

CHAPTER 4

DATA ANALYSIS AND RESULTS

4.1 Introduction

This chapter will cover the data analysis and results of the study, following the completion of the data collecting process. Specifically, this analysis will determine the results which will confirm the rejection or acceptance of the research questions, objectives and hypothesis proposed in this study. Overall, this chapter will consist of two major sections which represent the progression of the analysis. The first section is SPSS and the second section is structural equation modeling (SEM). The former displays the demographic information of the respondents as well as the results of a demographic profile. Following that, this chapter will go into the exploratory factor analysis and the results of the confirmatory factor analysis, such as the measurement model's goodness of fit, validity, and reliability. Finally, there will be a description of the structural model validity and mediation analysis.

4.2 Data Analysis

Data analysis is the most vital part of any research. Data analysis summarizes the collected data which involves the interpretation of data gathered through the use of analytical and logical reasoning to determine patterns, relationships or trends. In this current study, SPSS version 23.0 was used to generate the data gained from the respondents. As per Osborn and Costello (2009), SPSS was used to produce descriptive analysis, frequencies, mean, mod, and standard deviation, as well as exploratory factor analysis (EFA). Later, the study's variables, hypothesis, and model were evaluated using structural equation modelling (SEM) with AMOS 23.0. This study sought to determine the role of work-family enrichment in mediating the relationship between lifelong learning and job satisfaction among teachers in Malaysian public primary and secondary schools.

4.3 Response Rate

It has become a main concern of every single researcher to obtain the highest possible response rate during the questionnaire distribution. For this study, the researcher distributed 350 questionnaires to the respondents in total. However, after approximately a four-month duration of the survey session, the amount of return questionnaires was 234. A total of 234 responses were useable and used for a subsequent analysis giving a response rate of 67 percent.

4.3.1 The Demographic and Profiling of the Samples

This section will highlight and elaborate the profile of the respondents from the study. Altogether, there eight aspects were measured, namely gender, marital status, age, level of education, grade / position, teaching experience, working spouse, and number of children.

The profiles of the respondents who responded to the questionnaire are as follows. The respondents that took part in this survey consisted of 234 in which 74.4 percent (N = 174) of the employees were female and 25.6 percent (N = 60) of the employees were male. Table 4.1 illustrates the demographic profile of employees. Out of the 234 respondents who completed the survey, 81.6 percent (N = 191) of them were married, 17.1 per cent (N = 40) were single and only 1.3 percent were divorced. Most of the employees were between 36 to 40 years old (36.3%), 27.8 percent were 41 to 45 years old, 27.4 percent were 21 to 35 years old (N = 64), and 6.8 percent were 46 years old and above (N = 16). The majority of the students who responded to the survey have academic qualification of master's (59.0 percent, N = 138), and Phd (41.0 percent, N = 96). For grade of employees, most of them were grade DG 44 (67.9 percent, N=159), followed by grade DG 48 (15.4 percent,, N=36), DG 41 (14.2 percent, N=33), and DG52 (2.6 percent,, N=6). Besides working experience, almost all of the respondents have 11 to 20 years of experience (64.1per cent, N = 150), 1 to 10 years (26.9 per cent, N = 63), and 21 to 30 years (9.0 percent, N=21). The majority of the employees have a working spouse (82.1 percent, n=192), while 17.9 percent do not have a working spouse (17.9 percent, N=42). A total of 53.4 percent of employees in this study have 0 - 2 children (N=125), 35 percent have 3 - 4 children (N=82), 9.4 percent have 5 - 6 children (N=22), and 2.1 percent have more than 6 children (N=5).

Table 4.1: Demographic and Profile Details of Samples

Characteristic	Frequency	Percentage
Gender		
Male	60	25.6
Female	174	74.4
Marital Status		
Single	40	17.1
Married	191	81.6
Divorced	3	1.3
Age		
25 - 30 years	4	1.7
31 - 35 years	64	27.4
36 - 40 years	85	36.3
41 - 45 years	65	27.8
>46 years	16	6.8
Level Education		
Master	138	59.0
PhD	96	41.0
Gred		
DG 41	33	14.1
DG 44	159	67.9
DG 48	36	15.4
DG52	6	2.6
Teaching Experience		
1 - 10 years	63	26.9
11 - 20 years	150	64.1
21 - 30 years	21	9.0
Working Spouse		
Yes	192	82.1
No	42	17.9
Children		
0 - 2	125	53.4
3 - 4	82	35.0
5- 6	22	9.4
>6	5	2.1

In general, to conclude the attributes of the sample in this study which consists of N=234 from the HLP programme for public primary and secondary school teachers in year 2018, it can be seen that the majority were married females between the ages of 36 and 40 years old pursuing their studies at a master's level, had been employed for 11-20 years, and have 0-2 children.

4.4 Data Screening

According to Gallagher et al. (2008), data scanning is needed to prevent issues later in the study and to increase the model result. It is critical in any research to have an appropriate model and a seamless data analysis procedure. As a result, data screening is expected for this purpose. Furthermore, Pallant (2013) acknowledges that data screening is done to ensure that data has been transferred effectively by looking at unreliable responses and reviewing missing and outlier data to ensure that the data has a normal distribution. Furthermore, the use of multivariate demands that data be error-free, which is often regarded as a crucial first step when dealing with multivariate statistical techniques (Hair et al., 2010). Data screening proceed via SPSS software both descriptive and frequency command to detect any missing and out of range values.

Furthermore, in order to meet the assumptions of AMOS, various approaches to data screening were used, such as identifying missing values, outliers, performing normality checks, linearity, and homoscedasticity, and testing for multicollinearity between constructs.

4.4.1 Missing Data

In a quantitative research, the missing data is a common problem in sample surveys (Tabachnick & Fidell, 2007). According to Dong and Peng (2013), a missing rate of 15 to 20% was typical in psychological and educational studies. Meanwhile, as per Schafer and Graham (2002) and Tabachnick and Fidell (2007), missing values occur where data consist of different codes to signify a lack of response, resulting in partial loss of information that may cause error in the analysis.

Missing data is a major problem in research because it may have an effect on the findings (Sekaran & Bougie, 2010). As a consideration, the researcher directly administered the data completion in this study in order to eliminate and reduce the issue of missing data. This approach was also suggested by Sekaran (2003). In this study, the researcher found nine questionnaires with missing data at the end of the survey session. The distribution revealed that there were 12 cases of missing values from nine questionnaires distributed across the data. Eliminating the incomplete responses or cases with missing values can produce the accurate estimate of standard error, confidence interval and p-values of the analysis (Allison, 2002). Furthermore, based on the suggestions of Hair et al. (2010), when the missing values are still more than 50 percent and the study fulfills the sample size requirement, researchers are advised to delete the case respondents.

However, the researcher has analysed the overall questionnaires with a frequency command in SPSS and found that those nine questionnaires with missing values contained less than five percent of the missing values. To deal with missing values, alternative handling of missing data such as weighted or deletion may be used (Tabachnick & Fidell, 2007). Therefore, the researcher made a decision to treat the missing values by filling those values with the mean values.

Mean substitution provides some advantages. Its main advantage is that it produces internally consistent sets of results true correlation matrices. Another advantages of this method is it do not reduce the sample size. Applying mean imputation replaces all missing values. Since the number of respondents' response to the questionnaire slightly small, hence all the nine incomplete questionnaires were retained for data analysis purposes.

4.4.2 Analysis of Assumptions

The estimation model as well as the statistical predictions were put to the test. There are multivariate normality, outliers, linearity, homoscedasticity and multicollinearity (Hair et al. 2010). It is critical that these assumptions be checked since any deviation could jeopardize the validity of the results.

4.4.2.1 Multivariate Normality

The most fundamental assumption in multivariate analysis is normality, referring to the shape of data distribution for an individual variable and its normal distribution (Hair et al., 2010). To decide if the results followed the normality presumption, the synchronization between the respondent's answers and the data creation in the statistical analysis were analysed.

To measure the normality of study variables involved in this study, the sample size will give a potential influence to the normality. Sample size has the effect of increasing statistical power by reducing sampling error. It results in a similar effect here, in that larger sample sizes reduce the detrimental effects of non-normality.

According to Hair et al. (2010), In small samples of 50 or fewer observations, and especially if the sample size is less than 30 or so, significant departures from normality can have a substantial impact on the results. For sample sizes of 200 or more, however, these same effects may be negligible. Hence, as this study had obtained 234 samples, we could assume that the effect of non-normality to this study was reduced.

The study needs to assess for normality distribution of all items measuring the construct before modeling the structural model and executing SEM. Since SEM employs the parametric statistical approach of modeling, the study needs to assess the normality distribution of all items measuring their respective constructs. According to Awang (2015), Aziz et al. (2016), Yusuf et al. (2017), Awang et al. (2015, 2018) and Mohamad et al. (2016, 2017, 2018, 2019), the study only needs to show that the values of skewness for all items do not depart from normality since the Maximum Likelihood Estimator (MLE) algorithm is robust to skewed data (Awang, 2015; Awang et al., 2018). Thus, skewness values between -2.0 and 2.0 are considered normally distributed if the sample size is less than 200, and skewness values between -3.0 and 3.0 are permissible as normally distributed if the sample size is greater than 200. The skewness value in this study ranges from -0.646 to -2.702 which is acceptable as normality distribution. The evaluation of normality distribution for all items is displayed in Table 4.2.

Table 4.2: The Assessment of Normality for all Measuring Items

Variable	min	max	skew	c.r.	kurtosis	c.r.
Communication	4.000	28.000	-1.307	-8.165	3.170	9.899
Nature_Work	4.000	28.000	-1.522	-9.505	4.581	14.305
Co_Worker	4.000	28.000	-1.296	-8.090	3.412	10.655
Operating_Condition	4.000	28.000	-.849	-5.302	2.356	7.357
Contigent_Rewards	4.000	28.000	-.807	-5.038	1.648	5.144
Fringe_Benefits	4.000	28.000	-.882	-5.506	1.734	5.415
Supervision	4.000	28.000	-1.045	-6.525	1.415	4.419
Promotion	4.000	28.000	-.646	-4.035	.881	2.752
Pay	4.000	28.000	-.947	-5.914	1.802	5.625
Capital	3.000	21.000	-1.475	-9.211	3.949	12.332
Affect	3.000	21.000	-1.523	-9.512	4.025	12.570
Development	3.000	21.000	-1.881	-11.750	5.553	17.339
Career	8.000	42.000	-1.951	-12.187	7.580	23.667
Knowledge & Skill	6.000	42.000	-2.702	-16.873	10.897	34.027
Motivation	6.000	42.000	-2.701	-16.869	10.438	32.594
Multivariate					126.234	42.753

The study concludes that the data distribution does not deviate from normality and complies to the requirements for parametric statistical analysis such as correlation, regression, and structural equation modelling based on the skewness values given in Table 4.2. (Awang et al., 2016; Afthanorhan et al., 2019).

4.4.2.2 Outliers

In concern to SEM analysis, a multivariate next assumption is concerned with outliers. This assumption is vital because it can influence the parameter estimates (Randall & Richard, 1996; Schumacher & Lomax, 1996). As explained by Byrne (2010), outliers are data points that are significantly different from the standard and the majority of the respondents. Outliers, according to Cox (2017), can occur as a result of data entry error.

Besides that, Kline (2005) explained that outliers are cases that have extreme scores on two or more variables. Based on this scenario, the data will be discarded to allow for the creation of valid and usable data in order to achieve successful SEM analysis progress and outcomes. The standardized z-score and Mahalanobis distance can be used to quantify outliers in this process.

In order to measure and detect outliers, the Mahalanobis distance was used. In this study, there were nine questionnaires that had cases of outliers. The details can be referred to in Table 4.3. The cases are considered outliers if the squared Mahalanobis distance value exceeds the critical chi-square value which in this case is 37.697, using an alpha level of 0.001 as suggested by Tabachnick and Fidell (2007). Table 4.3 shows there were nine cases (observation numbers 234, 233, 224, 231, 230, 9, 229, 228 and 82).

Table 4.3: Mahalanobis Distance Value

Observation number	Mahalanobis d-squared	p1	p2
234	115.892	.000	.000
233	89.920	.000	.000
224	55.011	.000	.000
231	53.550	.000	.000
230	50.702	.000	.000
9	42.272	.000	.000
229	42.107	.000	.000
228	39.648	.001	.000
82	39.626	.001	.000

Before making any decision in deleting, the researcher examined based on the Cook's Distance value in identifying cases (Pallant, 2013), whether those cases have any unjustified influence on the results. The cases with Cook's Distance values of larger than 1 found a potential problem (Tabachnick & Fidell 2007). The value of Cook's Distance in Table 4.4 shows its maximum value is 0.350 (less than 1). Therefore, the nine cases were retained because there are no major problems (Pallant 2013).

Table 4.4: Cook's Distance Table

	Minimum	Maximum	Mean	Std. Deviation	N
Cook's Distance	0.000	0.350	0.006	0.027	234

4.4.2.3 Linearity and Homoscedasticity

The residual analysis that emerged from the regression analysis was used to test linearity. Furthermore, the findings of the homoscedasticity test, as shown by scatter plot diagrams of standardized residuals, revealed the presence of homoscedasticity in the set of independent variables and the variance of the dependent variable. According to Pallant (2011), the linearity and homoscedasticity assumptions are tested by analysing a scatterplot of the standardised residuals. Figure 4.1 shows that the scatterplot, linearity, and homoscedasticity criteria have been fulfilled (Pallant, 2011) since the scores are clustered in the center (along with the 0 points). As a result, this data achieves linearity and homoscedasticity.

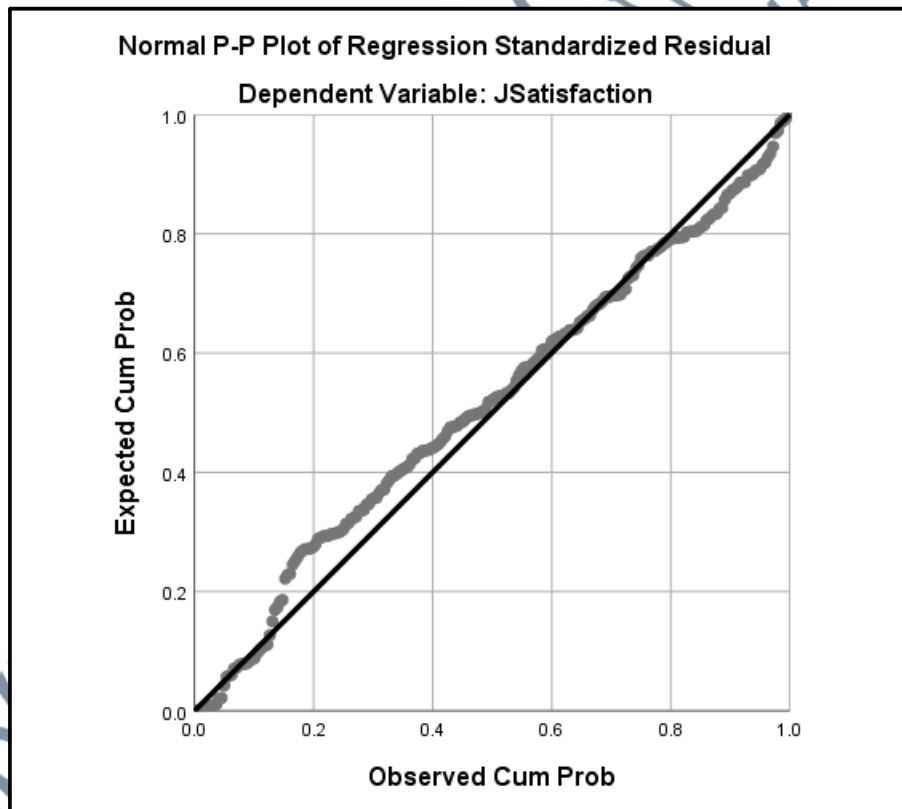


Figure 4.1: Scatterplot of the Standardized Residuals

4.4.2.4 Multicollinearity

The structural model for collinearity is important to be examined before assessing the structural model. Hair et al. (2010) explained that it appears when there are high correlations among predictor variables, which leads to inconsistent and unreliable assessments.

Multicollinearity, means that the component variable items are nearly identical (Zainudin, 2015). SEMs are an effective and powerful tool for dealing with multicollinearity in groups of predictor variables. Multicollinearity occurs when two or more variables are not independent of one another. It is a matter of degree which can be diagnosed. When variables are used as predictors and their interdependence is strong enough, model effects are insufficient and inaccurate. The significance of multicollinearity is to specify the correlation value between variables, which must not exceed 0.90 and must be free of redundant items (Zainudin, 2015).

Table 4.5: Correlation Matrix

	Lifelong Learning	Work Family Enrichment	Job Satisfaction
Lifelong Learning	-		
Work Family Enrichment	.834	-	
Job Satisfaction	.649	.652	-

Multicollinearity happens when the correlation matrix between any two variables is incredibly high; if it is greater than 0.90, each variable is redundant (Zainudin, 2015).

Table 4.5 shows that lifelong learning, work-family enrichment and job satisfaction correlate between values 0.649 to 0.834. Table 4.5 reveals that the correlation between variables is less than 0.90 and that there are no redundant items between variables.

4.5 Factor Analysis Results and Interpretations

The aim of this study is to investigate the relationship between the independent variables and the dependent variables, as well as the mediating effect of work-family enrichment on the relationship between the independent variables and the dependent variables. Details in the following section are on the use of Confirmatory Factor Analysis (CFA) and Exploratory Factor Analysis (EFA) to analyse the relationship between the sample variables. In particular, the positive and negative effects, correlation, loading, as well as the significance of the variables and the relationship between them, are all explained in detail.

Factor analysis can be divided into two types which are exploratory factor analysis (EFA) and confirmatory factor analysis (CFA). Exploratory factor analysis (EFA) is used to determine the existence of the constructs influencing a group of responses (DeCoster, 1998). Meanwhile, confirmatory factor analysis (CFA) is used to determine whether a given group of constructs influences responses in the intended way. According to Osborne and Costello (2009), the most widely used factor analysis is oblique rotation, which involves the use of principal axis factoring, screen plots, and various test runs to evaluate the number of meaningful variables in a dataset. In this study, before performing the analysis, a variety of graphs and tests had to be completed. Among the first tests to be run are the KMO test and the initial eigenvalues for the items.

Firstly, KMO and Bartlett's Test were carried out. The KMO test informs whether enough items are the predictor by each factor. The KMO measures the sampling sufficiency and adequacy which varies less than 1 and more than 0. Varies closer to 1 value is considered as accurate for an acceptable factor analysis accuracy.

Based on Kaiser (1974), KMO value of 0.50 is the minimum accepted, values between 0.70 and 0.80 are acceptable, and a value above 0.90 is superb.

Table 4.6: KMO and Bartlett's Test

KMO and Bartlett's Test		
Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		.952
Bartlett's Test of Sphericity	Approx. Chi-Square	14952.950
	df	1953
	Sig.	.000

Based on Table 4.6, it shows the KMO and Bartlett's Test results for all constructs involved in this study. The KMO is 0.952, which is higher than the 0.5 cutoff value (Kline, 1994; Tabachnick & Fidell, 2007; Hair et al., 2010; George & Mallery, 2001). As a matter of fact, the KMO value denotes a high degree of adequacy of the items in evaluating the study's variables. Meanwhile, for the Bartlett test according to Osborne (2014), the significant value should be less than 0.05. The Bartlett test value for this study is less than 0.05, hence indicating that the variables are correlated highly enough to provide a reasonable basis for factor analysis in this case.

4.6 Exploratory Factor Analysis (EFA)

An Exploratory Factor Analysis (EFA) plays a vital role in this study to examine the interrelationship among the items of lifelong learning, work-family enrichment and job satisfaction which are used to reveal the cluster of items that have adequate ordinary variation to justify their grouping together as a factor. In significance, this process compresses a group of items into a smaller set of combination factors for lifelong learning (motivation, knowledge and skill, career development), work-family

enrichment (development, affect, capital) and job satisfaction (pay, promotion, supervision, fringe benefit, contingent rewards, operating condition, co-worker, nature work, communication) with a minimum loss of information, and hence laid the foundation of structural equation modelling (Hair et al., 2006).

In this study, an exploratory factor analysis (Principal Component Analysis) with Varimax rotation was conducted on the items for the constructs of lifelong learning, work-family enrichment and job satisfaction. Applying a principal component analysis with Varimax rotation was deemed an appropriate approach for exploring the interrelationship among a set of items. The principal component analysis approach was used in this research to classify (extract) the number of underlying factors. The number of underlying factors was determined using the degree of eigenvalue. The Kaiser-Meyer-Olkin test was used to determine the suitability of data for factor analysis prior to conducting the principal component analysis.

In the Measure of Sampling Adequacy (KMO) and Bartlett's Test of Sphericity value, if the outcomes follow the requirements of both tests, the next step is to decide the number of factors that should be used to better reflect the set of items' interrelationships. After determining the number of factors, these factors are rotated using the Varimax rotation to evaluate the loading pattern of each item on the factors. To achieve easier and more interpretable factor solutions, the Varimax Rotation method was used (Hair et al., 2017).

According to Hair et al. (2010), three guidelines that must be utilised in factor analysis, that are: (1) identifying factors that have eigenvalue greater than 1, (2) items should have factor loading greater than 0.50 which is considered necessary for practical significance, and (3) no item cross-loading greater than 0.50. An item that has factor loading less than 0.5 must be deleted from the analysis and at least there must be

minimum three items in a factor (Hair et al., 2010). The sub-sections that follow include a thorough discussion of the EFA result for the constructs of lifelong learning, work-family enrichment and job satisfaction.

4.6.1 Lifelong Learning

The lifelong learning construct consists of three dimensions known as motivation, knowledge and skill, and career development. Below are the results of EFA for each dimension, followed by the EFA result for the lifelong learning variable.

The Kaiser-Meyer-Olkin Measure of Sampling Adequacy (KMO) value for three dimensions, namely motivation, knowledge and skill, and career development, was greater than the prescribed value of 0.6, as seen in Table 4.7. The Bartlett's Test of Sphericity χ^2 , $p < 0.001$, was also statistically significant. Both findings suggest that the data obtained are ideal for factor analysis. The principal component analysis result (Table 4.7) showed the existence of only one component with an eigenvalue greater than one, explaining more than 50.00 percent of the variance.

All of the items have factor loadings greater than the required significant loading of 0.50. Additionally, the three dimensions demonstrate adequate consistency with Cronbach's alpha values of motivation, knowledge and skill, and career development greater than 0.5 (Hair et al., 2006) for the items to obtain internal reliability (Table 4.7). This demonstrates that the three dimensions evaluated the same underlying construct, and therefore those dimensions were maintained for confirmatory factor analysis.

Table 4.7: The Summary Result for Exploratory Factor Analysis (EFA) – Lifelong Learning

No	Items	FACTOR		
		1	2	3
Variable: LIFELONG LEARNING				
Motivation				
LLB1	I learn because I can enjoy the hobbies and daily activities.	.716		
LLB2	I learn because I can contribute to the society.	.909		
LLB3	I learn because I can understand myself better.	.830		
LLB4	I learn because I can help others.	.887		
LLB5	I learn because I feel that I am a self - learner.	.749		
LLB6	I learn because I love learning for its utmost importance.	.788		
Knowledge and Skills				
LLC1	I learn because I can improve my working skills to facilitate my work.		.870	
LLC2	I learn because I can know what is happening around the world.		.773	
LLC3	I learn because I can get along well with others.		.807	
LLC4	I learn because I can relate something to my next generations.		.893	
LLC5	I learn because I can enhance basic daily skills, namely reading, writing and counting.		.808	
LLC6	I learn because when I learn something new, I will try to focus in depth, rather than look it as an overall picture.		.722	
Career Development				

No	Items	FACTOR		
		1	2	3
LLD1	I learn because I can obtain a degree or certificate for career development or getting a side income.			.551
LLD2	I learn because I am able to manage my career well.			.818
LLD3	I learn because for my personal and spiritual development.			.810
LLD4	I learn because for fun and enjoyment learning something new for my career development.			.843
LLD5	I learn because I prefer to have other plans on my learning.			.576
LLD6	I learn because I always think on my own learning and how to improve it.			.733
Eigenvalue		3.995	3.997	3.209
Percentage of Variance (%)		66.585	66.283	53.485
KMO Measure of Sampling Adequacy		.884	.895	.789
Approximate Chi-Square		863.721	792.871	534.233
Sig		.000	.000	.000
Cronbach Alpha		.860	.894	.805

4.6.2 Work Family Enrichment (WFE)

The work-family enrichment consists of three dimensions known as development, affect, and capital. Below are the results of EFA for each dimension, followed by the EFA result for the work-family enrichment variable

The three dimensions which are development, affect, and capital, as referred in Table 4.8, shows the Kaiser-Meyer-Olkin Measure of Sampling Adequacy (KMO) value exceeded the recommended value of 0.6. The Bartlett's Test of Sphericity χ^2 , $p < 0.001$. It was denoted that both results reached statistical significance and the collected data was appropriate for factor analysis. The principal component analysis result (Table 4.8) showed the existence of only one component with an eigenvalue greater than one, explaining more than 50.00 percent of the variance respectively.

The Cronbach's alpha value for the three dimensions, which are development, affect, and capital, must be greater than 0.5 (Hair et al., 2006) for the items to achieve internal reliability. Moreover, the factor loading of all the items beyond the minimum significant loading 0.50. This demonstrates that the three dimensions are evaluating the same underlying construct, and hence the three dimensions are maintained for confirmatory factor analysis.

Table 4.8: The Summary Result for Exploratory Factor Analysis (EFA) – Work Family Enrichment

No	Items	FACTOR		
		1	2	3
Variable: WORK FAMILY ENRICHMENT				
Development				
WFEA1	Helps me to understand different opinions and be a good member in the family.	.953		
WFEA2	Helps me to gain knowledge and be a good member in the family.	.959		
WFEA3	Helps me to acquire skills and be a good member in the family.	.949		
Affect				
WFEB1	Puts me in a great mood and be a good member in the family.		.957	
WFEB2	Makes me happy, thus helping me be a good member in the family.		.969	
WFEB3	Making me a cheerful person and be a good member in the family.		.950	
Capital				
WFEC1	Helps me to be a fulfilled person and be a good member in the family.			.947
WFEC2	Provides me with a sense of accomplishment and this helps me be a better family member.			.965
WFEC3	Provides me with a sense of success and this help me be a better family member.			.934
Eigenvalue		2.729	2.756	2.699
Percentage of Variance (%)		90.964	91.882	89.957
KMO Measure of Sampling Adequacy		.774	.766	.749
Approximate Chi-Square		694.949	754.829	671.621
Sig		.000	.000	.000
Cronbach Alpha		.950	.955	.944

4.6.3 Job Satisfaction

The job satisfaction construct consists of nine dimensions known as pay, promotion, supervision, fringe benefit, contingent rewards, operating condition, co-worker, nature work, and communication. Below are the results of EFA for each dimension, followed by the EFA result for the job satisfaction variable.

The nine dimensions, which are pay, promotion, supervision, fringe benefit, contingent rewards, operating condition, co-worker, nature work, and communication, as referred in Table 4.9 shows the Kaiser-Meyer-Olkin Measure of Sampling Adequacy (KMO) value exceeded the recommended value of 0.6 and the Bartlett's Test of Sphericity χ^2 , $p < 0.001$. It was denoted that both results reached statistical significance and the collected data was appropriate for factor analysis. The principal component analysis result (Table 4.9) highlights the existence of only one component with an eigenvalue greater than one, indicating more than 50.00 percent of the variance.

All of the items have factor loadings greater than the required significant loading of 0.50, and three items were removed because of low factor loading. Three items were removed (JSB3, JSD4, JSF4). Then, the nine dimensions showed satisfactory consistency with a Cronbach's alpha value of pay, promotion, supervision, fringe benefit, contingent rewards, operating condition, co-worker, nature work, and communication which is the value must be greater than 0.5 (Hair et al., 2006) for the items to attain internal reliability (Table 4.9). This demonstrates that the nine elements assessed the same underlying construct, and therefore the nine elements were held for confirmatory factor analysis.

Table 4.9: The Summary Result for Exploratory Factor Analysis (EFA) – Job Satisfaction

No	Items	FACTOR								
		1	2	3	4	5	6	7	8	9
Variable: JOB SATISFACTION										
Pay										
JSA1	I feel that I am paid with the salary based on the work that I do.	.874								
JSA2	Work promotions are fair and enough.	.890								
JSA3	I feel appreciated by this organisation when I look at the payment that I get.	.816								
JSA4	I am satisfied with the opportunity to get a better salary.	.873								
Promotion										
JSB1	Opportunities for me to get a job promotion are high.		.917							
JSB2	Whoever performs well in his or her job will get a better opportunity for a promotion.		.901							
JSB3	Others will advance quickly to another section.		Removed (0.463)							

No	Items	FACTOR								
		1	2	3	4	5	6	7	8	9
JSB4	I am satisfied with the opportunities in job promotion.		.917							
Supervision										
JSC1	My supervisor is very competent during work.			.885						
JSC2	The supervisor is being fair to me.			.955						
JSC3	My supervisor shows much interest in the feelings of his or her subordinates.			.943						
JSC4	I like my supervisor.			.909						
Fringe Benefits										
JSD1	I am happy with the benefits that I get.				.907					
JSD2	The benefits received is like the ones offered by other organisations.				.903					
JSD3	The benefits package is fair.				.942					
JSD4	There are other benefits that we do not have, of which we should have.				Removed (0.487)					
Contingent Rewards										
JSE1	When I do my job, I am rewarded with the right appreciation that I supposed to receive.					.907				

No	Items	FACTOR								
		1	2	3	4	5	6	7	8	9
JSG3	I am happy with my working colleagues.							.928		
JSG4	There are too little of bickering and fighting at work.							.701		
Nature Work										
JSH1	I feel that my work is meaningful.								.923	
JSH2	I like doing things that I do at work.								.904	
JSH3	I am proud of doing my job.								.922	
JSH4	My work is fun.								.886	
Communication										
JSI1	There is good communication in this organisation.									.892
JSI2	The goal for this organisation is clear to me.									.946
JSI3	I always feel that I know what is going on in the organisation.									.910
JSI4	My job task is fully explained.									.913
Eigenvalue		2.986	2.492	3.412	2.524	3.193	1.912	2.625	3.304	3.352
Percentage of Variance (%)		74.640	83.074	85.300	84.133	79.826	63.746	65.624	82.596	83.795
KMO Measure of Sampling Adequacy		.833	.750	.840	.733	.849	.590	.710	.860	.857

No	Items	FACTOR								
		1	2	3	4	5	6	7	8	9
Approximate Chi-Square		514.655	424.974	914.226	466.657	661.379	178.899	436.208	748.279	808.272
Sig		.000	.000	.000	.000	.000	.000	.000	.000	.000
Cronbach Alpha		.886	.800	.942	.905	.915	.714	.797	.928	.933

4.7 Structural Equation Modeling (SEM)

In recent years, Structural Equation Modeling (SEM) as an analysis software has grown in popularity among researchers. This study used the structural equation modeling (SEM) statistical package AMOS 23 software and IBM SPSS statistic 23. Through SEM, the research's input and population are generalizable on the status concerning the study domain. Besides that, Zaefarian et al. (2013) stated that SEM can offer statistical assistance in the resolution of problems and issues in a wide range of industries. Moreover, if the aim is to achieve accurate results and outcomes, using more than one method of data analysis is permissible (Zabkar, 2010; Zaefarian et al. 2013).

SEM is also able to determine any measurement errors and test the measurement model using the EFA and structural model simultaneously through a path analysis (Kline, 2010). This became the main reason of choosing SEM compared to factor analysis or multiple regression analysis in SPSS. In terms of visual ease, model viewing, option for modification, and quality graphics for publication purposes, SEM is the best over other multivariate analysis techniques.

Furthermore, the SEM was employed to analyse the latent constructs and perform analysis on the causal links between latent constructs. Besides, the SEM is efficient when applied to analysis, including (i) estimate variance and covariance, (ii) test hypotheses, (iii) conventional linear regression, and (iv) confirmatory factor analysis (Jöreskog, 1993). Simultaneously, the SEM will evaluate the unidimensionality, reliability, and validity of every individual construct (Hair et al., 2010; Kline, 2010). The SEM executes an overall test of model fit as well as individual parameter estimation tests at the same time, allowing for the best model fit to the data.

In the meantime, Gefen et al. (2000) claimed that SEM is a set of multivariate statistical methods used to examine the direct and indirect relationships between one or more independent latent variables and one or more dependent variables. Furthermore, in addition to evaluating the structural model, SEM also tests the overall fit of a model (Gefen et al. 2000; Hai). Ringle et al. (2012) agreed that SEM has the ability to assess both the measurement model as well as structural model. The researcher used SEM to evaluate both the hypothesised structural linkages among variables and the linkages that exist between a variable and its corresponding measures.

As per Byrne (2016), SEM allows statistical testing of a hypothesised model in a simultaneous analysis of the whole model for the purpose of evaluating how well the model is compatible with the data. Furthermore, Hoyle (2012) stated that SEM is a broad statistical technique that researchers may use to test hypotheses about the relationships between observed and latent variables. Awang (2015) has mentioned that SEM has proven to be able to overcome limitations in previous approaches, especially for simultaneous interrelationship analysis among constructs in a model.

4.8 The Confirmatory Factor Analysis (CFA)

The study adopted the two-step approach of modelling and analysing the structural model, namely Confirmatory Factor Analysis (CFA) and Structural Equation Modeling (SEM). Thus, prior to modelling the structural model and executing Structural Equation Modeling (SEM), the study needed to validate all the measurement models of latent constructs for Unidimensionality, Validity, and Reliability (Awang, 2014, 2015). This validation procedure is called Confirmatory Factor Analysis (CFA). According to Asnawi et al. (2019), the measurement model of latent constructs needs

to pass three types of validity, namely Construct Validity, Convergent Validity, and Discriminant Validity.

As for reliability, it is adequate for the study to assess Composite Reliability (CR) since it replaced the traditional method of computing the Cronbach's alpha for analysis using Structural Equation Modeling (SEM) (Awang et al., 2018). This particular latent construct is considered valid if its' fitness indexes achieved the three Model Fit categories, namely Absolute Fit, Incremental Fit and Parsimonious Fit (Awang et al., 2018). The fitness indexes and their respective threshold values are given in Table 4.10.

Table 4.10: The three categories of model fit and their level of acceptance

Name of Category	Name of Index	Level of Acceptance
Absolute Fit Index	RMSEA	RMSEA < 0.1 (ideal < 0.08)
	GFI	GFI > 0.85 (ideal if > 0.90)
Incremental Fit Index	AGFI	AGFI > 0.85 (ideal if > 0.9)
	CFI	CFI > 0.85 (ideal if > 0.9)
	TLI	TLI > 0.85 (ideal if > 0.9)
	NFI	NFI > 0.85 (ideal if > 0.90)
Parsimonious Fit Index	Chisq/df	Chi-Square/ df < 5.0 (ideal if < 3.0)

***The indexes in bold are recommended since they are frequently reported in literatures (Source: Awang et al. (2018)).

The framework for this study consists of one exogenous construct, one mediator constructs and one endogenous construct. The theoretical framework for this study and the path of interest where the hypotheses to be tested is presented in Figure 4.2.

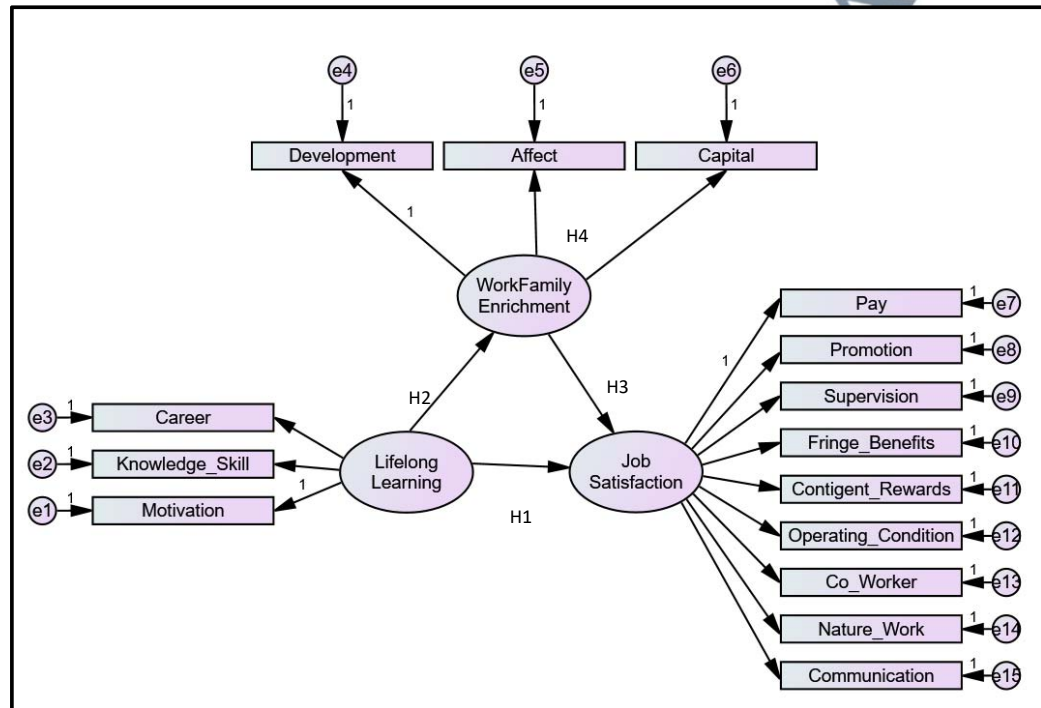


Figure 4.2: The Research Framework Showing the hypothesis to be tested in the Study

The hypotheses statement for every path in Figure 4.2 and the method of statistical analysis to be employed are listed in Table 4.11 and Table 4.12 respectively.

Table 4.11: The Direct Effect Hypothesis and Method of Analysis

	Direct Effect Hypothesis	Method of Analysis
H1	Lifelong learning significant affect teachers' job satisfaction.	Path Analysis in SEM
H2	The lifelong learning is significantly affect work family enrichment.	Path Analysis in SEM
H3	Work-to-family enrichment (WFE) is positively related to job satisfaction.	Path Analysis in SEM

Table 4.12: The Hypothesis Testing for Mediators and Method of Analysis

	The Hypothesis for testing Mediators	Method of Analysis
H4	Work-to-family enrichment mediates the relationship between lifelong learning and teachers' job satisfaction.	Path Analysis in SEM and Bootstrapping

Every construct involved in the framework (Figure 4.2) is measured using a certain number of items in the questionnaire resulted from the Exploratory Factor Analysis (EFA) procedure based on data from the pilot study (Bahkia et al., 2019; Shkeer & Awang, 2019; Rahlin et al., 2019,). The items measuring every construct are presented in Figure 4.3.

This study adopted the two-step approach of Structural Equation Modelling (SEM) as proposed by Awang et al. (2018). The two analysis steps are:

- i. **The measurement models.** The measurement models of all latent constructs need to be validated for validity and reliability through the Confirmatory Factor Analysis (CFA) procedure.

- ii. **The structural model.** Once the CFA is complete, the study develops the structural model. The Structural Equation Modelling (SEM) procedure will be employed to estimate the parameters involved and test the hypothesis regarding interrelationships among the constructs in the model.

Both procedures (CFA and SEM) are carried out using IBM-SPSS-AMOS 23.0. Figure 4.3 illustrates the research framework, which indicates that all constructs in the model are assessed using a certain number of measuring items in the questionnaire. Each item is scored on an interval scale ranging from 1 (strongly disagree) to 7 (strongly agree) with the given statement (Awang et al., 2016; Bahkia et al., 2019; Rahlin et al., 2019). This study decided to conduct the CFA procedure at once for all constructs involved in the study. This particular method of CFA is called Pooled-CFA for the measurement model.

Prior to modeling the structural model and executing SEM, the researcher needs to prove the validity and reliability for all constructs. The measurement model of the construct would be assessed for construct validity, convergent validity and discriminant validity (Awang et al., 2018; Rahlin et al., 2019; Asnawi et al., 2019). The construct validity would be assessed using a set of fitness indexes, the convergent validity would be assessed using the Average Variance Extracted (AVE), and the discriminant validity would be assessed through the Discriminant Validity index summary (Awang et al., 2015, 2017, 2018; Rahlin et al., 2019 and Afthanorhan et al., 2018, 2019). The discriminant validity or synonym as multi-collinearity assessment is required to ensure no items are redundant and no constructs are highly correlated. For any two exogenous structures that are strongly correlated (correlation greater than 0.85), a huge problem known as Multi-collinearity occurs (Awang, 2015; Awang et al., 2018).

As for reliability assessment, the study computed the Composite Reliability (CR) which is closely associated with Cronbach's alpha in the first-generation method of analysis. The CR was computed based on the factor loading for each measuring item which is more efficient compared to Cronbach's alpha that is computed based on the measurement score for each item. The value for CR should exceed 0.6 for the construct to achieve Composite Reliability (Afthanorhan et al., 2018, 2019; Awang et al., 2018; Rahlin et al., 2019a; Mahfouz et al., 2019; Shkeer & Awang, 2019).

Nowadays, the Pooled-CFA is more efficient and highly suggested method for assessing the measurement model. This method is more preferred since it could address the issue of identification problem (Afthanorhan et al., 2014). Furthermore, Awang (2015) acknowledge that if possible, the measurement model for all constructs involved in a particular study should be assessed together at once. Thus, in this current study, the Pooled-CFA would be executed for the Pooled-Measurement Model presented in Figure 4.4.

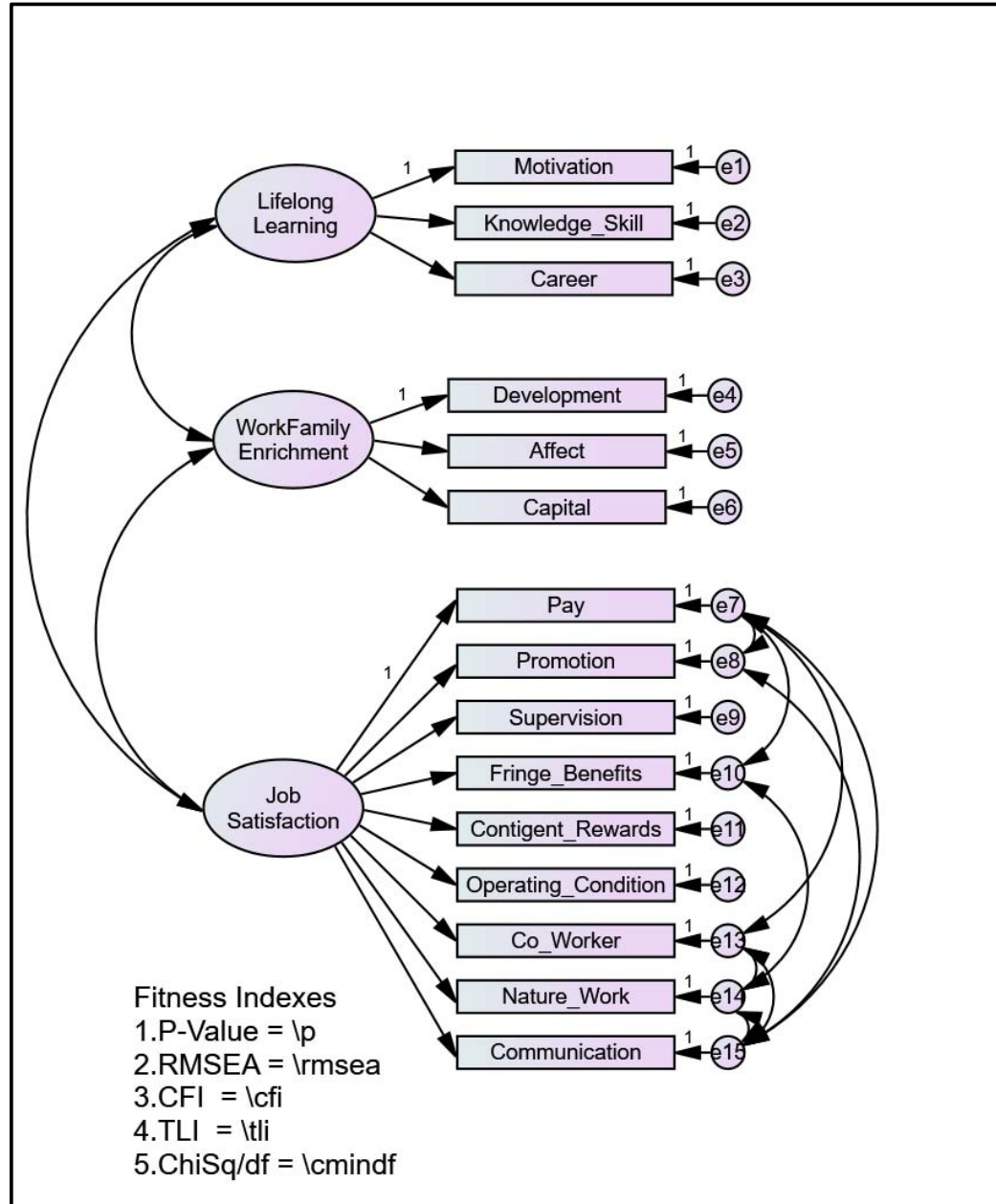


Figure 4.3: The Latent Constructs in the Study and their respective measuring Items

The CFA output for the model in Figure 4.3 is presented in Figure 4.4. The results display the fitness indexes for all constructs (combined) in the model, the factor loading for each measuring item in the model, and the correlation between constructs in the model. The fitness indexes must fulfil the threshold values as seen in Table 4.10, the factor loading for each item must be at least 0.6, and the correlation coefficient of any two constructs must not surpass 0.85. (Afthanorhan et al., 2017, 2019). All dimensions are greater than 0.60.

If the correlation of any two constructs exceeds 0.85, the issue of multi-collinearity arises (highly correlated). In Figure 4.4, none of the correlation values (at the double-headed arrow) between constructs were found to be greater than 0.85. As a result, the multi-collinearity issue is avoided.

After obtaining the pooled-CFA results, the analysis may begin the validation procedure for construct validity, convergent validity, discriminant validity, and composite reliability (Awang et al., 2018; Asnawi et al., 2019).

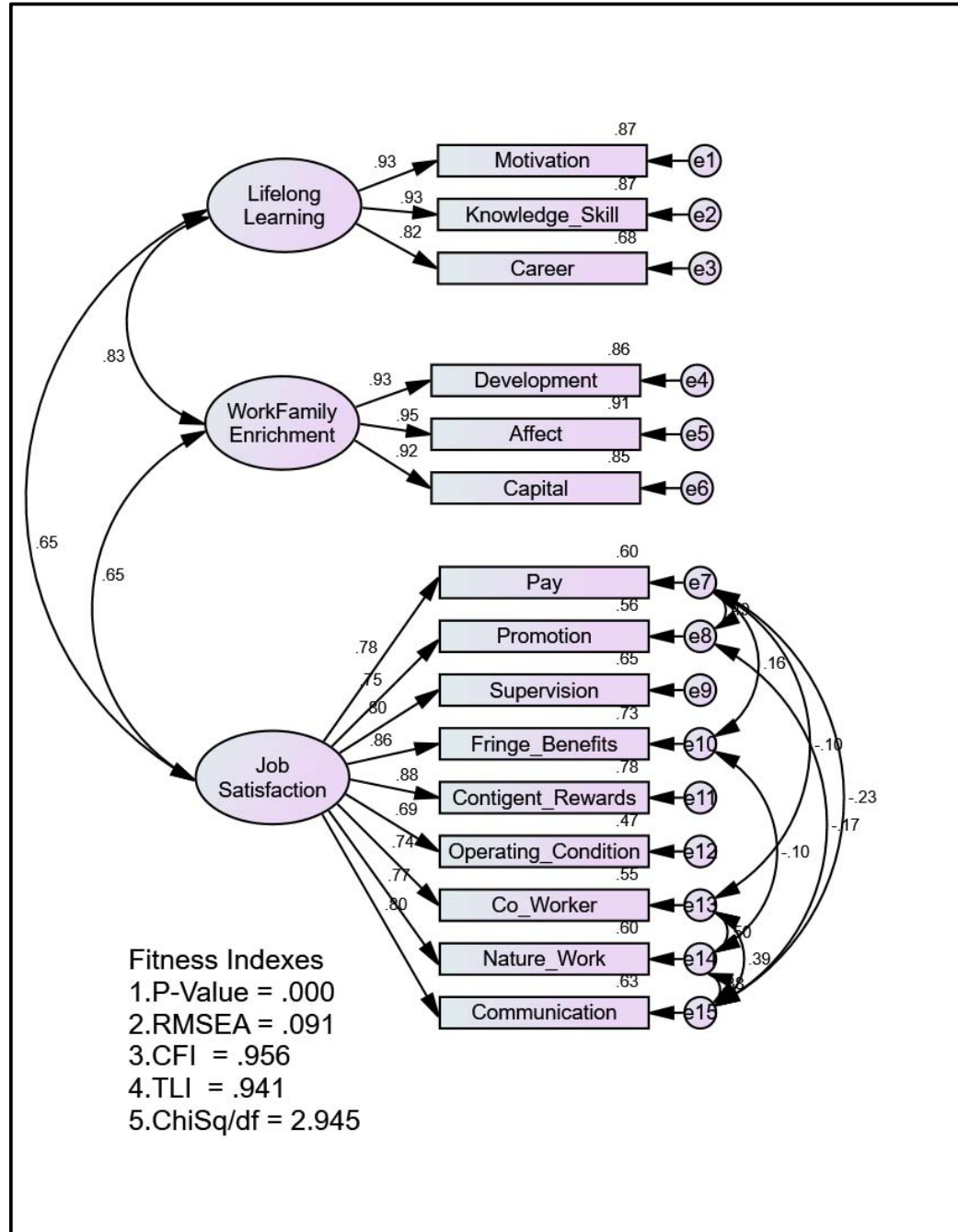


Figure 4.4: The Results for Pooled-CFA for the Measurement Model of Constructs

4.8.1 The Assessment for Construct Validity

The fitness indexes in Figure 4.4 have met the threshold values as stated in Table 4.10. The Absolute Fit category, namely RMSEA, is 0.091 (achieved the threshold of less than 1.00), the Incremental Fit category, namely CFI, is 0.956 (achieved the threshold of greater than 0.85), and the Parsimonious Fit category, namely the ratio of Chisq/df is 2.945 (achieved the threshold of 3.0). Thus, the measurement model of all latent constructs in Figure 4.4 have achieved the requirement for Construct Validity (Awang, 2015; Awang et al., 2018).

4.8.2 The Assessment for Convergent Validity and Composite Reliability

For the assessment of Convergent Validity, the study needs to compute the Average Variance Extracted (AVE). The construct achieved Convergent Validity if its AVE exceeds the threshold value of 0.5 (Awang, 2014; 2015). As for assessing the Composite Reliability, the study needs to compute the CR and its value should exceed the threshold value of 0.6 for this reliability to be achieved (Kashif et al., 2015, 2016). The AVE and CR for all constructs are computed and presented in Table 4.13.

Table 4.13: The Average Variance Extracted (AVE) and Composite Reliability (CR)

Construct	Items	Factor Loading	CR (above 0.6)	AVE (above 0.5)
Lifelong Learning	Motivation	0.93	0.837	0.801
	Knowledge Skill	0.93		
	Career	0.82		
Work Family Enrichment	Development	0.93	0.953	0.871
	Affect	0.95		
	Capital	0.92		
Job Satisfaction	Pay	0.78	0.936	0.620
	Promotion	0.75		
	Supervision	0.80		
	Fringe Benefits	0.86		
	Contingent Rewards	0.88		
	Operating Condition	0.69		
	Co-Worker	0.74		
	Nature Work	0.77		
	Communication	0.80		

With reference to the Average Variance Extracted (AVE) and Composite Reliability (CR) values in Table 4.13, the study found all AVE and CR exceeded their threshold values of 0.5 and 0.6 respectively (Awang et al., 2015, 2018;). Thus, the study can conclude that the Convergent Validity and Composite Reliability for all latent constructs in the model have been achieved.

4.8.3 The Assessment of Discriminant Validity among Constructs

Another kind of validity for the model must be assessed in the analysis, namely discriminant validity. The aim of the discriminant validity assessment is to ensure that there are no duplicate or redundant constructs in the model. When any pair of constructs in the model is strongly or highly correlated, the model has a redundant construct. To assess the discriminant validity, one needs to develop the discriminant validity index summary as shown in Table 4.14. The diagonal values in bold are the square root of the AVE of the respective constructs while other values are the correlation coefficient between the pair of the respective constructs.

Table 4.14: The Discriminant Validity Index Summary for all Constructs

	Lifelong Learning	Work Family Enrichment	Job Satisfaction
Lifelong Learning	.894		
Work Family Enrichment	.834	.933	
Job Satisfaction	.649	.652	.787

Referring to Table 4.14, the discriminant validity of the respective construct is achieved if the square root of its AVE exceeds its correlation value with other constructs in the model (Awang, 2015; Awang et al., 2015, 2018). In other words, the discriminant validity is achieved if the diagonal values (in bold) are higher than any other value in its row and its column. The tabulated values in Table 4.14 meet the threshold of discriminant validity. Thus, the study concludes that the discriminant validity for all constructs in this study is achieved.

4.9 The Structural Model and Structural Equation Modeling (SEM)

Once the CFA process report is finished or done, and all values exceed the necessary validity and reliability thresholds, the researcher will assume that the measurement models for all latent constructs involved in the model have been validated (Awang, 2015; Awang et al., 2018). In the next step, the researcher must then assemble these constructs into the structural model in order to perform Structural Equation Modeling (SEM).

The constructs should be ordered from left to right, starting with the exogenous constructs on the far left, then the mediator constructs in the centre, and finally the endogenous construct on the far right (Awang, 2015; Awang et al., 2018; Afthanorhan et al., 2019).

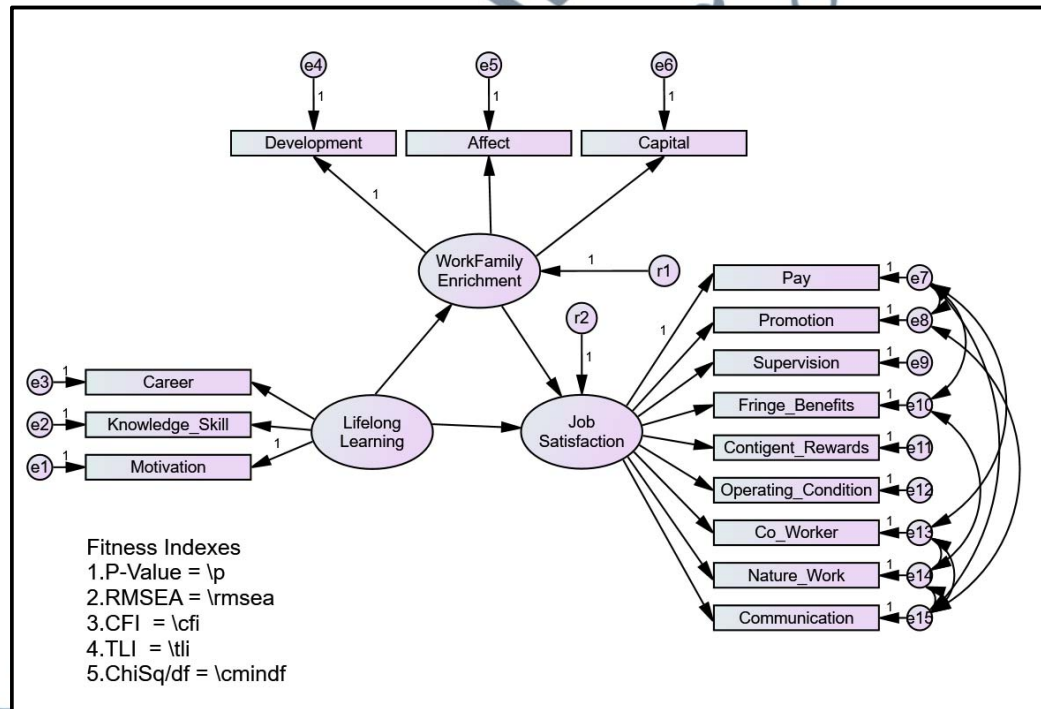


Figure 4.5: The Structural Model for this Study

The output resulted from executing the SEM is given in Figure 4.6 for the estimated Standardized Regression Path Coefficients. Meanwhile, Figure 4.7 displays the estimated Regression Path Coefficients between constructs.

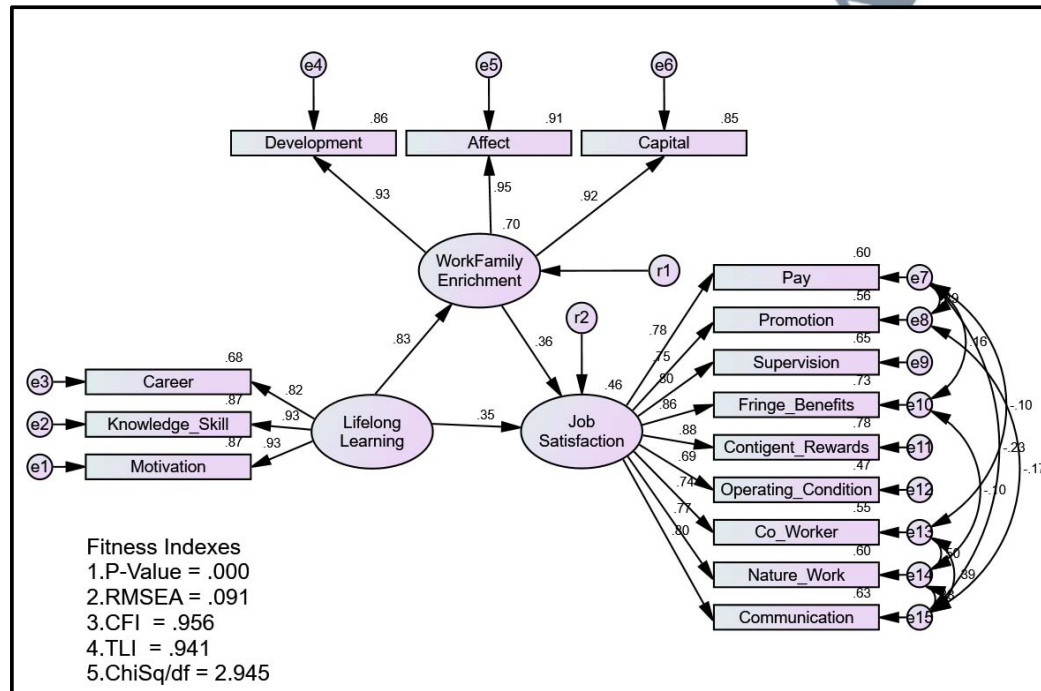


Figure 4.6: The Estimated Standardized Regression Path Coefficient

The explanation and clarification regarding the performance of R^2 (coefficient of multiple determination) of the model (obtained from Figure 4.6) is explained in Table 4.15.

Table 4.15: The Coefficient of Multiple Determination or R^2 and its implication in this study

Endogenous Construct	R^2	Conclusion
Work Family Enrichment	0.70	The lifelong learning manage to explain about 70 percent of the work family enrichment
Job Satisfaction	0.46	The work family enrichment and lifelong learning manage to explain about 46 percent of the job satisfaction

The estimated regression path coefficients for all constructs are presented in Figure 4.7.

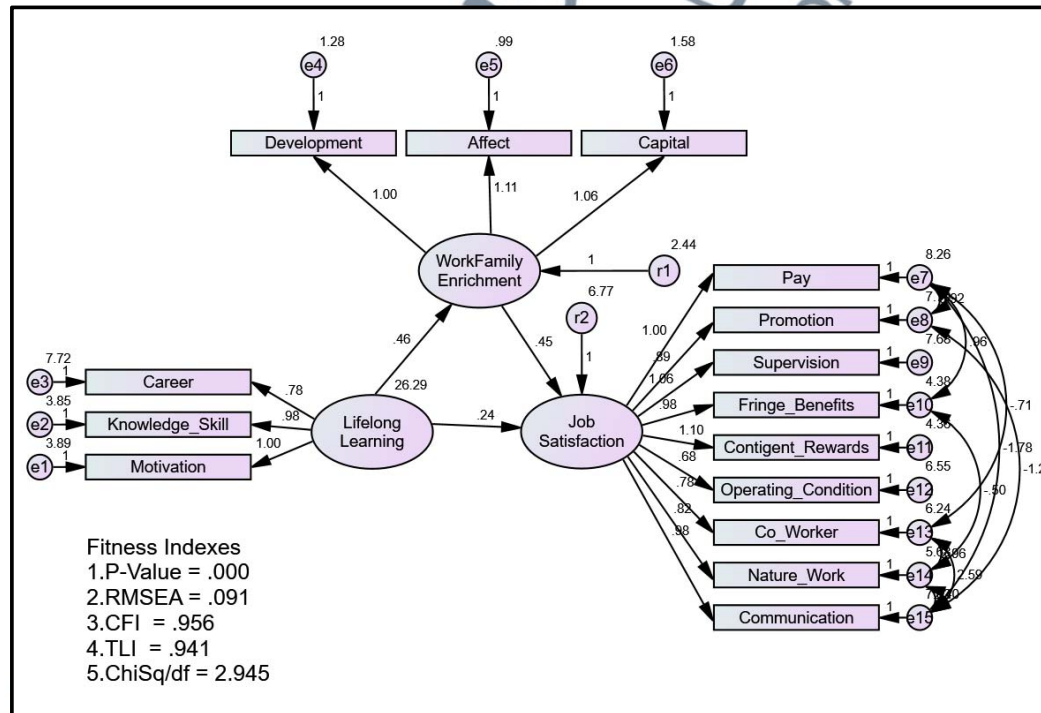


Figure 4.7: The Estimated Regression Path Coefficient between the constructs

Next, the output for estimated regression path coefficient (beta) from the exogenous constructs on endogenous construct extracted from Figure 4.7 is shown in Table 4.16.

Table 4.16: The Regression Path Coefficient obtained from Figure 4.7.

Exogenous	Endogenous	Beta	Conclusion
Lifelong Learning	Work Family Enrichment	.461	When Work Family Enrichment goes up one unit, Lifelong Learning goes up by 0.461 unit
Work Family Enrichment	Job Satisfaction	.454	When Job Satisfaction goes up one unit, Work Family Enrichment goes up by 0.454 unit
Lifelong Learning	Job Satisfaction	.239	When Job Satisfaction goes up one unit, Lifelong Learning goes up by 0.239 unit

The Regression weights or Regression coefficient in the model together with their respective significance is presented in Table 4.17.

Table 4.17: The Regression Weight/ Coefficient and Their Significance

			Estimate	S.E.	C.R.	P	Result
Work Family Enrichment	<---	Lifelong Learning	.461	.028	16.307	.000	Significant
Job Satisfaction	<---	Work Family Enrichment	.454	.138	3.301	.000	Significant
Job Satisfaction	<---	Lifelong Learning	.239	.077	3.103	.002	Significant

Using the results in Table 4.17, the hypothesis is carried out in Table 4.18.

Table 4.18: The Hypothesis Testing for Direct Effect Relationships

	Direct Effect Hypothesis	P	Result
H1	Lifelong learning significantly affect teachers' job satisfaction.	.002	Supported
H2	The lifelong learning is significantly affect work family enrichment.	.000	Supported
H3	Work-to-family enrichment (WFE) is positively related to job satisfaction.	.000	Supported

As seen in Table 4.19, hypothesis testing for the mediation effects of a mediator construct in the model is conducted separately.

Table 4.19: The Hypothesis Testing for Mediators and Method of Analysis

	The Hypothesis for testing Mediators	Method of Analysis
H4	Work-to-family enrichment mediates the relationship between lifelong learning and teachers' job satisfaction	Path Analysis in SEM and Bootstrapping

The fourth hypothesis testing is carried out in Table 4.20 and Figure 4.8

Table 4.20: Testing the Mediator in the Lifelong Learning – Work Family Enrichment – Job Satisfaction

	The Hypothesis for testing a Mediator
H4	Work-to-family enrichment mediates the relationship between lifelong learning and teachers' job satisfaction

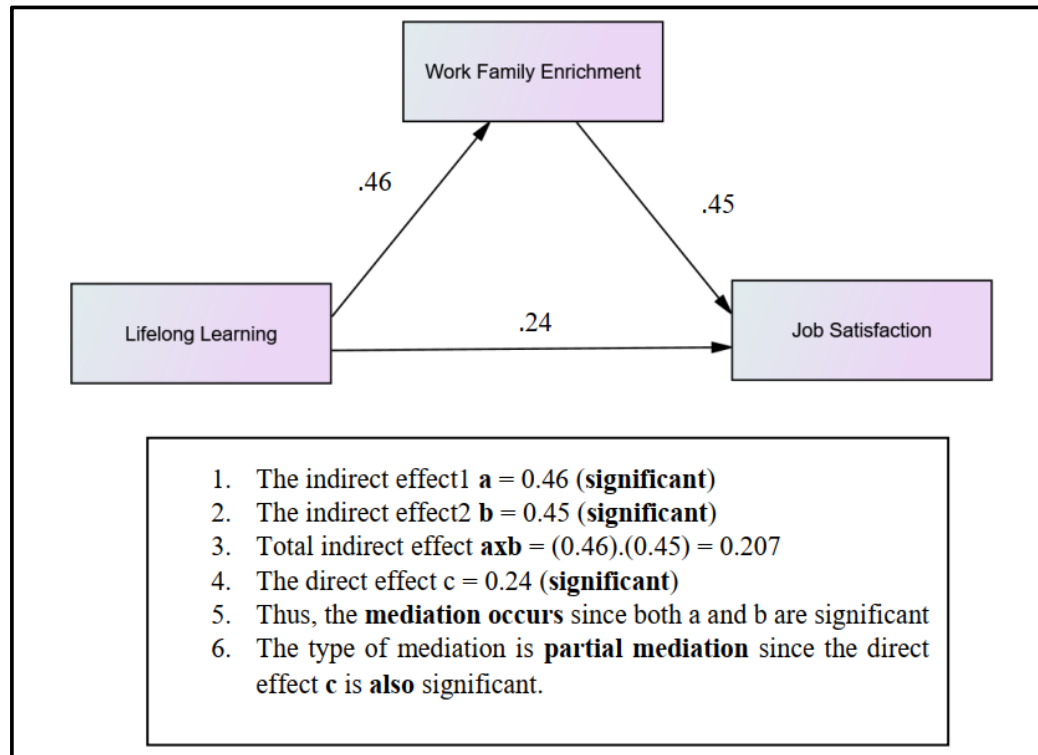


Figure 4.8: Testing the Following Relationship Lifelong Learning – Work Family Enrichment – Job Satisfaction

In order to reconfirm the above mediation test results, this study employed the ML (Maximum Likelihood) bootstrapping procedure using 1000 bootstrap samples with percentile confidence interval 95% and bias-corrected confidence interval 95%. The result is presented in the following table. The bootstrap results are summarized in Table 4.21.

Table 4.21: The Bootstrapping Procedure for Confirming Mediation Test in Figure 4.9.

	Indirect Effect	Direct Effect
Bootstrapping Results	0.506	0.589
Bootstrapping P-Value	0.010	0.021
Result	Significant	Significant
Mediation Type	Partial Mediation since the direct effect is significant	

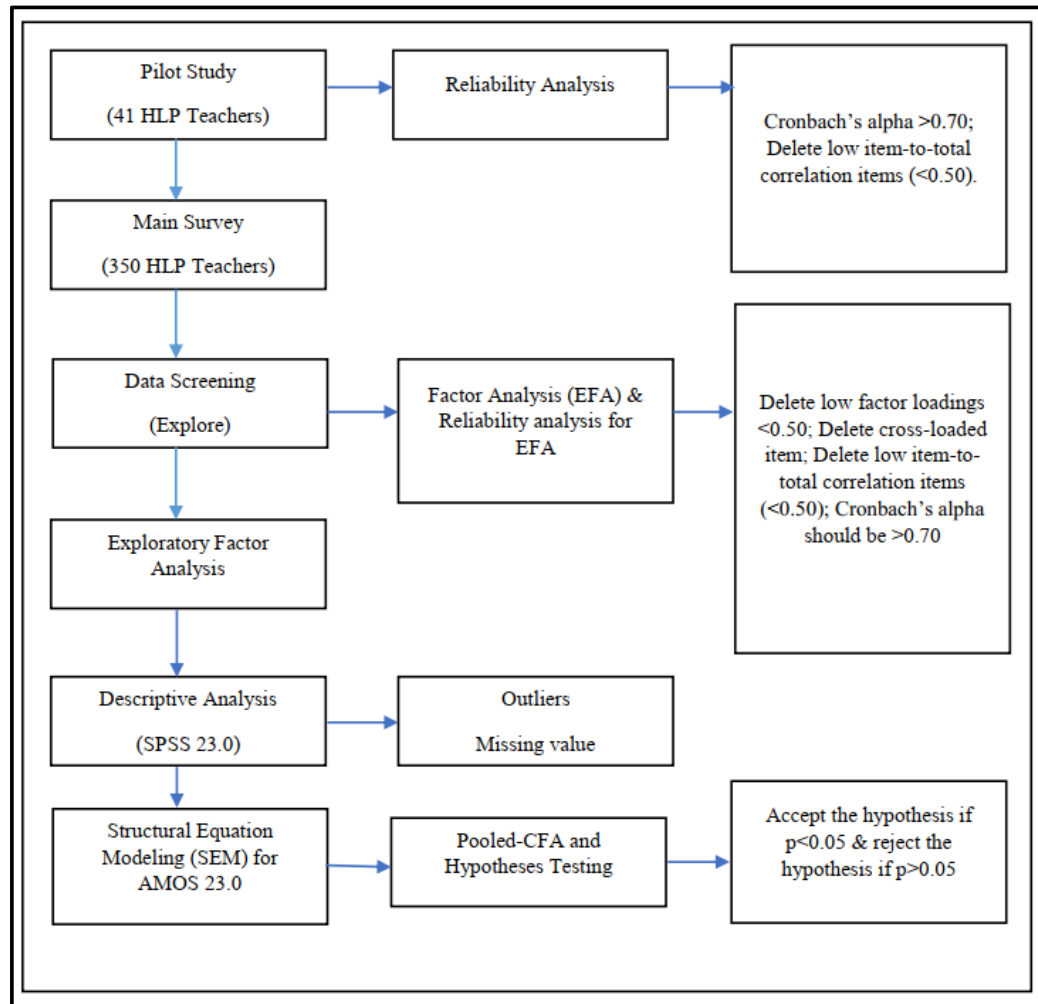


Figure 4.9: Summary for Data Preparation, Screening and Analysis

4.10 Summary

This study examined the roles of work-family enrichment in mediating the relationship between lifelong learning and job satisfaction. Accordingly, the data were gathered from public universities in Malaysia among the teachers in public primary and secondary schools participating in the 'Hadiah Latihan Persekutuan' (HLP) programme.

The study proposed four goals, four research questions, and four theories to be evaluated and tested. The analyses were carried out, and the results provided answers to the questions, allowing the researcher to meet the defined objectives of the study. Tests were performed and conducted to confirm the approval or rejection of the hypotheses, while the conceptual framework's property was also equally determined.

This chapter explored the respondent's attributes or characteristics as well as certain basic values about the data to guarantee the research's validity and reliability. The details on data cleaning to achieve optimum data that could promote the purpose of this research were then provided. Following that, EFA was performed to ensure that the items accurately represented the variables. Certain questionnaire items were removed as needed. This was done to ensure that the variables suit the proposed model.

CFA was performed to improve the model's goodness of this study. In the end, items were edited to conform to the research requirements and to fit with the construct model analysis. The hypotheses were tested and supported by the results. In addition, the effect of the mediator variable, which is work-family enrichment, was assessed and verified. The following chapter will summarize the study's findings and results. This is followed by a detailed discussion on the study's recommendations as well as contributions.