

The Impact of Usability and Reliability on ChatGPT Satisfaction among Gen Z and Gen Y

Mirjana PEJIĆ BACH¹, Mirko PALIĆ¹, Vanja ŠIMIČEVIĆ²

¹ University of Zagreb, Faculty of Economics and Business, Croatia, mpejic@efzg.hr, mirkopalic@gmail.com

² University of Libertas, Zagreb, Croatia, vsimicevic@libertas.hr

Background/Purpose: ChatGPT's rapid diffusion has transformed large-language-model (LLM) technology from a specialist tool into a mainstream companion for study and work. However, empirical evidence on what drives user satisfaction outside medical settings remains scarce. Focusing on future business and management professionals in Croatia, this study examines how perceived ease of use and perceived reliability shape satisfaction with ChatGPT and whether those effects differ between Generation Z (18–25 years) and Generation Y (26–35 years).

Methodology: An online survey administered in August 2024 yielded 357 valid responses. The measurement model met rigorous reliability and validity criteria (CFI = 0.96, SRMR = 0.04).

Results: Structural-equation modelling showed that, in the pooled sample, ease of use ($\beta = 0.42$) and reliability ($\beta = 0.46$) jointly explained 72 % of satisfaction. Multi-group analysis revealed a generational split: both predictors were significant for Gen Z. However, only reliability remained significant for Gen Y. Gaussian graphical models corroborated these findings, indicating a densely interconnected attitude network for younger users and a reliability-centred network for older users.

Conclusion: The study extends technology-acceptance research to the management domain, underscores the moderating role of generation and illustrates the value of combining SEM with network analytics. Insights inform designers and educators aiming to foster informed, responsible and gratifying engagement with generative AI.

Keywords: Artificial intelligence, Large language models (LLM), Marketing, User satisfaction, Croatia, ChatGPT

1 Introduction

The advent of ChatGPT at the close of 2022 accelerated public exposure to large-language-model technology, pushing conversational AI from the laboratory into the mainstream almost overnight. Croatian users mirrored this global surge: within a few months, ChatGPT had become a routine aide for homework, report writing and everyday

fact-finding. Such dialogic media settings differ fundamentally from the broadcast era in which audiences were passive receivers; they oblige users to craft prompts, evaluate algorithmic output and often remix it into new content. As Seyoung and Park (2023) argue, navigating this landscape demands more than basic digital literacy. It calls for technical fluency—knowing how to interact with and tweak the system—and for cognitive-ethical competence: the capacity to scrutinise accuracy, detect bias, deploy re-

sults creatively and do so responsibly.

Whether people continue to embrace a tool that imposes these cognitive costs ultimately hinges on the quality of the experience it provides. Prior work on technology acceptance shows that satisfaction with interactive systems is shaped most directly by two beliefs: the perceived ease with which the system can be used and the perceived reliability of the information it supplies. However, the relative weight of these beliefs may not be constant across demographic segments. Generation Z, whose media habits were forged in a mobile-first environment, might value frictionless, always-on access. In contrast, Generation Y could concentrate more on trustworthy performance once basic usability thresholds are met. In addition, everyday knowledge of how large-language models work—or the absence of such knowledge—may colour people's impressions of ChatGPT's intelligence and, by extension, their willingness to rely on it.

A notable gap exists in the literature that addresses these issues. Most empirical studies to date have centred on healthcare professionals, exploring, for example, how physicians or medical students assess AI chatbots for diagnostic assistance and patient education. Research that probes the perceptions of business and management professionals—or, more broadly, the university-educated cohorts who will deploy generative AI in organisational settings—remains sparse. This paper helps close that gap by examining Croatian users whose primary exposure to ChatGPT occurs in academic business courses, entrepreneurial projects and managerial tasks, rather than in clinical practice.

The present study, therefore, sets out to map usage patterns, to quantify how ease of use and reliability shape satisfaction, and to discover whether these relationships diverge between Generation Z (18- to 25-year-olds) and Generation Y (26- to 35-year-olds). Beyond conventional structural equation modelling (SEM), it visualises item-level interdependencies through network analysis, offering a complementary view of the attitude system that underpins user experience.

Our contribution is threefold. First, the study extends the technology-acceptance conversation from clinical to managerial domains, documenting how future business practitioners in a small European market appropriate ChatGPT. Second, it demonstrates that generational differences do, in fact, modulate the satisfaction process: younger users balance usability and reliability, whereas older users hinge largely on reliability. Third, by combining covariance-based SEM with graphical least absolute shrinkage and selection operator (LASSO) networks, the paper supplies both confirmatory evidence and an exploratory map that highlights the specific questionnaire items, such as perceived time-saving and task efficiency, that act as bridges between constructs. These insights offer actionable guidance for AI developers seeking to fine-tune

large-language-model interfaces and for educators' intent on cultivating informed, critical and ethically grounded use of generative AI.

Empirically, the research draws on a cross-sectional online survey administered in Croatia during August 2024. A convenience sample of 357 valid respondents, recruited mainly through WhatsApp and Facebook channels linked to the University of Zagreb, completed a battery of closed-ended items that measured perceived ease of use, reliability, satisfaction, usage frequency, perceived intelligence and knowledge of large-language-model technology. The measurement model was validated via confirmatory factor analysis; structural paths were estimated with multi-group SEM, and the graphical structure of item correlations was explored with network analysis. Together, these methods provide a nuanced, generationally informed picture of what makes ChatGPT satisfying or disappointing for business-oriented users in a rapidly evolving AI ecosystem.

2 Literature review

2.1 Language Models Based on Artificial Intelligence

Artificial Intelligence (AI) has experienced significant development since its inception, fundamentally altering how humans interact with machines. AI-based language models, a critical subset of AI, have revolutionised natural language processing (NLP), leading to more sophisticated conversational agents, virtual assistants, and machine translation systems. These advancements have transformed numerous industries, from customer service to content creation, making AI an indispensable tool for modern technology. So, it is worthwhile to explore the theoretical foundation, historical development, core principles, and the advantages and challenges associated with AI-driven language models further.

The definition of AI has evolved as the field has progressed. The American Psychological Association (Neisser et al, 1996) describes intelligence as the capacity to comprehend complex ideas, adapt to change, learn from experience, and engage in logical reasoning. The foundational goal of AI has been to replicate or exceed these human cognitive abilities in machines. The Dartmouth Research Project in 1955 marked the beginning of AI as an academic discipline, with the aim of creating machines capable of autonomous problem-solving and decision-making (McCarthy et al., 2006). Scholars categorise AI into three main types: narrow AI, which is specialised for specific tasks; general AI, a hypothetical construct capable of human-like cognition across multiple domains; and super AI, which could surpass human intelligence and autonomously improve its capabilities (Bringsjord, 2011; McLean, 2021).

While narrow AI dominates contemporary applications, researchers continue to explore the potential for developing more advanced AI systems.

The historical development of AI can be traced back to the early 20th century, when scientists and philosophers speculated about the possibility of machine intelligence. Alan Turing's seminal work, "Computing Machinery and Intelligence" (1950), introduced key concepts like machine learning and the Turing Test, which assesses a machine's ability to exhibit human-like intelligence (Michie, 1993). The 1960s witnessed the emergence of early neural network models and symbolic reasoning approaches, laying the foundation for modern AI systems. Joseph Weizenbaum's ELIZA, developed in 1966, was the first chatbot to mimic human conversation by using pattern recognition and scripted responses (Sharma et al., 2017). Although ELIZA lacked true comprehension, it demonstrated AI's potential in natural language processing.

AI research has experienced periods of rapid progress and stagnation, commonly referred to as "AI winters." The 1980s saw a surge in interest with the rise of expert systems—AI programs designed to emulate human decision-making in specialised fields. However, limitations in computational power and funding constraints slowed progress. The resurgence of AI in the 21st century has been fuelled by breakthroughs in deep learning, data availability, and enhanced computing power. IBM's Deep Blue famously defeated world chess champion Garry Kasparov in 1997, demonstrating AI's ability to process vast amounts of information and make strategic decisions (Feng-Hsiung, 1999). The development of Google's AlphaGo, which surpassed human players in the complex game of Go in 2015, marked another milestone in AI's evolution (Haenlein & Kaplan, 2019).

Several fundamental principles underpin AI research, including machine learning, deep learning, and neural networks. Machine learning involves training algorithms to detect patterns and make predictions based on data without explicit programming (Bolf, 2021; Pejić Bach et al., 2023). Deep learning, a subset of machine learning, utilises artificial neural networks with multiple layers to refine predictions through iterative learning, enabling AI to recognise speech, generate text, and classify images with high accuracy (Singh, 2023). NLP enables AI to understand and generate human language, facilitating applications like speech recognition and sentiment analysis (Yegnanarayana, 1994). Artificial neural networks, modelled after the human brain, consist of interconnected nodes that process information similarly to biological neurons. These networks have evolved significantly, transitioning from basic perceptron models to sophisticated architectures capable of handling complex linguistic and cognitive tasks (McCulloch & Pitts, 1943).

AI offers numerous advantages across various industries, including enhanced efficiency, automation of repet-

itive tasks, and data-driven decision-making. AI-powered chatbots and virtual assistants provide real-time customer support, minimising the need for human intervention and optimising service delivery (Sotala, 2012). In healthcare, AI models assist in medical diagnoses and predictive analytics, improving patient outcomes. The finance sector leverages AI for fraud detection and risk assessment, streamlining complex financial processes (Bolf, 2021). AI's ability to process vast datasets enables businesses to personalise user experiences and refine marketing strategies (Banjac & Palić, 2020). However, AI also presents substantial challenges. The development and maintenance of AI systems require significant financial investments, often limiting accessibility to larger corporations (Girdhar, 2022). Ethical concerns, such as bias in AI decision-making and data privacy risks, remain key issues that must be addressed (Hua et al., 2024). AI bias can result from skewed training data, leading to discriminatory outcomes in hiring, lending, and law enforcement applications (Isada, 2024). Additionally, fears of widespread job displacement due to automation highlight the socio-economic impact of AI adoption.

To maximise AI's benefits while mitigating its risks, ongoing research focuses on improving model transparency, reducing bias, and implementing robust governance frameworks. As AI continues to evolve, the importance of ethical considerations and responsible deployment will shape its integration into society. Future advancements in AI-driven language models are expected to enhance human-computer interactions, making AI more adaptive, context-aware, and capable of generating nuanced responses. These innovations will further bridge the gap between machine intelligence and human communication, solidifying AI's role as a transformative force in technology and beyond.

2.2 User Satisfaction and Key Factors in AI-Based Large Language Models

User satisfaction is a crucial factor in evaluating AI-driven services, as it determines long-term engagement, trust, and adoption rates. AI-based large language models (LLMs) such as Chatgpt, Google Bard, Claude, Deepseek, Grok, and other NLP-powered systems have gained widespread use, offering users intelligent, responsive, and context-aware interactions. However, their effectiveness in delivering high-quality service experiences is contingent on several key factors. These include perceived service quality, trust, personalisation, and perceived benefits. This paper provides a comprehensive exploration of the concept of satisfaction, key determinants influencing user perceptions, and their implications for AI-based language models.

The concept of customer satisfaction has been exten-

sively studied in service marketing and consumer behaviour research. Kotler and Armstrong (2017) emphasise that beyond delivering products, businesses must ensure that their services align with customer expectations. Satisfaction is a psychological state that arises when the perceived performance of a product or service meets or exceeds user expectations (Churchill & Surprenant, 1982). If expectations are not met, users experience dissatisfaction, which can lead to disengagement or negative word-of-mouth. Crosby et al. (1990) argue that satisfaction plays a vital role in fostering long-term customer relationships, enhancing retention rates, and encouraging brand loyalty. In the context of AI, satisfaction is influenced by various cognitive and emotional factors, such as trust in the system's accuracy, perceived efficiency, and the relevance of AI-generated responses.

One of the most important determinants of satisfaction in AI-based services is perceived service quality. Zeithaml (1988) defines perceived quality as the user's subjective assessment of a service's overall excellence. AI-generated services, including LLMS, must deliver high levels of accuracy, coherence, and contextual relevance to be perceived as valuable. The SERVQUAL model, developed by Parasuraman et al. (1988), identifies five key dimensions of perceived service quality: reliability, responsiveness, assurance, empathy, and tangibility. In the context of AI, reliability refers to an AI model's ability to generate accurate and meaningful responses consistently. Responsiveness is the system's ability to understand and quickly address user queries. Assurance relates to the confidence users have in the system's credibility and correctness, while empathy involves the AI's capacity to recognise and adapt to user-specific needs. Although AI lacks human emotions, advancements in sentiment analysis and personalisation algorithms have improved AI's ability to deliver context-aware responses that enhance user engagement.

Trust is another critical element in determining user satisfaction. Users must feel confident that AI-generated responses are accurate, unbiased, and free from manipulation (Ou et al., 2024). Trust in AI systems depends on multiple factors, including transparency in how AI processes information, data privacy assurances, and the ability of AI to acknowledge errors. When AI models provide misleading or incorrect information, user trust diminishes, potentially leading to abandonment of the service. Research suggests that users are more likely to trust AI when they understand how it functions and when it demonstrates consistent accuracy in its outputs (Papenmeier et al., 2022). Additionally, users trust AI more when they perceive it as fair and free from bias (Adeiza et al., 2022). Bias in AI-generated language models has been a growing concern, as AI systems trained on biased datasets may produce skewed or discriminatory outputs. Addressing these concerns through explainable AI (XAI) techniques and fairness-enhancing algorithms can improve user trust and

overall satisfaction (Singh, 2023).

Personalisation plays a crucial role in enhancing the user experience with AI-based language models. AI systems that adapt their responses based on user preferences, history, and context are more likely to deliver relevant and engaging interactions. Personalisation in AI services involves learning from previous interactions, adjusting responses to align with individual preferences, and offering tailored recommendations. Studies show that AI-driven systems that incorporate personalised experiences lead to higher levels of user satisfaction and retention (Rust & Oliver, 1994). However, personalisation also raises privacy concerns, as AI systems require extensive user data to optimise their interactions. Balancing personalisation with data privacy is a challenge that AI developers must address to maintain user trust while delivering customised experiences.

Perceived benefits are another key determinant of satisfaction, referring to the extent to which users find AI-generated interactions useful, efficient, and valuable in their daily tasks (Zeithaml, 1988). Users expect AI to provide quick, relevant, and insightful responses that add value to their interactions. When users perceive AI as helpful, they are more likely to integrate it into their workflows, leading to higher engagement and long-term adoption (Uren & Edwards, 2023). The perceived usefulness of AI varies depending on the context; for example, in customer service, users value AI's ability to provide instant responses and resolve issues efficiently, while in content creation, users appreciate AI's capacity to generate high-quality text with minimal effort. However, if AI-generated content lacks depth, coherence, or originality, users may perceive it as redundant or unreliable, reducing satisfaction (Heskett et al., 1997).

Empirical research on AI satisfaction suggests that multiple factors contribute to the user experience, including ease of use, cognitive effort, and the system's ability to handle complex queries effectively (Adeiza et al., 2022). User experience (UX) research has shown that frustration arises when AI systems fail to understand user intent or generate responses that are irrelevant or misleading. To improve user satisfaction, AI developers must continuously refine models to enhance accuracy, contextual awareness, and conversational fluency. Ethical considerations, including AI fairness, transparency, and adaptability, also play a role in shaping user perceptions of AI reliability and usefulness (Singh, 2023).

As AI technology advances, organisations leveraging AI-driven services must prioritise improving user experience by optimising model performance, addressing ethical concerns, and ensuring responsible AI deployment. Future AI developments should focus on reducing algorithmic bias, enhancing personalisation capabilities, and providing clear explanations for AI-generated decisions. The ongoing refinement of AI-based language models will

contribute to more effective, trustworthy, and engaging interactions, ultimately driving greater user satisfaction and adoption.

2.3 Relationship between age and attitude towards artificial intelligence

Numerous studies support the relationship between age and attitudes towards artificial intelligence (AI), demonstrating how age influences perceptions, acceptance, and willingness to embrace AI technologies (Pejić Bach & Marić, 2025). However, most of these studies have been conducted in relation to health care and medical research. Generally, younger individuals tend to exhibit more favourable attitudes towards AI compared to older generations.

Research indicates that younger respondents often have higher trust levels in AI systems, which correlates with their familiarity and comfort with new technologies (Ongena et al., 2021; York et al., 2020). In contrast, older demographics frequently display scepticism and apprehension towards AI applications, especially in healthcare settings where concerns about decision-making and precision in AI capabilities arise (Fritsch et al., 2022).

Furthermore, attitudes towards AI vary significantly across different age groups due to generational differences in technology exposure and inherent learning curves. Yigitcanlar et al. highlight how individual factors such as age and knowledge about AI significantly shape public perception, with older individuals generally being less informed about AI developments. Middle-aged populations have shown mixed responses, exhibiting scepticism towards adopting AI, particularly chatbots and similar technologies, due to perceived complexities and usability challenges (Wang et al., 2024).

While educational attainment plays a role in shaping attitudes towards AI, with higher education levels correlating with greater acceptance and trust, age often serves as a primary barrier. Shevtsova et al. noted that older participants (aged 40-60 and above) exhibited awareness of and positive attitudes towards AI technologies; however, this was not uniform across all older individuals, with some demonstrating reluctance (Shevtsova et al., 2024). Additionally, factors such as gender, experience with technology, and anxiety about AI modulate these attitudes (Sindermann et al., 2022; Alkhalifah et al., 2024). Research by Sindermann et al. illustrates that personal characteristics and previous interactions with AI systems create a reciprocal influence, complicating the dynamics of acceptance among varying age groups (Sindermann et al., 2022).

In summary, understanding the relationship between age and attitudes toward AI requires a multifaceted examination of demographic variables, personal experiences, and educational background. As age increases, the inclina-

tion to trust and accept AI technologies often diminishes, influenced by individual experiences and societal perceptions regarding technology and its intersection with daily life (Kauttonen et al., 2025; Zhang et al., 2023).

3 Methodology

3.1 Research questions and hypotheses

Guided by the aim of explaining why Croatian users embrace Chatgpt, the study poses two overarching research questions. RQ1 asks: Which experiential beliefs most strongly predict overall satisfaction with ChatGPT? Building on the Technology Acceptance Model and service quality theory, we propose that two beliefs—perceived ease of use and perceived reliability—serve as the primary antecedents. Accordingly, we advance H1: Perceived ease of use exerts a positive influence on satisfaction, and H2: Perceived reliability exerts a positive influence on satisfaction. RQ2 asks: Do these relationships differ across generations that have grown up with distinct digital habits? Drawing on generational theory, we expect younger “digital natives” to weigh usability and reliability more evenly, whereas older “digital adapters” may lean more heavily on reliability once basic usability is assured. Hence we specify H1a and H1b—replicating H1 and H2, respectively, for Generation Z (18–25 years)—and H2a and H2b for Generation Y (26–35 years). Specifically, we predict that both paths will be significant among Generation Z, but among Generation Y, only the reliability path will remain significant. Testing this hypothesis set allows us to isolate the universal drivers of satisfaction while detecting demographic nuances that can inform tailored design and outreach strategies.

3.2 Research instrument

After an introductory greeting, respondents were briefly informed about the purpose of the study and notified that completing the questionnaire would take about five minutes. It was also emphasised that participation in the research was entirely voluntary and anonymous, and that the collected data would be presented exclusively in an aggregated format. An elimination question was included to determine whether the respondent had ever used ChatGPT, with further participation allowed only for those who answered affirmatively. The research instrument was a survey questionnaire composed of 11 closed-ended questions divided into three sections (Table 1). All 11 items employ a five-point Likert format ranging from 1 (“strongly disagree”) to 5 (“strongly agree”).

Table 1: Research instrument

Latent variable	Code	Items
Perceived ease of use (PEOU)	PEOU1	I appreciate the ability to start interacting with ChatGPT regardless of location and time.
	PEOU2	ChatGPT saves me time by providing quick access to information.
	PEOU3	Interacting with ChatGPT justifies the time and effort spent to get the information I want.
	PEOU4	I find ChatGPT easy to use.
Reliability (REL)	REL1	ChatGPT has provided a wide range of information related to my questions, including detailed explanations and relevant examples.
	REL2	ChatGPT service offers greater efficiency in finding information compared to using other tools.
	REL3	ChatGPT provides me with exactly the level of service and quality of information that I expected.
	REL4	ChatGPT helps me to complete many tasks and efficiently.
Satisfaction (SAT)	SAT1	I am satisfied with the overall experience of using ChatGPT.
	SAT2	I plan to continue using ChatGPT in the future.
	SAT3	I would recommend others to use ChatGPT.

Source: Authors' work

Perceived ease of use (PEOU) is gauged using four items that prompt respondents to judge, first, how effortlessly they can initiate a ChatGPT session regardless of time or place and, second, whether the system demonstrably saves them time by delivering information rapidly. Two additional statements ask participants to weigh the overall cost-benefit of the interaction and give a direct appraisal of how easy the tool is to handle. Reliability (REL) is likewise assessed with four indicators. Respondents reflect on the breadth and depth of explanations received, the comparative efficiency of ChatGPT vis-à-vis alternative information sources, the extent to which the service meets their prior expectations of quality, and its usefulness in completing everyday tasks with minimal friction. User satisfaction (SAT) is measured by three items: a global affective evaluation of the experience, an intention to continue using the system, and a willingness to recommend it to others. At the end of the questionnaire, sociodemographic data of respondents were collected.

3.3 Sample and data collection

The overall sample consisted of 357 respondents who had at least used ChatGPT once. A total of 357 people completed the survey. Roughly two-thirds of them were women, while a little over one-third were men. The group was quite young overall: about seven out of every ten respondents were between 18 and 25 years old, and the remaining three out of ten were 26 to 35; no one older

than 35 took part. Educational backgrounds ranged from secondary to postgraduate levels. Just under half of the participants had finished high school without yet earning a university degree. Around one quarter held a bachelor's degree, and almost the same proportion had completed a master's programme. Only one respondent reported having a PhD. In sum, the sample represents a predominantly young, female-leaning population with education spanning from high school through master's studies.

3.4 Statistical analysis

To examine the hypothesised relationships among perceived ease of use, reliability and user satisfaction, we applied structural-equation modelling (SEM) in JASP 0.18, which relies on the Lavaan package for maximum-likelihood estimation. Screening showed no extreme multivariate outliers (Mahalanobis distance, $p > .001$). Univariate skewness and kurtosis fell within ± 2 , allowing the use of (robust) ML estimation; nonetheless, we adopted the Satorra–Bentler correction to guard against any residual non-normality.

All eleven survey items were specified as reflective indicators of their respective latent constructs. A confirmatory factor analysis (CFA) was first run on the pooled sample to verify factorial validity. Internal consistency was judged with both Cronbach's α and composite reliability (CR); values ≥ 0.70 were deemed acceptable. Convergent validity was inspected through standardised factor load-

Table 2: Sample structure and demographics

Characteristic	Modalities	n	%	Cumulative %
Gender	Female	224	62.7/	62.7
	Male	133	37.3%	100.0
Age	18-25	248	69.5	69.5
	26-35	109	30.5	100.0
Education	High School	173	48.5	48.5
	Bachelor	96	26.9	75.4
	Master	87	24.4	99.7
	PhD	1	0.3	100.0
	Total	357	100.0	

Source: Authors' work

Table 3: The main purpose of using ChatGPT

	Total	18-25	26-35	Chi-square
Purpose	n=357	n=248	n=109	59.648**
Help with learning/education	44.5%	56.5%	17.4%	
Writing/editing text	28.3%	23.4%	39.4%	
Translation/language support	7.6%	4.4%	14.7%	
Seeking information relevant to work	3.9%	2.0%	8.3%	
Health advice/therapeutic purposes	0.6%	0.4%	0.9%	
For asking simple questions (e.g., "What time is it?", "What is the capital of Italy?", etc.)	2.5%	2.0%	3.7%	
Entertainment/chatting	5.3%	4.8%	6.4%	
Analysis of large amounts of data	3.9%	3.6%	4.6%	
Recommendations (books, movies, restaurants, etc.)	2.5%	2.8%	1.8%	
Other purposes	0.8%	0.0%	2.8%	
Total	100.0%	100.0%	100.0%	

Note: *** statistically significant at 1%

Source: Authors' work

ings (target ≥ 0.70) and average variance extracted (AVE ≥ 0.50). After psychometric adequacy was confirmed, we proceeded to the structural step.

The posited paths from perceived ease of use and reliability to satisfaction were estimated simultaneously. Model-level fit was evaluated with multiple indices to offset the limitations of any single statistic: the comparative fit index (CFI ≥ 0.90 for good fit), the Tucker–Lewis index (TLI ≥ 0.90), the root-mean-square error of approximation (RMSEA ≤ 0.06 , 90 % CI reported) and the standardised root-mean-square residual (SRMR ≤ 0.08). Predictive power was gauged with the squared coefficient of determination (R^2) for satisfaction as the endogenous construct.

Given the age split in the sample, we tested the structural model separately for the 18–25 and 26–35 cohorts, using the multigroup approach.

JASP's Network Analysis module was used to estimate a Gaussian Graphical Model by applying the graphical LASSO to the research items of variables PEOU, SAT and REL. The procedure decreases small partial correlations to zero and chooses the optimal amount of regularisation with the EBIC (engl. Extended Bayesian Information Criterion) tuning rule. The resulting graph displays only those conditional associations that survive this penalisation.

4 Results

4.1 Attitudes towards ChatGPT among Generation Z and Generation Y

Results presented in Table 1 indicate that most respondents say they turn to ChatGPT for study-related tasks: overall 45 % use it primarily to support learning, and this motive dominates in the 18-to-25 cohort (57 %) but drops sharply among 26-to-35-year-olds (17 %). In contrast, older users rely on the tool mainly for writing or editing text (39 % versus 23 % in the younger group) and are more likely to seek translation help or job-related information. Smaller shares across both ages mention casual queries, entertainment, data analysis or recommendations, and only a handful cite health advice or “other” reasons. The chi-square value ($\chi^2 = 59.65$, $p < .01$) confirms that the pattern of purposes differs significantly between the two age brackets.

Roughly one user in four opens ChatGPT only occasionally: 28 % of the total sample report logging in less

than once a month, with no great age difference at that lowest tier of engagement (Table 4). Beyond that point, however, the two cohorts diverge. Younger respondents (18–25) are in the mid-range: they are more likely to say they use the tool “once a month” or “several times a month,” and fewer reach the higher-frequency categories. Older respondents (26–35) lean in the opposite direction: a fifth of them enter ChatGPT several times a week, and one in six does so many times a day—more than double the proportion seen in the younger group. These contrasting usage patterns yield a chi-square of 16.58 ($p < .01$), confirming that frequency of interaction with ChatGPT varies significantly by age.

About three-quarters of all respondents—77 %—say they regard ChatGPT as an intelligent system, while roughly one in four do not (Table 5). This perception is virtually identical in both age groups (77.4 % among 18- to 25-year-olds versus 77.1 % among 26- to 35-year-olds). The very small, non-significant chi-square value ($\chi^2 = 0.014$) confirms that age makes no observable difference to this judgement.

Table 4: Frequency of ChatGPT use

Usage frequency	Total n=357	18-25 n=248	26-35 n=109	Chi-square
Rarely (less than once a month)	27.5%	28.6%	24.8%	16.583**
Once a month	7.0%	9.3%	1.8%	
Several times a month	24.6%	26.6%	20.2%	
Weekly	9.8%	9.3%	11.0%	
Several times a week	18.5%	16.9%	22.0%	
Once a day	2.8%	2.4%	3.7%	
Multiple times a day	9.8%	6.9%	16.5%	
Total	100.0%	100.0%	100.0%	

Note: *** statistically significant at 1%

Source: Authors' work

Table 5: Considering ChatGPT as intelligent

Consider ChatGPT intelligent	Total n=357	18-25 n=248	26-35 n=109	Chi-square
No	22.7%	22.6%	22.9%	0.014 [†]
Yes	77.3%	77.4%	77.1%	
Total	100.0%	100.0%	100.0%	

Note: [†] not statistically significant

Source: Authors' work

Table 6: Level of LLM knowledge

Knowledge LLM	Total n=357	18-25 n=248	26-35 n=109	Chi-square
Yes	30.0%	27.0%	36.7%	3.381 [†]
No	70.0%	73.0%	63.3%	
Total	100.0%	100.0%	100.0%	

Note: [†] not statistically significant
Source: Authors' work

Table 7: Relationship between the level of LLM knowledge and considering ChatGPT as intelligent

Knowledge LLM	Consider ChatGPT intelligent			Chi-square
	No	Yes	Total	
Yes	39,5%	26,9%	29,8%	4.749*
No	60,5%	73,1%	70,2%	
Total	100,0%	100,0%	100,0%	

Note: * statistically significant at 5%
Source: Authors' work

Table 8: Chi-Square test

Model	x ²	df	p
Baseline model	2548.617	55	
Factor model	144.264	41	< .001

Source: Authors' work

Table 9: Fit indices

Fit indices	Value
Comparative Fit Index (CFI)	0.959
Tucker-Lewis Index (TLI)	0.944
Root mean square error of approximation (RMSEA)	0.084
Standardised root mean square residual (SRMR)	0.039
Goodness of fit index (GFI)	0.988

Source: Authors' work

Only three out of ten respondents say they already know what a large-language model (LLM) is, while the remaining seven out of ten admit they do not (Table 6). Self-reported familiarity is somewhat higher among the 26- to 35-year-olds (37 %) than among the 18- to 25-year-olds (27 %), but the gap is small, and the chi-square test ($\chi^2 = 3.38$) shows it is not statistically reliable. In other words, most users interact with ChatGPT without being able to describe the underlying technology, and this lack of technical knowledge is shared across both age groups.

A user's grasp of what a large-language model is shapes the way they judge Chatgpt's intelligence, as indicated by Table 7. Among participants who say they understand LLMs, only about 27 % call the system intelligent, whereas the figure rises to 73 % for those who lack that technical knowledge. Conversely, knowledgeable users make up a larger share of the "not intelligent" camp. The association is modest but statistically reliable ($\chi^2 = 4.75$, $p < .05$), indicating that deeper familiarity with the technology tends to temper perceptions of ChatGPT's intelligence.

4.2 Measurement model

The chi-square test pits the hypothesised three-factor solution against an independence (baseline) model in which all items are assumed uncorrelated (Table 8). The baseline model shows an enormous misfit ($\chi^2 = 2,548.62$, $df = 55$), whereas the factor model cuts the discrepancy to $\chi^2 = 144.26$ with 41 degrees of freedom. Although the χ^2 statistic for the factor model is still significant (reflecting its sensitivity to sample size), the reduction of more than 2,400 chi square units demonstrates that the latent variable structure explains the observed covariances far better than

a null model.

Most descriptive indices meet or exceed conventional benchmarks (Table 9). The CFI (.959) and GFI (.988) signal a very good fit ($\geq .95$), and the TLI (.944) is just below the same threshold, still considered acceptable. Residual-based measures are also favourable: the SRMR (.039) is comfortably under the .08 criterion, while the RMSEA (.084) is in the "reasonable fit" band (.05–.08) but slightly above the ideal .06 cutoff, suggesting mild room for improvement in model parsimony.

All 11 indicators load strongly and significantly on their intended factors; standardised loadings range from .685 (REL3) to .879 (SAT2), well above the 0.70 rule of thumb (Table 10), which confirms that each item is a reliable reflection of its latent construct.

Coefficient ω and Cronbach's α exceed .84 for every scale, indicating high reliability (Table 11). Average variance extracted (AVE) surpasses .50 for PEOU (.621), REL (.592) and SAT (.728), confirming that, within each construct, the indicators share more variance with the latent factor than with measurement error.

All heterotrait-monotrait (HTMT) ratios fall below the conservative .85 threshold (largest = .806 between SAT and PEOU). Thus, the three constructs are statistically distinguishable despite being moderately correlated (Table 12).

Taken together, the chi-square comparison, fit indices, strong loadings, high reliability, adequate AVE, and satisfactory HTMT ratios provide a coherent body of evidence that the three-factor measurement model is reliable and valid for capturing perceived ease of use, reliability, and satisfaction with ChatGPT.

Table 10: Factor loadings

Factor	Indicator	Estimate	Std. Error	z-value	p	Std. Est. (all)
PEOU	PEOU1	0.746	0.046	16.091	< .001	0.755
	PEOU2	0.808	0.042	19.292	< .001	0.854
	PEOU3	0.723	0.043	16.835	< .001	0.780
	PEOU4	0.684	0.042	16.227	< .001	0.761
REL	REL1	0.750	0.044	17.157	< .001	0.792
	REL2	0.800	0.049	16.476	< .001	0.769
	REL3	0.704	0.050	14.034	< .001	0.685
	REL4	0.846	0.046	18.311	< .001	0.828
SAT	SAT1	0.730	0.041	18.007	< .001	0.812
	SAT2	0.792	0.039	20.416	< .001	0.879
	SAT3	0.805	0.040	19.924	< .001	0.865

Source: Authors' work

Table 11: Reliability and average variance extracted

	Coefficient ω	Coefficient α	AVE
PEOU	0.869	0.866	0.621
REL	0.849	0.853	0.592
SAT	0.892	0.886	0.728
Total	0.942	0.926	

Source: Authors' work

Table 12: Heterotrait-monotrait ratio

	PEOU	REL	SAT
PEOU	0.788		
REL	0.734	0.769	
SAT	0.806	0.800	0.853

Source: Authors' work

Table 13: SEM results for the total sample and Gen Z vs. Gen Y

Group	Outcome	Predictor	Estimate	Std. Error	z-value	R-squared	Hypothesis
Total	SAT	PEOU	0.424	0.072	5.871**	0.716	H1 ✓ (+1%)
		REL	0.457	0.072	6.321**		H2 ✓ (+1%)
18-25	SAT	PEOU	0.435	0.082	5.297**	0.722	H1a ✓ (+1%)
		REL	0.469	0.080	5.872**		H1b ✓ (+1%)
26-35		PEOU	0.262	0.178	1.469†	0.715	H2a ∅
		REL	0.522	0.171	3.057**		H2b ✓ (+1%)

Note: ** statistically significant at 1%; † not statistically significant

Source: Authors' work

4.3 Structural equation modelling

The structural-equation results in [Table 13] show that, in the full sample, both perceived ease of use (PEOU) and perceived reliability (REL) make sizable, statistically significant contributions to satisfaction with ChatGPT ($\beta = 0.42$ and 0.46 , respectively). Together, they explain about 72 % of the variance in satisfaction.

When the analysis is split by generation, the pattern diverges slightly. For the Gen Z cohort (18–25 years), both paths remain strong and significant ($\beta \approx 0.44$ for PEOU and 0.47 for REL), again accounting for roughly 72 % of satisfaction. For the Gen Y group (26–35 years), reliability is still a significant driver ($\beta = 0.52$). However, ease of

use drops to a weaker, non-significant role ($\beta = 0.26$, $p > .05$). Thus, older respondents appear to base their satisfaction chiefly on how dependable ChatGPT's answers are. In contrast, younger users weigh usability and reliability almost equally. Overall, hypotheses H1, H2, H1a and H1b are supported, while H2a is not, and H2b is confirmed.

4.4 Network analysis

Figure 1 visualises the partial-correlation network among the eleven survey items for the whole sample. Each node represents one questionnaire statement, coloured by its latent construct (pink = Perceived Ease of Use, green = Reliability, blue = Satisfaction). Lines indicate regularised

partial correlations that remain after controlling for all other items: thicker, darker blue lines mark stronger positive linkages, whereas the virtual absence of red lines means no appreciable negative associations survived the graphical LASSO penalty.

Three observations stand out. First, nodes cluster almost perfectly by their theoretical section, confirming that items within the same construct share stronger conditional ties with each other than with items from other constructs. Second, the densest within-cluster edges appear in the Satisfaction trio—especially between “continue using” (SAT2) and “recommend to others” (SAT3)—highlighting their conceptual closeness. Third, two cross-cluster bridges emerge: PEOU2 (“saves me time”) connects to REL2 (“more efficient than other tools”), and REL4 (“helps me complete tasks”) links to SAT1 (“overall satisfied”). These bridges reflect the pathways later captured in the SEM: ease of use and reliability channel their influence into satisfaction via efficiency and task accomplishment.

Centrality analysis (strength) shows SAT2 and REL4 as the most influential nodes in the network, suggesting that intentions to keep using the service and perceptions of task efficiency play pivotal roles in holding the entire attitude system together.

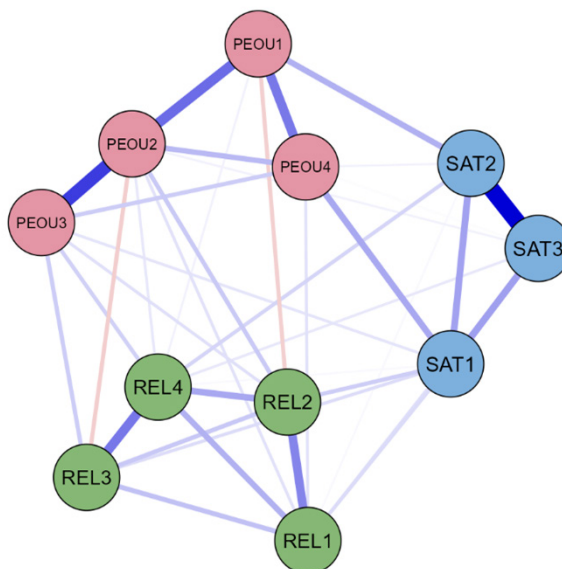
To complement the multigroup SEM, Gaussian graphical models were estimated separately for the 18–25-year-old respondents (Generation Z) and the 26–35-year-old respondents (Generation Y).

In each subsample, the eleven manifest variables were treated as continuous, and the networks were obtained with

the graphical LASSO, selecting the optimal tuning parameter by the EBIC rule ($\gamma = 0.50$). The resulting graphs, depicted in Figure 2, visualise regularised partial correlations: edge width conveys absolute strength, blue hues denote positive relations, and the sparse red lines indicate residual negative associations that survived regularisation. Nodes are colour-coded by their theoretical domain—pink for perceived ease of use (PEOU1–4), green for reliability (REL1–4) and blue for satisfaction (SAT1–3).

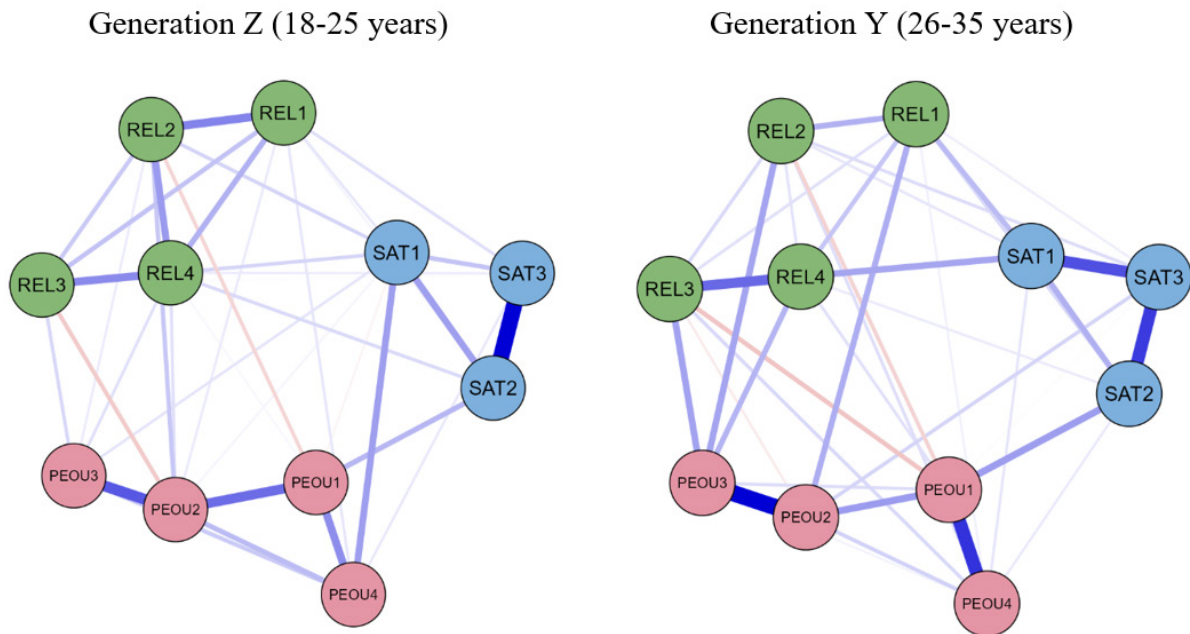
For Generation Z, the topology is highly interwoven. Although the three construct clusters are still discernible, several medium-to-strong bridges link the ease-of-use nodes directly to the satisfaction cluster. The most prominent of these connects the “time-saving” item PEOU2 to both SAT1 (overall experience) and SAT2 (intention to continue). Centrality indices corroborate the visual impression: SAT2 and PEOU2 show the highest strength centrality, indicating that usability cues and future-use intentions form the main hubs through which information flows in the younger cohort. This pattern aligns with the SEM finding that, for Gen Z, perceived ease of use contributes to satisfaction almost as strongly as perceived reliability.

By contrast, Generation Y’s network is more compartmentalised. Most cross-construct connections are weak or absent after regularisation, and the reliability items form the densest sub-graph. REL4 (“helps me complete tasks”) emerges as the key bridge to satisfaction, maintaining a thick edge to SAT1; links emanating from ease-of-use items are much thinner. Consequently, strength centrality ranks REL4 and SAT3 well above the usability nodes,



Source: Authors' work

Figure 1: Network analysis of a total sample



Source: Authors' work

Figure 2: Network analyses of Generation Z and Generation Y

mirroring the SEM result in which reliability, but not ease of use, significantly predicts satisfaction among older respondents.

Taken together, the generation-specific networks reinforce the multigroup SEM conclusions. Where Gen Z's perceptions of ChatGPT resemble a highly integrated attitude system in which usability and reliability jointly feed into satisfaction, Gen Y's perceptions resemble a modular system whose satisfaction component is supplied chiefly by reliability cues. These structural differences suggest that design and communication strategies aimed at younger users should emphasise both frictionless interaction and dependable output. In contrast, strategies for slightly older users may generate greater returns by foregrounding the system's trustworthiness and task efficacy.

5 Conclusion

This study set out to answer two questions: RQ1, which experiential beliefs drive satisfaction with Chatgpt, and RQ2, whether those drivers differ across generations, and to test the associated hypotheses (H1–H2b). The evidence affirms that perceived ease of use and perceived reliability are the principal antecedents of satisfaction, thereby supporting H1 and H2 for the total sample. However, the multi-group analysis reveals a generational inflexion.

Among Generation Z (18–25 years), both ease of use and reliability significantly shape satisfaction, validating H1a and H1b. Among Generation Y (26–35 years), only reliability retains explanatory power, leading to the acceptance of H2b and the rejection of H2a.

These findings are best understood against the backdrop of generational digital literacy. Gen Z has grown up with ubiquitous, intuitive technology; for them, the enjoyment of an AI assistant is tightly linked to how engaging the interface feels. Gen Y is equally tech-savvy but has accumulated more professional experience; accordingly, it places greater weight on the trustworthiness and consistency of information. Designers targeting Gen Z should prioritise interactive, visually rich and highly personalised features, while for Gen Y, marketing messages and product roadmaps should foreground data security and transparent sourcing.

The results also speak to product-development strategy. Segmenting the user base by generation and tailoring feature sets to the specific expectations of each cohort can raise adoption and retention rates. Further research that probes the psychological reasons behind these generational preferences would help refine such segmentation. Finally, user-education programmes should mirror these needs: Gen Z may benefit from tutorials that showcase creative

prompt engineering and playful use cases. Gen Y may prefer guidance on evaluating output quality, integrating citations and enforcing ethical safeguards.

These patterns carry several practical implications. Developers hoping to retain younger audiences must continue to streamline prompts, reduce latency and integrate conversational cues that signal effortlessness, while also safeguarding output quality. For older, professionally focused users, investments in source transparency, factual accuracy and task-specific guidance are likely to yield greater returns. Educators, meanwhile, should recognise that most students still lack a working knowledge of LLM technology and thus need structured training not only in prompt engineering but also in critical appraisal and ethical deployment of generative AI.

The study also expands the empirical reach of technology-acceptance research, which has so far concentrated on clinical contexts, by demonstrating that the same constructs operate—and operate differently—among future managers and entrepreneurs. Methodologically, it shows the value of coupling covariance-based SEM with graphical LASSO networks to obtain both confirmatory and exploratory insight.

Future work should replicate the model in organisational field studies, track longitudinal adoption trajectories and probe additional moderators such as task complexity or domain expertise. Additionally, it would be interesting to investigate whether the ChatGPT tool is more useful for natural science or social science professional users, and what the differences are between these two groups of users. Limitations remain and should be considered when taking into account the results of this research. The convenience sample, reliance on self-report and cross-sectional design restrict generalisability and causal inference. Even so, the present findings offer a grounded starting point for designing, teaching and governing conversational AI in the management arena.

References

- Adeiza, A., Abdullahi, M., Abdelfattah, F., & Fawehinmi, O. (2022). Mediating mechanism of customer satisfaction on customer relationship management implementation and customer loyalty among consolidated banks. *Supply Chain Management*, 10(3), 819–832. <https://doi.org/10.5267/j.uscm.2022.3.012>
- Alkhalifah, J., Bedaiwi, A., Shaikh, N., Seddiq, W., & Meo, S. (2024). Existential anxiety about artificial intelligence (AI)- is it the end of the era of humanity or a new chapter in the human revolution: questionnaire-based observational study. *Frontiers in Psychiatry*, 15. <https://doi.org/10.3389/fpsy.2024.1368122>
- Banjac, I., & Palić, M. (2020). “Analysis of Best Practice of Artificial Intelligence Implementation in Digital Marketing Activities”, 5th International Scientific and Professional Conference CRODMA 2020, October 23rd, Faculty of Organization and Informatics, University of Zagreb. Paper published in: Gregurec, I., Kvošca, V., Keglević Kozjak, S., Cvetko, L. (ur.): Book of Papers: 5th International Scientific and Professional Conference (CRODMA 2020), str. 1 – 12, (ISSN 2459-7953). objavljen i u CroDiM; *International Journal of Marketing Science*, Vol. 4, No. 1, 2021. pp. 45-56.
- Bhbosale, S., Pujari, V., & Multani, Z. (2020). Advantages and disadvantages of artificial intelligence. *International Interdisciplinary Research Journal*, 13(1), 227–230.
- Bolf, N. (2021). Umjetne neuronske mreže. *Kemija u industriji: Časopis kemičara i kemijskih inženjera Hrvatske*, 70(9–10), 591–593.
- Bratić, D., Sačer, S., & Palić, M. (2020). Implications of Artificial Intelligence in Marketing Activities on Multimedia Platforms. in: Šimurina, J., Načinović Braje, I., & Pavić, I. (ur.) Proceedings of FEB Zagreb 11th International Odyssey Conference on Economics and Business. doi:10.22598/odyssey/2020.2
- Bringsjord, S. (2011). Psychometric artificial intelligence. *Journal of Experimental & Theoretical Artificial Intelligence*, 23(3), 271–277. <https://doi.org/10.1080/0952813X.2010.502314>
- Churchill, G. A., & Surprenant, C. (1982). An investigation into the determinants of customer satisfaction. *Journal of Marketing Research*, 19(4), 491–504. <https://doi.org/10.1177/002224378201900410>
- Crosby, L. A., Evans, K. R., & Cowles, D. (1990). Relationship quality in services selling: An interpersonal influence perspective. *Journal of Marketing*, 54(3), 68–81. <https://doi.org/10.1177/002224299005400306>
- Feng-Hsiung, H. (1999). IBM's Deep Blue Chess grandmaster chips. *IEEE Micro*, 19(2), 70–81. <https://doi.org/10.1109/40.755469>
- Fritsch, S., Blankenheim, A., Wahl, A., Hetfeld, P., Maaßen, O., Deffge, S., ... & Bickenbach, J. (2022). Attitudes and perception of artificial intelligence in healthcare: a cross-sectional survey among patients. *Digital Health*, 8, 205520762211167. <https://doi.org/10.1177/20552076221116772>
- Girdhar, A. (2022). History of Artificial Intelligence – A Brief History of AI. *Appy Pie*. <https://www.appypie.com/history-of-artificial-intelligence>
- Haenlein, M., & Kaplan, A. (2019). A brief history of artificial intelligence: On the past, present, and future of artificial intelligence. *California Management Review*, 61(4), 1–10. <https://doi.org/10.1177/0008125619864925>
- Heskett, J. L., Sasser, W. E., & Schlesinger, L. A. (1997). The service profit chain. *Free Press*.
- Hua, X., Hasan, N. A. M., De Costa, F., & Qiao, W. (2024). Opportunities or Challenges? The Interplay between

- Artificial Intelligence and Corporate Social Responsibility Communication. *Business Systems Research: International journal of the Society for Advancing Innovation and Research in Economy*, 15(1), 131-157. <https://doi.org/10.2478/bsrj-2024-0007>
- Isada, F. (2024). Inter-Organisational Collaboration Networks in The Introduction Phase of Generative AI. *ENTRENOVA-ENTERprise REsearch INNOVation*, 10(1), 481-488. <https://doi.org/10.54820/entrenova-2024-0037>
- Kauttonen, J., Rousi, R., & Alamäki, A. (2025). Trust and acceptance challenges in the adoption of AI applications in health care: quantitative survey analysis. *Journal of Medical Internet Research*, 27, e65567. <https://doi.org/10.2196/65567>
- Kotler, P., & Armstrong, G. (2017). Principles of marketing (17th ed.). *Pearson Education Limited*.
- McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (2006). A proposal for the Dartmouth Summer Research Project on Artificial Intelligence, *AI Magazine*, 27(4), 12. <https://doi.org/10.1609/aimag.v27i4.1904>
- McCulloch, W. S., & Pitts, W. (1943). A logical calculus of the ideas immanent in nervous activity. *The Bulletin of Mathematical Biophysics*, 5(4), 115-133. <https://doi.org/10.1007/BF02478259>
- McLean, S. (2021). The risks associated with Artificial General Intelligence: A systematic review. *Journal of Experimental & Theoretical Artificial Intelligence*, 35(5), 649-663.
- Michie, D. (1993). Turing's test and conscious thought. *Artificial Intelligence*, 60(1), 1-22. [https://doi.org/10.1016/0004-3702\(93\)90032-7](https://doi.org/10.1016/0004-3702(93)90032-7)
- Neisser, U., Boodoo, G., Bouchard, T. J., Jr., Boykin, A. W., Brody, N., Ceci, S. J., Halpern, D. F., Loehlin, J. C., Perloff, R., Sternberg, R. J., & Urbina, S. (1996). Intelligence: Knowns and unknowns. *American Psychologist*, 51(2), 77-101. <https://doi.org/10.1037/0003-066X.51.2.77>
- Ongena, Y., Yakar, D., Haan, M., & Kwee, T. (2021). Artificial intelligence in screening mammography: a population survey of women's preferences. *Journal of the American College of Radiology*, 18(1), 79-86. <https://doi.org/10.1016/j.jacr.2020.09.042>
- Ou, M., Zheng, H., Zeng, Y., & Hansen, P. (2024). Trust it or not: Understanding users' motivations and strategies for assessing the credibility of AI-generated information. *New Media & Society*, 14614448241293154. <https://doi.org/10.1177/14614448241293154>
- Papenmeier, A., Kern, D., Englebienne, G., & Seifert, C. (2022). It's complicated: The relationship between user trust, model accuracy and explanations in AI. *ACM Transactions on Computer-Human Interaction (TOCHI)*, 29(4), 1-33. <https://doi.org/10.1145/3495013>
- Parasuraman, A., Zeithaml, V. A., & Berry, L. L. (1988). SERVQUAL: A multiple-item scale for measuring consumer perceptions of service quality. *Journal of Retailing*, 64(1), 12-40.
- Pejić Bach, M., Topalović, A., & Turulja, L. (2023). Data mining usage in Italian SMEs: an integrated SEM-ANN approach. *Central European Journal of Operations Research*, 31(3), 941-973. <https://doi.org/10.1007/s10100-022-00829-x>
- Pejić-Bach, M., & Marić, J. (2025). Unlocking artificial intelligence for all: Navigating the complexities of the artificial intelligence digital divide. In *Trust in Generative Artificial Intelligence* (pp. 36-48). Routledge.
- Rust, R. T., & Oliver, R. L. (1994). Service quality: Insights and implications. *Journal of Service Research*, 7(1), 65-82.
- Seyoung, L. & Park, G. (2023.), Exploring the Impact of ChatGPT Literacy on User Satisfaction, The Mediating Role of User Motivations, *Cyberpsychology, Behavior, and Social Networking*, 26(12), 913-912. <https://doi.org/10.1089/cyber.2023.0312>
- Sharma, V., Goyal, M., & Malik, D. (2017). An intelligent behaviour shown by chatbot system. *International Journal of New Technology and Research*, 3(4), 52-54.
- Shevtsova, D., Ahmed, A., Boot, I., Sanges, C., Hudecek, M., Jacobs, J., ... & Vrijhoef, H. (2024). Trust in and acceptance of artificial intelligence applications in medicine: mixed methods study. *Jmir Human Factors*, 11, e47031. <https://doi.org/10.2196/47031>
- Sindermann, C., Yang, H., Elhai, J., Yang, S., Quan, L., Li, M., ... & Montag, C. (2022). Acceptance and fear of artificial intelligence: associations with personality in a German and a Chinese sample. *Discover Psychology*, 2(1). <https://doi.org/10.1007/s44202-022-00020-y>
- Singh, H. (2023). Deep learning 101: Beginners guide to neural network. *Analytics Vidhya*. <https://www.analyticsvidhya.com/blog/2021/03/basics-of-neural-network/>
- Sotala, K. (2012). Advantages of artificial intelligences, uploads, and digital minds. *International Journal of Machine Consciousness*, 4(1), 275-291. <https://doi.org/10.1142/S1793843012400161>
- Uren, V., & Edwards, J. S. (2023). Technology readiness and the organizational journey towards AI adoption: An empirical study. *International Journal of Information Management*, 68, 102588. <https://doi.org/10.1016/j.ijinfomgt.2022.102588>
- Wang, A., Zhou, Y., Ma, H., Tang, X., Li, S., Pei, R., ... & Piao, M. (2024). Preparing for aging: understanding middle-aged user acceptance of AI chatbots through the technology acceptance model. *Digital Health*, 10. <https://doi.org/10.1177/20552076241284903>
- Yegnanarayana, B. (1994). Artificial neural networks for pattern recognition. *Sadhana*, 19(2), 189-238. <https://doi.org/10.1007/BF02811896>
- Yiğitcanlar, T., Degirmenci, K., & Inkinen, T. (2022). Drivers behind the public perception of artificial intel-

ligence: insights from major Australian cities. *Ai & Society*, 39(3), 833-853. <https://doi.org/10.1007/s00146-022-01566-0>

York, T., Jenney, H., & Jones, G. (2020). Clinician and computer: a study on patient perceptions of artificial intelligence in skeletal radiography. *BMJ Health & Care Informatics*, 27(3), e100233. <https://doi.org/10.1136/bmjhci-2020-100233>

Zeithaml, V. A. (1988). Consumer perceptions of price, quality, and value: A means-end model and synthesis of evidence. *Journal of Marketing*, 52(3), 2-22. <https://doi.org/10.1177/002224298805200302>

Zhang, C., Schießl, J., Plöchl, L., Hofmann, F., & Gläser-Zikuda, M. (2023). Acceptance of artificial intelligence among pre-service teachers: a multigroup analysis. *International Journal of Educational Technology in Higher Education*, 20(1). <https://doi.org/10.1186/s41239-023-00420-7>

Mirjana Pejić Bach is a full professor at the Department of Informatics, Faculty of Economics in Zagreb. She holds a PhD in systems dynamics modelling from the Faculty of Economics, University of Zagreb. She was trained at the MIT Sloan School of Management in system dynamics and at the Olivia Group in data mining. Mirjana is the leader and collaborator of numerous projects in which she cooperates with Croatian companies and international organisations, mainly through European Union projects and the bilateral research framework. Her research areas include the strategic application of information

technology in business, data science, simulation modelling, research methodology, and both qualitative and quantitative methods, with a focus on multivariate statistics and modelling structural equations.

Mirko Palić, PhD, is a full professor at the Marketing Department, Faculty of Economics and Business, University of Zagreb, Croatia. He is the head of the postgraduate program in Sales Management, and his teaching and research interests include marketing channels, personal selling, retail, and marketing management. His recent focus has been directed toward the research and application of artificial intelligence in marketing and higher education. He is the author of more than 70 scientific papers and several university textbooks. Before joining the University, he worked in managerial positions in several prominent Croatian companies.

Vanja Šimičević is full professor at the Libertas International University, Zagreb, Croatia. She received PhD from the University of Zagreb, Faculty of Economics and Business in the area of quantitative economics. Her major area of research is focused on applications of quantitative methods, particularly modelling techniques and analysis, in the field of business, economic and other social sciences, using a wide range of methods and techniques. Her scientific work is a significant contribution to the application of modern quantitative methods in social sciences.
