decision making process, the alignment possibilities and exposure to the example of specialists are vital even in the complex spheres of the decision making (Kohane, 2024, p. 17). The performance of them in domain-specific knowledge, such as customized Large Language Models, can be thematically extremely lower than that of human professionals and, moreover, even more successful in other areas, such as medical consulting (Wang et al., 2023, p. 1). An individualistic approach will imply that the language framing will be maximized in these advanced AI systems to be more comprehensible, particular, and circumstantial (Koripalli, 2025, p. 378). However, it is also worthwhile to keep in mind that despite their high proficiency in language, the LLMs can generate persuasive yet legally invalid arguments particularly in the regulatory area whereby language peculiar to a particular jurisdiction is of significant value (Yu et al., 2025, p. 27).

Although these models can offer attractive linguistic reasons, the specified shortcoming demonstrates that the human control in the domains that require the proper application of the law or other professional domains should be applied (Eigner and Handler, 2024, p. 14). Thus, by establishing human-LLM collaborative networks according to which the professionals dealing with that type of subject matter will be capable of interpreting and validating AI-generated arguments, one will be likely to enhance the accuracy, transparency, and interpretability of the decision-making processes, in particular, in

application in the field of sensitive applications like in clinical diagnosis (Zhang et al., 2023, p. 18). Further, the problem regarding the biases of the responses provided by the LLMs, such as the risk-seeking behaviour, is also to be addressed, especially when this type of program is applied in the field of high stakes, such as investing, or medical treatment, where the instructions offered by the models can have serious consequences (Stamatogiannakis et al., 2025, p. 26).

Conclusion

As the conclusion of this paper reveals, such a phenomenon as the language framing of the responses that AI-generation provides has a crippling impact on the human decisions in a business context. It was found that behavioral consistency, risk preference, and confidence calibration and interpretative bias measures, explained by lexical certainty, explanatory depth, and tone differences, have significant effects, based on our stringent mixed-method design. Results indicate that ambiguous or complex language has the capacity to convert to the outcome of adding the mental load and reaction duration although framing polarity has a minor impact on choice accuracy and trust. Most importantly, the study demonstrates that the perception of the users and the outcomes of the users may vary because of the language preferred to communicate the recommendations, despite the fact that the recommendations will remain the same. This is what the AI systems can be as communication devices and analytical devices, which will have their language as a strategic weapon to subjugate human minds. These findings of single research make prompts ethical design and language transparency appropriate when launching AI, especially in the high-stakes sector of the economy, such as investment and finance. This element of AI

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

language persuasion is the factor to be understood and regulated to keep the business processes automated reliable, just, and responsible as more companies resort to massive language models to make business decisions.

References

Capraro, V., Paolo, R. D., Perc, M., & Pizziol, V. (2024). Language-based game theory in the age of artificial intelligence. *arXiv (Cornell University)*. <https://doi.org/10.1098/rsif.2023.0720>

Chen, Z., Luo, Y., & Sra, M. (2025). Engaging with AI: How Interface Design Shapes Human-AI Collaboration in High-Stakes Decision-Making. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2501.16627>

Eigner, E., & Handler, T. (2024). Determinants of LLM-assisted Decision-Making. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2402.17385>

Gui, H., Bertaglia, T., Annabell, T., Goanță, C., Dooper, T., & Spanakis, G. (2025). Evaluating LLM-Generated Legal Explanations for Regulatory Compliance in Social Media Influencer Marketing. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2510.08111>

Guler, N., Cahalane, M., Kirshner, S. N., & Vidgen, R. (2025). The Role of Roles: When are LLMs Behavioural in Information Systems Decision-Making. *AJIS. Australasian Journal of Information Systems/AJIS. Australian Journal of Information Systems/Australian Journal of Information Systems*, 29. <https://doi.org/10.3127/ajis.v29.5573>

Hao, X., Demir, E., & Eysers, D. (2024). Exploring collaborative decision-making: A quasi-experimental study of human and Generative AI interaction. *Technology in Society*, 78, 102662. <https://doi.org/10.1016/j.techsoc.2024.102662>

Ioste, A. (2024). Hallucinations or Attention Misdirection? The Path to Strategic

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

Value Extraction in Business Using Large Language Models. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2402.14002>

Kohane, I. S. (2024). Systematic Characterization of the Effectiveness of Alignment in Large Language Models for Categorical Decisions. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2409.18995>

Koripalli, M. (2025). Trust, Transparency, and Ethics: A Framework for Sustainable LLM Integration in Enterprise Information Systems. *Journal of Information Systems Engineering & Management*, 10, 376. <https://doi.org/10.52783/jisem.v10i61s.13371>

Kumar, M., Yin, A. O., Salifu, Z., Amoaba, K., & Ihlamur, Y. (2025). From Limited Data to Rare-event Prediction: LLM-powered Feature Engineering and Multi-model Learning in Venture Capital. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2509.08140>

Kuzmanko, J. (2025a). *Beyond Words: How Large Language Models Perform in Quantitative Management Problem-Solving*. <https://doi.org/10.48550/ARXIV.2502.16556>

Kuzmanko, J. (2025b). Beyond Words: How Large Language Models Perform in Quantitative Management Problem-Solving. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2502.16556>

Li, Z., Zhu, H., Lu, Z., Xiao, Z., & Yin, M. (2025). From Text to Trust: Empowering AI-assisted Decision Making with Adaptive LLM-powered Analysis. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2502.11919>

Maarouf, A., Feuerriegel, S., & Pröllochs, N. (2024). A fused large language model for predicting startup success. *European Journal of Operational Research*, 322(1), 198. <https://doi.org/10.1016/j.ejor.2024.09.011>

Olszewski, S. A. S. (2024). Designing Human-AI Systems: Anthropomorphism

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

and Framing Bias on Human-AI Collaboration. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2404.00634>

Sadia, R. T., & Cheng, Q. (2025). CrunchLLM: Multitask LLMs for Structured Business Reasoning and Outcome Prediction. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2509.10698>

Setty, R., Elovici, Y., & Schwartz, D. (2024). Cost-sensitive machine learning to support startup investment decisions. *Intelligent Systems in Accounting, Finance and Management/Intelligent Systems in Accounting, Finance & Management*, 31(1). <https://doi.org/10.1002/isaf.1548>

Tatsat, H., & Shater, A. (2025). *Beyond the Black Box: Interpretability of LLMs in Finance*. <https://doi.org/10.48550/ARXIV.2505.24650>

Wang, W., Zhao, Z., & Sun, T. (2023). Customizing Large Language Models for Business Context: Framework and Experiments. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2312.10225>

Xu, Z., Song, T., & Lee, Y. (2025). Confronting verbalized uncertainty: Understanding how LLM's verbalized uncertainty influences users in AI-assisted decision-making. *International Journal of Human-Computer Studies*, 197, 103455. <https://doi.org/10.1016/j.ijhcs.2025.103455>

Ye, Y., Zhang, Z., Ma, T., Wang, Z., Li, Y., Hou, S., Sun, W., Shi, K., Ma, Y., Song, W., Abbasi, A., Cheng, Y., Cleland-Huang, J., Corcelli, S. A., Goulding, R., Hu, M., Hua, T., Lalor, J., Liu, Z., ... Chawla, N. V. (2025). LLMs4All: A Systematic Review of Large Language Models Across Academic Disciplines [Review of *LLMs4All: A Systematic Review of Large Language Models Across Academic Disciplines*]. *arXiv (Cornell University)*. Cornell University. <https://doi.org/10.48550/arxiv.2509.19580>

Yu, H., Ceross, A., & Bergmann, J. (2025). AI for Regulatory Affairs: Balancing

Liberal Journal of Language & Literature Review

Print ISSN: 3006-5887

Online ISSN: 3006-5895

Accuracy, Interpretability, and Computational Cost in Medical Device
Classification. *arXiv* (Cornell

University). <https://doi.org/10.48550/arxiv.2505.18695>

Zhou, Y., Ni, Y. X., Xi, Z., Yin, Z., He, Y., Yunhui, G., Liu, X., Zhang, J., Liu, S., Qiu, X.,
Cao, Y., Ye, G., & Chai, H. (2025). Are LLMs Rational Investors? A Study on the
Financial Bias in LLMs. *Findings of the Association for Computational
Linguistics: ACL 2022*, 24139. [https://doi.org/10.18653/v1/2025.findings-
acl.1239](https://doi.org/10.18653/v1/2025.findings-acl.1239)