



POLİTEKNİK DERGİSİ

JOURNAL of POLYTECHNIC

ISSN: 1302-0900 (PRINT), ISSN: 2147-9429 (ONLINE)

URL: <http://dergipark.org.tr/politeknik>



Artificial intelligence (AI) use in data analysis: A comparison of ChatGPT and SmartPLS outputs in PLS-SEM analysis

*Veri analizinde yapay zeka kullanımı: PLS-SEM
analizinde ChatGPT ve SmartPLS çıktılarının
karşılaştırılması*

Yazar(lar) (Author(s)): Yavuz TORAMAN¹, Orçun Muhammet ŞİMŞEK²

ORCID¹: 0000-0002-5196-1499

ORCID²: 0000-0001-8028-3394

To cite to this article: Toraman Y., and Şimşek O. M., “Artificial İntelligence (AI) Use in Data Analysis: A Comparison of ChatGPT and SmartPLS Outputs in PLS-SEM Analysis”, *Journal of Polytechnic*, *(*) : *, (*).

Bu makaleye şu şekilde atıfta bulunabilirsiniz: Toraman Y., ve Şimşek O. M., “Artificial İntelligence (AI) Use in Data Analysis: A Comparison of ChatGPT and SmartPLS Outputs in PLS-SEM Analysis”, *Politeknik Dergisi*, *(*) : *, (*).

Erişim linki (To link to this article): <http://dergipark.org.tr/politeknik/archive>

DOI: 10.2339/politeknik.1739414

Artificial Intelligence (AI) Use in Data Analysis: A Comparison of ChatGPT and SmartPLS Outputs in PLS-SEM Analysis

Highlights

- ❖ ChatGPT demonstrates high consistency with SmartPLS in conducting PLS-SEM analyses
- ❖ Key model constructs show comparable reliability and validity across both tools.
- ❖ ChatGPT can serve as a valid alternative for PLS-SEM under guided user input.
- ❖ The study highlights AI's expanding role in empirical social science methodologies.
- ❖ Limitations include lack of graphical output and dependency on sequential commands.

Graphical Abstract

These results suggest that ChatGPT may be a complementary and cost-effective tool for researchers, especially those without access to licensed software. Future research should explore its usability with other statistical techniques and artificial intelligence models.

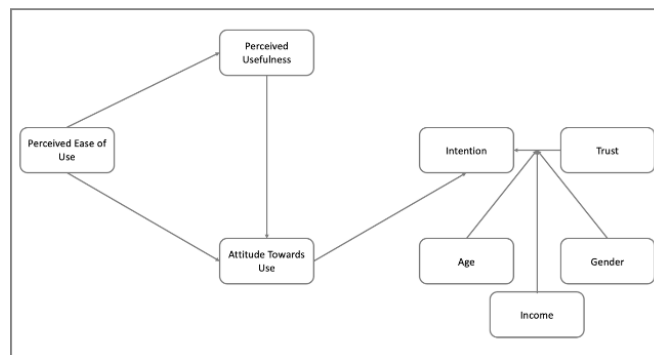


Figure. Research Model

Aim

This study aims to evaluate the usability and reliability of ChatGPT, a large language model, to conduct Partial Least Squares Structural Equation Modeling (PLS-SEM) analysis.

Design & Methodology

Search model based on the Technology Acceptance Model was tested using both SmartPLS and ChatGPT software through a Python-based statistical coding.

Originality

The originality of the research lies in testing the reliability and validity of AI technology, which is used in many fields in the final process, and presenting the results in simple and understandable language so that they can be used as a research tool.

Findings

Both tools identified the same significant and nonsignificant relationships, and similar mediation effects were found.

Conclusion

Although AI technologies have encountered some problems during their development, successful results have been achieved at this point. This situation shows that AI technologies will find their place in different sectors, especially in education, in the future.

Declaration of Ethical Standards

Ethical committee reviews were conducted by the Istanbul Nisantasi University (reference number: 2024/16).

Artificial Intelligence (AI) Use in Data Analysis: A Comparison of ChatGPT and SmartPLS Outputs in PLS-SEM Analysis

Araştırma Makalesi / Research Article

Yavuz TORAMAN^{1*}, Orçun Muhammet ŞİMŞEK²

¹Vocational School, Foreign Trade, İstanbul Nişantaşı University, Türkiye

²Faculty of Economics, Administrative And Social Sciences, Social Work, İstanbul Nişantaşı University, Türkiye

(Geliş/Received : 10.07.2025 ; Kabul/Accepted : 29.10.2025 ; Erken Görünüm/Early View : 02.11.2025)

ABSTRACT

This study aims to evaluate the usability and reliability of ChatGPT, a large language model, to conduct Partial Least Squares Structural Equation Modeling (PLS-SEM) analysis. For this purpose, a research model based on the Technology Acceptance Model was tested with 523 participants using both SmartPLS and ChatGPT. Key statistical indicators such as internal consistency, convergent and discriminant validity, VIF values, path coefficients, and R-square values were compared. The results show a high level of consistency between ChatGPT and SmartPLS outputs throughout all stages of the analysis. Furthermore, the statistical validity of SmartPLS and ChatGPT outputs is supported by examining their Pearson Correlation (r), Mean Absolute Error (MAE), and Root Mean Square Error (RMSE) values. Both tools identified the same significant and nonsignificant relationships, and similar mediation effects were found. Although ChatGPT lacks visualization capabilities and requires step-by-step guidance, it successfully reproduced statistically valid results. These results suggest that ChatGPT may be a complementary and cost-effective tool for researchers, especially those without access to licensed software. Future research should explore its usability with other statistical techniques and artificial intelligence models.

Anahtar Kelimeler: SmartPLS, artificial intelligence, statistical analysis, ChatGPT, PLS-SEM.

Veri Analizinde Yapay Zeka Kullanımı: PLS-SEM Analizinde ChatGPT ve SmartPLS Çıktılarının Karşılaştırılması

ÖZ

Bu çalışma, Kısmi En Küçük Kareler Yapısal Eşitlik Modellemesi (PLS-SEM) analizi yapmak için geniş bir dil modeli olan ChatGPT'nin kullanılabilirliğini ve güvenilirliğini değerlendirmeyi amaçlamaktadır. Bu amaçla, Teknoloji Kabul Modeline dayalı bir araştırma modeli, hem SmartPLS hem de ChatGPT kullanılarak 523 katılımcıyla test edilmiştir. İç tutarlılık, güvenilirlik ve geçerlilik, VIF değerleri, yol katsayıları ve R-kare değerleri gibi temel istatistiksel göstergeler karşılaştırılmıştır. Sonuçlar, analizin tüm aşamalarında ChatGPT ve SmartPLS çıktıları arasında yüksek düzeyde tutarlılık olduğu gözlemlenmiştir. Ayrıca SmartPLS ve ChatGPT çıktılarının Pearson Correlation (r), Mean Absolute Error (MAE) ve Root Mean Square Error (RMSE) değerleri incelenerek istatistiksel olarak desteklenmiştir. Her iki araç da aynı anlamlı ve anlamsız ilişkileri tespit etmiş ve benzer aracılık etkileri bulunmuştur. ChatGPT görselleştirme yeteneklerinden yoksun olmasına ve adım adım rehberlik gerektirmesine rağmen, istatistiksel olarak geçerli sonuçları başarılı bir şekilde yeniden üretmiştir. Bu sonuçlar ChatGPT'nin özellikle lisanslı yazılımlara erişimi olmayan araştırmacılar için tamamlayıcı ve uygun maliyetli bir araç olabileceğini göstermektedir. Gelecekteki araştırmalar bu aracın diğer istatistiksel teknikler ve yapay zeka modelleri ile kullanılabilirliğini araştırmalıdır.

Keywords: SmartPLS, yapay zeka, istatistiksel analiz, ChatGPT, PLS-SEM.

1. INTRODUCTION

Statistical methods provide valuable approaches for defining relationships between variables. Each method operates based on different statistical foundations. Common methods in the literature, such as correlation, regression, factor analysis, and structural equation modeling (SEM), focus on testing different assumptions and offer different perspectives [1], [2], [3].

Each method is known to make significant contributions to the literature and ease the work of researchers who aim to uncover the complex relationships between variables.

In particular, analysis methods like structural equation modeling, which focus directly on complex relationships, are widely studied in the literature [4]. SEM offers two common approaches: covariance-based and partial least squares. Programs like AMOS allow the development of complex models using covariance-based structural equation modeling, while programs like SmartPLS work with partial least squares-based structural equation modeling [5], [6], [7]. Although all these methods focus on relationships between variables, due to the ability to

*Sorumlu Yazar (Corresponding Author)

e-posta : yavuz.toraman@nisantasi.edu.tr

work with smaller sample sizes, the SmartPLS software package has been focused on in this study.

The foundation of PLS-SEM was laid by Swedish econometrician Herman Wold in the 1970s through his work on new models and methods in social sciences [8], [9]. Subsequently, less complex models and less data were emphasized and these studies evolved into PLS-SEM after the 2000s [10]. Unlike covariance-based structural equation modeling, partial least squares-based structural equation modeling (PLS-SEM) is used in many fields due to its fewer rigid rules. The ease of use and relatively flexible structure of PLS-SEM made it a frequently used method in social sciences in a short period of time. Today, there are different software programs used specifically for PLS-SEM [11], [12]. For example, software and algorithms based on tools like WarpPLS, ADANCO, MATLAB (PLS-GUI), R, and Python are found in the literature. However, due to its user-friendly interface, the SmartPLS software package is often preferred. In summary, the PLS-SEM literature highlights both the methodological flexibility and the practical utility of SmartPLS in social sciences. Having established this foundation, the next step is to examine how emerging technologies—particularly artificial intelligence—may interact with and potentially complement these statistical tools.

Building on the discussion of PLS-SEM, this study also situates itself within the growing literature on artificial intelligence (AI). AI, a relatively new field based on software and algorithms, has spread to every field as the Renaissance of the 21st century, making human life easier. Artificial intelligence, which is considered to have originated from the 1955 Dartmouth Conference, has reached a practical form in the last five years [13], [14]. Today, there are numerous artificial intelligence chatbots, such as ChatGPT, Google Gemini, and Bing Chat (Microsoft), and these AI tools are used for various purposes depending on the user [15]. ChatGPT, as a more capable AI application compared to its peers, serves a large number of users [16]. Although ChatGPT is currently used in fields such as healthcare, chemistry, tourism, psychology, and education, it is expected to be used in more complex tasks in various fields in the future, in line with its increasing capabilities [15]. While these AI tools bring about questions regarding their reliability, the present study aims to provide evidence to eliminate the uncertainties in statistical processes. Thus, after reviewing the methodological foundations of PLS-SEM and the rapid growth of AI applications, our study connects these two streams by directly comparing ChatGPT with SmartPLS. In this context, analyses performed with SmartPLS were tested with ChatGPT, and the findings were compared. This study emphasizes the use of ChatGPT in statistical analyses and demonstrates its compatibility with SmartPLS, a widely used and reliable statistical program. What differentiates this study from previous research is that, to the best of our knowledge, it is one of the first attempts to systematically compare the outputs of a large language

model (ChatGPT) with those of a specialized statistical software (SmartPLS) in the context of PLS-SEM. While prior studies have mainly focused on the theoretical development of PLS-SEM or the application of SmartPLS in various fields, there has been little to no research on whether AI-powered tools can replicate and validate such statistical processes. By addressing this gap, the present study contributes to the literature by positioning ChatGPT not merely as a conversational AI but as a potential complementary tool for empirical research methods.

2. STRUCTURAL EQUATION MODELING (SEM)

Linear regression, logistic regression, analysis of variance, and multiple regression are statistical methods that primarily examine the relationships between variables. Moreover, the influence of more than one independent variable on the dependent variable or some mediating and moderating roles can be adopted by researchers other than examining only the relationship between two variables. For such analyses, Structural Equation Modelling (SEM) stands out in the current statistical literature. Beyond the logic of linear regression, SEM provides the opportunity to analyze multiple relationships between more than one independent variable and the dependent variable or multiple relationships with mediating and moderating variables [11], [18].

There are two basic methods in SEM. The first is covariance-based SEM (CB-SEM) and the second is partial least squares SEM (PLS-SEM). The main approach in CB-SEM is to test a theoretical hypothesis with large data sets. In CB-SEM, hypotheses are confirmed or rejected depending on how accurately the proposed model reproduces the covariance matrix of the data set from which it is constructed [4]. On the other hand, PLS-SEM focuses on explaining the variation in the dependent variable. The absence of the assumption of normal distribution and its suitability for small sample sizes facilitate exploratory studies. The assumptions of CB-SEM are relaxed in PLS-SEM. The assumptions relaxed in PLS-SEM are particularly related to distributional assumptions [8]. The PLS-SEM will be discussed in detail in the next section.

2.1. Partial Least Squares Structural Equation Modeling (PLS-SEM)

PLS-SEM is often preferred because it is more flexible in terms of distributional assumptions compared to other techniques [17]. Moreover, the ability to work with smaller samples and test complex models has brought PLS-SEM (Partial Least Squares Structural Equation Modeling) to the forefront [18], [19], [20]. In this way, PLS-SEM has been mostly used in social sciences research [21], [22], [23], [24], [25]. When the measurement approaches are analyzed, PLS-SEM can perform two types of measurements: reflective and formative. Due to the nature of the current study, the focus of this study is based on reflective measurement.

This measurement model is based on validity and reliability analyses such as Discriminant Validity, Internal Consistency Reliability and Convergent Validity. Two stages are taken into consideration to carry out the model analyses. Firstly, validity and reliability outputs are examined and then the research model is analyzed. Like other SEMs, the aim here is to ensure the validity and reliability required for the application of the research model. Internal Consistency Reliability is calculated using Cronbach's Alpha, Composite Reliability, and rho_A reliability coefficients. These values must be equal to or greater than ≥ 0.70 . For Convergent Validity, the Outer Loading and Average Variance Extracted (AVE) coefficients are calculated. Outer Loadings should be equal to or greater than ≥ 0.70 , and the AVE coefficient should be equal to or greater than ≥ 0.50 [4], [11], [16]. For Discriminant Validity, the Fornell-Larcker Criterion, Cross Loadings, and Heterotrait-Monotrait Ratio (HTMT) coefficients are calculated. Among these, HTMT and Fornell-Larcker are frequently preferred analyses. In the present study, the correlation analysis calculated according to the Fornell-Larcker Criterion will be used. Fornell-Larcker is calculated by taking the square root of the AVE coefficients. Therefore, in the correlation table, it is expected that the correlation of each variable with itself, shown at the top of the column, will be higher than the correlations with other variables. Additionally, the Variance Inflation Factor (VIF) values should be less than 5. Since this study also includes a comparison between SmartPLS and ChatGPT, the Excess Kurtosis and Skewness values, which are not typically considered in PLS-SEM, have been included in the research. After statistically valid results are obtained from the reliability and validity analyses, hypothesis testing can be conducted. In this section, hypothesis tests were performed using the Bootstrapping method in PLS-SEM, while in ChatGPT, the tests were conducted using the "statsmodels.OLS" Ordinary Least Squares (OLS) model in an open-access Python database. Finally, the R-square and R-square adjusted coefficients of the dependent variable are obtained from the analysis. In this section, the R^2 of the dependent variable is categorized as follows:

$0.25 < R^2 < 0.50 \Rightarrow$ Weak
 $0.50 \leq R^2 \leq 0.70 \Rightarrow$ Moderate
 $0.70 < R^2 \Rightarrow$ Strong

These value ranges can be defined in this way [26].

3. ARTIFICIAL INTELLIGENCE (AI)

Artificial Intelligence (AI) is an approach aimed at endowing machines with human-like cognitive functionalities. The goal is to channel cognitive abilities such as decision-making, learning, reasoning, and problem-solving from human intelligence to machines, which is conceptualized as artificial intelligence. By establishing causal relationships and drawing on previous experiences, artificial intelligence (AI) enables computers to simulate human intelligence [27], [12]. AI

is a breakthrough for humanity as it can enhance the efficiency of human tasks across various domains and potentially replace them [28]. It is remarkable that non-human, lifeless machines can communicate and engage in logical reasoning, similar to the biological capacities of humans [29]. However, like human intelligence, the source of learning for AI is data, which makes the process more understandable. Through algorithms and statistics, data functions like neurons in the human brain, establishing connections between each other and forming the process of thinking in machines.

Strong language models, like the Generative Pre-trained Transformer (GPT) series, have been developed as a result of recent advances in artificial intelligence. These models are characterized by being pre-trained, which enables them to achieve outstanding success in numerous NLP tasks, including language translation, summarizing large texts, and providing quick answers to questions. Today, large language models (LLMs) such as ChatGPT (GPT-3.5, GPT-4 and GPT-5), Google Gemini and Bing Chat (Microsoft) are among these models. Among them, ChatGPT has particularly demonstrated its capabilities in various fields, from human-computer interaction to numerous areas of scientific research. This potential has made ChatGPT an object of wide interest and broadened its horizon of usability in various application areas.

4. PRESENT STUDY

The purpose of this paper is to highlight the role of large language models such as ChatGPT in statistical analyses as part of this expanding horizon. To shed light on ChatGPT's capabilities in various fields, we subjected ChatGPT to PLS-SEM-based analyses and compared the results with analyses performed with PLS-SEM-based SmartPLS software. In this context, we sought to answer the research question RQ1: Is ChatGPT reliable for use in PLS-SEM analyses? We tested the research model presented in Figure 1. Our research model is built on relationships between variables based on the Technology Acceptance Model. According to this model, electronic money was considered a new technology, and the factors affecting its acceptance were examined. In this context, the following paths between variables were analyzed:

Perceived Ease of Use (PEOU) $\rightarrow$ Perceived Usefulness (PU)

Perceived Ease of Use (PEOU) $\rightarrow$ Attitude Towards Use (AT)

Perceived Usefulness (PU) $\rightarrow$ Attitude Towards Use (AT)

Attitude Towards Use (AT) $\rightarrow$ Intention (I)

Trust (T) $\square$ Intention (I)

Perceived Ease of Use (PEOU) $\rightarrow$ Perceived Usefulness (PU) $\rightarrow$ Attitude Towards Use (AT)

Income x Trust (T) $\rightarrow$ Intention (I)

Gender x Trust (T) $\rightarrow$ Intention (I)

Age x Trust (T) $\rightarrow$ Intention (I)

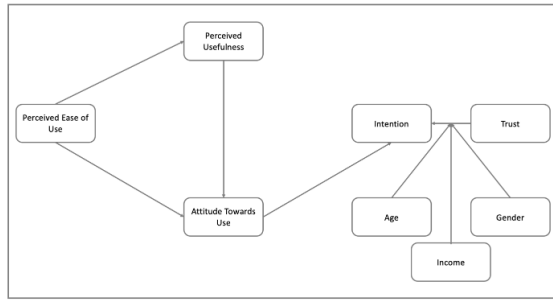


Figure 1. Research Model

5. MATERIALS AND METHODS

5.1. Study Design and Participants

In this study, partial least squares structural equation modeling (PLS-SEM) method is used to examine the capabilities of ChatGPT in different domains. The performance of ChatGPT is evaluated by PLS-SEM based analysis using Python programming language. In addition, the results obtained are compared with the analyses performed with SmartPLS software using the PLS-SEM method. In the implementation process, the data uploaded to ChatGPT were subjected to PLS-SEM analysis with Python; model estimation and validation were performed. Then, the same structural models were analyzed using SmartPLS software and performance evaluation was performed by comparing ChatGPT outputs with SmartPLS results. This method provided a comprehensive and comparative evaluation of ChatGPT's capabilities in terms of statistical modeling.

The 523 participants who voluntarily supported the study were reached between May 1 and June 1, 2024, through the online survey platform Survey Monkey (<https://tr.surveymonkey.com>, accessed on June 2, 2024). 51.1% of participants are aged 18-21; 48.9% were aged 22-30. 42.4% of the participants were female, 57.6% were male. In terms of income level, the majority of the participants were in the low-income level (44.5%); 30.6% middle income, and 24.9% had high income. Participants were informed about the purpose of the study, confidentiality, and the reliability of the data collected, and their consent was obtained. No personal information was requested, IP addresses were not recorded, and participants' privacy was protected. They were allowed to start, complete, or withdraw from the survey at any time. It was determined that the survey took approximately 3 minutes to complete. The study protocol was reviewed and approved by İstanbul Nişantaşı University (reference number:2024/16). The Declaration of Helsinki's guiding principles were followed when conducting the survey.

5.2. Data Analysis

SmartPLS software (version 4.1) was used to perform the study's normality, reliability, descriptive statistics, correlation, and structural equation modeling. A 95% confidence interval based on 1000 resamples produced

by the bias-corrected bootstrapping method was used to test the model's direct and indirect effects. To observe the indirect effects, specific indirect effects were also investigated. Additionally, ChatGPT (version 4o Plus) was used to perform the same analyses, and the consistency of the results was evaluated by comparing them.

5.3. Measures

In this study, a questionnaire consisting of demographic questions and Perceived Ease of Use, Perceived Usefulness, Attitude Towards, and Intention Factors within the scope of the Technology Acceptance Model, as well as the Trust variable were used as data collection tools.

5.4. Demographic Information Form

The first part of the questionnaire consists of demographic questions to obtain demographic information of the participants. This form includes questions on age, gender, and income levels.

5.5. Technology Acceptance Model (TAM)

The Technology Acceptance Model was developed by Davis (Davis, 1986). The model consists of a combination of the dimensions Perceived Ease of Use, Perceived Usefulness, Attitude Towards Use, and Intention. Perceived Ease of Use includes 4 items, Perceived Usefulness 4 items, Attitude Towards Use 3 items, and Intention 3 items. A 5-point Likert scale ranging from 'Strongly Agree' = 5 to 'Strongly Disagree' = 1 was used for the responses.

5.6. Trust

During the development of the Technology Acceptance Model, the variable Trust was tested within the model as a new factor by Pavlou (Pavlou, 2001, 2003). The variable consists of 3 items. A 5-point Likert scale ranging from 'Strongly Agree' = 5 to 'Strongly Disagree' = 1 was used for the responses.

6. RESULTS

In the current study, the values of Excess Kurtosis and Skewness, which allow for the examination of normality, are presented in Table 1 for SmartPLS and in Table 2 for ChatGPT. Excess Kurtosis and Skewness values within the range of -1.5 to +1.5 indicate that the assumption of normality is statistically met [30]. In this study as well, the analysis results for both SmartPLS and ChatGPT show that the assumption of normality is met. The dataset was uploaded to ChatGPT, and the prompt "Calculate the Excess Kurtosis and Skewness values for the current dataset" was given in order to obtain these values.

Subsequently, the item-level VIF (Variance Inflation Factor) values were examined to determine whether there was a multicollinearity issue. Although VIF values below 10 and 5 are generally accepted, values below 3 are considered an indicator that there is no multicollinearity problem [4].

Table 1. Construct validity and reliability (SmartPLS)

Constructs	Items	Excess Kurtosis	Skewness	VIF	Factor Loading	Cronbach's Alpha	CR	AVE
Attitude Towards Use	AT1	0.114	-0.187	2.464	0.896	0.880	0.926	0.807
	AT2	0.151	-0.240	2.223	0.876			
	AT3	-0.170	-0.217	2.879	0.922			
Intention	I1	0.058	-0.180	2.561	0.902	0.894	0.934	0.825
	I2	0.142	-0.358	2.937	0.917			
	I3	0.150	-0.263	2.604	0.906			
Perceived Usefulness	PU1	0.151	-0.206	2.124	0.848	0.889	0.923	0.751
	PU2	0.226	-0.373	2.836	0.890			
	PU3	0.265	-0.378	2.906	0.888			
Perceived Ease of Use	PEOU1	0.274	-0.399	2.023	0.839	0.907	0.935	0.781
	PEOU2	0.329	-0.398	2.701	0.876			
	PEOU3	0.574	-0.436	2.470	0.868			
Trust	PEOU4	0.226	-0.305	3.240	0.904	0.867	0.918	0.789
	T1	0.287	-0.363	2.824	0.887			
	T2	-0.019	-0.131	2.020	0.880			
	T3	0.133	-0.208	2.570	0.902			
		0.326	-0.222	2.360	0.883			

Note. VIF: Variance Inflation Factor, CR: Composite Reliability, AVE: Average Variance Extracted

In the current study, the VIF values are presented in Table 1 and Table 2. Upon reviewing the item-level VIF values in both SmartPLS and ChatGPT results, it was found that only PEOU3 had a value above 3, specifically 3.240. Both ChatGPT and SmartPLS produced the same VIF value. In this context, the VIF values of the study are statistically acceptable, and it is noteworthy that both software tools yielded identical results. To obtain these values, the prompt “Calculate multicollinearity using VIF” was given to ChatGPT. To determine the validity and reliability of the variables, Factor Loadings, Cronbach's Alpha, Composite Reliability (CR), and Average Variance Extracted (AVE) values were calculated.

Table 2. Construct validity and reliability (ChatGPT)

Constructs	Items	Excess Kurtosis	Skewness	VIF	Factor Loading	Cronbach's Alpha	CR	AVE
Attitude Towards Use	AT1	0.094	-0.186	2.464	0.895	0.880	0.926	0.807
	AT2	0.131	-0.239	2.223	0.881			
	AT3	-0.185	-0.216	2.879	0.919			
Intention	I1	0.040	-0.180	2.561	0.901	0.894	0.934	0.825
	I2	0.122	-0.356	2.937	0.919			
	I3	0.130	-0.262	2.604	0.904			
Perceived Usefulness	PU1	0.131	-0.205	2.124	0.848	0.889	0.923	0.751
	PU2	0.205	-0.372	2.836	0.889			
	PU3	0.243	-0.376	2.906	0.891			
Perceived Ease of Use	PEOU1	0.252	-0.397	2.023	0.838	0.907	0.935	0.781
	PEOU2	0.306	-0.396	2.701	0.878			
	PEOU3	0.547	-0.434	2.470	0.869			
Trust	PEOU4	0.205	-0.304	3.240	0.902	0.867	0.918	0.790
	T1	0.265	-0.361	2.824	0.886			
	T2	-0.036	-0.130	2.020	0.869			
	T3	0.113	-0.207	2.570	0.904			
		0.303	-0.221	2.360	0.893			

Note. VIF: Variance Inflation Factor, CR: Composite Reliability, AVE: Average Variance Extracted

In this context, to evaluate the constructs in terms of reliability and validity and to test the hypotheses, threshold values are accepted as 0.70 for factor loadings, Cronbach's Alpha and CR, and 0.50 for AVE [11].

Table 1 and Table 2 present the values of factor loadings, Cronbach's Alpha, CR and AVE. Examining the values,

it was found that the results of the analysis of both SmartPLS and ChatGPT produced statistically significant results and that the values were identical or very similar

Table 3. Discriminant validity analysis-Fornell-Larcker (SmartPLS)

Items	1	2	3	4	5	6	7	8
1 Attitude Towards Use	0.898							
2 Gender	-0.137	1.000						
3 Intention	0.810	-0.210	0.908					
4 Income	-0.034	0.156	-0.132	1.000				
5 Age	0.042	-0.000	0.025	0.204	1.000			
6 Perceived Ease of Use	0.680	-0.181	0.721	-0.146	0.048	0.884		
7 Perceived Usefulness	0.763	-0.143	0.786	-0.099	0.066	0.847	0.867	
8 Trust	0.651	-0.164	0.700	-0.134	0.040	0.678	0.678	0.888

In the correlation analysis of the study, the results were obtained based on the Fornell-Larcker Criteria. Correlation analysis is presented in Table 3 and Table 4. In correlation analysis, the correlation coefficient of a variable with itself is expected to be high. By examining the results in both tables, it was determined that they are statistically significant [11]. ChatGPT follows a different sequence of analysis steps than SmartPLS. In SmartPLS, the model is first subjected to validity and reliability analysis, followed by bootstrapping to test direct, indirect, and moderating effects. ChatGPT, on the other hand, performs each analysis step separately, following the sequence of validity and reliability analysis, then

direct, indirect, and moderating effect analyses. Due to this sequential approach, while demographic factors were included in the correlation analysis in SmartPLS, they were not included in ChatGPT's correlation analysis. Although SmartPLS and ChatGPT conducted correlation analyses with different contents due to their structural differences, both produced generally similar results (Tables 3 and 4). In both analyses, the highest correlation was found between PEOU and PU. To obtain the values in this section, ChatGPT was commanded with the prompt "Calculate discriminant validity according to Fornell-Larcker Criterion".

Table 4. Discriminant validity analysis-Fornell-Larcker (ChatGPT)

Items	1	2	3	4	5
1 Perceived Usefulness	0.867				
2 Perceived Ease of Use	0.846	0.884			
3 Attitude Towards Use	0.762	0.679	0.898		
4 Trust	0.675	0.677	0.649	0.889	
5 Intention	0.785	0.722	0.810	0.697	0.908

After the success in reliability and validity analyses, the paths were tested. The test results of the direct and moderation analysis conducted with the SmartPLS

program are shown in Table 5 and the results of the mediating variable analysis are shown in Table 6.

Table 5. Path coefficients (SmartPLS)

Paths	Path coefficient	t statistics	p statistics
Perceived Ease of Use → Perceived Usefulness	0.847	32.637	0.000*
Perceived Ease of Use → Attitude Towards Use	0.118	1.160	0.246
Perceived Usefulness → Attitude Towards Use	0.663	6.957	0.000*
Attitude Towards Use → Intention	0.617	12.441	0.000*
Trust → Intention	0.293	4.414	0.000*
Income x Trust → Income	-0.013	0.311	0.756
Gender x Trust → Income	-0.029	0.487	0.626
Age x Trust → Income	0.012	0.382	0.702

Note: * significant is 0.001 level.

In the research model, 4 out of the 8 paths were found to be significant, while the other 4 were not. Firstly, Perceived Ease of Use (PEOU) had a positive effect on Perceived Usefulness (PU) ($\beta = 0.847$, $p < 0.05$). PEOU had no effect on Attitude Toward Use (AT) ($\beta = 0.118$, $p > 0.05$). PU had a positive effect on AT ($\beta = 0.663$, $p < 0.05$).

AT had a positive effect on Intention to Use (I) ($\beta = 0.617$, $p < 0.05$). Trust (T) had a positive effect on I ($\beta = 0.293$, $p < 0.05$). Furthermore, demographic characteristics did not have a moderating effect on the relationship between T and I.

Table 6. Specific indirect effects (SmartPLS)

Path	Path coefficient	t statistics	p statistics
Perceived Ease of Use → Perceived Usefulness → Attitude Towards Use	0.561	6.979	0.000*

Note: * significant is 0.001 level.

Table 6 indicates the presence of a significant indirect relationship between the variables. According to the results, PU plays a mediating role in the relationship between PEOU and AT.

In Tables 7 and 8, the paths in the model were analyzed using ChatGPT. To obtain these values, the prompt “Calculate path coefficients and specific indirect effects” was given to ChatGPT.

Similarly, ChatGPT, like SmartPLS, found 4 out of the 8 paths to be significant and 4 to be non-significant in

Table 7. PEOU had a positive effect on PU ($\beta = 0.852$, $p < 0.05$). PEOU had no effect on AT ($\beta = 0.132$, $p > 0.05$). PU had a positive effect on AT ($\beta = 0.723$, $p < 0.05$). AT had a positive effect on I ($\beta = 0.613$, $p < 0.05$). T had a positive effect on I ($\beta = 0.317$, $p < 0.05$). Users' demographic characteristics did not have a moderating effect on the relationship between T and I.

When the relationship between PEOU and AT through PU was examined using ChatGPT, PU was similarly found to have a mediating role, as shown in Table 8.

Table 7. Path coefficients (ChatGPT)

Paths	Path coefficient	p statistics
Perceived Ease of Use → Perceived Usefulness	0.852	0.000*
Perceived Ease of Use → Attitude Towards Use	0.132	0.074
Perceived Usefulness → Attitude Towards Use	0.723	0.000*
Attitude Towards Use → Intention	0.613	0.000*
Trust → Intention	0.317	0.000*
Income x Trust → Income	-0.028	0.560
Gender x Trust → Income	-0.078	0.735
Age x Trust → Income	0.002	0.741

Note: * significant is 0.001 level.

Table 8. Specific indirect effects (ChatGPT)

Path	Path coefficient	p statistics
Perceived Ease of Use → Perceived Usefulness → Attitude Towards Use	0.616	0.000*

Note: * significant is 0.001 level.

The R^2 values of the study for both SmartPLS and ChatGPT are presented in Table 9. The explained variance of the dependent variable I was found to be 0.718 for SmartPLS and 0.706 for ChatGPT. These values indicate a strong explanatory power [26], [31], [32]. To obtain these values, the prompt “Calculate results of R-square and R-square adjusted” was given to ChatGPT.

When all tables are examined, both SmartPLS and ChatGPT produced close and similar results. Due to ChatGPT's lack of visualization capabilities for model analyses, the results can only be presented in a summarized visual form, as shown in Figure 2, through SmartPLS.

Finally, it must be statistically verified and proven that the analyses conducted using both the traditional SmartPLS software package and ChatGPT, an artificial intelligence (AI) application currently used in many fields, within the scope of PLS-SEM are very similar and almost identical.

Therefore, the SmartPLS and ChatGPT outputs obtained will be tested using Pearson Correlation (r) Mean, Absolute Error (MAE), and Root Mean Square Error (RMSE) analyses. Excess Kurtosis, Skewness, VIF, and Factor Loading were examined with 68 parameters, Cronbach's Alpha, CR, and AVE with 15 parameters, Path Coefficients with 8 parameters, and finally R^2 and $Radj^2$ with 3 parameters.

Table 9. Results of R-square and R-square adjusted (SmartPLS and ChatGPT)

Variables	SmartPLS		ChatGPT	
	R ²	Radj ²	R ²	Radj ²
Attitude Towards Use	0.587	0.584	0.585	0.581
Intention	0.718	0.710	0.706	0.703
Perceived Usefulness	0.717	0.716	0.726	0.724

Note: R²: R-square, Radj²: R-square adjusted.

Pearson Correlation Coefficient (r)

The Pearson Correlation coefficient (r) indicates the strength and direction of the linear relationship between two variables [33], [34]. The formula is presented below.

$$r = \frac{\sum_{i=1}^n (S_i - \bar{S})(C_i - \bar{C})}{\sqrt{\sum_{i=1}^n (S_i - \bar{S})^2} \sqrt{\sum_{i=1}^n (C_i - \bar{C})^2}}$$

- S_i = the value generated by SmartPLS,
- C_i = the corresponding value generated by ChatGPT,
- $\bar{S}$ and $\bar{C}$ = the mean values of SmartPLS and ChatGPT results,
- n = number of comparable outputs.

In this study, the outputs of the traditional software package (SmartPLS) and AI (ChatGPT) were compared. The similarity of the statistical outputs obtained from the two different programs was supported by the Pearson Correlation coefficient (r). As seen in Table 10, the r coefficient supports that the coefficients in the table (1, 2, 5, 7, and 9) show parallel trends. The r coefficient, which indicates the strength and direction of the relationship between variables, takes values between +1 and -1. Furthermore, when the r value is 0, it indicates that there is no linear relationship between the variables [34]. When examining the r values in Table 10, it was found that they ranged from 0.9984 to 0.99999. These values, falling between 0.90 and 0.99, indicate the presence of a very strong positive relationship [33], [34].

Mean Absolute Error (MAE)

MAE is an error metric that measures the magnitude of the difference between two data sets (SmartPLS and ChatGPT). MAE was used to determine the average difference between the outputs. The formula for MAE is presented below (33).

$$MAE = \frac{1}{n} \sum_{i=1}^n |S_i - C_i|$$

- S_i = SmartPLS output for the parameter,
- C_i = ChatGPT output for the same parameter,
- n = number of comparable outputs.

In this study, the outputs of the traditional software package (SmartPLS) and AI (ChatGPT) were compared.

The difference between the statistical outputs obtained from the two different programs was supported by the MAE value. The smaller the MAE value, the more accurate the outputs of SmartPLS and ChatGPT are. An MAE value of 0 indicates that there is no difference between the two groups (33).

As presented in Table 10, the low error rate in the differences (MAE) between the coefficients in the table (1, 2, 5, 7, and 9) supports the similarity of the outputs. MAE shows a deviation of 0.008 and 0.0142 between SmartPLS and ChatGPT outputs. When MAE values are examined, Tables (5 and 7) show very good agreement with 0.0142, while other Tables show perfect agreement [35].

Root Mean Square Error (RMSE)

RMSE is an error measure calculated by taking the square root of the average of the squares of the differences between two outputs. Therefore, it is more sensitive and an error measure compared to MAE. In RMSE, the differences between the outputs of SmartPLS and ChatGPT were detected as in MAE [36]. The formula for RMSE is presented below.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (S_i - C_i)^2}$$

- S_i = SmartPLS output for the parameter,
- C_i = ChatGPT output for the parameter,
- n = number of comparable outputs.

This study compares the outputs of a traditional software package (SmartPLS) and AI (ChatGPT). The difference between the statistical outputs obtained from the two different programs is supported by the RMSE value. The smaller the RMSE value, the more it supports the accuracy of the SmartPLS and ChatGPT outputs. An RMSE value of 0 indicates that there is no difference between the two groups [36]. As presented in Table 10, the low error rate in the differences between the coefficients (RMSE) in the table (1, 2, 5, 7, and 9) supports the similarity of the outputs. The highest value of RMSE being 0.235 is among the important outputs of the study. $RMSE \leq 0.5$ indicates a statistically perfect fit between the outputs [36].

Table 10. Results of Pearson Correlation, Report Mean Absolute Error (MAE) or Root Mean Square Error (RMSE) (SmartPLS and ChatGPT Table 1,2,5,7 and 9)

Table	Parameters Included	n	r	MAE	RMSE
1 vs 2	Excess Kurtosis, Skewness, VIF, Factor Loading	68	0.998	0.008	0.012
1 vs 2	Cronbach's Alpha, CR and AVE	15	0.99999	0.00007	0.00026
5 vs 7	Path Coefficients	8	0.9984	0.0142	0.0235
9	R ² and Radj ²	3	0.9999	0.0017	0.0019

n: Parameters

Overall, when Table 10 is examined, it is evident that the analyses conducted using SmartPLS and ChatGPT within the PLS-SEM framework exhibit a high degree of

consistency, which is further substantiated by statistical measures such as r, MAE, and RMSE values.

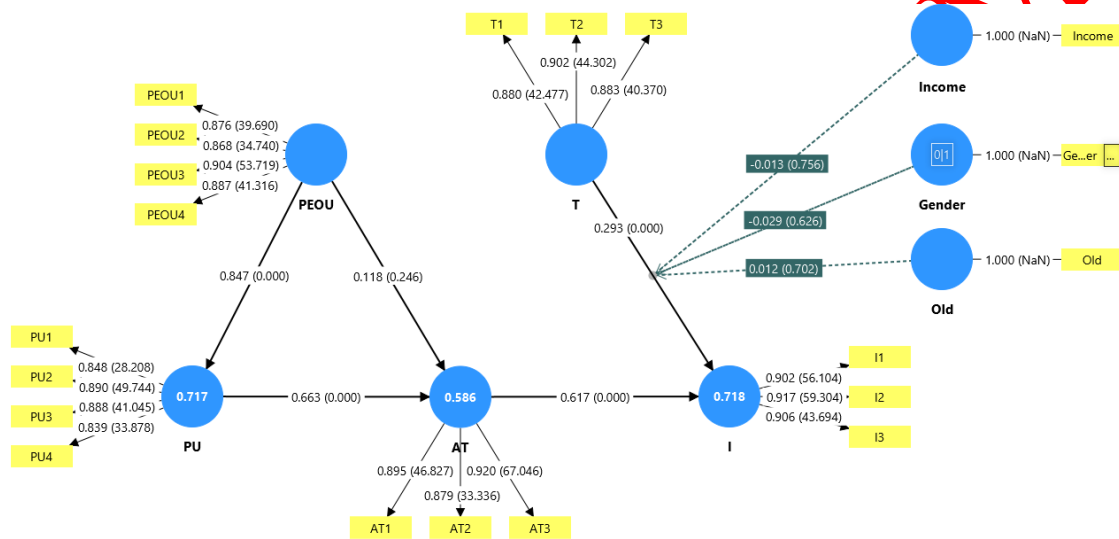


Figure 2. Visualization of the Path Analysis Results of the Study (SmartPLS)

7. DISCUSSION

This study aimed to evaluate whether ChatGPT, a widely used artificial intelligence tool, can reliably perform PLS-SEM analyses by comparing its outputs with those of the SmartPLS software. The findings indicated a high level of consistency between the two tools across key indicators such as reliability, validity, path coefficients, and explained variances (R² values). With nearly identical values for factor loadings, Cronbach's alpha, composite reliability (CR), average variance extracted (AVE), and discriminant validity based on the Fornell-Larcker criterion, both instruments specifically identified four statistically significant and four nonsignificant paths. Additionally, both approaches found that the relationship between perceived ease of use (PEOU) and attitude towards use (AT) was significantly mediated by perceived usefulness (PU). These results suggest that there are reasonable grounds to claim that ChatGPT, when used within specific guidelines, has the potential to yield results that can be regarded as statistically valid. This underlines the study's originality, as it directly addresses an unexplored gap in the literature: the lack of systematic evaluations of whether AI-based tools can serve as reliable alternatives to licensed statistical

software in PLS-SEM. By doing so, the study provides a novel contribution that extends both the methodological and practical boundaries of AI-assisted research.

PLS-SEM is successful in modelling complex relationships between variables due to its ability to work with smaller samples and its capacity to make fewer distributional assumptions, and the findings of the current study are in line with previous studies that emphasize the flexibility and effectiveness of PLS-SEM [11], [8]. While SmartPLS has been recognized as a standard tool used in the social sciences for almost a quarter of a century, the current study extends the potential of such large language models (LLMs) in the field of statistical analysis by showing that ChatGPT can provide similar statistical outputs when properly guided through the right commands and prompts. In line with previous literature on the use of AI in different disciplines such as tourism, health and education, this study fills an important gap in how AI can directly contribute to empirical research methods [15], [16].

Despite this success in statistical analyses, ChatGPT needs some improvements. While SmartPLS provides the opportunity to include all variables together in the analysis thanks to its visual modelling capability,

ChatGPT applies the analysis by following the prompts entered sequentially by the users. Although it has successfully calculated basic statistical values such as VIF, path coefficients and R-squared, these steps must be given in a sequential manner and with clear commands. This means that ChatGPT may be more efficient for researchers with a basic grasp of statistical logic than for beginners who are not familiar with the current analysis method.

Considering the practical contributions of the current study, researchers who have difficulty in accessing licensed software packages such as SmartPLS may consider ChatGPT as a more accessible and less costly alternative, especially when working with Python-based statistical programs. ChatGPT's ability to follow prompts, process data and apply advanced statistical analyses can increase efficiency and save time in the research process. Therefore, ChatGPT should go beyond being an assistant for empirical processes compatible with expanding areas of use.

Although the present study is based on ChatGPT in large language models (LLM), future studies can include other artificial intelligence tools such as Google Gemini or Microsoft Bing Chat for similar statistical analysis tasks and compare the agreement between large language models. In addition, since ChatGPT currently lacks visualization capabilities, future research may explore whether these alternative AI tools can overcome this limitation by providing graphical outputs (e.g., path diagrams or model fit visuals). Highlighting such tools as potential solutions could further enhance the practical usability of AI-based statistical analyses. In conclusion, the present study has shown that ChatGPT can be a reliable and useful tool for PLS-SEM analyses. Although not as extensive as full-featured statistical software, ChatGPT can be considered as a promising complement to existing analysis techniques in the digital age due to its capacity to provide valid and similar results.

8. LIMITATIONS AND SUGGESTIONS FOR FUTURE STUDIES

Although the present study provides valuable insights into the use of ChatGPT as a statistical analysis tool, it is not without limitations. First, the analysis conducted with ChatGPT required manual prompting for each step of the process. Unlike specialized statistical software such as SmartPLS, ChatGPT does not support automated or integrated workflows. This limitation may pose challenges for novice researchers who are unfamiliar with statistical procedures or Python-based programming environments. Second, ChatGPT lacks built-in visualization capabilities for modeling results. While numerical outputs were successfully generated and interpreted, the absence of graphical representations such as path diagrams or model fit visuals may hinder the interpretability of complex models, especially for visual learners or stakeholders who rely on visual summaries. Third, the current study relied on a single dataset and

focused solely on the PLS-SEM method. Future research should therefore test ChatGPT's performance across different datasets with varying sample sizes and characteristics. Such replications will be crucial to evaluate the robustness, generalizability, and stability of ChatGPT's outputs in PLS-SEM analyses. Although the results were highly consistent between ChatGPT and SmartPLS, the generalizability of these findings to other statistical methods (e.g., confirmatory factor analysis, mediation-moderation analysis in CB-SEM, or logistic regression) remains untested. Moreover, the analysis assumes that the researcher has prior knowledge of both the theoretical model and the statistical prompts needed to guide ChatGPT effectively. Without adequate prompting, ChatGPT may misinterpret the task or provide incomplete results. Furthermore, potential biases and errors inherent in AI-generated outputs should be acknowledged. Since ChatGPT is trained on large-scale datasets, it may reflect or amplify existing biases in the training data, which can influence statistical reasoning and interpretation. Although no systematic biases were observed in the present study, future research should more explicitly assess whether AI-based analyses introduce distortions that could affect validity and reliability. Another critical consideration is reproducibility. Because ChatGPT is non-deterministic, identical prompts may sometimes yield slightly different outputs. In this study, reproducibility checks showed that the results remained highly stable across multiple runs. However, reproducibility challenges remain an inherent concern, and future studies should systematically test repeated analyses, different datasets, and varied conditions to evaluate the robustness of AI-assisted outputs. On the other hand, for future research, it is recommended that similar studies be conducted using alternative statistical techniques such as exploratory factor analysis, multivariate regression, or hierarchical linear modeling. This would help determine whether ChatGPT's statistical reasoning is equally robust across a variety of quantitative methods. In addition, due to the non-deterministic nature of ChatGPT, identical prompts may occasionally produce slight variations in the outputs. While our checks indicated that the statistical results remained highly consistent, this characteristic should be considered when evaluating the reliability of AI-based analyses. Additionally, testing the performance of other AI-powered language models, such as Google Gemini or Microsoft Bing Chat, may offer comparative insights into the reliability and utility of various platforms. Researchers are also encouraged to explore the integration of ChatGPT with coding environments (e.g., Jupyter Notebooks or Google Colab) to automate sequential statistical analyses. Such integration could enhance usability and improve efficiency in research workflows. Lastly, future studies could focus on developing standardized prompt templates or user-friendly interfaces to make advanced statistical analysis more accessible to a wider range of users. Moreover, considering the current visualization limitation of

ChatGPT, researchers are encouraged to examine whether alternative AI-based tools (e.g., Google Gemini, Microsoft Bing Chat) can provide effective graphical outputs such as path diagrams or model fit visuals. Including such tools in future studies may help overcome this constraint and broaden the applicability of AI-powered statistical analyses.

AUTHORS' CONTRIBUTIONS

Yavuz TORAMAN: Contributed to the design, manufacturing, and testing of the machine presented in this study and participated in the writing of the manuscript.

Orçun Muhammet ŞİMŞEK: Contributed to the design, manufacturing, and testing of the machine presented in this study and participated in the writing of the manuscript.

DECLARATION OF ETHICAL STANDARDS

Ethical committee reviews were conducted by the Istanbul Nisantasi University (reference number: 2024/16).

CONFLICT OF INTEREST

There is no conflict of interest in this study.

REFERENCES

- [1] T. Hayes and S. Usami, 'Factor Score Regression in the Presence of Correlated Unique Factors', *Educational and Psychological Measurement*, 80(1): 5–40, Feb. (2020).
- [2] J. Miles, 'Structural Equation Modelling', *Historical Social Research*, vol. 23, p. 159171, (1998).
- [3] C. M. Musil, S. L. Jones, and C. D. Warner, 'Structural equation modeling and its relationship to multiple regression and factor analysis', *Res. Nurs. Health*, 21(3):271–281, (1998).
- [4] J. F. Hair, G. T. M. Hult, C. M. Ringle, M. Sarstedt, N. P. Danks, and S. Ray, 'An Introduction to Structural Equation Modeling', in *Partial Least Squares Structural Equation Modeling (PLS-SEM) Using R*, in Classroom Companion: Business, Cham: Springer International Publishing, pp. 1–29., (2021).
- [5] G. Cepeda-Carrion, J.-G. Cegarra-Navarro, and V. Cillo, 'Tips to use partial least squares structural equation modelling (PLS-SEM) in knowledge management', *JKM*, 23(1):67–89, (2024).
- [6] A. Moosavi and R. bt Juhari, 'SPSS vs. AMOS: A Comparative Approach to Analyzing Moderators in Behavioral Research', *IJSSHR*, 7(8):6159–6164, Aug. (2024).
- [7] B. K. Sarker, D. K. Sarker, S. R. Shaha, D. Saha, and S. Sarker, 'Why Apply SPSS, SmartPLS and AMOS: An Essential Quantitative Data Analysis Tool for Business and Social Science Research Investigations', *IJRISS*, vol. VIII, no. IX, pp. 2688–2699, (2024).
- [8] M. Sarstedt, C. M. Ringle, and J. F. Hair, 'Partial Least Squares Structural Equation Modeling', in *Handbook of Market Research*, C. Homburg, M. Klarmann, and A. Vomberg, Eds., Cham: Springer International Publishing, pp. 1–40., (2017).
- [9] H. Wold, 'Path Models with Latent Variables: The NIPALS Approach', in *Quantitative Sociology*, Elsevier, pp. 307–357., (1975).
- [10] M. Sarstedt, 'Der Knacks and a Silver Bullet', in *The Great Facilitator*, B. J. Babin and M. Sarstedt, Eds., Cham: Springer International Publishing, pp. 155–164., (2019).
- [11] J. F. Hair, G. T. M. Hult, C. M. Ringle, and M. Sarstedt, *A primer on partial least squares structural equation modeling (PLS-SEM)*, Second edition. Los Angeles London New Delhi Singapore Washington DC Melbourne: SAGE, (2017).
- [12] Alnıpak, Serdar, and Yavuz Toraman, "“I am ChatGPT. would you accept me to assist you in logistics management?” An Empirical Study From Türkiye." *Logforum* 21.3, 11., (2025).
- [13] J. Ma, P. Wang, B. Li, T. Wang, X. S. Pang, and D. Wang, 'Exploring User Adoption of ChatGPT: A Technology Acceptance Model Perspective', *International Journal of Human-Computer Interaction*, 41(2):1431–1445, (2025).
- [14] J. McCarthy, M. L. Minsky, N. Rochester, and C. E. Shannon, 'A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence, august 31, 1955', *AI Magazine*, 27(4):12–12, (2006).
- [15] K. Uludag, 'The Use of AI-Supported Chatbot in Psychology', in *Advances in Medical Technologies and Clinical Practice*, K. Uludag and N. Ahmad, Eds., IGI Global, pp. 1–20., (2024).
- [16] M. Cascella, J. Montomoli, V. Bellini, and E. Bignami, 'Evaluating the Feasibility of ChatGPT in Healthcare: An Analysis of Multiple Clinical and Research Scenarios', *J Med Syst*, 47(1):33, (2023).
- [17] S. Akter, S. Fosso Wamba, and S. Dewan, 'Why PLS-SEM is suitable for complex modelling? An empirical illustration in big data analytics quality', *Production Planning & Control*, 28(11-12):1011–1021, (2017).
- [18] J. Hair, C. L. Hollingsworth, A. B. Randolph, and A. Y. L. Chong, 'An updated and expanded assessment of PLS-SEM in information systems research', *IMDS*, 117(3):442–458, (2017).
- [19] D. S. Kumar and K. Purani, 'Model specification issues in PLS-SEM: Illustrating linear and non-linear models in hospitality services context', *JHTT*, 9(3):338–353, (2018).
- [20] M. Rönkkö and J. Evermann, 'A Critical Examination of Common Beliefs About Partial Least Squares Path Modeling', *Organizational Research Methods*, 16(3):425–448, (2013).
- [21] R. Calderon Jr., G. Kim, C. Ratsameemonthon, and S. Pupanead, 'Assessing the Adaptation of a Thai Version of the Ryff Scales of Psychological Well-Being: A PLS-SEM Approach', *PSYCH*, 11(7):1037–1053, (2020).
- [22] E. W. L. Cheng, S. K. W. Chu, and C. S. M. Ma, 'Students' intentions to use PBWorks: a factor-based PLS-SEM approach', *ILS*, 120(7/8):489–504, Jul. (2019).
- [23] E. Erwin, Y. K. M. Suade, and N. Alam, 'Social Media Micro-enterprise: Utilizing Social Media Influencers, Marketing Contents and Viral Marketing Campaigns to Increase Customer Engagement', in *Proceedings of the International Conference of Economics, Business, and Entrepreneur (ICEBE 2022)*, vol. 241, Nairobi, Yuliansyah, H. Jimad, R. Perdana, G. E. Putrawan, and T. Y. Septiawan, Eds., in Advances in Economics, Business

- and Management Research, vol. 241., Paris: Atlantis Press SARL, pp. 578–593., (2023).
- [24] G. T. M. Hult, J. F. Hair, D. Proksch, M. Sarstedt, A. Pinkwart, and C. M. Ringle, 'Addressing Endogeneity in International Marketing Applications of Partial Least Squares Structural Equation Modeling', *Journal of International Marketing*, 26(3):1–21, Sep. (2018).
- [25] J. Riou, H. Guyon, and B. Falissard, 'An introduction to the partial least squares approach to structural equation modelling: a method for exploratory psychiatric research', *Int J Methods Psych Res*, 25(3): 220–231, (2016).
- [26] J. Henseler, C. M. Ringle, and R. R. Sinkovics, 'The use of partial least squares path modeling in international marketing', in *Advances in International Marketing*, vol. 20, R. R. Sinkovics and P. N. Ghauri, Eds., Emerald Group Publishing Limited, 277–319., (2009).
- [27] K. Keskinbora and F. Güven, 'Artificial Intelligence and Ophthalmology', *tjo*, 50(1):37–43, (2020).
- [28] Y. K. Dwivedi *et al.*, 'Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy', *International Journal of Information Management*, vol. 57, p. 101994, (2021).
- [29] H. Gil De Zúñiga, M. Goyanes, and T. Durotoye, 'A Scholarly Definition of Artificial Intelligence (AI): Advancing AI as a Conceptual Framework in Communication Research', *Political Communication*, 41(2):317–334, (2024),
- [30] B. G. Tabachnick and L. S. Fidell, *Using multivariate statistics*, 6. ed., International ed. in Always learning. Boston Munich: Pearson, (2013).
- [31] KOÇ USTALI, Nesrin, et al. "Küresel Risk Yönetim İndeksi Değerlendirmesi: Gri Tabanlı Topsis Yöntemi Uygulaması." *Journal of Polytechnic* 27(5), (2024).
- [32] Toraman Y. Space Logistics and Risks: A Study on Spacecraft. *Politeknik Dergisi.*, 28(2):573-584., (2025).
- [33] De Winter, Joost CF, Samuel D. Gosling, and Jeff Potter. "Comparing the Pearson and Spearman correlation coefficients across distributions and sample sizes: A tutorial using simulations and empirical data." *Psychological methods* 21.3, 273., (2016).
- [34] Benesty, Jacob, et al. "Pearson correlation coefficient." *Noise reduction in speech processing*. Berlin, Heidelberg: Springer Berlin Heidelberg, 1-4., (2009).
- [35] Willmott, Cort J., and Kenji Matsuura. "Advantages of the mean absolute error (MAE) over the root mean square error (RMSE) in assessing average model performance." *Climate research* 30.1, 79-82., (2005).
- [36] Chai, Tianfeng, and Roland R. Draxler. "Root mean square error (RMSE) or mean absolute error (MAE)?—Arguments against avoiding RMSE in the literature." *Geoscientific model development* 7.3 1247-1250., (2014).