

Title	An exploratory study of factors influencing the use of Generative AI for decision#making: a missing link
Authors	Heavin, Ciara;Phillips-Wren, Gloria;Jefferson, Theresa
Publication date	2026-06-30
Original Citation	Heavin, C, Phillips-Wren, G & Jefferson, T 2026, 'An exploratory study of factors influencing the use of Generative AI for decision#making: a missing link', Journal of Decision Systems, vol. 35, no. 1, 2688525, pp. 1-9. https://doi.org/10.1080/12460125.2026.2688525
Type of publication	Article (Peer reviewed)
Link to publisher's version	10.1080/12460125.2026.2688525
Rights	cc_by
Download date	2026-08-30 09:41:39
Item downloaded from	https://hdl.handle.net/10468/19056



UCC

University College Cork, Ireland
Coláiste na hOllscoile Corcaigh



An exploratory study of factors influencing the use of Generative AI for decision-making: a missing link

Ciara Heavin, Gloria Phillips-Wren & Theresa Jefferson

To cite this article: Ciara Heavin, Gloria Phillips-Wren & Theresa Jefferson (2026) An exploratory study of factors influencing the use of Generative AI for decision-making: a missing link, Journal of Decision Systems, 35:1, 2688525, DOI: [10.1080/12460125.2026.2688525](https://doi.org/10.1080/12460125.2026.2688525)

To link to this article: <https://doi.org/10.1080/12460125.2026.2688525>



© 2026 The Author(s). Published by Informa UK Limited, trading as Taylor & Francis Group.



Published online: 30 Jun 2026.



[Submit your article to this journal](#)



Article views: 57



[View related articles](#)



[View Crossmark data](#)

An exploratory study of factors influencing the use of Generative AI for decision-making: a missing link

Ciara Heavin^a, Gloria Phillips-Wren^b and Theresa Jefferson^b

^aDepartment of Business Information Systems, Cork University Business School, University College Cork, Cork, Ireland; ^bSellinger School of Business and Management, Loyola University Maryland, Baltimore, MD, USA

ABSTRACT

Generative AI (GenAI), particularly large language models (LLMs), is increasingly embedded in decision-making processes, making it essential to understand the drivers of its actual use. Despite their accessibility, LLMs are prone to hallucinations and bias, which contribute to user skepticism. This study examines the factors shaping perceptions of LLMs in personal and professional decision-making contexts, focusing on perceived usefulness, user trust, and personal values. An anonymous survey of 215 professionals was analysed using descriptive statistics, hierarchical cluster analysis, and regression analysis. Findings reveal that, contrary to expectations, trust and perceived usefulness alone do not sufficiently explain LLM adoption. Instead, personal values emerge as a significant yet underexplored determinant influencing usage decisions. These results extend current GenAI adoption research by incorporating the role of individual values. The study highlights the importance of context-sensitive, stakeholder-aware strategies and underscores the need for further research on value alignment across diverse cultural and regulatory environments.

ARTICLE HISTORY

Received 2 March 2026
Accepted 27 April 2026

KEYWORDS

Generative AI; decision making; trust; personal values; usefulness; professionals; ChatGPT

1. Background

Professionals face growing tension between the transformative potential of large language models (LLMs) and concerns about their opaque ‘black box’ nature, which undermines trust and adoption (Hassija et al., 2024; Shin, 2021). As LLMs become embedded in high-stakes decision-making across sectors such as finance, healthcare, law and education (Dwivedi et al., 2023), adoption has accelerated rapidly, rising from 33% to 65% within a year of ChatGPT’s 2022 launch Tavakoli et al. (2024). Yet uncertainties about accuracy, bias and ethical alignment persist (Rana et al., 2024). Traditional decision support systems are ill-suited for today’s complex, data-rich environments (Jarrahi, 2018; Shim et al., 2002), and while LLMs offer context-aware insights from vast datasets (Sedkaoui & Benaichouba, 2024), organisations are still determining effective implementation strategies (Dencik et al., 2023). As LLMs move into professional workflows, transparency, trust and accountability become essential (Do Khac et al., 2025; Rana et al., 2024), highlighting the need to understand other factors that shape professionals’ willingness to rely on LLMs for decision support. While trust in AI has been explored in recent IS literature, less is known about other factors that shape how professionals use LLMs for decision support. Thus, we pose the question:

Research Question: What factors influence a user’s actual use of Generative AI for decision-making tasks?

2. Artificial intelligence and trust

Trust is defined as ‘willingness to be vulnerable to another party based on expectations of their actions, irrespective of monitoring ability’ (Mayer et al., 1995, p. 712). Mayer et al.’s (1995) seminal framework identifies three trust dimensions: ability (competence to perform), benevolence (motivation to act in the trustor’s interest) and integrity (adherence to acceptable principles). In technology contexts, trust encompasses beliefs about system reliability, accuracy and ethical integrity (Gefen et al., 2003; McKnight et al., 2011;

CONTACT Ciara Heavin  c.heavin@ucc.ie  Department of Business Information Systems, Cork University Business School, University College Cork, Cork, Ireland

© 2026 The Author(s). Published by Informa UK Limited, trading as Taylor & Francis Group.

This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited. The terms on which this article has been published allow the posting of the Accepted Manuscript in a repository by the author(s) or with their consent.

Zhang et al., 2026). Trust directly predicts technology adoption and sustained use, particularly in high-stakes domains where errors carry significant consequences (McKnight et al., 2011). In healthcare and finance, where LLMs increasingly support clinical diagnoses and financial advice, trust becomes non-negotiable (Jarrahi, 2018).

However, AI trust is uniquely complex, involving technical, ethical and social dimensions (Habbal et al., 2024). The EU AI Act emphasises trustworthiness as foundational to AI development (European Commission, 2024). Critically, ‘incorrect levels of trust may cause misuse, abuse or disuse of technology’ (Jacovi et al., 2021, p. 624). Overtrust leads to uncritical acceptance of flawed outputs, while undertrust prevents beneficial adoption. Users often anthropomorphise AI systems, projecting human-like qualities and shifting from reliance to trust relationships (Jacovi et al., 2021). Anthropomorphism complicates trust calibration, as users may attribute understanding or intentionality to stochastic systems operating on probabilistic patterns. Such misattributions can further obscure the distinction between genuine model competence and erroneous outputs, creating fertile ground for misleading responses.

LLM hallucinations arise when models generate false or fabricated information with unwarranted confidence (Zhang et al., 2023). These hallucinations occur when outputs diverge from factual or expected content leading to inaccurate statements, invented references or flawed reasoning (Phillips-Wren & Virvou, 2025). They are not evidence of internal cognition, but rather statistical artefacts produced through probabilistic token prediction (Zhang et al., 2023). As a result, users must critically evaluate LLM-generated outputs (Phillips-Wren & Virvou, 2025).

LLMs also embed values and biases that may differ from human or culturally diverse value systems (Hadar-Shoval et al., 2024). For example, ChatGPT tends to reflect Western, Anglo-American perspectives (Cao et al., 2023; Tuna et al., 2024), a limitation acknowledged by OpenAI. Such cultural biases raise concerns about fairness across global user groups, and models trained on large-scale human-generated datasets may reproduce existing societal biases, potentially generating discriminatory content. A further challenge emerges as LLM-generated material increasingly becomes part of the training ecosystem. This introduces a feedback loop in which models learn from their own outputs, risking performance degradation and a reduction in the diversity of human knowledge captured in training dataset (Choudhury & Chaudhry, 2024). These issues can contribute to unreliable or unsound model decisions and may erode trust in artificial reasoning systems (Jacovi et al., 2021).

2.1. Personal values and technology adoption

Studies acknowledge that personal values influence human behaviour (Lee & Lyu, 2016). Values are ‘enduring beliefs or motivational goals relating to desired modes of conduct or end-states of existence’ and ‘direct people to behave in a way that fulfills their values, and drive them to change their behavior in order to diminish the inconsistency between their value priorities and behavior’ (Lee & Lyu, 2016, p. 324).

Personal values influence technology adoption. For example, the adoption of e-government services was shown to be mediated by two personal values related to its advantages, citizens’ time consciousness and environmental concern. In this case, personal values yielded a positive effect on adoption (Belanche et al., 2012). As a counterpoint, the adoption of self-service technologies was impacted negatively by personal values due to the need for human interaction and influenced by self-efficacy with SSTs. In this paper, we investigate the role of personal values in the adoption of LLMs for decision making.

2.2. From traditional DSS to AI-Augmented decision-making

DSS have evolved from early interactive tools for semi-structured problems (Power, 2002) to sophisticated AI-driven systems. The IFIP Working Group 8.3 (1981) defined DSS as ‘approaches for applying information-systems technology to increase the effectiveness of decision-makers’. Early DSS drew on Simon’s three-stage model (intelligence, design, choice); subsequent work shifted towards data-driven paradigms powered by analytics, machine learning and AI (Davenport & Harris, 2017; Pomerol, 1997). LLMs represent a paradigm shift, offering context-aware insights that exceed traditional analytics (Brynjolfsson et al., 2018; Javaid, 2024) and yet are easy to use. Their promise lies not in replacing human judgement but in creating human-AI

symbiosis (Jarrahi, 2018). Machines excel at large-scale data processing and pattern recognition, while humans contribute contextual understanding, ethical nuance and strategic vision.

Drawing on established IS theories, three factors shape intentions to use LLMs:

- **Perceived Usefulness** – The belief that LLMs improve decision quality or efficiency (Davis, 1989; Venkatesh & Davis, 2000). When users see clear benefits such as greater accuracy or faster analysis, adoption likelihood increases (Dwivedi et al., 2023; Gill et al., 2022).
- **Trust and Reliability** – Confidence in the model's accuracy, transparency and lack of bias (Glikson & Woolley, 2020; McKnight et al., 2011). Given risks of hallucinations, users rely on trustworthy performance, explainability and ethical output.
- **Value Compatibility** – Alignment with users' ethical principles, professional norms and organisational standards (Belanche et al., 2012; Turja et al., 2020). Adoption is more likely when LLMs fit moral frameworks, especially in domains with high ethical stakes (Oniani et al., 2023; Schwartz, 2012).

While existing research establishes trust as critical for technology adoption (Do Khac et al., 2025; McKnight et al., 2011) and documents LLM risks such as hallucinations and bias (Håkansson & Phillips-Wren, 2024), we did not find a study that integrates these streams to examine how the three factors, shaped by awareness of LLM limitations, influence professionals' adoption for decision-making tasks. This integrated perspective advances IS scholarship by bridging trust theory, values alignment, technology acceptance and emerging AI research to inform both theory and practice in AI-augmented decision-making.

3. Methodology

The aim of this research is to explore how GenAI and LLMs are used among professionals for decision-making purposes. To answer this research question and to ensure reliability and validity, an anonymous survey was designed using questions adapted from existing research (Choudhury & Shamszare, 2023; Turja et al., 2020). Survey questions were designed to assess how frequently professionals use GenAI technologies and to gauge their attitudes towards them, including trust, perceived usefulness and concerns such as alignment with personal values. We used a cross-sectional approach and multiple measurement items to operationalise the focal constructs (Hulland et al., 2018).

This anonymous survey enabled us to access a diverse audience, gathering a broad set of data pertaining to GenAI use. Professionals, part-time graduate students, undergraduate business and data science students were offered the opportunity to participate in this survey which was made available by the research team in Qualtrics over a ten-month period between October 2023 and July 2024. Following data cleaning, a total of 225 responses were obtained with 215 responses deemed complete and useful for data analysis purposes. Participant responses to survey items were captured using a scale of 1–5, with 1 corresponding to strongly disagree and 5 corresponding to strongly agree. The sample consisted mainly of younger respondents (60% aged 18–24) and was predominantly male (59%), with most participants using ChatGPT (74%) at least occasionally. LLM engagement was generally low-intensity, with 59% using them only 1–2 times per week, and participants represented a wide range of industries, especially finance/accounting (31%) and IT/operations (29%). Most used LLMs for personal purposes (65%), followed by school-related (40%) and work-related tasks (44%), with many selecting multiple use cases.

4. Results

Two primary analyses were undertaken with the survey data. SAS Enterprise Miner was used to conduct both analyses. A review of the correlations between variables indicated that there were 2 instances where correlation was greater than 0.75. Following the removal of these 2 variables, all remaining inputs had correlations less than 0.75 and values less than 5, thereby satisfying the risk of multicollinearity (Montgomery et al., 2021). The first analysis aimed to understand the characteristics of the study population and to identify distinct dimensions measured by the survey. The second analysis explored relationships between the identified dimensions with a focus on the perspectives of the sample group towards using GenAI to support decision-making tasks.

Table 1. Cluster characteristics.

Cluster	Number of items	Mean 1-R ²	Proportion Variance Explained	Example items	Interpretative Level
1	3	0.14	0.72	ChatGPT (or similar LLM) is trustworthy in the sense that it is dependable and credible.	Trust
2	3	0.26	0.59	Using ChatGPT (or similar LLM) runs counter to my own values.	Values
3	3	0.17	0.85	I think ChatGPT (or similar LLM) is useful to help me do my job.	Usefulness

4.1. Identification of distinct dimensions

Determining the appropriate input variables to include in a predictive model is often challenging due to redundancy that exists among the independent variables. Including redundant variables leads to models that overfit due to correlations between input variables. Therefore, to explore the underlying structure of questionnaire items and reduce dimensionality, we conducted a hierarchical cluster analysis using the Variable Clustering node in SAS Enterprise Miner (Aggarwal & Kosian, 2011). The Variable Clustering node within SAS Enterprise Miner provides an alternative to principal component analysis (PCA) by grouping variables instead of transforming them into unobservable components (Nelson, 2001). This technique has been successfully demonstrated by Lee et al. (2008), Woolston et al. (2012), Sanche and Lonergan (2006) and Khattab et al. (2020). Variable Clustering decreases variable redundancy, removes collinearity and assists in revealing the underlying structure of the input variables. This method iteratively partitions the variables into mutually exclusive clusters, ensuring that each cluster is highly correlated internally while maintaining low correlations with variables from other clusters. For each cluster, a cluster component is created. The cluster component is a linear combination of the variables assigned to the cluster. The cluster components are used as substitutes for the original input variables, thereby reducing the number of variables in a predictive model with ‘little loss of information’ (SAS, 2022).

The analysis was performed on a set of 9 items designed to measure survey respondents’ perceptions concerning the use of LLMs. (The survey also contained one item used to capture the respondents’ perceptions concerning the use of LLMs for decision making, which will be discussed in more detail in the next section). The clustering criterion was based on the $1-R^2$ ratio, set at a maximum threshold of 0.75, which allows each cluster to retain variables that explain at least 25% of shared variance.

The analysis resulted in the identification of three distinct dimensions, based on 9 survey questions, that identify the underlying dimensions of perceptions concerning the use of LLMs. Each of the clusters contained items with high internal cohesion, as indicated by low 1-R² values. Table 1 summarises the characteristics of each cluster, including the mean 1-R².

The internal consistency of each cluster was assessed using Cronbach’s alpha. The α values ranged from 0.84 to 0.88, suggesting that the items within each cluster are measuring cohesive constructs (Hulland et al., 2018). The inter-cluster correlations ranged from 0.56 to 0.71, indicating that each of the three dimensions is relatively distinct.

Each cluster was labelled based on the themes represented by the items within each cluster.

- (1) Trust: Cluster 1 grouped items associated with feelings of trust associated with the use of LLMs, including integrity, trustworthiness, reliability and security.
- (2) Values: Cluster 2 grouped items related to a user’s personal values and the use of LLMs, including runs counter to values, does not fit with worldview and is not appropriate considering values.
- (3) Usefulness: Cluster 3 grouped items related to the usefulness of LLMs, including usefulness in job and work-related tasks.

The three distinct dimensions represented by the above clusters suggest that these are key areas for understanding individuals’ perceptions concerning the use of LLMs.

4.2. Relationships among constructs related to decision-making

Linear regression analysis was conducted to determine the relationship between the dimensions of Trust, Values and Usefulness and the use of LLMs to support decision-making. The survey question, ‘I think it is

a good idea to use ChatGPT (or similar LLM) to support the decisions I make' was used for the dependent variable. The composite constructs, discussed in Section 4.1, of Trust, Values and Usefulness, were used as the predictor variables.

A multiple linear regression model was specified as follows:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \epsilon$$

where:

- Y = Decision Making
- X₁ = Values composite variable
- X₂ = Trust composite variable
- X₃ = Usefulness composite variable
- $\beta_0, \beta_1, \beta_2, \beta_3$ = Regression coefficients
- ϵ = Error term

Table 2 presents the coefficients from the regression analysis.

Based on the *p*-value of the F-test (<0.0001), the null hypothesis was rejected, indicating that collectively the independent variables contribute significantly to the model. The *R*² value of 0.47 suggests the model explains 48% of the variance concerning using LLMs to support decision making. All variables made significant independent contributions to the model.

As can be seen in Table 3, Trust is selected as the first variable and improves the model (over the intercept only model) by 37%. The next variable to enter the model is Usefulness, this improves the existing model by 7.5%. The last variable to enter the model is Values, which improves the existing model by 1.7%.

5. Discussion

5.1. Interpretation of results

The cluster analyses suggest that Trust, Values and Usefulness are distinct dimensions. The 1-*R*² values for all three clusters are within the acceptable range for meaningful clustering. The high proportion of variance explained by the first eigenvalue (>0.59) indicates that the items within each cluster are well aligned with the underlying factor represented by the cluster.

The low 1-*R*² values for Cluster 1 (Trust), Clusters 2 (Values) and Cluster 3 (Usefulness) are indicative of the strong relationships among the items within each cluster. Cluster 2 (Values) is focused on 'counter to values' with questions including runs counter to values, does not fit with worldview and is not appropriate considering values. Although we were surprised at the magnitude of Values in the prediction model, we note that AI has been characterised by some authors as a threat to society, taking people's jobs, and a concern for security and privacy (Khan et al., 2023). Cluster 3 (Usefulness) is focused on questions related

Table 2. Regression coefficients.

Variable	Coefficient (β)	p-value
Values	-0.1897	0.0404
Trust	0.4251	<.0001
Usefulness	0.2357	0.0053
Intercept (β_0)	-0.0388	0.0053

**According to the model, decision-making is negatively related to Values ($p < 0.04$) and positively related to Trust ($p < 0.01$) and Usefulness ($p < 0.01$). It is noted that the variable Values is coded in the survey as 'counter to my values', and the coding explains the negative relationship in the model.

Table 3. Adjusted R-square values and model improvement.

Model	Predictors included	<i>R</i> ²	Adjusted <i>R</i> ²	Model Improvement
Model 1	X ₁ (Trust)	0.371	0.3654	Model improves by 37% (over intercept alone) when Trust is added.
Model 2	X ₁ , X ₂ (Usefulness)	0.451	0.4406	Model improves by 7.5% when Useful is added to the model.
Model 3	X ₁ , X ₂ , X ₃ (Values)	0.473	0.4576	Model improves by 1.7% when Values is added to the model.

to usefulness in job and work-related tasks, and it is clear from the results that users recognise the technology as useful.

The higher $1-R^2$ values for Cluster 2 (Values) suggest that there is more variability within the cluster items. The corresponding survey questions are broader in scope reflecting the characteristics of the dimensions. Cluster 1 (Trust) included survey questions such as integrity, trustworthiness, reliability, and security. In particular, the lack of security of AI systems is often highlighted as a threat (Habbal et al., 2024), so users may give qualified answers to trust GenAI.

Exploring the relationship between the three dimensions yielded a linear regression model showing that the three dimensions of Trust, Values and Usefulness are statistically significant predictors of the use of GenAI for decision-making. The results show that Trust and Usefulness have a positive effect on decision-making, while Values has a negative effect due to its representation as 'counter to values'. While usefulness has long been recognised as essential to an individual's use of decision support technologies (Shim et al., 2002), and trust is a current research topic in GenAI especially regarding hallucinations (e.g. Christensen et al., 2024), it is surprising to see that personal values had a strong effect in the model. We have called Values (i.e. counter to personal values) as a missing link in understanding the intention to use GenAI and LLMs for decision-making tasks. Although these technologies offer potential benefits to decision-making such as data manipulation, analysis, interpretation and alternatives, an identifiable group feels that the use of GenAI is counter to their personal values.

6. Conclusion

Despite the widespread adoption and significance of GenAI and specific LLMs, our understanding of factors that shape professionals' use of these advanced technologies for decision support remains limited. Several limitations of the study are acknowledged. The sample size is modest, drawn culturally from the US and Ireland, and focused on professional and emerging professional groups that are expected to be familiar with technology systems including GenAI. However, we posit that it is representative of business groups who use technology and are likely to extend that use to GenAI technologies for multiple purposes including decision making. Another limitation concerns the measurement of the primary dependent variable. This research relied on a single survey item to assess decision-making: 'I think it is a good idea to use ChatGPT (or similar LLM) to support the decisions I make'. This approach may be considered 'vague' as it does not distinguish between different types of professional decisions, such as strategic, operational or ethical. In addition, the dependent variable does not address contextual complexity, a characteristic of modern decision environments.

The study is exploratory and leads to future research questions. Future iterations of the survey should employ a multi-item scale that specifies the nature and stakes of the decision-making task. Additionally, as the adoption of GenAI continues to increase, longitudinal research is needed to determine if user 'trust' and 'value alignment' change over time as users become more comfortable with these systems. As a new technology, GenAI offers opportunities for decision support and potential pitfalls such as the inclusion of bias and underrepresentation of alternate viewpoints that the research community should examine.

Consistent with existing empirical accounts for adoption of AI and other emerging technologies, trust in AI is one key factor influencing professionals' adoption use (Choudhury & Shamszare, 2023; Do Khac et al., 2025; Kaur et al., 2022; Petzel & Sowerby, 2025). However, while GenAI can enhance efficiency and provide data-driven insights, its opaque logic and potential for bias can erode managerial confidence and trust (Phillips-Wren & Virvou, 2025). Professionals and organisations must navigate the tension between leveraging GenAI's capabilities and maintaining human oversight to ensure accountable and value-aligned decisions.

A unique aspect of this study is the identification of personal values as a crucial factor in understanding the use of AI and GenAI for decision support. Sedkaoui and Benaichouba (2024) advocate for new studies on GenAI and LLMs, highlighting the need to investigate 'the long-term societal and economic implications of GenAI, as well as its impact on human values and cultural expression, which will be crucial for guiding responsible innovation'. While much of the existing literature focuses on the technical and cognitive dimensions of GenAI adoption (Berente et al., 2021; Choudhury & Shamszare, 2023), this research suggests that a values-driven approach could yield deeper insights into user

behaviour and technology acceptance, particularly beyond existing Western values and cultural frameworks. These perspectives have largely overlooked how deeply personal and collective values shape individuals' willingness to engage with emerging tools such as GenAI. By explicitly recognising the interplay between users' ethical priorities and cultural norms, researchers can uncover richer patterns of technology appropriation and resistance. Future studies could build on this finding by exploring a 'values by design' approach, incorporating personal and organisational values directly into the design and implementation of AI to foster social inclusion, trust, and just, broad, equitable innovation.

The interplay between personal values and the use of GenAI for decision making also points to a broader gap in AI design. Current AIs lack the ability to account for the diverse personal values of their users, which may hinder their effectiveness and trustworthiness in professional decision-making contexts. Addressing this issue requires thoughtful integration of values into the design process, ensuring that AI is not only technically sound but also contextually relevant and ethically grounded. This research makes a significant theoretical contribution by highlighting the importance of personal values in GenAI adoption. Practically, it suggests that organisations should consider these factors when implementing GenAI, as aligning technology with user values may enhance the trust and effectiveness of decision support in the future.

Author contributions

CRedit: **Ciara Heavin:** Conceptualization, Data curation, Formal analysis, Methodology, Writing – original draft, Writing – review & editing; **Gloria Phillips-Wren:** Conceptualization, Data curation, Formal analysis, Methodology, Writing – original draft, Writing – review & editing; **Theresa Jefferson:** Data curation, Formal analysis, Methodology, Writing – original draft, Writing – review & editing.

Disclosure statement

No potential conflict of interest was reported by the author(s).

Generative AI and AI-assisted technologies were not used in any capacity during the preparation of this manuscript.

ORCID

Ciara Heavin  <http://orcid.org/0000-0001-8237-3350>

Gloria Phillips-Wren  <http://orcid.org/0000-0002-0461-5577>

Theresa Jefferson  <http://orcid.org/0009-0001-4950-5216>

Data statement

Survey data are publicly available at <https://drive.google.com/drive/folders/10i93c6i8yib4bsofrAj3xpF4ypnpRriN?usp=sharing>

References

- Aggarwal, V., & Ksian, S. (2011). *Feature selection and dimension reduction techniques in SAS*. EXL Service.
- Belanche, D., Casaló, L.V., & Flavián, C. (2012). Integrating trust and personal values into the technology acceptance model: The case of e-government services adoption. *Cuadernos de Economía y Dirección de la Empresa*, 15(4), 192–204. <https://doi.org/10.1016/j.cede.2012.04.004>
- Berente, N., Gu, B., Recker, J., & Santhanam, R. (2021). Managing artificial intelligence. *MIS Quarterly*, 45(3), 1433–1450. <https://doi.org/10.25300/MISQ/2021/16274>
- Brynjolfsson, E., Mitchell, T., & Rock, D. (2018, May). What can machines learn and what does it mean for occupations and the economy? AEA Papers and Proceedings (Vol. 108. pp. 43–47). Nashville, TN: American Economic Association.
- Cao, Y., Zhou, L., Lee, S., Cabello, L., Chen, M., & Hershcovich, D. (2023). Assessing cross-cultural alignment between ChatGPT and human societies: An empirical study. In *Proceedings of the first workshop on cross-cultural considerations in NLP (C3NLP)* (pp. 53–67).
- Choudhury, A., & Chaudhry, Z. (2024). Large language models and user trust: Consequence of self-referential learning loop and the deskilling of health care professionals. *Journal of Medical Internet Research*, 26, e56764. <https://doi.org/10.2196/56764>

- Choudhury, A., & Shamszare, H. (2023). Investigating the impact of user trust on the adoption and use of ChatGPT: Survey analysis. *Journal of Medical Internet Research*, 25, e47184. <https://doi.org/10.2196/47184>
- Christensen, J., Hansen, J.M., & Wilson, P. (2024). Understanding the role and impact of generative AI hallucination within consumers' tourism decision-making processes. *Current Issues in Tourism*, 28(4), 1–16. <https://doi.org/10.1080/13683500.2023.2300032>
- Davenport, T., & Harris, J. (2017). *Competing on analytics: Updated, with a new introduction: The new science of winning*. Harvard Business Press.
- Davis, F.D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- Dencik, J., Goehring, B., & Marshall, A. (2023). Managing the emerging role of generative AI in next-generation business. *Strategy & Leadership*, 51(6), 30–36.
- Do Khac, L.T., Leyer, M., Wichmann, J., & Richter, A. (2025). Divisional perspectives on trust in generative AI: An empirical study of senior decision-makers. *Enterprise Information Systems*, 19(7–8), 2481557. <https://doi.org/10.1080/17517575.2025.2481557>
- Dwivedi, Y.K., Kshetri, N., Hughes, L., Slade, E.L., Jeyaraj, A., Kar, A.K., Wright, R., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M.A., Al-Busaidi, A.S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wirtz, J. (2023). Opinion paper: “So what if ChatGPT wrote it?” multidisciplinary perspectives on generative conversational AI for research, practice, and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- European Council of the European Union. (2024). <https://consilium-europa.libguides.com/c.php?g=690732&p=4948483>
- Gefen, D., Karahanna, E., & Straub, D.W. (2003). Trust and TAM in online shopping: An integrated model. *MIS Quarterly*, 27(1), 51–90. <https://doi.org/10.2307/30036519>
- Gill, S.S., Xu, M., Ottaviani, C., Patros, P., Bahsoon, R., Shaghaghi, A., Uhlig, S., Stankovski, V., Wu, H., Abraham, A., Singh, M., Mehta, H., Ghosh, S.K., Baker, T., Parlikad, A.K., Lutfiyya, H., Kanhere, S.S., Sakellariou, R., Dustdar, S., ... Brandic, I. (2022). AI for next generation computing: Emerging trends and future directions. *Internet of Things*, 19, 100514. <https://doi.org/10.1016/j.iot.2022.100514>
- Glikson, E., & Woolley, A.W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Habbal, A., Ali, M.L., & Abuzaraida, M.A. (2024). AI trust, risk, and security management (AI TRISM): Frameworks, applications, challenges, and future research directions. *Expert Systems With Applications*, 240, 122442. <https://doi.org/10.1016/j.eswa.2023.122442>
- Hadar-Shoval, D., Asraf, K., Mizrahi, Y., Haber, Y., & Elyoseph, Z. (2024). Assessing alignment of LLMs with human values for mental health: A cross-sectional study using Schwartz's theory. *JMIR Mental Health*, 11, e55988. <https://doi.org/10.2196/55988>
- Håkansson, A., & Phillips-Wren, G. (2024). Generative AI and large language models: Benefits, drawbacks, future, and recommendations. *Procedia Computer Science*, 246, 5458–5468. <https://doi.org/10.1016/j.procs.2024.09.689>
- Hassija, V., Chamola, V., Mahapatra, A., Singal, A., Goel, D., Huang, K., Scardapane, S., Spinelli, I., Mahmud, M., & Hussain, A. (2024). Interpreting black-box models: A review on explainable artificial intelligence. *Cognitive Computation*, 16(1), 45–74. <https://doi.org/10.1007/s12559-023-10179-8>
- Hulland, J., Baumgartner, H., & Smith, K.M. (2018). Marketing survey research best practices: Evidence and recommendations from a review of JAMS articles. *Journal of the Academy of Marketing Science*, 46(1), 92–108. <https://doi.org/10.1007/s11747-017-0532-y>
- Jacovi, A., Marasović, A., Miller, T., & Goldberg, Y. (2021). Formalizing trust in artificial intelligence. *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency* (pp. 624–635). Virtual Event, Canada: ACM.
- Jarrah, M.H. (2018). Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577–586. <https://doi.org/10.1016/j.bushor.2018.03.007>
- Javaid, H.A. (2024). The future of financial services: Integrating AI for smarter, more efficient operations. *MZ Kolbjørnsrud Journal of Artificial Intelligence*, 1(2), 1–7.
- Kaur, D., Uslu, S., Rittichier, K.J., & Durresi, A. (2022). Trustworthy artificial intelligence: A review. *ACM Computing Surveys*, 55(2), 1–38. <https://doi.org/10.1145/3491>
- Khan, B., Fatima, H., Qureshi, A., Kumar, S., Hanan, A., Hussain, J., & Abdullah, S. (2023). Drawbacks of artificial intelligence and their potential solutions in the healthcare sector. *Biomedical Materials & Devices*, 1(2), 731–738. <https://doi.org/10.1007/s44174-023-00063-2>
- Khattab, M.F., Al-Muqdad, S.W., Abo, R.K., & Merkel, B.J. (2020). Variable reduction for water quality investigation using VARCLUS technique: A case study of Mosul Dam Lake, Northern Iraq. *Water Resources*, 47(6), 1005–1011. <https://doi.org/10.1134/S0097807820060093>
- Lee, H.J., & Lyu, J. (2016). Personal values as determinants of intentions to use self-service technology in retailing. *Computers in Human Behavior*, 60, 322–332. <https://doi.org/10.1016/j.chb.2016.02.051>
- Lee, T., Duling, D., Liu, S., & Latour, D. (2008). *Two-stage variable clustering for large data sets*. SAS Global Forum SAS Institute Inc.
- Mayer, R.C., Davis, J., & Schoorman, F.D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.2307/258792>

- McKnight, D.H., Carter, M., Thatcher, J.B., & Clay, P.F. (2011). Trust in a specific technology: An investigation of its components and measures. *ACM Transactions on Management Information Systems*, 2(2), Article 1.1–25. <https://doi.org/10.1145/1985347.1985353>
- Montgomery, D.C., Peck, E.A., & Vining, G.G. (2021). *Introduction to linear regression analysis* (6th ed.). John Wiley & Sons.
- Nelson, B.D. (2001). Variable reduction for modeling using PROC VARCLUS. Proceedings of the Twenty-Sixth Annual SAS® Users Group International Conference, SAS Institute Inc., Cary, NC.
- Oniani, D., Hilsman, J., Peng, Y., Poropatich, R.K., Pamplin, J.C., Legault, G.L., & Wang, Y. (2023). Adopting and expanding ethical principles for generative AI from military to healthcare. *NPJ Digital Medicine*, 6(1), 225. <https://doi.org/10.1038/s41746-023-00965-x>
- Petzel, Z.W., & Sowerby, L. (2025). Prejudiced interactions with large language models (LLMs) reduce trustworthiness and behavioral intentions among members of stigmatized groups. *Computers in Human Behavior*, 165, 108563. <https://doi.org/10.1016/j.chb.2025.108563>
- Phillips-Wren, G., & Virvou, M. (2025). Issues and trends in generative AI technologies for decision making. *Intelligent Decision Technologies*, 19(2), 574–584. <https://doi.org/10.1177/18724981251320551>
- Pomerol, J.C. (1997). Artificial intelligence and human decision making. *European Journal of Operational Research*, 99(1), 3–25. [https://doi.org/10.1016/S0377-2217\(96\)00378-5](https://doi.org/10.1016/S0377-2217(96)00378-5)
- Power, D.J. (2002). *Decision support systems: Concepts and resources for managers*. Greenwood Publishing.
- Rana, N.P., Pillai, R., Sivathanu, B., & Malik, N. (2024). Assessing the nexus of generative AI adoption, ethical considerations, and organizational performance. *Technovation*, 135, 103064. <https://doi.org/10.1016/j.technovation.2024.103064>
- Sanche, R., & Loneragan, K. (2006, March). Variable reduction for predictive modeling with clustering. *Casualty Actuarial Society Forum*, 89–100.
- SAS. (2022). Variable clustering node. <https://documentation.sas.com/doc/en/emref/15.2/p19e837tepmjz0n1hjt2gdk3sqfg.htm>
- Schwartz, S.H. (2012). An overview of the Schwartz theory of basic values. *Online Readings in Psychology and Culture*, 2(1), 11. <https://doi.org/10.9707/2307-0919.1116>
- Sedkaoui, S., & Benaichouba, R. (2024). Generative AI as a transformative force for innovation: A review of opportunities, applications and challenges. *European Journal of Innovation Management*, 29(3), 677–701. <https://doi.org/10.1108/EJIM-02-2024-0129>
- Shim, J.P., Warkentin, M., Courtney, J.F., Power, D.J., Sharda, R., & Carlsson, C. (2002). Past, present, and future of decision support technology. *Decision Support Systems*, 33(2), 111–126. [https://doi.org/10.1016/S0167-9236\(01\)00139-7](https://doi.org/10.1016/S0167-9236(01)00139-7)
- Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI. *International Journal of Human-Computer Studies*, 146, 102551. <https://doi.org/10.1016/j.ijhcs.2020.102551>
- Tavakoli, A., Harreis, H., Rowshankish, K., & Bogobowicz, M. (2024). *Charting a path to the data- and AI-driven enterprise of 2030*. McKinsey & Company. <https://www.mckinsey.com/capabilities/tech-and-ai/our-insights/charting-a-path-to-the-data-and-ai-driven-enterprise-of-2030>
- Tuna, M., Schaaff, K., & Schlippe, T. (2024). Effects of language- and culture-specific prompting on ChatGPT. *Proceedings of the 2nd International Conference on Foundation and Large Language Models (FLLM)* (pp. 73–81). Dubai, United Arab Emirates: IEEE. <https://doi.org/10.1109/FLLM63129.2024.10852463>
- Turja, T., Aaltonen, I., Taipale, S., & Oksanen, A. (2020). Robot acceptance model for care (RAM-care): A principled approach to the intention to use care robots. *Information & Management*, 57(5), 103220. <https://doi.org/10.1016/j.im.2019.103220>
- Venkatesh, V., & Davis, F.D. (2000). A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management Science*, 46(2), 186–204. <https://doi.org/10.1287/mnsc.46.2.186.11926>
- Woolston, A., Tu, Y.K., Baxter, P.D., & Gilthorpe, M.S. (2012). A comparison of different approaches to unravel the latent structure within metabolic syndrome. *PLOS ONE*, 7(4), e34410. <https://doi.org/10.1371/journal.pone.0034410>
- Zhang, Y., Li, Y., Cui, L., Cai, D., Liu, L., Fu, T., Huang, X., Zhao, E., Zhang, Y., Chen, Y., Luu, A.T., Bi, W., Shi, F., & Shi, S. (2023). Siren's song in the AI ocean: A survey on hallucination in large language models. *Computational Linguistics*, 51(4), 1373–1418.
- Zhang, Y., Tian, J., & Deng, F. (2026). In generative AI we trust: An exploratory study on dimensionality and structure of user trust in ChatGPT. *Interacting With Computers*, 38(1), 58–76. <https://doi.org/10.1093/iwc/iwaf029>