

Understanding User Experience in Voice Assistants: The Impact of Functional Suitability and Robust Speech Recognition

David Šajina, Andrej Korica, and Tihomir Orehovački

Faculty of Informatics, Juraj Dobrila University of Pula, Zagrebačka 30, 52100 Pula, Croatia

ABSTRACT

This study investigates user experience in the context of voice assistants, with a particular focus on the role of speech interaction quality. While prior research has emphasized multiple dimensions of system quality, voice-based interaction introduces unique challenges related to the accurate interpretation of spoken input. To address this, a research model comprising four constructs—Robust Speech Recognition, Functional Suitability, Completeness, and Attitude Toward Use—was developed and empirically tested. Data were collected from 50 voice assistant users and analyzed using partial least squares structural equation modeling (PLS-SEM). The results indicate that robust speech recognition significantly influences functional suitability, completeness, and attitude toward use, highlighting its central role in shaping user evaluations. In contrast, functional suitability did not exhibit a significant effect on attitude toward use, suggesting that users may form attitudes primarily based on immediate interaction quality rather than overall system performance. The findings contribute to the understanding of voice assistants by demonstrating that robust speech recognition is a key driver of perceived functional suitability, completeness, and attitude toward use. They further suggest that traditional assumptions from technology acceptance research may not fully apply in voice-based interaction contexts. From a practical perspective, the results emphasize the importance of improving speech recognition performance under real-world conditions, as this may strengthen perceived functional suitability and completeness while fostering more favorable attitudes toward use.

Keywords: Voice assistants, Robust speech recognition, Functional suitability, Completeness, Attitude toward use, User experience, PLS-SEM

INTRODUCTION

Voice assistants represent a rapidly evolving class of intelligent systems that enable users to interact with technology through natural language. By allowing hands-free and intuitive interaction, these systems have become widely adopted across various domains, particularly in smart home environments, mobile devices, and increasingly in educational contexts (Terzopoulos and Satratzemi, 2020). Their growing popularity can be attributed to the convenience they offer, as users can perform tasks without direct physical

interaction with devices. This is especially beneficial in situations where manual interaction is impractical, such as while driving, as well as for users with diverse functional needs. From a technological perspective, voice assistants rely on speech recognition and natural language processing to transform spoken input into executable commands, with advances in machine learning enabling improved performance across diverse accents, dialects, and usage conditions (Yang et al., 2014; Khanzode and Sarode, 2020).

Despite these advancements, ensuring a high-quality user experience remains a critical challenge. Prior research has examined various aspects of voice assistants, including system functionality, usability, and user satisfaction. Comparative studies have evaluated the performance of different voice assistants across tasks such as information retrieval and task execution, highlighting differences in response quality and interaction design (López et al., 2018; Berdasco et al., 2019). Other studies have focused on user experience measurement, proposing standardized scales to assess usability and perceived system quality (Zwakman et al., 2020). Additionally, research has shown that factors such as language proficiency and accent can significantly influence user satisfaction, indicating the importance of robust speech processing capabilities (Pal et al., 2019). Further research has examined user experience and performance aspects of intelligent personal assistants, particularly in educational contexts, highlighting the importance of usability, satisfaction, and task performance as key dimensions shaping user evaluations (Babić et al., 2018; Orehovački et al., 2019). While these studies provide valuable insights, they primarily conceptualize system quality as a multidimensional construct without explicitly isolating the role of speech interaction quality.

Understanding user experience in voice assistants therefore requires a focused examination of core system capabilities that directly shape interaction outcomes. Voice user interfaces (VUIs) have transformed human–computer interaction by enabling natural language communication, where system performance depends heavily on speech processing and interaction quality (McTear et al., 2016; Følstad and Brandtzaeg, 2017). In such systems, the ability to accurately interpret spoken input and generate appropriate responses represents a fundamental determinant of perceived system effectiveness and user satisfaction.

In response to this gap, this study examines how speech recognition quality influences user experience by shaping key system-level perceptions in voice-based interaction. By focusing on real-world usage conditions, the study aims to provide clearer insight into how core system capabilities shape user perceptions and evaluations in such environments.

RESEARCH DESIGN

Constructs

The proposed research model consists of four reflective constructs: Robust Speech Recognition (RSR), Functional Suitability (FUS), Completeness (CPL), and Attitude Toward Use (ATU). These constructs capture key dimensions of user experience in voice assistants, with particular emphasis on speech interaction quality and perceived system performance.

Robust Speech Recognition (RSR) refers to the ability of a voice assistant to accurately interpret spoken input despite grammatical inaccuracies, pronunciation variability, and environmental noise.

Functional Suitability (FUS) represents the extent to which a voice assistant effectively performs intended tasks under real-world conditions.

Completeness (CPL) refers to the degree to which a voice assistant provides comprehensive and relevant outputs that fully address user requests.

Attitude Toward Use (ATU) reflects users' overall evaluative judgment regarding the desirability of using a voice assistant.

Hypotheses Development

System quality represents a fundamental determinant of user experience in information systems. The DeLone and McLean IS success model conceptualizes system quality as a key antecedent of user satisfaction and system use, emphasizing that technical performance directly shapes user perceptions (DeLone and McLean, 2003). In the context of voice assistants, speech recognition accuracy constitutes a critical dimension of system quality, as it directly mediates interaction between users and the system.

Voice interaction differs from traditional interfaces in that it relies on natural language input, which is inherently variable, ambiguous, and sensitive to environmental conditions. This introduces additional complexity, as systems must process variations in pronunciation, grammar, and background noise. Prior research demonstrates that speech recognition errors disrupt interaction flow, increase cognitive effort, and reduce perceived system reliability (Luger and Sellen, 2016; Porcheron et al., 2018). These disruptions not only affect immediate interaction but also influence broader perceptions of system performance.

From a causal perspective, accurate input processing represents a prerequisite for successful task execution. When user input is correctly interpreted, systems are able to perform intended actions more reliably and efficiently. Conversely, recognition errors propagate through the interaction process, limiting the system's ability to execute tasks. Empirical studies on conversational agents confirm that improvements in speech understanding significantly enhance perceived usability and system effectiveness (McTear, 2020; Følstad et al., 2018). Based on this theoretical reasoning, the following hypothesis is formulated:

H1. Robust Speech Recognition (RSR) positively influences Functional Suitability (FUS).

Beyond enabling task execution, input accuracy also determines the quality of system outputs. Information quality theory suggests that system usefulness depends on the completeness and relevance of the information provided to users (Petter et al., 2013). In voice assistants, output generation is directly dependent on how accurately the system interprets user intent.

When speech recognition performs reliably, systems can retrieve relevant resources, suggest appropriate alternatives, and generate more comprehensive responses. In contrast, failures in natural language understanding reduce the scope and relevance of outputs, often resulting in incomplete or incorrect responses. Prior research indicates that such failures significantly reduce the coverage and usefulness of system outputs (Zamora, 2017; Gnewuch et al., 2017), while interaction effectiveness has been shown to influence perceived output quality (Ashktorab et al., 2019; Kiseleva et al., 2016).

From a process perspective, speech recognition represents the initial stage in an information processing pipeline, where errors propagate to subsequent stages and affect output completeness. In line with these arguments, the following hypothesis is proposed:

H2. Robust Speech Recognition (RSR) positively influences Completeness (CPL).

The relationship between system performance and user attitudes is well established in technology acceptance research. The Technology Acceptance Model (TAM) posits that perceived usefulness is a primary determinant of user attitudes (Marangunić and Granić, 2015). Systems that effectively support task completion reduce effort and increase perceived value, which contributes to more favorable evaluations.

In voice assistants, reliable task execution signals system effectiveness and reduces uncertainty associated with interaction outcomes. This is particularly important in voice-based environments, where users rely heavily on system responses due to the absence of visual feedback. Empirical evidence confirms that system effectiveness and performance quality play a central role in shaping user attitudes and behavioral responses (Dwivedi et al., 2021).

Furthermore, consistent performance contributes to the development of trust and habitual use, as users are more likely to adopt systems that reliably meet their expectations. Therefore, it is hypothesized that:

H3. Functional Suitability (FUS) positively influences Attitude Toward Use (ATU).

In addition to indirect effects, interaction quality may also directly influence user attitudes. In voice assistants, speech recognition is the primary interaction mechanism, making it one of the most salient and immediately observable system characteristics. Users form rapid evaluations based on interaction fluency, responsiveness, and perceived system competence.

High-quality interaction reduces cognitive effort and enhances the overall user experience, while recognition errors disrupt conversational flow and increase frustration. Prior research shows that interaction quality strongly influences user satisfaction and perceived system competence (Cowan et al., 2017; Clark et al., 2019), while interaction breakdowns caused by recognition errors reduce user acceptance (Luger and Sellen, 2016; Yang et al., 2020).

From an experiential perspective, interaction quality represents a direct driver of user evaluation, independent of task outcomes. Consequently, the following hypothesis is advanced:

H4. Robust Speech Recognition (RSR) positively influences Attitude Toward Use (ATU).

Data Collection and Analysis

Data were collected using a structured questionnaire distributed via Google Forms. The instrument included multiple items for each construct: Robust Speech Recognition (RSR), Functional Suitability (FUS), and Attitude Toward Use (ATU) were each measured using four items, while Completeness (CPL) was measured using three items. All items were assessed on a five-point Likert scale ranging from 1 (strongly disagree) to 5 (strongly agree). The questionnaire also captured demographic characteristics such as gender, age, participant status, type of voice assistant used, and usage frequency.

To examine the proposed relationships, partial least squares structural equation modeling (PLS-SEM) was applied (Hair et al., 2022). This approach is suitable for exploratory research and models with relatively small samples. Following the inverse square root method (Kock and Hadaya, 2018), the minimum required sample size was estimated at 45, and the obtained sample size satisfies this threshold. All analyses were conducted using SmartPLS 4.1.1.8 (Ringle et al., 2022).

RESULTS

Sample Characteristics

The final sample consisted of 50 respondents and was predominantly male (74%), with 26% female participants. In terms of age distribution, the majority of respondents belonged to Generation Z (78%), while 22% were classified as Generation Y, indicating a relatively young and digitally experienced sample.

Regarding participants' status, 50% were students, 24% were employed, 18% were pupils, and 8% were unemployed. With respect to the type of voice assistant used, Google Assistant was the most frequently reported platform (48%), followed by Amazon Alexa (30%), while 22% of respondents reported using other voice assistants. In terms of usage frequency, the majority of participants reported using voice assistants occasionally (78%), while only 22% indicated frequent use. This suggests that most respondents can be characterized as casual users, which may have implications for how they evaluate system characteristics.

Measurement Model Assessment

The measurement model was evaluated by examining indicator reliability, internal consistency reliability, convergent validity, and discriminant validity (Hair et al., 2022). Indicator reliability was assessed using standardized outer loadings. Following commonly applied purification criteria (Hair et al., 2022), indicators with loadings below 0.708 were excluded from the final model. Based on this threshold, item CPL1 was removed from further analysis. For the retained indicators, standardized loadings ranged from 0.782 to 0.948, indicating that the constructs explained between 61.15% and 89.87% of the variance in their respective items.

Internal consistency reliability was assessed using Cronbach's alpha, composite reliability (rho_C), and the consistent reliability coefficient (rho_A). All constructs demonstrated good reliability, with values ranging from 0.721 to 0.959, exceeding the recommended threshold of 0.70 (Hair et al., 2022). Convergent validity was examined via Average Variance Extracted (AVE), and all constructs surpassed the 0.50 criterion, supporting adequate convergent validity (Hair et al., 2022). Table 1 reports these reliability and convergent validity results.

Table 1: Construct reliability and validity.

Construct	Cronbach's α	rho_A	rho_C	AVE
ATU	0.853	0.873	0.900	0.693
CPL	0.721	0.937	0.867	0.767
FUS	0.943	0.949	0.959	0.855
RSR	0.889	0.924	0.923	0.750

Discriminant validity was evaluated using the Heterotrait–Monotrait ratio (HTMT). All HTMT values were below the recommended cutoff of 0.85 for conceptually distinct constructs, indicating adequate separation between constructs (Hair et al., 2022). An overview of the HTMT results is provided in Table 2.

Table 2: HTMT ratios.

	ATU	CPL	FUS	RSR
ATU				
CPL	0.199			
FUS	0.490	0.284		
RSR	0.549	0.351	0.619	

Structural Model Assessment

Following the establishment of measurement model adequacy, the structural model was evaluated with respect to collinearity, explanatory power, hypothesis testing, and effect sizes (Hair et al., 2022). Collinearity among predictor constructs was assessed using the Variance Inflation Factor (VIF). The obtained VIF values ranged from 1.000 to 1.532, indicating no issues with multicollinearity (see Table 3).

Table 3: Collinearity statistics (VIF).

	ATU	CPL	FUS	RSR
ATU				
CPL				
FUS	1.532			
RSR	1.532	1.000	1.000	

The explanatory power of the model was examined using the coefficient of determination (R^2), with adjusted R^2 values reported to account for model complexity (Hair et al., 2022). Interpretation followed established benchmarks from software quality evaluation research, where values of 0.15, 0.34, and 0.46 indicate weak, moderate, and substantial explanatory power, respectively (Orehovački, 2013). The results indicate that robust speech recognition explained 33.4% of the variance in functional suitability, reflecting moderate explanatory power. In contrast, robust speech recognition and functional suitability jointly accounted for 25.5% of the variance in attitude toward use, indicating weaker explanatory power. Additionally, robust speech recognition explained 7.9% of the variance in completeness, suggesting very low explanatory capacity.

Hypotheses were tested using a bootstrapping procedure with 5,000 resamples and two-tailed significance testing (Hair et al., 2022). The results revealed that robust speech recognition had a significant positive effect on functional suitability ($\beta = 0.589$, $p < 0.001$), completeness ($\beta = 0.313$, $p < 0.05$), and attitude toward use ($\beta = 0.372$, $p < 0.05$), thereby supporting H1, H2, and H4. Conversely, the effect of functional suitability on attitude toward use was not statistically significant ($\beta = 0.223$, $p = 0.114$), indicating that H3 was not supported. The hypothesis testing results are summarized in Table 4.

Table 4: Hypothesis testing results.

Hypothesis	β	t	p	Decision
H1. RSR $\rightarrow$ FUS	0.589	6.487	< 0.001	Supported
H2. RSR $\rightarrow$ CPL	0.313	2.269	0.023	Supported
H3. FUS $\rightarrow$ ATU	0.223	1.582	0.114	Not Supported
H4. RSR $\rightarrow$ ATU	0.372	3.199	0.001	Supported

Effect sizes were evaluated using f^2 , where values of 0.02, 0.15, and 0.35 correspond to small, medium, and large effects, respectively (Hair et al., 2022). The findings indicate that robust speech recognition exerted a large effect on functional suitability ($f^2 = 0.532$), while its effects on attitude toward use ($f^2 = 0.127$) and completeness ($f^2 = 0.108$) were small.

DISCUSSION

Study findings highlight the central role of robust speech recognition in shaping user experience with voice assistants. The results indicate that speech recognition capabilities significantly influence both functional suitability and completeness, confirming that accurate interpretation of spoken input represents a foundational prerequisite for effective system performance. In voice-based interaction, where communication relies entirely on natural language, input processing directly determines the success of subsequent system responses.

The strong effect of robust speech recognition on functional suitability suggests that users primarily evaluate voice assistants based on their ability to correctly execute commands under realistic conditions. In environments characterized by noise, pronunciation variability, and imperfect language use, systems that maintain high recognition accuracy are perceived as more reliable and effective. This reinforces the view that interaction quality in voice interfaces is closely intertwined with perceived system functionality.

The significant relationship between robust speech recognition and completeness further indicates that accurate input interpretation contributes not only to task execution but also to the quality and coverage of system outputs. When user intent is correctly identified, the system is better positioned to retrieve relevant resources and provide comprehensive responses. Conversely, recognition errors may propagate through the interaction process, resulting in incomplete or less relevant outputs.

Robust speech recognition also demonstrated a direct influence on attitude toward use, highlighting its role as a salient and immediately observable system characteristic. Users form rapid evaluations based on interaction fluency and responsiveness, suggesting that reduced cognitive effort and smoother interaction contribute to more favorable perceptions of the system.

In contrast, the relationship between functional suitability and attitude toward use was not supported. This finding suggests that users—particularly those with limited or occasional usage—may form attitudes primarily based on immediate interaction experiences rather than on a broader evaluation of system performance. In such contexts, speech recognition quality appears to overshadow downstream system characteristics. Additionally, the relatively low explanatory power observed for attitude toward use indicates that other factors, such as trust, perceived privacy, or hedonic aspects, may play a more prominent role in shaping user evaluations.

From a theoretical perspective, the findings suggest that speech-driven interaction introduces specific dynamics that differ from traditional information systems contexts. In voice-based environments, user evaluations appear to be shaped more strongly by immediate interaction quality than by broader assessments of system functionality, indicating that established relationships in technology acceptance research may not fully apply.

From a practical standpoint, the results point to the need for design strategies that prioritize interaction robustness under real-world conditions. This includes improving system performance in noisy environments, accommodating diverse speech patterns, and ensuring reliable interpretation of imperfect input. Such improvements may strengthen perceived functional suitability and completeness while fostering more favorable attitudes toward use.

Despite these contributions, several limitations should be acknowledged. The study is based on a relatively small and demographically homogeneous sample, which may limit the generalizability of the findings. The predominance of younger and predominantly occasional users suggests that the results

primarily reflect everyday interaction scenarios rather than advanced usage contexts. Furthermore, the relatively low explanatory power observed for certain constructs indicates that additional variables not included in the model may influence user experience.

CONCLUSION

This study provides evidence that user experience in voice assistants is strongly shaped by core interaction capabilities, particularly those related to speech-based input processing. Rather than being driven primarily by broader system characteristics, user evaluations appear to depend on how effectively systems support natural and seamless interaction in real-world conditions.

These findings contribute to the growing body of research on voice interfaces by highlighting the need to reconsider traditional assumptions from technology acceptance literature in the context of speech-driven interaction. The results suggest that interaction quality may play a more prominent role than previously assumed, particularly for users with limited or occasional experience.

Future research should extend the proposed model by incorporating additional factors, such as trust, perceived privacy, and hedonic motivation, in order to better capture the complexity of user evaluations. Furthermore, studies based on larger and more diverse samples, as well as longitudinal and context-specific approaches, could provide deeper insights into how user experience evolves across different usage scenarios.

REFERENCES

- Ashktorab, Z., Jain, M., Liao, Q. V., Weisz, J. D. (2019). Resilient Chatbots: Repair Strategy Preferences for Conversational Breakdowns, *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, pp. 1–12.
- Babić, S., Orehovački, T., Etinger, D. (2018). Perceived User Experience and Performance of Intelligent Personal Assistants Employed in Higher Education Settings, *Proceedings of the 41st International Convention on Information and Communication Technology, Electronics and Microelectronics (MIPRO)*, pp. 830–834.
- Berdasco, A., López, G., Diaz, I., Quesada, L., Guerrero, L.A. (2019). User Experience Comparison of Intelligent Personal Assistants: Alexa, Google Assistant, Siri and Cortana, *Proceedings*, 31, 51.
- Clark, L., Pantidi, N., Cooney, O., Doyle, P., Garaialde, D., Edwards, J., Spillane, B., Gilmartin, E., Murad, C., Munteanu, C., Wade, V., Cowan, B. R. (2019). What Makes a Good Conversation?: Challenges in Designing Truly Conversational Agents, *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, pp. 1–12.
- Cowan, B. R., Pantidi, N., Coyle, D., Morrissey, K., Clarke, P., Al-Shehri, S., Earley, D., Bandeira, N. (2017). “What Can I Help You With?”: Infrequent Users’ Experiences of Intelligent Personal Assistants, *Proceedings of the 19th International Conference on Human-Computer Interaction with Mobile Devices and Services*, pp. 1–12.

- DeLone, W.H., McLean, E.R. (2003). The DeLone and McLean Model of Information Systems Success: A Ten-Year Update, *Journal of Management Information Systems*, 19(4), pp. 9–30.
- Dwivedi, Y.K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., Duan, Y., Dwivedi, R., Edwards, J., Eirug, A., Galanos, V., Ilavarasan, P.V., Janssen, M., Jones, P., Kar, A.K., Kizgin, H., Kronemann, B., Lal, B., Lucini, B., Medaglia, R., Le Meunier-FitzHugh, K., Le Meunier-FitzHugh, L.C., Misra, S., Mogaji, E., Sharma, S.K., Singh, J.B., Raghavan, V., Upadhyay, N. (2021). Artificial Intelligence (AI): Multidisciplinary Perspectives on Emerging Challenges, Opportunities, and Agenda for Research, Practice and Policy, *International Journal of Information Management*, 57, 101994.
- Følstad, A., Brandtzaeg, P.B. (2017). Chatbots and the New World of HCI, *Interactions*, 24(4), pp. 38–42.
- Følstad, A., Nordheim, C.B., Bjørkli, C.A. (2018). What Makes Users Trust a Chatbot for Customer Service? An Exploratory Interview Study. In: Bodrunova, S. (eds) *Internet Science, INSCI 2018, Lecture Notes in Computer Science*, Volume 11193, pp. 194–208.
- Gnewuch, U., Morana, S., Maedche, A. (2017). Towards Designing Cooperative and Social Conversational Agents for Customer Service, *Proceedings of the 38th International Conference on Information Systems (ICIS)*, pp. 1–13.
- Hair, J.F., Hult, G.T.M., Ringle, C.M., Sarstedt, M. (2022). *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*. 3rd edn. Thousand Oaks, CA: Sage.
- Khanzode, K.C.A., Sarode, R.D. (2020). Advantages and Disadvantages of Artificial Intelligence and Machine Learning: A Literature Review, *International Journal of Library and Information Science (IJLIS)*, 9(1), pp. 30–36.
- Kiseleva, J., Williams, K., Hassan Awadallah, A., Crook, A.C., Zitouni, I., Anastasakos, T. (2016). Understanding User Satisfaction with Intelligent Assistants, *Proceedings of the 2016 ACM on Conference on Human Information Interaction and Retrieval*, pp. 121–130.
- Kock, N., Hadaya, P. (2018). Minimum Sample Size Estimation in PLS-SEM: The Inverse Square Root and Gamma-Exponential Methods, *Information Systems Journal*, 28(1), pp. 227–261.
- López, G., Quesada, L., Guerrero, L.A. (2018). Alexa vs. Siri vs. Cortana vs. Google Assistant: A Comparison of Speech-Based Natural User Interfaces. In: Nunes, I. (eds) *Advances in Human Factors and Systems Interaction, AHFE 2017, Advances in Intelligent Systems and Computing*, Volume 592, pp. 241–250.
- Luger, E., Sellen, A. (2016). “Like Having a Really Bad PA”: The Gulf Between User Expectation and Experience of Conversational Agents. *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems*, pp. 5286–5297.
- Marangunić, N., Granić, A. (2015). Technology Acceptance Model: A Literature Review from 1986 to 2013, *Universal Access in the Information Society*, 14(1), pp. 81–95.
- McTear, M. (2020). *Conversational AI: Dialogue Systems, Conversational Agents, and Chatbots*. San Rafael: Morgan & Claypool.
- McTear, M., Callejas, Z., Griol, D. (2016). *The Conversational Interface: Talking to Smart Devices*. Cham: Springer.
- Orehovački, T. (2013). *Methodology for Evaluating the Quality in Use of Web 2.0 Applications*. PhD thesis. Varaždin: University of Zagreb, Faculty of Organization and Informatics.

- Orehovački, T., Etinger, D., Babić, S. (2019). The Antecedents of Intelligent Personal Assistants Adoption. In: Nunes, I. (eds) *Advances in Human Factors and Systems Interaction, AHFE 2018, Advances in Intelligent Systems and Computing*, Volume 781, pp. 76–87.
- Pal, D., Arpnikanondt, C., Funilkul, S., Varadarajan, V. (2019). User Experience with Smart Voice Assistants: The Accent Perspective, *Proceedings of the 10th International Conference on Computing, Communication and Networking Technologies (ICCCNT)*, pp. 1–6.
- Petter, S., DeLone, W.H., McLean, E.R. (2013). Information Systems Success: The Quest for the Independent Variables, *Journal of Management Information Systems*, 29(4), pp. 7–62.
- Porcheron, M., Fischer, J.E., Reeves, S., Sharples, S. (2018). Voice Interfaces in Everyday Life, *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*, pp. 1–12.
- Ringle, C.M., Wende, S., Becker, J.-M. (2022). *SmartPLS 4*. Oststeinbek: SmartPLS GmbH.
- Terzopoulos, G., Satratzemi, M. (2020). Voice Assistants and Smart Speakers in Everyday Life and in Education, *Informatics in Education*, 19(3). pp. 473–490.
- Yang, Q., Steinfeld, A., Rosé, C., Zimmerman, J. (2020). Re-examining Whether, Why, and How Human-AI Interaction Is Uniquely Difficult to Design, *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, pp. 1–13.
- Yang, Y.F., Zhang, J., Zeng, D.H. (2014). The Design and Realization of Remote Voice Control System Based on the Intelligent Home, *Applied Mechanics and Materials*, 530–531, pp. 1112–1118.
- Zamora, J. (2017). I’m Sorry, Dave, I’m Afraid I Can’t Do That: Chatbot Perception and Expectations. *Proceedings of the 5th International Conference on Human Agent Interaction*, pp. 253–260.
- Zwakman, D. S., Pal, D., Triyason, T., Arpnikanondt, C. (2020). Voice Usability Scale: Measuring the User Experience with Voice Assistants, *Proceedings of the 2020 IEEE International Symposium on Smart Electronic Systems (iSES)*, pp. 308–311.