

## Comparison of the Various Methods Used in Solving Missing Data Problems\*

Zekeriya Nartgun<sup>1</sup> and Merve Sahin Kursad<sup>2</sup>

<sup>1</sup>*Abant Izzet Baysal University, Faculty of Education, Bolu, Turkey*

<sup>2</sup>*Ankara University, Faculty of Education, Ankara, Turkey*

*E-mail: <sup>1</sup><nartgun@yahoo.com>, <sup>2</sup><sahinmerv@gmail.com>*

**KEYWORDS** Missing Data. Research. Reliability. Scale. Validity.

**ABSTRACT** The aim of this study is to compare various methods used in solving missing data problems in terms of the validity and reliability of scale results under some conditions such as missing completely at the random mechanism, normal distribution, unidimensionality, a large sample size and different missing data rates. In the study the PISA 2012 Math Work Ethics Scale results of Turkey sample were used. In the research process estimated values of validity and reliability from original complete and new complete data sets of various methods were compared. The findings reveal that the estimated values of validity and reliability from the new complete data sets of the expectation maximization, regression imputation and multiple imputation methods are the most similar values estimated from the complete data set.

### INTRODUCTION

Missing data is a significant problem encountered by researchers in a lot of research. It is most often present in the fields of social (Vansteelandt et al. 2010) and behavioral sciences (Ginkel et al. 2010). Mechanical errors such as skipping questions in long questionnaires or not recording all the data in experimental studies, research on sensitive topics such as sexual behavior (Field 2009) and unanswered questions due to lack of attention (Finch and Margraf 2008) are some of primary causes of missing data problems. Missing data in any research means a lack of knowledge and loss of information. As statistical analysis package programs are prepared for complete data sets, encountering missing data affects study results in various ways (Bal 2003).

Missing data affect both descriptive and inferential statistics estimated from that data. Thus, there may be bias and estimation errors depending on the pattern and amount of missing data (Demir and Parlak 2012; Demir 2013). In particular, when the amount of missing data is high, the generalizability of the research findings and the statistical inferences are affected negatively as well as the validity and reliability of the results (McKnight et al. 2007). Therefore, researchers are expected to prevent missing data problems at the beginning or to correct them when they encounter them (McKnight et al. 2007).

In order to decide whether the presence of missing data is a significant problem for research,

and if it is, how to overcome it, one needs to understand what causes missing data and through what kind of pattern it emerges. Missing completely at random (MCAR), missing at random (MAR) and nonignorable (NI) are the missing data mechanisms patterns. Determining the mechanism is very important in determining the method that should be used to solve it (Allison 2003).

There are different methods that can be used to deal with missing data problems. While some methods require the analysis to continue with missing data, some methods require the missing data to be deleted, and others require data to be input to replace the missing data (Little 1988; Downey and King 1998; Bal 2003; Carpita and Manisera 2011). The most preferred methods for the missing data problem are listwise deletion and pairwise deletion, which involve leaving the variable with the missing data out of analysis. However, it has been shown that these methods cause problems such as bias, sample loss, etc. (Oguzlar 2001; Allison 2009; Satıcı and Kadilar 2009; Van Der Ark and Vermunt 2010; Cumming 2013), and as a result they decrease the level of representation of the population by the sample (Little 1988; Demir and Parlak 2012). In order to reduce bias, to make effective parameter predictions and to ensure greater statistical power, in recent years researchers have increasingly suggested the use of other methods such as expectation maximization, regression imputation, multiple imputation etc. instead of deletion methods

(Enders 2013). However, it should not be forgotten that each of these methods used for dealing with missing data has a different effect on the conclusions to the research (Cheema 2012). As missing data impacts on the conclusions in different ways depending on the conditions (Ginkel et al. 2007), it is important to choose the appropriate method by considering those conditions such as the mechanism of the missing data, the sample size, the structure of the data distribution, the structure of the data gathering tool, the missing data rate etc. (Cheema 2012; Kose and Oztemur 2014; Akbas and Tavsancil 2015; Nartgun 2015).

### **The Purpose of Research**

Using the knowledge that missing data affects the results of research in different ways, as mentioned above, this study aims to compare the different missing data methods used in solving the missing data problem in terms of the validity and reliability of the scale results by examining their effectiveness under certain conditions such as missing completely at the random mechanism (MCAR), normal distribution, unidimensionality, large sample size and different missing data rates. Finding out the effectiveness of these different missing data methods in the above-mentioned conditions is considered the most important aspect of this study.

Approximate value imputation methods - series mean, mean of nearby points, median of nearby points, linear interpolation, linear trend at points - listwise deletion, expectation maximization, regression imputation and multiple imputation are the methods compared in this study.

## **METHODOLOGY**

### **Research Model**

This study is based on fundamental research in which the effectiveness of various missing data methods used in the solution of missing data problems are compared under different conditions. Fundamental research are theoretical or experimental studies which are carried out to add to the existing knowledge (Karasar 2007).

### **Sample, Data Collection Tools and Data**

In this study, the PISA 2012 "Math Work Ethics" scale (nine items and the four-option Likert-

type scale) and the PISA 2012 Turkey sample were used. In the research process, a complete data set of 1,000 students was taken from the PISA 2012 Turkey sample randomly, and 5 percent, 10 percent and 20 percent of the data were deleted from the complete data set to generate missing data sets under the MCAR data mechanism. After completing these missing data and creating new complete data sets using approximate value imputation methods - series mean, mean of nearby points, median of nearby points, linear interpolation, linear trend at points - listwise deletion, expectation maximization, regression imputation and multiple imputation, analysis of validity and reliability were carried out.

In this study, the values estimated for the complete data set for each missing data rate, within the context of the validity and reliability of the scale, were compared with the values estimated for the new complete data sets. The estimated values for the complete data set were used as reference values for comparison with the values estimated from the new complete data sets.

### **Data Analysis Process**

In the first phase of data analysis, the normality of the distribution of the data and the single factor construct of scale were tested. The skewness and kurtosis coefficients for the complete data set were found to be 0.458 and 0.321, respectively. That these coefficients lie between the values of -1 and +1 has been taken as an indicator of the normality of the distribution of the data (Huck 2012). To test the factor construct of scale, principal component analysis was used. This analysis showed that all the items on the scale had high factor loadings in the first factor, and the eigenvalue for this single factor were 5.24. The variance explained by this single factor was determined as 58.20 percent of the total variance. An eigenvalue of 1 or more and a minimum explained variance of 30 percent are sufficient values for a scale to be accepted as a single factor construct (Buyukozturk 2012).

In the second phase of the data analysis process, 5 percent, 10 percent and 20 percent of the data from the complete data set of 1,000 students were deleted to create missing data sets. Although randomness was taken into account in the deletion of the data, a Little's MCAR test was applied to the data to ensure randomness. To this end, the Little's MCAR test, which is calculated by the EM algorithm of the SPSS 20.0

program, was used. The Little's MCAR test yielded a  $p$  value of 0.80 for the 5 percent, 0.23 for the 10 percent and 0.80 for the 20 percent missing data rates. Since these  $p$  values are greater than 0.05, the data are accepted as being suitable for the MCAR mechanism (IBM 2012).

Once the single factor construct, normality of data and the results of little's MCAR test had been examined, the missing data sets of 5 percent, 10 percent and 20 percent were converted into the new complete data sets using different missing data methods. These methods are approximate value imputation methods - series mean, mean of nearby points, median of nearby points, linear interpolation, linear trend at points - listwise deletion, expectation maximization, regression imputation and multiple imputation. Analyses and comparisons regarding the validity of the complete data set and the new complete data sets were carried out with the values estimated from exploratory factor analysis. Furthermore, analyses and comparisons regarding reliability were carried out using Cronbach alpha reliability coefficients and Fisher's  $z$  statistics. The values related to validity and reliability obtained through the analysis of the complete data set of 1,000 students were used as reference values for comparison.

## RESULTS

### Findings Related to Validity

The results of the validity analyses for the complete data set and the new completed data sets, for each missing data rates, are given separately, below in Tables 1, 2 and 3.

The item factor loadings estimated for the complete data set of 1,000 students vary between 0.72 and 0.80. The eigenvalue for this data set is 5.24, while the explained variance value is 58.20. These values have been taken as references for the comparison of the values obtained from the new completed data sets, which are created by using different missing data methods.

When the values given in Table 1 are examined, it can be seen that the estimated factor loadings of all the items from all the new complete data sets are the same as those estimated from the complete data set except for item four for the new data sets of series mean and the linear trend at point methods, and items one and three for the new data set of listwise deletion.

It can be seen that in Table 1 the eigenvalues vary between 5.22 and 5.27. The eigenvalues estimated for the new complete data sets of the regression imputation and multiple imputation methods are the same as the eigenvalue estimat-

**Table 1: Results of validity analysis for the complete data set (sample size of 1,000) and the new completed data sets created by using different missing data methods for the condition of the 5% missing value rate**

| Item | Complete data | Series mean | Mean of nearby points | Median of nearby points | Linear interpolation | Linear trend at point | Listwise deletion | Expectation maximization | Regression imputation | Multiple imputation |
|------|---------------|-------------|-----------------------|-------------------------|----------------------|-----------------------|-------------------|--------------------------|-----------------------|---------------------|
| 1    | .75           | .75         | .75                   | .75                     | .75                  | .75                   | .76               | .75                      | .75                   | .75                 |
| 2    | .80           | .80         | .80                   | .80                     | .80                  | .80                   | .80               | .80                      | .80                   | .80                 |
| 3    | .77           | .77         | .77                   | .77                     | .77                  | .77                   | .78               | .77                      | .77                   | .77                 |
| 4    | .76           | .75         | .76                   | .76                     | .76                  | .75                   | .76               | .76                      | .76                   | .76                 |
| 5    | .77           | .77         | .77                   | .77                     | .77                  | .77                   | .77               | .77                      | .77                   | .77                 |
| 6    | .80           | .80         | .80                   | .80                     | .80                  | .80                   | .80               | .80                      | .80                   | .80                 |
| 7    | .77           | .77         | .77                   | .77                     | .77                  | .77                   | .77               | .77                      | .77                   | .77                 |
| 8    | .72           | .72         | .72                   | .72                     | .72                  | .72                   | .72               | .72                      | .72                   | .72                 |
| 9    | .72           | .72         | .72                   | .72                     | .72                  | .72                   | .72               | .72                      | .72                   | .72                 |
| E    | 5.24          | 5.22        | 5.22                  | 5.23                    | 5.22                 | 5.22                  | 5.27              | 5.25                     | 5.24                  | 5.24                |
| EV   | 58.20         | 58.00       | 58.02                 | 58.02                   | 58.01                | 58.00                 | 58.52             | 58.31                    | 58.18                 | 58.21               |

E= Eigenvalue; EV = Explained Variance

ed for the complete data set (5.24). The lowest and most distant eigenvalues from that of the complete data set were estimated for the new complete data sets of series mean (5.22), the mean of nearby points (5.22), the linear interpolation (5.22), and the linear trend at point methods (5.22), while the highest eigenvalue was obtained for the new complete data set of the listwise deletion method (5.27).

Table 1 shows that the new complete data set of the multiple imputation method (58.21) has the closest explained variance value to that estimated for the complete data set (58.20). The greatest explained variance value was estimated for the new complete data set of the listwise deletion method (58.52), while the lowest value was estimated for the new complete data sets of series mean and the linear trend at point methods (58.00).

When the item factor loadings given in Table 2 are examined, it can be seen that the new complete data set of the expectation maximization method and the complete data set have the same factor loadings except for items 1, 2, 3 and 9. Three items' factor loadings are different from the complete data set for the new complete data sets using the following methods: median of nearby points (items 1, 4 and 7), regression imputation (items 2, 3 and 9) and multiple imputation (items 3, 4 and 9). The most similar results to the

complete data set were estimated for the new complete data sets of the series mean and the linear trend at point methods. For these two methods only the fourth items are different from those of the complete data set.

An examination of the eigenvalues given in Table 2 shows that the values vary between 5.17 and 5.28. The estimated eigenvalues for the new complete data set of the linear trend at point (5.22) and multiple imputation (5.26) methods are the closest to those estimated for the complete data set. The lowest and most distant eigenvalue from that of the complete data set was estimated for the new complete data set of the linear interpolation method (5.17), while the highest eigenvalues were estimated for the new complete data sets of listwise deletion (5.28) and expectation maximization (5.28) methods.

Table 2 shows that the new complete data set of the multiple imputation method (58.42) has the closest explained variance value to that estimated for the complete data set. The greatest explained variance value was estimated for the new complete data set of the listwise deletion method (58.65), while the lowest value was estimated for the new complete data set of the linear interpolation method (57.48).

When the item factor loadings given in Table 3 are examined, it can be seen that the estimated values for the new complete data sets of the ex-

**Table 2: Results of validity analysis for the complete data set (sample size of 1,000) and the new completed data sets created by using different missing data methods for the condition of the 10% missing value rate**

| <i>Item</i> | <i>Complete data</i> | <i>Series mean</i> | <i>Mean of nearby points</i> | <i>Median of nearby points</i> | <i>Linear interpolation</i> | <i>Linear trend at point</i> | <i>Listwise deletion</i> | <i>Expectation maximization</i> | <i>Regression imputation</i> | <i>Multiple imputation</i> |
|-------------|----------------------|--------------------|------------------------------|--------------------------------|-----------------------------|------------------------------|--------------------------|---------------------------------|------------------------------|----------------------------|
| 1           | .75                  | .75                | .75                          | .74                            | .75                         | .75                          | .75                      | .76                             | .75                          | .75                        |
| 2           | .80                  | .80                | .80                          | .80                            | .80                         | .80                          | .80                      | .81                             | .81                          | .80                        |
| 3           | .77                  | .77                | .77                          | .77                            | .77                         | .77                          | .78                      | .78                             | .78                          | .78                        |
| 4           | .76                  | .75                | .75                          | .75                            | .75                         | .75                          | .76                      | .76                             | .76                          | .75                        |
| 5           | .77                  | .77                | .77                          | .77                            | .76                         | .77                          | .77                      | .77                             | .77                          | .77                        |
| 6           | .80                  | .80                | .80                          | .80                            | .80                         | .80                          | .81                      | .80                             | .80                          | .80                        |
| 7           | .77                  | .77                | .76                          | .76                            | .77                         | .77                          | .77                      | .77                             | .77                          | .77                        |
| 8           | .72                  | .72                | .72                          | .72                            | .72                         | .72                          | .72                      | .72                             | .72                          | .72                        |
| 9           | .72                  | .72                | .72                          | .72                            | .72                         | .72                          | .72                      | .73                             | .73                          | .73                        |
| E           | 5.24                 | 5.21               | 5.19                         | 5.19                           | 5.17                        | 5.22                         | 5.28                     | 5.28                            | 5.27                         | 5.26                       |
| EV          | 58.20                | 57.93              | 57.67                        | 57.68                          | 57.48                       | 57.94                        | 58.65                    | 58.64                           | 58.52                        | 58.42                      |

E= Eigenvalue; EV = Explained Variance

**Table 3: Results of validity analysis for the complete data set (sample size of 1,000) and the new completed data sets created by using different missing data methods for the condition of the 10% missing value rate**

| <i>Item</i> | <i>Complete data</i> | <i>Series mean</i> | <i>Mean of nearby points</i> | <i>Median of nearby points</i> | <i>Linear interpolation</i> | <i>Linear trend at point</i> | <i>Listwise deletion</i> | <i>Expectation maximization</i> | <i>Regression imputation</i> | <i>Multiple imputation</i> |
|-------------|----------------------|--------------------|------------------------------|--------------------------------|-----------------------------|------------------------------|--------------------------|---------------------------------|------------------------------|----------------------------|
| 1           | .75                  | .74                | .75                          | .75                            | .75                         | .74                          | .76                      | .75                             | .75                          | .75                        |
| 2           | .80                  | .80                | .80                          | .80                            | .80                         | .80                          | .81                      | .81                             | .81                          | .81                        |
| 3           | .77                  | .77                | .77                          | .77                            | .76                         | .77                          | .78                      | .78                             | .77                          | .78                        |
| 4           | .76                  | .75                | .75                          | .75                            | .75                         | .75                          | .77                      | .76                             | .75                          | .76                        |
| 5           | .77                  | .76                | .76                          | .75                            | .75                         | .76                          | .78                      | .77                             | .77                          | .77                        |
| 6           | .80                  | .79                | .79                          | .79                            | .79                         | .79                          | .81                      | .80                             | .80                          | .80                        |
| 7           | .77                  | .76                | .76                          | .76                            | .76                         | .76                          | .77                      | .77                             | .76                          | .76                        |
| 8           | .72                  | .72                | .72                          | .72                            | .72                         | .72                          | .72                      | .72                             | .72                          | .72                        |
| 9           | .72                  | .72                | .72                          | .72                            | .72                         | .72                          | .74                      | .73                             | .72                          | .72                        |
| E           | 5.24                 | 5.16               | 5.13                         | 5.12                           | 5.12                        | 5.16                         | 5.34                     | 5.28                            | 5.23                         | 5.24                       |
| EV          | 58.20                | 57.32              | 57.04                        | 56.92                          | 56.86                       | 57.33                        | 59.32                    | 58.71                           | 58.12                        | 58.20                      |

E= Eigenvalue; EV = Explained Variance

pectation maximization, regression imputation and multiple imputation methods are closest to that of the complete data set. Three items' factor loadings are different from the complete data set for the new complete data sets of the expectation maximization (items 2, 3 and 9), regression imputation (items 2, 4 and 7) and multiple imputation (items 2, 3 and 7) methods. Four items' factor loadings are different from the complete data set for the new complete data sets of the median of nearby points (items 4, 5, 6 and 7) and the mean of nearby points (items 4, 5, 6 and 7), and five items' factor loadings are different from the complete data set for the new complete data sets of the series mean (items 1, 4, 5, 6 and 7), the linear interpolation (items 3, 4, 5, 6 and 7) and the linear trend at points (items 1, 4, 5, 6 and 7) methods. Finally, only two items (items 7 and 8) of the new complete data set of listwise deletion are similar to that estimated for the complete data set.

It can be seen that in Table 3 the eigenvalues vary between 5.12 and 5.34. The eigenvalue estimated for the new complete data sets of the multiple imputation (5.24) method is the same as the eigenvalue estimated for the complete data set (5.24). The lowest and the most distant eigenvalues from that of the complete data set were estimated for the new complete data sets of the median of nearby points (5.12) and linear interpolation

(5.12), while the highest eigenvalue was obtained for the new complete data set of the listwise deletion method (5.34).

Table 3 shows that the new complete data set of the multiple imputation method (58.20) has the closest explained variance value to that estimated for the complete data set. The greatest explained variance value was estimated for the new complete data set of the listwise deletion method (59.32), while the lowest value was estimated for the new complete data set of the linear interpolation (56.86) method.

In terms of single factor construct, taken as a whole, all missing data methods give the similar results at all missing data rates. These findings can be interpreted as all missing data methods operate in the same way at this context.

For all the missing data rates, although similar item factor loadings were observed in terms of the small and large size of the factor loadings, there is a fall in the resemblance of the item factor loadings estimated for the new complete data sets as the rate of the missing data increased. These results can be interpreted that researchers should be much more careful in their research process when they encounter the big missing data rates.

In terms of explained variance and eigenvalues, the estimated values for the new complete data sets are varies to that of complete data set

at different missing data rates. These results can be interpreted as the use of some methods are more accurate than the others.

### Findings Related to Reliability

The results of the reliability analysis (Cronbach alpha reliability coefficients and Fisher's z test) of the complete data set and the new completed data sets are given in Table 4 for each of the 5 percent, 10 percent and 20 percent missing data rates.

The Cronbach alpha coefficient estimated for the complete data set of 1,000 students is 0.910. This value has been taken as reference for the comparisons of values obtained from the new completed data sets, which are created by using different missing data methods.

As shown in Table 4, for the 5 percent missing data rate the estimated Cronbach alpha reliability coefficients of the new complete data sets vary between 0.909 and 0.911. The lowest Cronbach alpha coefficients ( $\alpha = 0.909$ ) were estimated for the new complete data sets of the approximate value imputation methods and regression imputation method, while the new complete data sets of the expectation maximization and multiple imputation methods have the same coefficients ( $\alpha = 0.910$ ) as that of the complete data set. On the other hand, the highest coefficient ( $\alpha = 0.911$ ) was estimated for the new complete data set of the listwise deletion method.

For the 10 percent missing data rate, the estimated Cronbach alpha reliability coefficients of the new complete data sets vary between 0.907 and 0.911. While the new complete data sets of the regression imputation and multiple imputation

methods have the same coefficient ( $\alpha = 0.910$ ) as that of the complete data set, the highest coefficient ( $\alpha = 0.911$ ) was estimated for the new complete data sets of the listwise deletion and expectation maximization methods. Although the new complete data sets of the listwise deletion and expectation maximization methods yielded higher coefficients compared to the complete data set, these coefficients are still quite close to that of the complete data set. The lowest coefficient compared to that of the complete data set was estimated for the new complete data sets of the mean of nearby points, the median of nearby points and the linear interpolation ( $\alpha = 0.907$ ) methods.

For the 20 percent missing data rate, the estimated Cronbach alpha reliability coefficients of the new complete data sets vary between 0.904 and 0.913. While the lowest coefficient ( $\alpha = 0.904$ ) was estimated for the new complete data sets of the median of nearby points and linear interpolation methods, the highest coefficient value ( $\alpha = 0.913$ ) was estimated for the new complete data sets of the listwise deletion method. While the coefficient ( $\alpha = 0.910$ ) estimated for the new complete data set of the multiple imputation method is the same as the coefficient estimated for the complete data set, the new complete data sets of the expectation maximization ( $\alpha = 0.911$ ) and regression imputation ( $\alpha = 0.909$ ) methods have the closest values to the coefficient of the complete data set. Finally, it can be seen that the coefficients estimated for the new complete data sets of the approximate value imputation method are the lowest compared to that estimated for the complete data set.

The results of the Fisher's z test that was carried out to determine whether there is a sig-

**Table 4: Results of reliability analysis for the complete data set (sample size of 1,000) and the new completed data sets created by using different missing data methods for the condition of the 5%, 10% and 20% missing value rates**

| Missing data method      | 0% missing<br>(Complete Data)<br>Cronbach $\alpha$ | 5%                              | 10%                             | 20%                             |
|--------------------------|----------------------------------------------------|---------------------------------|---------------------------------|---------------------------------|
|                          |                                                    | Cronbach $\alpha$<br>(Fisher z) | Cronbach $\alpha$<br>(Fisher z) | Cronbach $\alpha$<br>(Fisher z) |
| Series mean              | .910                                               | .909 (.000)                     | .908 (.000)                     | .906 (.650)                     |
| Mean of nearby points    |                                                    | .909 (.000)                     | .907 (.650)                     | .905 (.650)                     |
| Median of nearby points  |                                                    | .909 (.000)                     | .907 (.650)                     | .904 (.650)                     |
| Linear interpolation     |                                                    | .909 (.000)                     | .907 (.650)                     | .904 (.650)                     |
| Linear trend at point    |                                                    | .909 (.000)                     | .908 (.000)                     | .906 (.650)                     |
| Listwise deletion        |                                                    | .911 (.000)                     | .911 (.000)                     | .913 (-.650)                    |
| Expectation maximization |                                                    | .910 (.000)                     | .911 (.000)                     | .911 (.000)                     |
| Regression imputation    |                                                    | .909 (.000)                     | .910 (.000)                     | .909 (.000)                     |
| Multiple imputation      |                                                    | .910 (.000)                     | .910 (.000)                     | .910 (.000)                     |

nificant difference between the Cronbach alpha reliability coefficient of the complete data set and the Cronbach alpha reliability coefficients of the new complete data sets created by using different missing data solving methods shows that the  $z$  values calculated for all the new complete data sets, under the conditions of the 5 percent, 10 percent and 20 percent missing data rates, lie within the limits of -1,96 and + 1,96 (Akhun 1994; Kenny 1987), which means that there is no significant difference between the Cronbach alpha reliability coefficients estimated for the complete data set and for the new complete data sets.

In terms of Cronbach alpha reliability coefficient, taken as a whole, all missing data methods give the similar results at all missing data rates. These findings can be interpreted as all missing data methods operate in the same way at this context.

## DISCUSSION

The comparisons between the complete data set and the new complete data sets for validity show that for the different missing data rates, no matter which missing data method was used, the single factor construct of the complete data set was retained. This finding is consistent with some studies in the relevant literature (Cokluk and Kayri 2011; Chen et al. 2012).

When the item factor loadings are taken in the context of the exploratory factor analysis, a fall in the resemblance of the item factor loadings estimated for the new complete data sets was observed as the rate of the missing data increased. However, it was seen that for all the missing data rates, similar item factor loadings were observed in terms of the small and large size of the factor loadings. However, under the condition of the high missing data rate in particular, the similarity of the item factor loadings estimated for the complete data set and the new complete data sets of the approximate value imputation and listwise deletion methods decreased. On the other hand, the estimated item factor loadings of the new complete data sets of the expectation maximization, regression imputation and multiple imputation methods are the closest to those estimated for the complete data set under all conditions, especially when the missing data rate is high.

In terms of explained variance and eigenvalues, the estimated values for the new complete data sets of the approximate value imputation methods are low compared to the estimated val-

ues for the complete data set. This finding is consistent with the findings of Cokluk and Kayri (2011) and Akbas and Tavsancil (2015). The lowest values are obtained from the new complete data set of the linear interpolation method of the approximate value imputation methods. On the other hand, the highest values are obtained from the new complete data set of listwise deletion for all the missing data rates. The estimated explained variance and eigenvalues for the new complete data set of the multiple imputation method are the closest values to that of the complete data set. As well as the multiple imputation method, similar values were also estimated for the new complete data sets of the expectation maximization and regression imputation methods. These findings are similar to the findings of some studies in the literature (Graham et al. 1994; Bernaards and Sijtsma 2000; Musil et al. 2002; Streiner 2002; Bal 2003; Enders 2004; Peng et al. 2006; Cheema 2012).

When analyzed as a whole, under the conditions of all the missing data rates, it is seen that the estimated values of the Cronbach alpha reliability coefficients for all the new complete data sets are very close to or the same as that estimated for the complete data set. However, it can be seen that under the conditions of high missing data rates, the coefficients estimated for the new complete data sets of the approximate value imputation methods were slightly low compared to that estimated for the complete data set. These results are consistent with the results of some studies in the field literature (Cokluk and Kayri 2011; Demir 2013; Nartgün 2015). Under the conditions of low missing data rates, the estimated coefficients of the new complete data set of listwise deletion are very similar to that of the complete data set, but when the missing data rates increased, the estimated coefficients were higher than that estimated for the complete data set. This finding is also consistent with the finding of Demir (2013). The estimated coefficients of the new complete data sets of the expectation maximization, regression imputation and multiple imputation methods are the closest coefficients to that estimated for the complete data set. Among them the estimated coefficients of the new complete data set of the multiple imputation method is especially noteworthy. The multiple imputation method performing better than the other methods is supported by some studies in the field (Granberg-Rademacker 2007;

Leite and Beretvas 2010; Young et al. 2011). While some descriptive differences were observed in the reliability coefficients for the different missing data methods under different missing data rates, according to Fisher's z test there is no statistically significant difference between the coefficients estimated for the complete data set and the new complete data sets.

According to the literature, it can be seen that the relative superiority of the different missing data methods changes depending on factors such as the missing data rate, the sample size, the missing data mechanism etc. In the present study some missing data methods designed to solve the missing data problems are compared under the conditions of the missing completely at random mechanism (MCAR), normal distribution, unidimensionality, large sample size and different missing data rates. As a result of this study, the values estimated for the new complete data sets of the different missing data methods in terms of validity and reliability are similar to or the same as some cases with that estimated for the complete data set. The findings of this study may be explained by using the fact that the missing data structure was the MCAR (Bernaards and Sijtsma 2000; Buyukozturk 2012; Buhi et al. 2008; Graham 2009).

### CONCLUSION

The single factor construct of complete data set and the new complete data sets, no matter which missing data method was used, remained same for all missing data rates.

When the item factor loadings are taken in the context of the exploratory factor analysis, a fall in the resemblance of the item factor loadings estimated for the new complete data sets of different methods was observed as the rate of the missing data increased.

In terms of explained variance and eigenvalues, the estimated values for the new complete data sets of different methods are varies to that of complete data set at different missing data rates.

When analysed as a whole, under the conditions of all the missing data rates, it is seen that the estimated values of the Cronbach alpha reliability coefficients for all the new complete data sets of different methods are very close to or the same as that estimated for the complete data set.

### RECOMMENDATIONS

The following suggestions may be considered based on the aims and findings of the present study:

1. Under the conditions considered in this study the multiple imputation, expectation maximization or regression imputation methods may be preferable to create new complete data sets when the missing data rates are high.
2. Under the conditions considered in this study, when the missing data rates are low, all the methods except for listwise deletion can be used to create new complete data sets.

### NOTE

\*This article was presented at the 1st International Conference on Lifelong Education and Leadership, in Olomouc, Czech on October 29-31, 2015.

### REFERENCES

- Akbas U, Tavsancil E 2015. Investigation of psychometric properties of scales with missing data techniques for different sample sizes and missing data patterns. *Journal of Measurement and Evaluation in Education and Psychology*, 6 (1): 38-57.
- Akhun I 1994. *Statistical Formulas and Tables*. Ankara: Faculty of Education, Hacettepe University Publications.
- Allison PD 2003. Missing data techniques for structural equation modeling. *Journal of Abnormal Psychology*, 112(4): 545-557. DOI: 10.1037/0021-843X.112.4.545.
- Allison PD 2009. Missing Data. In: *Sage University Paper Series on Quantitative Applications in the Social Sciences*. London: Sage Publication, pp. 72-89.
- Bal C 2003. *The Solution of Missing Value Problem in Multigroup Data Sets and an Application in Health*. PhD Thesis, Unpublished, Eskisehir: Osmangazi University.
- Bernaards CA, Sijtsma K 2000. Influence of imputation and EM methods on factor analysis when item nonresponse in questionnaire data is non-ignorable. *Multivariate Behavioral Research*, 35(3): 321-364. DOI: 10.1207/S15327906MBR3503\_03.
- Buhi ER, Goodson P, Neilands TB 2008. Out of sight, not out of mind: Strategies for handling missing data. *American Journal of Health Behavior*, 32(1): 83-92.
- Buyukozturk S 2012. *Data Analysis Manuel for Social Sciences*. Ankara: Pegem Academy.
- Carpita M, Manisera M 2011. On the imputation of missing data in surveys with Likert-type scales. *Journal of Classical*, 28: 93-112. DOI: 10.1007/s00357-011-9074-z.

- Cheema J 2012. *Handling Missing Data in Educational Research Using SPSS*. PhD Thesis, Unpublished. USA: George Mason University.
- Chen SF, Wang S, Chen YC 2012. A simulation study using EFA and CFA programs based the impact of missing data on test dimensionality. *Expert Systems with Applications*, 39: 4026-4031.
- Cumming P 2013. Missing data and multiple imputation. *Clinical Review and Education*, 167(7): 656-661.
- Cokluk O, Kayri M 2011. The effects of methods of imputation for missing values on the validity and reliability of scales. *Educational Sciences: Theory and Practise*, 11(1): 289-309.
- Demir E, Parlak B 2012. Missing data problems in educational researches in Turkey. *Journal of Measurement and Evaluation in Education and Psychology*, 3(1): 230-241.
- Demir E 2013. Item and test parameters estimations for multiple choice tests in the presence of missing data: The case of SBS. *Journal of Educational Sciences Research*, 3(2): 47-68.
- Downey RG, King CV 1998. Missing data in likert ratings: A comparison of replacement methods. *The Journal of General Psychology*, 125(2): 175-191. DOI: 10.1080/00221309809595542.
- Enders CK 2004. The impact of missing data on sample reliability estimates: Implications for reliability reporting practices. *Educational and Psychological Measurement*, 64(3): 419-436. DOI: 10.1177/0013164403261050.
- Enders CK 2013. Dealing with missing data in developmental research. *Child Development Perspectives*, 7(1): 27- 31.
- Field A 2009. *Discovering Statistics Using SPSS*. London: Sage Publication.
- Finch H, Margraf M 2008. Imputation of Categorical Missing Data: A Comparison of Multivariate Normal and Multinomial Methods. From <<http://www.mwsug.org/proceedings/2008/stats/MWSUG-2008-S05.pdf>> (Retrieved on 20 November 2014).
- Ginkel JRV, Van der Ark LA, Sijta K, Vermunt JK 2007. Two-way imputation: A Bayesian method for estimating missing scores in tests and questionnaires, and an accurate approximation. *Computational Statistics and Data Analysis*, 51: 4013-4027. DOI: 10.1016/j.csda.2006.12.022.
- Ginkel JRV, Sijta K, Van der Ark LA, Vermunt JK 2010. Incidence of missing item scores in personality measurement, and simple item-score imputation. *Methodology*, 6(1): 17-30. DOI: 10.1027/1614-2241/a000003.
- Graham JW, Hofer SM, Piccinin AM 1994. Analysis with missing data in drug prevention research. US National Library of Medicine National Institutes of Health, *NIDA Research Monographs*, pp.13-64.
- Graham JW 2009. Missing data analysis: Making it work in the real world. *Annual Review of Psychology*, 60: 549-576.
- Granberg Rademacker JS 2007. A comparison of three approaches to handling incomplete state level data. *State Politics and Policy Quarterly*, 7(3): 325-338.
- Huck SW 2012. *Reading Statistics and Research*. 6<sup>th</sup> Edition. USA: Pearson.
- Karasar N 2007. *Scientific Research Methods. Concepts, Principles, Techniques*. Ankara: Nobel Publishing.
- Kose IA, Oztemur B 2014. Examining the effect of missing data handling methods on the parameters of t test and anova. *Abant Izzet Baysal University Journal of Education*, 14(1): 400-412.
- Kenny DA 1987. *Statistics for the Social and Behavioral Science*. Boston: Little, Brown and Company.
- Leite W, Beretvas SN 2010. The performance of multiple imputation for Likert-type items with missing data. *Journal of Modern Applied Statistical Methods*, 9(1): 64-74.
- Little RJA 1988. Missing data adjustments in large surveys. *Journal of Business and Economic Statistics*, 6(3): 287-296.
- McKnight PE, McKnight KM, Sidani S, Figueredo AJ 2007. *Missing Data: A Gentle Introduction*. United States of America: The Guilford Press.
- Musil CM, Warner CB, Yobas PK, Jones SL 2002. A comparison of imputation techniques for handling missing data. *Western Journal of Nursing Research*, 24(7): 815-829. DOI: 10.1177/019394502237390.
- Nartgun Z 2015. Comprison of various methods used in solving missing data problems in terms of psychometric features of scales and measurement results under different missing data conditions. *International Online Journal of Educational Sciences*, 7(4): 252-265. DOI: <http://dx.doi.org/10.15345/ijoes.2015.04.017>.
- Oguzlar A 2001. Missing Data Problems and Recommendations for Solutions in Surveys. *Paper presented in V National Econometry and Statistics Symposium*. Cukurova University, Adana, 19-22 September.
- Peng CYJ, Harwell M, Liou SM, Ehman LH 2006. Advances in missing data methods and implications for educational research. In: S Sawilowsky (Ed.): *Real Data Analysis*. Greenwich, CT: Information Age Publishing Inc, pp. 31-78.
- Satici E, Kadilar C 2009. Estimation of the population mean when some observations are missing. *Anadolu University Journal of Science and Technology*, 10(2): 549-556.
- Streiner DV 2002. The case of the missing data: Methods of dealing with dropouts and other research vagaries. *Research methods in Psychiatry*, 47: 68-75.
- Van der Ark LA, Vermunt JK 2010. New developments in missing data analysis. *Methodology*, 6(1): 1-2. DOI: 10.1027/1614-2241/a000001.
- Vansteelandt S, Carpenter J, Kenward MG 2010. Analysis of incomplete data using inverse probability weighting and doubly robust estimators. *Methodology*, 6(1): 37-48. DOI: 10.1027/1614-2241/a000 005.
- Young W, Weckman G, Holland W 2011. A survey of methodologies for the treatment of missing values within datasets: Limitations and benefits. *Theoretical Issues in Ergonomics Science*, 12(1): 15-43. DOI: 10.1080/14639220903470205.