

Addressing missing values in structural equation modeling with application to digital transformation survey data and its role in improving the quality of healthcare in Mosul

Safa Ahmad Jumaa, Omar Salim Ibrahim*

Department of Statistics and Informatics, College of Computer Science and Mathematics, University of Mosul, Iraq

Abstract The presence of missing data is a pivotal methodological challenge across various scientific disciplines, particularly in structural equation modeling (SEM), due to its direct impact on the accuracy of statistical estimation and the quality of model fit. This study aimed to evaluate and compare the performance of three common missing data handling methods: full information (FIML), pairwise deletion (PD), and total deletion (LD), within the framework of SEM using questionnaire data on digital transformation and its role in improving the quality of sustainable healthcare in Mosul.

The study employed a descriptive-analytical approach with a sample of 303 individuals. SEM analysis was performed using R software, relying on three estimation methods: maximum likelihood (ML), unweighted least squares corrected by mean and variance (ULSMV), and weighted least squares corrected by mean and variance (WLSMV). The effect of each processing method on estimation accuracy, fit quality, and statistical efficiency was evaluated.

To support the applied findings, a systematic simulation study was conducted comparing the performance of pairwise and whole deletion with the ULSMV and WLSMV estimation methods under sample sizes of 200, 400, 600, 800, 1000, and 1500 and loss ratios of 0.01 and 0.1, respectively, on a four-factor model, each factor comprising five observed variables. Standardized fit indices were used for evaluation. The results showed that the FIML method outperformed most evaluation criteria in terms of statistical accuracy, fit quality, and stability. The simulation also demonstrated that pairwise deletion outperformed whole deletion at low and medium loss ratios, with relative improvement as sample size increased. ULSMV and WLSMV showed similar performance at low loss, with ULSMV tending towards better stability in smaller samples. These results underscore the need to adopt advanced methods for handling missing data and avoid exclusive reliance on traditional deletion methods in SEM applications.

Keywords Structural Equation Modeling, Missing Data Handling, Elimination Methods, Whole Information Method, Digital Transformation in Healthcare

DOI: 10.19139/soic-2310-5070-4064

1. Introduction

The problem of missing data is one of the most important methodological challenges in quantitative research, as it directly affects the quality of statistical results and the validity of scientific conclusions, especially in studies that rely on questionnaires and standardized data collection tools. Research has confirmed that missing data, if left unaddressed or handled inappropriately, can lead to biased estimates and a loss of analytical power [5, 13]. In the context of structural equation modeling (SEM), which is used to analyze complex relationships between observed and latent variables, handling missing data becomes more important, as SEM models rely on multiple estimation of parameters and their relationships. Therefore, poor handling of missing values may weaken the reliability of results and indices of conformity [13]. Traditional methods for handling missing data include listwise deletion and pairwise deletion. While listwise deletion provides simplicity of application by deleting any case

*Correspondence to: Omar Salim Ibrahim (Email: omarsalim85@uomosul.edu.iq). Department of Statistics and Informatics, College of Computer Science and Mathematics, University of Mosul, Iraq.

containing a missing value, it reduces the sample size and may skew estimates if the missing data is not completely random. Pairwise deletion, on the other hand, allows the use of the largest amount of available data for each pair of variables, but it may lead to inconsistent estimates across different calculations [5]. Among the most advanced methods employed in SEM environments is the Full Information Maximum Probability (FIML) method, which estimates parameters using all available information without the need to eliminate entire cases. This preserves sample size and reduces bias compared to traditional elimination methods [8]. Recent studies have shown that probability-based methods like FIML often provide more accurate and stable estimates in SEM models, especially at medium to high percentages of missing values, compared to traditional elimination methods [13]. Based on this theoretical and methodological framework, this study aims to review and compare methods for handling missing data in structural equation modeling and apply them to data from a five-point Likert scale survey on improving sustainable healthcare. The goal is to evaluate the impact of each method choice on the accuracy of estimates, the quality of fit, and the strength of the scientific inferences.

2. Research Problem and Objective

Studies based on questionnaires face a common problem: missing data that affects the accuracy of statistical estimates and the validity of Structural Equation Modeling (SEM) results. Inappropriate processing of these missing values can lead to biased results and weak fit indicators. Since the current study relies on questionnaire data about digital transformation and its role in improving sustainable healthcare, collected from a sample of 303 individuals, missing values emerged that require systematic processing. The research problem lies in identifying the most appropriate method for processing the missing data and demonstrating the impact of different processing methods on the results of the structural model and its fit indicators. The study aims to diagnose the nature and mechanism of missing data in structural equation modeling, apply several processing methods, and compare their impact on the model's estimates and fit indicators. It also seeks to demonstrate the effect of processing methods on interpreting the relationship between digital transformation and improving sustainable healthcare, ultimately leading to more accurate and reliable results that can be used in applied studies.

3. Missing Data

Missing data represents a fundamental and comprehensive challenge in modern data science, significantly hindering analytical capabilities and decision-making processes across a wide range of disciplines, including healthcare, bioinformatics, social sciences, e-commerce, and industrial control systems [6]. The reasons for data loss can be summarized as follows: respondents refuse to answer certain items because they relate to personal issues or are perceived as threatening; because they address sensitive topics; because the content may be inappropriate or irrelevant to the respondents; because the scale is too long, causing fatigue or boredom; due to time constraints; loss during data entry after collection or during coding; due to wars and disasters; high costs; negligence; or a combination of these factors [7, 8]. It is essential to identify the causes of data loss through patterns and mechanisms to select the appropriate method for addressing this missing data [8].

4. Loss Mechanisms

Due to the incompleteness of the data, the statistical methods for analyzing this data also vary, as does the mechanism leading to data loss. Understanding this mechanism and identifying its nature greatly assists in selecting the appropriate analytical method. It even serves as the starting point for diagnosing the method whose results are closest to optimum for the studied data. The relationship between data loss for a specific variable and the values of that variable itself, or the values of other variables, is as follows [9]:

The loss of X values is independent of the values of other variables and of the missing values themselves.

The loss of X values depends on the value itself.

The missing X values depend on the values of other variables in the sample.

5. Types of Data Loss Mechanisms

5.1. Random Loss (MAR)

In this pattern, the lack of response on an item where the loss is present is directly or indirectly related to external variables. If the response on an item is missing, and this item measures the latent trait at the scale level as well as the other items, then the reason for the loss in this case is attributed to some external variables.

$$P(Y_{mis}/Y_{com}, X) = P(Y_{mis}/Y_{com}) \quad , \quad \forall X \quad (1)$$

5.2. Complete Random Loss (MCAR)

means that the item is missing due to randomness alone.

$$P(Y_{mis}/Y_{com}, X) = P(Y_{mis}) \quad , \quad \forall X \quad (2)$$

5.3. Non-random loss (MANAR)

means that the loss in the data is related to the measured variable itself. The data loss may be related to either the item's parameters (difficulty in distinguishing, guessing, or the individual's ability).

$$P(Y_{mis}/Y_{com}, X) = P(Y_{mis}/X) \quad , \quad \forall X \quad (3)$$

If the cause of the loss is independent of the missing value and the values of the other variables in the sample, then it can be said that the data is missing completely at random (MCAR). However, if the cause of the loss is related only to the values of the other variables and is independent of the missing value, then the data loss is random (MAR). Sometimes, the cause of the loss is due to the missing value itself and is independent of the values of the other variables. In this case, the data is not missing at random (Not MAR). When analyzing this type of data, a mechanism must be taken Loss is taken into account, but in the case of MCAR and MAR, the distribution can be neglected [10] .

6. Methods of Handling Missing Values

There are several methods that can be used to handle missing values, and these methods can be presented as follows [11].

These methods use missing values to display the data containing the missing values as complete data, but these methods of handling are criticized for often giving biased and inefficient results [13].

Methods Depend on Deletion: These methods involve deleting cases whose data contain missing values. This leads to a reduction in the sample size and the available data for calculating item parameters and the individual ability parameter. Examples include case deletion, the method of considering a missing answer as incorrect, and the method of discarding only missing answers. Case deletion generally addresses data loss by deleting data from individuals with missing values on some scale items, thus reducing the sample size. This, in turn, maximizes the standard error and reduces the significance level [17]. Two procedures are employed within this method:

6.1. Listwise Deletion

This method relies on completely deleting the case—which contains a missing value for one of the variables—from the dataset. Therefore, it is also called Complete Case Analysis. It simply relies on removing incomplete cases from the analysis [13]. In this method, data from any participant whose data contains a missing value for any item on the

measurement instrument is removed, and only the data from participants with complete responses on all items are analyzed. Generally, this method assumes that the missing data are randomly distributed throughout the sample of participants of interest [17].

6.2. *Pairwise Deletion*

also known as Available Case Analysis, is a method of partially deleting some data related to a specific analysis. It utilizes all variables that contain complete data and are related to a specific analysis. This differs from total deletion, which relies on removing all missing data from the database [13]. Both deletion methods, whether at the list level or the pair level, are effective when the missing data follows the MCAR hypothesis, where the loss is not related to observed or unobserved variables. In MCAR scenarios, the complete data can be considered a simple random sample of the original target population. However, when missing data follow specific loss mechanisms, such as MNAR, deletion methods can lead to biased parameter estimates and compromise the validity of the analysis. For example, suppose all missing data are from high-income groups, and these cases are deleted. In this case, the resulting dataset will be limited to representing low-income groups, introducing bias and potentially misleading results. Furthermore, using deletion methods reduces the sufficient sample size, which can lead to decreased statistical power and accuracy.

It is essential to exercise caution when using deletion methods and consider their limitations. These methods neglect valuable information in the missing values, potentially missing important insights. Alternative approaches, such as assignment methods, can be explored to more effectively handle missing data and mitigate the limitations associated with deletion methods [18, 41].

6.3. *Maximum Probability of Complete Information (FIML)*

Another advanced method for missing data analysis is Maximum Probability of Complete Information. In this method, missing values are not replaced or completed, but rather the missing data is handled within the analysis model. The model is estimated using the Maximum Probability of Complete Information method, and in this way, all available information is used to estimate the model. In FIML, population parameters are estimated based on the probability of yielding estimates from the sample data being analyzed [19].

Several traditional statistical methods exist for analyzing missing data, and the Maximum Probability of Complete Information (FIML) is one of the most widely used tools [22, 27]. FIML has been implemented in virtually all structural equation modeling software [21, 35], due to two important properties: its ability to produce consistent parameter estimates and the efficiency of these estimates. However, these properties may not apply in cases of randomly missing data [23]. The full-information maximum likelihood method uses all available data, including partially missing or fully observed data, to produce parameter estimates that maximize the probability function [24]. Given a set of parameters, the probability function essentially represents how well the statistical model fits the observed data. Studies have shown that the full information maximum likelihood (FIML) method can provide unbiased estimates for normally distributed data under the partial missing data mechanism [24].

The general maximization function is given as follows:

The key advantage of FIML is that it utilizes all missing data rather than deleting or compensating for them. This makes parameter estimates more efficient and less biased, especially in regression models or structural equation models [25]. Examples of models frequently estimated using FIML include structural equation models, multilevel models, and growth models. Multiple interpolation and maximizing the complete information yield similar results when outcome data are missing and the same information is incorporated into the multiple interpolation model as in maximizing the complete information [19]. When data is missing, FIML:

Determines the likeness function for each case based on the variables available in that case. Maximizes the log likelihood of parameters at the level of the complete and available data only. Estimates typical parameters (such as means, skews, and ranges) in a way that maximizes the overall probability without compensating for missing values. This means that FIML does not explicitly fill in missing values but rather uses the available information roughly from each record in your data [25].

7. Structural Equation Modeling (SEM)

Structural equation modeling is a statistical test that helps evaluate a set of regression equations simultaneously. Its aim is to explore the relationships between one or more independent variables and one or more dependent variables. SEM is more complex than traditional statistical analysis models such as regression. It combines various computer algorithms, mathematical models, and statistical methods applied to a dataset [1].

The primary goal of SEM is to test causal relationships between multiple variables within a single, integrated model, rather than analyzing the relationship of each variable individually. Instead of using simple regression analysis or individual statistical analyses, SEM estimates a set of equations simultaneously, allowing for the analysis of both direct and indirect effects between variables, including latent (unobserved) and measured (observed) variables. Another distinguishing feature of this approach is its ability to test complex theoretical hypotheses built upon robust conceptual frameworks. This makes it suitable for fields such as education, psychology, and the social sciences in general, where variables overlap and their relationships are intertwined [6]. Furthermore, SEM boasts several statistical and methodological advantages over traditional methods like regression and discrete factor analysis. These include:

The ability to estimate multiple and interrelated relationships between variables simultaneously, rather than analyzing the relationship between each pair of variables separately. The integration of latent variables—variables not directly observed—into the model, with their measurement measured via indicators (questions/measurement items). The separation of the measurement model from the construct model, allowing for the evaluation of the quality of the measures (their validity and reliability) before estimating the relationships between constructs. The precise handling of standard errors for each indicator, which reduces bias in the results compared to simpler methods. It enables the researcher to test theoretical hypotheses with direct practical application using advanced statistical software such as AMOS, LISREL, SmartPLS and R. [15]

8. Methods for Estimating Structural Equation Model Parameters

In statistical modeling, the term "estimator" refers to the method used to obtain an estimate of an unknown parameter. Although many estimators are available for researchers to choose from when performing confirmatory factor analysis (CFA), the common goal of all estimators is to obtain a set of parameter estimates that minimizes the variance between observed and reproduced relationships as a function of those estimates [2]. Reducing variance based on the final set of parameter estimates forms the basis for obtaining model-fit statistics to evaluate confirmatory factor analysis models [2]. Inaccurate or biased estimators can lead to erroneous decisions about the assumed factor model [8]. Commonly used estimators in confirmatory factor analysis can be categorized into different "sets" of estimation techniques. These three sets include maximum likelihood estimators, least squares estimators, and Bayesian estimators [6, 7, 40, 42]. Maximum likelihood and least squares estimators, along with their robust variants, are based on traditional iterative probability statistics (such as the relative frequency in many trials or the long-term repetition), and are therefore also known as iterative estimators [5].

The choice of estimator depends heavily on the statistical assumptions associated with the estimator class [6]. Normalized estimators (such as the maximum likelihood method and the generalized least squares method) share common assumptions that must be met to ensure accurate results [5]. These assumptions include:

a large sample size; the independence of the observations; the measurement of observed variables on a continuous scale and their following a multivariate normal distribution; the correct selection of the confirmatory factor analysis (CFA) model underlying the analysis of the observed variables [6, 8]. Natural estimators tend to yield satisfactory results when these assumptions are met, but they may perform poorly when these assumptions are violated [7]. In response to this limitation, powerful estimators such as MLR, ULSM, ULSCMV, WLSM, and WLSMCV have been developed to address violations of distributional assumptions related to the data. They typically improve performance when the assumptions are not fully met [8].

8.1. Maximum Likelihood Family (ML) Estimation

The Maximum Likelihood Family (ML) estimator is the most common and widely used natural estimator in factor analysis. The parameter estimates produced by an ML estimator represent the most probable values that explain the observed relationship between indicators [5]. It provides unbiased estimates (on For example, the expected value of a statistic equals the true value).

Maximum likelihood (ML) estimates are accurate, consistent (e.g., sampling error decreases with increasing sample size), and efficient (e.g., have the lowest sampling error compared to other estimates) when statistical assumptions are met [4, 5, 6]. However, ML estimates tend to produce inflated test statistics and low standard errors of parameter estimates in the case of non-normalized data [6, 7, 13]. these biases can lead to erroneous data-based decisions.

Social science researchers frequently encounter data measured on an ordinal scale (e.g., Likert responses). Ordinal data are discrete in nature and do not follow a normal distribution, which violates two fundamental assumptions necessary for normalized estimators to produce reliable parameter estimates and appropriate statistics [50]. Applying normality estimators to confirmatory factor analysis using ordinal data is possible, but requires caution. Studies have shown that normality theory estimators (such as the maximum likelihood method and the generalized least squares method) are only suitable for analyzing ordinal data with five or more response classes, where the assumptions of multivariable ordinal data, whether normal or continuous, can be approximated [2, 4, 8]. Normality estimators often produce unreliable results when the sample size is small or the number of ordinal classes is less than five [8, 9].

8.2. Unweighted Least Squares (ULSMV)

Unweighted Least Squares (ULS) is a form of Ordinary Least Squares (OLS) used in structural equation modeling. It minimizes the sum of squared differences between the variance-covariance matrix of the sample and that estimated by the model. It is advantageous for providing unbiased estimates in random samples and does not require a strict distributional assumption or a specific positive variance matrix, as in Maximum Likelihood Estimation (ML). However, it assumes that the variables are on the same scale [2, 7]. Nevertheless, the conventional test statistic in ULS does not follow a chi-squared asymptotic distribution when the model is correct, necessitating the use of strong corrections. Furthermore, neglecting strong corrections can underestimate standard errors and affect the accuracy of statistical inference [2, 8].

8.3. Weighted Least Squares (WLSMV) Method

The Weighted Least Squares (WLS) method in structural equation modeling (SEM) relies on descriptive statistics, such as multi-class correlation matrices, instead of raw data, unlike WLS in the general linear model framework. However, using a full WLS requires estimating and inverting the full asymptotic covariance matrix, making it impractical for large samples or when dealing with ordered categorical data. Therefore, the Diagonal Weighted Least Squares (DWLS) method was developed to reduce these limitations. It uses only the diagonal elements of the weight matrix, thus simplifying the computational requirements and making it suitable for ordinal data. It can also be combined with robustness corrections, such as Satorra-Bentler scaling, to produce robust estimates and corrected standard errors [5, 9, 13, 39].

9. Conformity Quality Indicators

9.1. Standardized Root Mean Residual (SRMR)

The standardized root mean residual (SRMR) is an absolute measure of fit, defined as the standard difference between the observed correlation and the expected correlation. It is a measure with a positive bias, and this bias is greater in studies with small N-values and low df scores. Because the root mean residual (RMR) is an absolute measure of fit, a value of zero indicates a perfect fit. The RRM does not affect the complexity of the model. A value less than 0.08 is generally considered a good fit [7].

9.2. Root Mean Square Residual Correlation (CRM)

CRM is another standard measure of total mismatch, focusing on correlation residuals rather than covariance, making it less affected by measure variations. Studies recommend using standard measures such as CRM and SRM for their ease of interpretation and comparison between models [4, 6, 2].

9.3. Expected Cross-Value Verification Index (ECVI)

Michael W. Browne and Robert Cudeck developed the Expected Cross-Value Verification Index (ECVI) to measure a model's ability to generalize. A new sample. The ECVI is similar to the Akaike Information Criterion (AIC) in its logic of balancing fit quality and model complexity, but it focuses on estimating the model's suitability for future samples. Its value is calculated based on the minimum fit function, the number of free parameters, and the sample size, with smaller values indicating a better model. When using the maximum likelihood fit function (ML), both the ECVI and AIC provide a similar ranking of models, so reporting only one is sufficient when selecting the optimal model [5, 13, 7].

9.4. Comparative Fit Index (CFI)

Its value ranges from zero to one, with an acceptable value of 0.90. The closer to one, the better, and a value below 0.90 indicates a model modification. This index indicates an improvement in the fit of the assumed theoretical model compared to the baseline model, meaning that all correlations between variables are zero, or that the variables are independent of each other.

$$CFI = 1 - \max[(x_t^2 - df_t), 0] / \max[(x_t^2 - df_t), (x_b^2 - df_b), 0] \quad (4)$$

x_t^2 For the assumed model, df_t the degree of freedom of the assumed model, and x_b^2 the chi-squared value of the underlying theoretical model, and df_b the degree of freedom of the underlying theoretical model. Kline, R. B. (2011) [5]

9.5. Tucker-Lewis Index (TLI) or Non-Normed Fit Index (NFI)

There is another index called the Tucker-Lewis Non-Normed Fit Index (NFI), sometimes referred to as the Non-Normed Fit Index. This index, in addition to the logic of comparison with a baseline model (the independent model or null model), involves a penalty function. This penalty function is applied when the model is complicated by adding freely estimable parameters to the assumed model without any benefit, i.e., without this addition leading to any improvement in the level of fit to the assumed model. This compensates for the effect of the assumed model's complexity. The equation for the Tucker-Lewis Index (TLI) appears as follows:

$$RMSEA = \sqrt{\frac{\widehat{F}_0}{d}} \quad , \quad 0 \leq RMSEA \leq 1 \quad (5)$$

The $\widehat{F}_0$ all model represents the chi-squared value of the independent or null model. The default model represents the degrees of freedom for the independent model, the Xurget model represents the chi-squared value for the assumed model, and the dfurget model represents the degrees of freedom for the assumed model.

While the CFI index has a defined range, the TLI index lacks a defined range of values or criteria, meaning some of its values fall outside the range of zero to one. Therefore, it is non-standard. However, its interpretation follows the same pattern as the CFI index; that is, TLI values exceeding 0.90 indicate a reasonable fit to the research model [5].

$$GFI = 1 - \frac{tr[w^{-1/2}(s - \widehat{\sum}_0)w^{-1/2}][w^{-1/2}(s - \widehat{\sum}_0)w^{-1/2}]}{tr[w^{-1/2}(s)w^{-1/2}][w^{-1/2}(s)w^{-1/2}]} \quad (6)$$

d_b & $\widehat{C}_b$ the degrees of freedom for the underlying model, while C and D represent the degrees of freedom for the assumed model.

9.6. Root Mean Square of Approach Error (RMSEA) Index

This index incorporates a model complexity correction to examine discrepancies, differences, or inconsistencies. The expected value in the study population from which the study sample was taken is determined using the correlation matrix. This indicator measures the degree of fit or inadequacy of the theoretical model's correlation matrix with the study population. The closer the value of this indicator is to zero, the better and more optimal the fit of the population as measured by the correlation matrix between the model and the data. A value greater than zero indicates a decrease in fit. A value of 0.005 indicates a high and good fit. Values between 0.08 and 0.009 are acceptable. However, a value of 1 indicates a discrepancy between the model and the data, necessitating model modification [2, 38, 43].

$$RMSEA = \sqrt{\frac{\widehat{F}_0}{d}}, \quad 0 \leq RMSEA \leq 1 \quad (7)$$

the $\widehat{F}_0$ variance function obtained from the population data for the structural model, rather than sampling from it, represents the value of the variance function obtained from the population data, and d represents the degree of freedom of the structural model.

9.7. The Goodness of Fit (GFI)

Index measures the amount of variance in the analyzed matrix by the model the researcher assumes to explain it. It is analogous to the multiple correlation square Eq. (2) in multiple regression analysis, and its value ranges from 0 to 1. A higher value within this range indicates a better fit between the model and the sample data. In the current factor analysis, its value was high, and this index is considered good fit [8]

$$GFI = 1 - \frac{tr[w^{-1/2}(s - \widehat{\Sigma}_0)w^{-1/2}][w^{-1/2}(s - \widehat{\Sigma}_0)w^{-1/2}]}{tr[w^{-1/2}(s)w^{-1/2}][w^{-1/2}(s)w^{-1/2}]} \quad (8)$$

S represents the variance and covariance matrix of the assumed model, Σ the variance and covariance matrix of the base model, and w represents the transition weight matrix resulting from the sum of the elements of matrices w and s .

9.8. Chi-square

The Chi-square statistic is generally defined as the sum of the squared deviations of independent values from the repeated values. In the context of structural analysis (SEM), it is applied to study the relationship between two variables to determine whether or not there is a relationship between them [5].

9.9. Adjusted Goodness of Fit Index The Adjusted Goodness of Fit Index (AGFI)

is an evaluation indicator that measures the degree of data conformity to the proposed model. Its value ranges from zero to one, and a value of 0.90 or higher indicates that the data conforms to the assumed structural model. The mathematical formula for the Adjusted Goodness of Fit Index (AGFI) is:

$$AGFI = 1 - (1 - GFI) \frac{d_b}{d} \quad (9)$$

the d_b degree of freedom of the original model, and $[d]$ represents the degree of freedom of the assumed model [11].

9.10. Nomed Fit Index (NFI)

The Nomed Fit Index (NFI) is an indicator that shows how well the data fits the assumed structural model. Its value ranges from zero to one. The closer the value is to 0.90 or greater, the more we accept the assumed model. The

mathematical formula for the Normed Fit Index is [44]

$$NFI = 1 - \frac{\hat{C}}{\hat{C}_b} \quad (10)$$

the $\hat{C}$ smallest value of the variance for the assumed model and the $\hat{C}_b$ smallest value of the variance for the underlying model.

9.11. Practical Aspect

The topic of digital transformation and its role in improving the quality of sustainable healthcare is a vital subject in contemporary healthcare management research. Studies have demonstrated that adopting digital transformation in healthcare systems can enhance service efficiency, improve health outcomes, and support institutional sustainability by integrating information and communication technologies into operational processes [5]. Research has also indicated that digital transformation not only improves clinical performance but also contributes to enhancing the overall patient experience, expanding access to healthcare services, and promoting sustainable management practices within healthcare institutions [11]. Digital transformation has become an urgent necessity in light of global developments affecting all aspects of life. There is a growing interest from the state in improving the level of healthcare services for all citizens as a fundamental and essential requirement. This necessitates the continuous development of healthcare services with the aim of health reform. Raising awareness of the importance of digital transformation as a modern technology in healthcare institutions, particularly medical centers, is crucial for maximizing the benefits of its services. Medical centers strive to develop their services, activate e-management, adopt technological methods, and enhance healthcare services to ensure cybersecurity and confidentiality for data and information related to the healthcare sector. The importance of this study stems from its focus on developing healthcare services in health centers. Digital transformation is a vital tool in developing work methods within organizations by relying on modern means, thus improving the level of medical service delivery within medical centers [2]. Digital transformation enables medical centers to transition from traditional to digital systems to develop healthcare services. It utilizes information technology to improve access to information and leverage digital advancements between medical centers and their service beneficiaries. Developing healthcare services through a digital system involves (providing digital security, developing digital skills, implementing digital transformation plans, and establishing the necessary infrastructure). Healthcare Services:

Updating healthcare services in medical centers in light of technological and social changes to better meet patients' health needs and aim to raise the level of healthcare services through (efficiency in providing healthcare services - ease of access to services - entitlement to services - transparency in the distribution of healthcare services). [2] Sustainable development is considered one of the most important issues of concern to countries worldwide, whether developed or developing. Today, development is no longer seen solely as economic growth; attention is now turning to human development efforts, as human beings are the fundamental instrument of all progress in society. Therefore, it is essential to increase attention to the human element [3]. Development is considered the only way for the world to overcome the problems it faces, as human beings are the foundation of the development process. Therefore, it requires the optimal investment of material and human resources and the participation of community members in its various fields to achieve progress and advancement, as well as social welfare for all segments of society [4].

9.12. Studying the Structural Structure of Variable X1

The latent variable (X1) represents the first explanatory variable (first axis) in this study, indicating the digital transformation in healthcare institutions. This refers to the adoption of digital technologies in healthcare service delivery to facilitate processes, improve operational efficiency, and enhance the quality of healthcare. The observed variables associated with it are as follows:

(X11) Our organization possesses modern digital devices and systems that support digital transformation in daily operations.

(X12) All patient data is documented electronically in a secure and organized manner.

Table 1. Distribution of health staff in the questionnaire according to gender, academic qualification, years of service and age

Gender				
Females 38.5%	Males 61.5%			
academic qualification				
High school 5%	diploma 39.9 %	Bachelor's 39.9 %	Master's 9.9 %	PhD 5.3 %
years of service				
1-5 year 34.2 %	6-10 year 14.4 %	11-15 year 21.8 %	16-20 year 13.8 %	20 year or more 15.8 %
Age				
20-29 year 23.4 %	30-39 year 39.9 %	40-49 year 26.4 %	50-59 year 8.3 %	60 year or more 2 %

(X13) An effective electronic linking mechanism exists between departments and other hospitals.

(X14) Our organization implements strict policies to protect patient data from breaches or leaks.

(X15) Staff receive regular training on the use of modern digital systems.

The structural structure of the variables can be represented as follows:

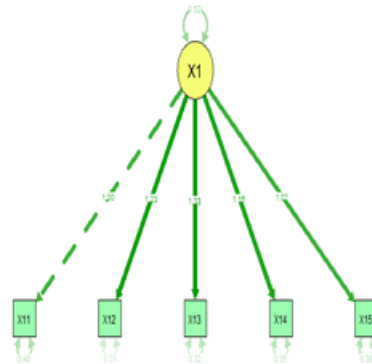


Figure 1. Structural structure of the variable (X1).

9.13. Structural Analysis of the (Yi) Variable

The response variable in this study (the second axis) is the quality of healthcare. This axis includes five observed variables:

(Y11) Digital systems have contributed to accelerating the patient service delivery process.

(Y12) Patient complaints are tracked electronically for rapid resolution.

(Y13) The quality of healthcare services has improved thanks to the use of digital systems.

(Y14) The application of technology has helped reduce medical errors.

(Y15) Patients feel confident in the accuracy of the healthcare procedures provided.

The following figure illustrates the nature of the relationship between the latent variable (Y1) and the observed variables (Yij).

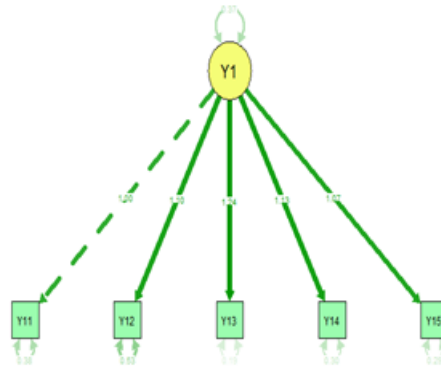


Figure 2. Structural structure of the variable (Y1).

9.14. Studying the structural framework of variable (X2)

which represents health sustainability and is the second explanatory latent variable in this study. Its dependent observed variables are:

- (X21) Digital transformation has contributed to reducing the use of paper and traditional resources.
- (X22) Electronic systems ensure service continuity even in emergencies.
- (X23) The organization strives to spread a culture of digital awareness among employees and beneficiaries.
- (X24) Digital transformation has helped reduce operational costs and improve efficiency.
- (X25) The organization bases its future plans on the principles of sustainable development.

The structural framework of these variables can be illustrated as follows:

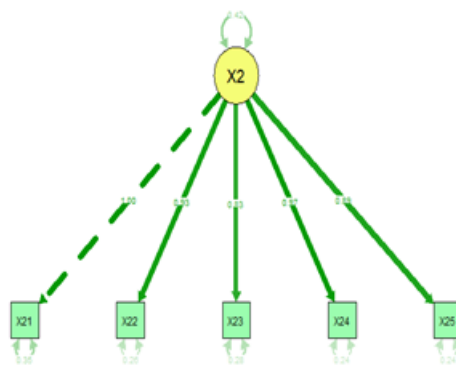


Figure 3. Structural structure of the variable (X2).

9.15. Studying the structural framework of variable (X3)

This variable represents the fourth axis, which is the organizational challenges of digital transformation. The variables indicative of it were as follows:

(X31) Weak technical skills among some employees.

(X32) Resistance to change from traditional staff.

(X33) Insufficient financial or technical support from senior management.

(X34) The need for continuous infrastructure updates.

(X35) The absence of clear legislation to regulate a sound digital transformation.

The relationship between the underlying variable (X3) and the variables (X3i) can be represented as follows:

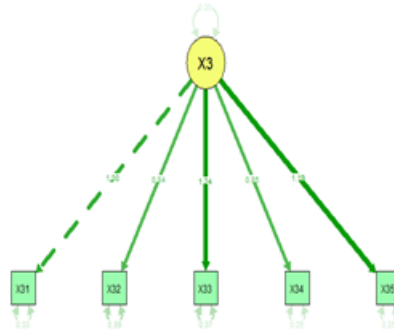


Figure 4. Structural structure of the variable (X3).

Table 2. Measurement errors for variables

Variances	Estimate	Std.Err	P-value
X11	0.339	0.033	0.000
X12	0.230	0.027	0.000
X13	0.301	0.034	0.000
X14	0.355	0.036	0.000
X15	0.368	0.035	0.000
Y11	0.377	0.035	0.000
Y12	0.528	0.049	0.000
Y13	0.236	0.025	0.000
Y14	0.269	0.027	0.000
Y15	0.243	0.025	0.000
X21	0.314	0.031	0.000
X22	0.221	0.023	0.000
X23	0.218	0.023	0.000
X24	0.249	0.025	0.000
X25	0.244	0.024	0.000
X31	0.337	0.034	0.000
X32	0.385	0.036	0.000
X33	0.340	0.039	0.000
X34	0.247	0.025	0.000
X35	0.236	0.030	0.000

The table of measurement errors for the variables shows the estimated values of the error variances for the indicators used in the model. The results indicate that all measurement error variances were statistically significant at the 0.05 level, with p-values for all indicators being 0.000. This confirms the presence of actual measurement errors, not due to statistical chance.

Furthermore, the standard error (Std. Err) values were relatively low, reflecting the accuracy and stability of the model's estimates and indicating the suitability of the data used for analysis. As for the estimated error variance, it fell within acceptable limits for social and human studies, demonstrating that the indicators measure the underlying variables with a reasonable degree of accuracy.

Therefore, it can be concluded that the measurement errors for the indicators did not reach levels that would necessitate the exclusion of any of them. This strengthens the reliability of the econometric model and its suitability for testing research hypotheses.

Table 3. Factor Overloads for Variables

Latent Variables	Estimate	Std.Err	P-value
X11	0.731	0.048	0.000
X12	0.859	0.047	0.000
X13	0.934	0.053	0.000
X14	0.822	0.051	0.000
X15	0.729	0.049	0.000
Y11	0.253	0.028	0.000
Y12	0.280	0.032	0.000
Y13	0.280	0.029	0.000
Y14	0.283	0.029	0.000
Y15	0.272	0.028	0.000
X21	0.614	0.044	0.000
X22	0.605	0.040	0.000
X23	0.593	0.039	0.000
X24	0.578	0.040	0.000
X25	0.534	0.039	0.000
X31	0.449	0.045	0.000
X32	0.366	0.045	0.000
X33	0.553	0.048	0.000
X34	0.374	0.038	0.000
X35	0.520	0.042	0.000

The factor loading table shows the values of the factor loads for the indicators on the underlying variables in the standard model. The results indicate that all factor loadings were statistically significant at the 0.05 level, with p-values of 0.000 for all indicators, demonstrating a substantial relationship between the indicators and their respective underlying variables.

The factor loadings ranged from 0.253 to 0.934. The indicators for the first variable (X11–X15) showed relatively high loadings, reflecting their strong contribution to measuring the underlying variable and indicating a high degree of convergent validity. The indicators for the second and third variables showed loadings ranging from moderate to statistically acceptable, which are within the limits accepted in social and humanistic studies.

The standard error (Std. Err) was relatively low, indicating the accuracy of the factor loading estimation and the stability of the model. Therefore, it can be concluded that all indicators demonstrated sufficient ability to represent the underlying variables, and no indicators exhibited weak overloads warranting exclusion.

Table 4. Correlation between underlying variables

Covariances	Estimate	Std.Err	P-value
X1 ~~			
X2	0.607	0.045	0.000
X3	0.089	0.070	0.204
X2 ~~			
X3	0.238	0.069	0.001

Descriptive statistics were also performed on the questionnaire items, revealing that the arithmetic means ranged between 1.66 and 2.32, indicating that the response level of the sample was within the average range. Furthermore, the standard deviation values showed a relatively low level, reflecting the homogeneity of the sample responses.

Table 5. Study of the regression relationship (Y) on (Xi) for the structural model

Regressions Y1 ~	Estimate	Std.Err	P-value
X1	0.860	0.147	0.000
X2	1.517	0.218	0.000
X3	0.118-	0.106	0.266

Additionally, the skewness and kurtosis values indicated that all items fell within statistically acceptable limits, suggesting that the data follow a distribution close to normality. Therefore, the data are suitable for subsequent statistical analyses such as CFA and SEM.

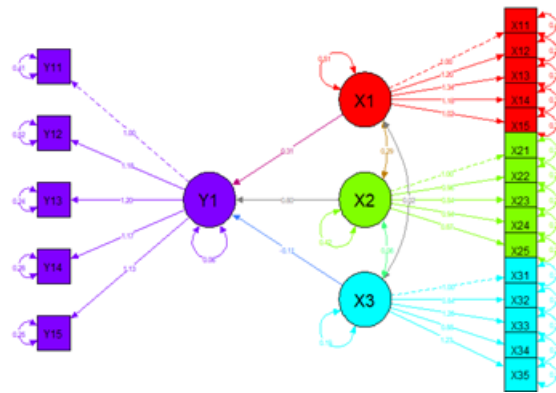


Figure 5. The structural model for studying hypotheses.

Table 6. Survey results without processing methods and with FIMIL, LISTWISE and PAIRWISE processing methods

Index	Unprocessed (for the General model)	FIMIL	LISTWISE	PAIRWISE
USED	303	302	223	302
Chisq	399.673	463.166	220.783	436.739
P-value	0.000	0.000	0.000	0.000
Df	164.000	164	164	164
Cfi	0.911	0.908	0.900	0.918
Tli	0.896	0.894	0.884	0.905
Nfi	0.858	0.866	0.847	0.876
Rmse	0.078	0.078	0.084	0.074
Rmr	0.038	0.042	0.048	0.042
Srmr	0.056	0.055	0.062	0.055
Crmr	0.059	0.058	0.066	0.058
Ecv	2.101	1.971	2.299	1.751
Gfi	0.853	0.954	0.841	0.871
Agfi	0.812	0.936	0.796	0.835

The structural model for studying the impact of digital transformation on improving sustainable healthcare was tested after addressing missing data using listwise deletion and pairwise deletion methods, compared to the whole-information method. In listwise deletion, all questionnaires with missing values were excluded, reducing the sample size to 223 individuals. The matching results showed CFI 0.900, TLI 0.884, and NFI 0.847 values, indicating acceptable but below-ideal levels. RMSEA 0.084 reflected a relatively high estimation error due to missing data,

while GFI 0.841 and AGFI 0.796 were relatively weak, indicating that listwise deletion negatively impacted model quality.

In pairwise deletion, all available information for each pair of variables was used without excluding all cases, maintaining the sample size at 302 individuals. The results showed a marked improvement in the fit indices, with CFI reaching 0.918, TLI 0.905, and NFI 0.876, while RMSEA decreased to 0.074, and the quality indices GFI increased to 0.871 and AGFI to 0.835. These results indicate that pairwise deletion is more efficient than whole deletion in representing the relationships between the dimensions of digital transformation, healthcare quality, and health sustainability. However, it remains less accurate compared to the Full Information Modeling (FIML), which relies on all available data without deletion. The comparison results suggest that FIML is the most suitable method for handling missing data within structural equation modeling, given its ability to maintain sample size and provide stable and unbiased estimates, compared to both whole and pairwise deletion methods. Although the Pairwise method showed better numerical values for some fit indices, the theoretical basis and statistical stability favor FIML for the final model analysis.

While Pairwise showed better numerical values for some fit indices, the theoretical basis and statistical stability favor FIML for the final model analysis.

Table 7. Methods for estimating ULSMV & WLSMV using the listwise & pairwise deletion processing methods

Index	ULSMV-LD	ULSMV-PD	WLSMV-LD	WLSMV-PD
USED	244	303	252	303
Chisq	167.383	186.131	146.185	190.178
P-value	0.412	0.114	0.838	0.079
Df	164.000	164.000	164.000	164.000
Cfi	0.996	0.983	1.000	0.978
Tli	0.996	0.980	1.024	0.974
Nfi	0.848	0.876	0.862	0.860
Rmse	0.009	0.021	0.000	0.023
Rmr	0.090	0.088	0.076	0.086
Srmr	0.045	0.043	0.039	0.044
Crmr	0.047	0.045	0.041	0.046
Ecvi	2.092	1.915	0.698	0.722
Gfi	0.986	0.987	0.982	0.979
Agfi	0.982	0.983	0.977	0.973

The results of comparing the WLSMV and ULSMV estimation methods using both types of deletion showed that all models achieved excellent fit indices. However, the WLSMV method with pairwise deletion (PW) offered the best balance between maintaining sample size and fit quality, as well as a lower ECVI value compared to ULSMV methods, indicating higher stability in the estimation and better generalizability of the results.

A simulation study was conducted to evaluate the performance of listwise deletion and pairwise deletion in structural equation modeling using the ULSMV and WLSMV estimation methods. The simulation included sample sizes of 200, 400, 600, 800, 1000, and 1500, and missing data percentages of 0.01, 0.1. This was applied to the same questionnaire model, with the simulation model built on four underlying factors, each comprising five observed variables, to examine the stability of the matching indices under different conditions of sample size, missing data percentage, processing method, and estimation.

10. Simulation

The results of the ULSMV estimation with total deletion showed that the model achieved excellent fit levels when the percentage of missing data was low (LD = 0.01). The CFI and TLI indices registered very high values, and the RMSEA and SRMR values dropped to optimal levels, indicating high model fit. Conversely, when the percentage of missing data increased (LD = 0.1), some fit indices deteriorated, particularly in small samples (N = 200). The NFI index decreased, while the RMR and ECVI values increased, suggesting that total deletion was affected by missing data. As the sample size increased, the indices gradually improved, indicating that the ULSMV method

Table 8. shows the conformity quality indicators for different sample sizes and loss ratios of 0.01, 0.1 for the ULSMV-LD estimation method

Loss ratios	N	Chisq Scaled	Cfi Scaled	Tli scaled	Nfi scaled	Rmsea scaled	Rmr nomean	Srmr bentler	Crmr nomean	Ecvi	Gfi	Agfi
0.01	200	59.666	1.000	1.032	0.075	0.000	0.070	0.043	0.046	0.913	0.990	0.985
	400	74.346	0.993	0.991	0.914	0.015	0.058	0.036	0.039	0.551	0.993	0.989
	600	54.516	1.000	1.018	0.953	0.000	0.040	0.025	0.027	0.310	0.996	0.994
	800	74.806	0.996	0.995	0.953	0.011	0.041	0.026	0.028	0.286	0.996	0.994
	1000	58.310	1.000	1.008	0.970	0.000	0.032	0.021	0.022	0.193	0.998	0.996
0.1	1500	58.772	1.000	1.005	0.979	0.000	0.027	0.017	0.018	0.132	0.998	0.997
	200	67.827	1.000	1.014	0.666	0.000	0.152	0.085	0.091	4.098	0.966	0.948
	400	69.318	0.998	0.997	0.690	0.007	0.105	0.069	0.074	1.954	0.975	0.961
	600	57.706	1.000	1.059	0.831	0.000	0.075	0.048	0.051	1.100	0.987	0.981
	800	60.112	1.000	1.037	0.852	0.000	0.072	0.045	0.048	0.951	0.989	0.983
	1000	52.176	1.000	1.050	0.902	0.000	0.055	0.036	0.039	0.652	0.993	0.989
	1500	69.263	1.000	0.999	0.907	0.003	0.057	0.037	0.039	0.559	0.992	0.988

becomes more stable in large samples despite the presence of missing data. However, it remains less efficient compared to modern methods for processing missing data.

Table 9. shows the conformity quality indicators for different sample sizes and loss ratios of 0.01, 0.1 for the ULSMV-PD estimation method

Loss ratios	N	Chisq Scaled	Cfi Scaled	Tli scaled	Nfi scaled	Rmsea scaled	Rmr nomean	Srmr bentler	Crmr nomean	Ecvi	Gfi	Agfi
0.01	200	62.221	1.000	1.021	0.880	0.000	0.067	0.042	0.045	0.830	0.991	0.986
	400	74.098	0.994	0.992	0.923	0.014	0.054	0.034	0.036	0.484	0.994	0.990
	600	52.476	1.000	1.016	0.964	0.000	0.036	0.022	0.024	0.254	0.997	0.996
	800	77.855	0.995	0.993	0.958	0.013	0.039	0.025	0.027	0.248	0.997	0.995
	1000	59.234	1.000	1.006	0.974	0.000	0.030	0.019	0.021	0.165	0.998	0.997
0.1	1500	68.800	1.000	1.000	0.979	0.000	0.027	0.017	0.019	0.123	0.998	0.997
	200	81.343	0.978	0.971	0.876	0.030	0.075	0.049	0.051	0.959	0.988	0.982
	400	93.070	0.980	0.973	0.927	0.030	0.057	0.048	0.039	0.524	0.993	0.989
	600	76.307	0.996	0.995	0.959	0.013	0.042	0.036	0.028	0.302	0.996	0.994
	800	99.851	0.986	0.982	0.957	0.024	0.042	0.026	0.028	0.272	0.996	0.994
	1000	89.529	0.993	0.991	0.970	0.017	0.035	0.026	0.024	0.201	0.997	0.997
	1500	87.810	0.995	0.994	0.979	0.013	0.029	0.023	0.019	0.134	0.998	0.997

The results of ULSMV estimation with pairwise deletion showed that the model achieved excellent fit levels at low missing data percentages (PD = 0.01), with the fit indices registering near-perfect values across different sample sizes. When the missing data percentage increased to (PD = 0.1), a slight decline in some indices was observed, particularly in small samples; however, the fit remained within acceptable limits and improved with increasing sample size. These results indicate that pairwise deletion with ULSMV is more flexible than whole deletion, but is negatively affected at high missing data percentages.

Simulation results using the WLSMV estimation method with total deletion showed that the model achieved excellent fit levels at low data loss (LD = 0.01), with most fit indices registering optimal values across different sample sizes. However, when the data loss increased to (LD = 0.1), a clear deterioration in some indices appeared, particularly in small samples (N = 200 and 400), where the NFI index decreased and the RMR and ECVI values increased, indicating some data loss due to total deletion. As the sample size increased, the fit indices gradually improved, suggesting that the ULSMV method becomes more stable in large samples, although it remains sensitive to data loss when using total deletion.

At high loss, pairwise deletion with WLSMV showed a marked deterioration in fit quality, especially in small samples, indicating significant information loss and increased estimation instability. Although increasing the

Table 10

Loss ratios	N	Chisq Scaled	Cfi Scaled	Tli scaled	Nfi scaled	Rmse scaled	Rmr nomean	Srmr bentler	Crmr nomean	Ecvi	Gfi	Agfi
0.01	200	61.551	1.000	.026	0.867	0.000	0.070	0.043	0.046	0.605	0.983	0.974
	400	76.262	0.990	0.987	0.906	0.017	0.058	0.036	0.039	0.341	0.988	0.981
	600	53.271	1.000	1.020	0.952	0.000	0.040	0.025	0.027	0.207	0.994	0.991
	800	74.289	0.996	0.995	0.950	0.011	0.042	0.026	0.028	0.176	0.993	0.990
	1000	59.109	1.000	1.008	0.967	0.000	0.032	0.021	0.022	0.129	0.996	0.994
	1500	58.193	1.000	1.006	0.977	0.000	0.027	0.017	0.018	0.085	0.997	0.995
	200	66.751	1.000	1.027	0.671	0.000	0.154	0.085	0.091	2.544	0.950	0.925
	400	70.998	0.984	0.979	0.669	0.018	0.106	0.069	0.074	1.301	0.957	0.935
	600	56.366	1.000	1.070	0.829	0.000	0.075	0.048	0.051	0.751	0.979	0.968
	800	59.146	1.000	1.044	0.847	0.000	0.073	0.045	0.048	0.622	0.980	0.969
0.1	1000	51.578	1.000	1.055	0.898	0.000	0.056	0.036	0.039	0.474	0.987	0.980
	1500	69.308	0.999	0.999	0.902	0.004	0.057	0.037	0.039	0.352	0.987	0.980

Table 11. shows the conformity quality indicators for different sample sizes and loss ratios of 0.01, 0.1 for the WLSMV-PD estimation method

PD-WLSMV	N	Chisq Scaled	Cfi Scaled	Tli scaled	Nfi scaled	Rmse scaled	Rmr nomean	Srmr bentler	Crmr nomean	Ecvi	Gfi	Agfi
0.01	200	63.927	1.000	1.016	0.872	0.000	0.067	0.042	0.045	0.555	0.983	0.975
	400	74.925	0.993	0.990	0.916	0.015	0.054	0.034	0.036	0.303	0.989	0.983
	600	50.877	1.000	1.018	0.963	0.000	0.036	0.022	0.024	0.172	0.995	0.993
	800	78.875	0.994	0.992	0.954	0.013	0.039	0.025	0.027	0.156	0.994	0.990
	1000	59.659	1.000	1.006	0.972	0.000	0.030	0.019	0.020	0.110	0.996	0.994
0.1	1500	67.353	1.000	1.001	0.978	0.000	0.027	0.017	0.018	0.078	0.997	0.995
	200	86.260	0.968	0.958	0.863	0.035	0.076	0.048	0.052	0.667	0.977	0.965
	400	97.264	0.974	0.966	0.919	0.032	0.057	0.036	0.039	0.353	0.987	0.980
	600	75.059	0.996	0.995	0.958	0.012	0.042	0.026	0.028	0.207	0.993	0.990
	800	100.314	0.985	0.980	0.954	0.024	0.042	0.026	0.028	0.181	0.992	0.988
	1000	90.453	0.992	0.989	0.967	0.018	0.035	0.022	0.024	0.137	0.995	0.992
	1500	86.450	0.995	0.994	0.978	0.013	0.029	0.018	0.019	0.089	0.996	0.994

sample size improved performance somewhat, the indices did not return to their optimal levels, limiting the method's reliability in cases of high loss.

The simulation studies for the four tables showed a clear impact of both the missing data handling method and the estimation method on the quality of the structural equation model fit. Pairwise deletion performed better and more stably than whole deletion in most simulation scenarios, maintaining higher values for comparative fit indices (CFI, TLI, NFI) and lower values for error indices (RMSEA, SRMR, RMR). This is attributed to its retention of a larger number of observations compared to whole deletion, which reduces the actual sample size and results in the loss of important information.

Regarding the estimation methods, ULSMV and WLSMV showed similar performance at low loss ratios, achieving excellent fit across different sample sizes. However, ULSMV appeared relatively more stable in small samples, while the two methods showed clear convergence in medium and large samples. In general, the simulation results confirm that the effect of the percentage of missing data, sample size, processing method, and estimation is a crucial factor in the quality of structural equation models, and that relying on pairwise deletion with ULSMV or WLSMV is appropriate at low and medium loss percentages, while caution should be exercised at high loss, especially with total deletion.

11. Conclusion

The proposed structural model for studying the impact of digital transformation on improving sustainable healthcare was tested using three methods for handling missing data: full information with maximum probability (FIML), listwise deletion, and pairwise deletion. The aim was to compare their efficiency in maintaining sample size and model fit quality. FIML utilizes all available information without excluding cases containing missing values, while listwise deletion removes any questionnaire containing at least one missing value. Pairwise deletion uses all available pairs of data without deleting the entire case.

The results showed that the total deletion method reduced the sample size to 223 items, which negatively impacted the quality of the fit. The CFI (Center for Financial Intake) value was 0.900, the TLI (Test for Interval) value was 0.884, and the NFI (Neutral Intake) value was 0.847. The RMSEA value increased to 0.084, while the GFI (Ground Intake) value was 0.841 and the AGFI (Aggregate Intake) value was below the optimal threshold. This indicates that the loss of a number of questionnaires affected the stability of the model's estimates and its ability to represent the relationships between the dimensions of digital transformation, healthcare quality, and health sustainability.

The double deletion method maintained the sample size at 302 units and resulted in a significant improvement in the fit indices, with CFI 0.918, TLI 0.905, and NFI 0.876 values. The RMSEA value decreased to 0.074, while GFI 0.871 and AGFI 0.835 values increased. This indicates that this method is more efficient than total deletion in representing the study data, although it may suffer from inconsistencies in the correlation matrix in some cases. In contrast, the full information method (FIML) achieved the best results in terms of fit quality and model stability. It maintained the full sample size of 302 units and recorded CFI 0.908, GFI 0.954, and AGFI 0.936 values, while the RMSEA value remained within acceptable limits 0.078, and error indices such as SRMR 0.055 decreased. These results indicate that FIML provided more accurate and less biased estimates of the relationships between digital transformation variables, healthcare quality, and health sustainability.

Therefore, it can be concluded that the full information method (FIML) is the most suitable for addressing the missing data in this study, followed by the pairwise deletion method. The listwise deletion method is the least efficient due to its negative impact on sample size and model fit quality. This reinforces the reliance on FIML results in testing the study's hypotheses regarding the role of digital transformation in improving sustainable healthcare.

REFERENCES

1. B. A. H. Rabie, *Digital Transformation and Development of Healthcare Services in Medical Centers*, The Journal Future of Social Sciences, vol. 18, no. 5, pp. 155–196, 2024.
2. M. B. Tighaza, *Exploratory and Confirmatory Factor Analysis*, Dar Al Masirah for Publishing and Distribution, Amman, 2012.
3. F. Lahoual, *The Problem of Applying Statistical Methods in Experimental Psychological and Educational Research at Algerian Universities: An Applied Study on a Sample of Doctoral Theses*, Al-Mi'yar Journal, vol. 28, no. 4, pp. 1072–1085, 2024.
4. D. B. Rubin, *Inference and Missing Data*, Biometrika, vol. 63, no. 3, pp. 581–592, 1976.
5. Q. P. Zhou, *Missing Value Imputation Methods for Parameter Estimates and Psychometric Properties of Likert Measures*, University of Maryland, Baltimore, 2001.
6. J. Fan, *An Interdisciplinary and Cross-Task Review on Missing Data Imputation*, Foundations and Trends in Signal Processing, vol. 20, no. 3, pp. 185–317, 2026.
7. M. J. Zyphur, C. V. Bonner, and L. Tay, *Structural Equation Modeling in Organizational Research: The State of Our Science and Some Proposals for Its Future*, Annual Review of Organizational Psychology and Organizational Behavior, vol. 10, no. 1, pp. 495–517, 2023.
8. M. K. M. Alawadi, Z. Meshkati, R. A. K. Alkurdi, Z. Serjuei, and P. Ahmadi, *Determining the Validity and Reliability of the Arabic Version of the Cognitive Flexibility Questionnaire Among Iraqi Volleyball Players: A Psychometric Study*, 2026.
9. S. Gemici, A. Bednarz, and P. Lim, *A Primer for Handling Missing Values in the Analysis of Education and Training Data*, International Journal of Training Research, vol. 10, no. 3, pp. 233–250, 2012.
10. Y. Zhou, S. Aryal, and M. R. Bouadjene, *Review for Handling Missing Data with Special Missing Mechanism*, arXiv preprint arXiv:2404.04905, 2024.
11. L. M. Collins, J. L. Schafer, and C. M. Kam, *A Comparison of Inclusive and Restrictive Strategies in Modern Missing Data Procedures*, Psychological Methods, vol. 6, no. 4, p. 330, 2001.
12. B. Muthén and B. O. Muthén, *Statistical Analysis with Latent Variables*, Wiley, New York, vol. 123, no. 6, pp. 55–112, 2009.
13. Y. Rosseel, *lavaan: A Brief User's Guide*, Index of yrosseel/lavaan, 2012.
14. D. Rodriguez, *Area Under the Curve as an Alternative to Latent Growth Curve Modeling When Assessing the Effects of Predictor Variables on Repeated Measures of a Continuous Dependent Variable*, Stats, vol. 6, no. 2, pp. 674–688, 2023.

15. K. H. Yuan, *Normal Distribution Based Pseudo ML for Missing Data: With Applications to Mean and Covariance Structure Analysis*, Journal of Multivariate Analysis, vol. 100, no. 9, pp. 1900–1918, 2009.
16. C. K. Enders and D. L. Bandalos, *The Relative Performance of Full Information Maximum Likelihood Estimation for Missing Data in Structural Equation Models*, Structural Equation Modeling, vol. 8, no. 3, pp. 430–457, 2001.
17. T. Lee and D. Shi, *A Comparison of Full Information Maximum Likelihood and Multiple Imputation in Structural Equation Modeling with Missing Data*, Psychological Methods, vol. 26, no. 4, p. 466, 2021.
18. J. E. Karr, C. N. Areshenkoff, P. Rast, S. M. Hofer, G. L. Iverson, and M. A. Garcia-Barrera, *The Unity and Diversity of Executive Functions: A Systematic Review and Re-analysis of Latent Variable Studies*, Psychological Bulletin, vol. 144, no. 11, p. 1147, 2018.
19. K. Al-Assaf, Z. Bahroun, and V. Ahmed, *Transforming Service Quality in Healthcare: A Comprehensive Review of Healthcare 4.0 and Its Impact on Healthcare Service Quality*, Informatics, vol. 11, no. 4, p. 96, 2024.
20. R. B. Kline, *Principles and Practice of Structural Equation Modeling*, Guilford Publications, 2023.
21. C. H. Li, *Confirmatory Factor Analysis with Ordinal Data: Comparing Robust Maximum Likelihood and Diagonally Weighted Least Squares*, Behavior Research Methods, vol. 48, no. 3, pp. 936–949, 2016.
22. F. N. Yilmaz, *Comparison of Different Estimation Methods Used in Confirmatory Factor Analyses in Non-normal Data: A Monte Carlo Study*, International Online Journal of Educational Sciences, vol. 11, no. 4, pp. 131–140, 2019.
23. V. Savalei and M. Rhemtulla, *The Performance of Robust Test Statistics with Categorical Data*, British Journal of Mathematical and Statistical Psychology, vol. 66, no. 2, pp. 201–223, 2013.
24. K. A. Bollen, *Structural Equations with Latent Variables*, John Wiley & Sons, 2014.
25. B. O. Muthén, L. K. Muthén, and T. Asparouhov, *Regression and Mediation Analysis Using Mplus*, Muthén & Muthén, Los Angeles, CA, 2017.
26. A. Maydeu-Olivares, D. Shi, and Y. Rosseel, *Assessing Fit in Structural Equation Models: A Monte-Carlo Evaluation of RMSEA versus SRMR Confidence Intervals and Tests of Close Fit*, Structural Equation Modeling: A Multidisciplinary Journal, vol. 25, no. 3, pp. 389–402, 2018.
27. D. Shi, T. Lee, and A. Maydeu-Olivares, *Understanding the Model Size Effect on SEM Fit Indices*, Educational and Psychological Measurement, vol. 79, no. 2, pp. 310–334, 2019.
28. M. W. Browne and R. Cudeck, *Alternative Ways of Assessing Model Fit*, Sociological Methods & Research, vol. 21, no. 2, pp. 230–258, 1992.
29. S. Gürbüz, N. Ding, and A. Bakker, *Research Methods for Business and Management: A Practical Guide to SPSS, Process Macro, Atlas.ti, and AMOS*, Taylor & Francis, 2026.
30. R. J. Little and D. B. Rubin, *Statistical Analysis with Missing Data*, John Wiley & Sons, 2019.
31. M. R. Marshall, S. Curd, J. Kennedy, D. Khatri, S. Lee, K. Pireva, et al., *Structural Equation Modelling to Identify Psychometric Determinants of Medication Adherence in a Survey of Kidney Dialysis Patients*, Patient Preference and Adherence, pp. 855–878, 2024.
32. S. J. Finney and C. DiStefano, *Nonnormal and Categorical Data in Structural Equation Modeling*, 2013.
33. J. T. Newsom, *Longitudinal Structural Equation Modeling: A Comprehensive Introduction*, Routledge, 2023.
34. D. L. Bandalos, *Measurement Theory and Applications for the Social Sciences*, Guilford Press, 2018.
35. O. S. Ibrahim and Z. Y. Algarni, *Hybrid V-Support Vector Regression as Statistical Methodology for Forecasting Daily Natural Gas Prices*, Industrial Engineering & Management Systems, vol. 24, no. 4, pp. 694–705, 2025.
36. O. S. Ibrahim and Z. Y. Algarni, *Improving Nonlinear Regression Model Estimation Based on Coati Optimization Algorithm*, Statistics, Optimization & Information Computing, vol. 14, no. 2, pp. 923–936, 2025.
37. R. T. Ahmed and O. Salim Ibrahim, *Comparison Between the FUZZY-ARFIMA Model and the Hybrid ARFIMA-FUZZY Model with Application to Agricultural Data in the City of Mosul*, Statistics, Optimization & Information Computing, vol. 13, no. 2, pp. 741–758, 2025.
38. Alkhateeb, Ahmed N., Alghoory, Ahmed M., Hadied, Zainab A., Al-Saqal, Omar E., and Algarni, Zaid Y., *Enhancing Surface Water Potability Assessment through a ν -SVM Hybrid Statistical Learning Model*, Statistics, Optimization & Information Computing, vol. 15, no. 6, pp. 5481–5494, 2026.
39. Ibrahim, O.S. and Mohammed, M.J., *A Proposed Method for Cleaning Data from Outlier Values Using the Robust RFCH Method in Structural Equation Modeling*, International Journal of Nonlinear Analysis and Applications, vol. 12, no. 2, pp. 2269–2293, 2021.
40. Snaan, R.A.S. and Shukur, O.B., *Using Bayesian Ridge Regression Model and ESN for Climatic Time Series Forecasting*, Statistics, Optimization & Information Computing, vol. 14, no. 3, pp. 1226–1243, 2025.
41. AlBazzaz, Z.M. and Shukur, O.B., *Using LSTM Network Based on Logistic Regression Model for Classifying Solar Radiation Time Series*, In: *Proceedings of the International Conference on Explainable Artificial Intelligence in the Digital Sustainability*, Springer Nature Switzerland, Cham, pp. 375–388, 2024.
42. Hasan, S.T. and Shukur, O.B., *Using Bayesian AR-ESN for Climatic Time Series Forecasting*, Statistics, Optimization & Information Computing, vol. 14, no. 6, pp. 3816–3832, 2025.
43. Alghoory, Ahmed M., Alkhateeb, Ahmed N., and Algarni, Zaid Y., *A Modified Jackknife Liu-Type Estimator for the Gamma Regression Models Data*, Statistics, Optimization & Information Computing, vol. 14, no. 5, pp. 2142–2154, 2025.