

연구논문

## 구조방정식모형을 이용한 청소년들의 안녕감에 대한 분석

### Analysis of Adolescents' Well-being Using Structural Equation Models

오영창<sup>a)</sup> · 박은식<sup>b)</sup> · 윤경희<sup>c)</sup>

Youngchang Oh · Eunsik Park · Kyounghee Yun

일반적으로 구조방정식모형은 다변량 정규분포를 가정한 최대우도법을 가장 많이 사용한다. 본 연구는 청소년들의 안녕감에 대한 구조방정식모형에서 정규성 가정에 위배되는 경우의 분석방법에 대해서 살펴보았다. 비정규 자료에 대한 구조방정식모형의 분석방법으로는 점근분포무관법(ADF)과 Satorra-Bentler scaled 통계량과 Bollen-Stine 부트스트랩에 의한 방법 등이 있다. 분석한 자료는 청소년들의 안녕감에 대한 자료로서 구조방정식모형의 적합도 지수를 전체 표본과 이를 무작위 추출한 표본으로 분석하였다. 표본의 크기에 따른 적합도 지수를 비교함으로써 정규성 가정이 필요한 최대우도법과 비정규성 자료에 대한 여러 방법들 중 어떠한 방법이 가장 적절한지를 살펴보았다.

**주제어:** 구조방정식모형, 점근분포무관법, Satorra-Bentler scaled 검정통계량, Bollen-Stine 부트스트랩, 비정규성

In general, the structural equation modeling uses the maximum likelihood method assuming multivariate normality. This study have investigated methods for the structural equation models of adolescents' well-being, when the normality assumption is violated, such as asymptotic

a) 전남대학교 통계학과 석사.

b) 교신저자(corresponding author): 전남대학교 통계학과 교수 박은식.  
E-mail: [espark02@jnu.ac.kr](mailto:espark02@jnu.ac.kr)

c) 전남대학교 생활환경복지학과 박사과정

distribution free method, analyses using Satorra-Bentler scaled statistic or Bollen-Stine bootstrap method. Observed samples and their random selections have been analysed in terms of goodness of fit indices of structural equation modeling. Based on these indices, we have recommended what method is the most suitable among several methods developed for non-normally distributed data as well as the maximum likelihood method.

Key words: structural equation model, asymptotic distribution free(ADF) method, Satorra-Bentler scaled test statistics, Bollen-Stine bootstrap, non-normality

## I. 서론

구조방정식모형(Structural Equation Model, SEM)은 변수들 간의 구조적 관계를 측정모형과 구조모형을 통하여 분석하는 것으로, 연구자에 따라서 공분산 구조 분석(covariance structure analysis) 또는 공분산 구조모형(covariance structure model)으로 불리기도 한다. 구조방정식 모형은 사회학, 교육학, 생물학, 경제학, 경영학, 의학 등 수많은 분야에서 사용하고 있는 통계적인 방법론이다(Tenko & George 2012).

구조방정식모형은 통계학뿐만 아니라 다른 많은 분야에서 더 많이 사용하고 있다. 따라서 통계적인 쟁점을 많이 무시하고 분석을 하는 경우가 많다. 예를 들면 데이터의 생성, 변수의 척도, 잠재변수 식별, 모형 식별 및 모수 추정, 모형 적합 및 검정 등 여러 가지 통계적 쟁점들에 대해 문제가 있다(김규성 2016). 간호학, 체육학, 교육학 등 여러 분야에서는 구조방정식모형을 이용한 논문들을 탐구하여 어떠한 통계적 쟁점에 대해서 문제가 있는지를 파악하고 연구하였다

(조형대 2011; 박일혁 외 2014; 김정희 외 2015).

일반적으로 구조방정식모형은 모수 추정방법에 있어서 최대우도법(Maximum Likelihood, ML)을 가장 많이 사용하는데, 다른 방법들에 비해서 모수 추정시에 적합도가 가장 좋게 나타나기 때문이다. 하지만 최대우도법은 데이터가 다변량 정규성을 만족해야 한다는 강한 가정이 존재한다. 그리고 데이터가 심각한 비정규성 분포를 가지고 있을 때 다음과 같은 문제점이 발생할 수 있다. 먼저, 최대우도 모수 추정치 값은 대표본에서는 상대적으로 정확할 수 있으나, 추정오차는 약 25%~50% 정도 낮아지는 경향이 있다. 그 결과, 귀무가설을 더 많이 기각하게 되고 이는 제1종 오류를 높인다. 둘째로는 모형의 적합도 통계량의 값이 지나치게 커지는 경향이 있다. 이는 모집단 내에서 실제 모형이 옳다고 할 귀무가설을 빈번하게 기각시킨다. 오류의 비율은 데이터나 모형에 따라 다르지만 기대되는 비율이 5% 일 때, 50%로 높아질 수도 있다(Kline, 2015). 따라서 데이터의 정규성이 만족되지 않는 경우에는 다른 방법으로 모수 추정을 해야 한다.

대표적으로 가장 많이 사용하는 방법으로 가중최소제곱법(Weighted Least Square), Satorra-Bentler의 scaled 통계량, Bollen-Stine 부트스트랩에 의한 방법이 있다. 통계프로그램인 SAS와 AMOS를 이용하여 분석을 실시하고 위에 제시한 세 가지 방법과 다변량 정규성 가정에 의한 ML방법에 의한 결과를 비교하고자 한다.

본 논문의 구성은 다음과 같다. 2절에서는 구조방정식모형의 적합도를 검정하는 통계량을 계산하는 방법들을 간략히 살펴보고자 한다. 3절에서는 청소년들의 안녕감에 대한 구조방정식 모형을 적합하고, 여러 적합도 검정방법들에 의한 결과를 비교하고자 한다. 4절에서는 논문을 결론짓고 마무리한다.

## II. 구조방정식모형의 적합도 검정 통계량

Hu & Bentler(1999)는 여러 적합도 검정 통계량을 비교하여 다중통계량(multiindex) 전략을 추천하며, 상황에 따른 적절한 cutoff를 제시하였다. Marsh, Hau, & Wen(2004)는 이와 같은 지침이 모든 상황에서 황금률은 아니라고 지적하고, 안전한 방법으

로 여러 통계량들을 비교한 후 전체적인 평가를 하는 것을 추천하였다. 구조방정식모형의 적합도 검정에 널리 사용되는 네 가지 방법을 간단히 살펴보고자 한다.

## 1. 최대우도법에 의한 검정 통계량

Joreskog(1970)은 최대우도법에 의한 적합함수의 추정에 대해서 다음과 같이 설명하였다.

$$F_{ML} = \log|\Sigma| + tr(S\Sigma^{-1}) - \log|S| - (p+q)$$

$S$ 는 모공분산행렬  $\Sigma$ 의 표본 추정치이며,  $p$ 와  $q$ 는 각각  $x$ 와  $y$ 의 관찰변수의 수이다.

최대우도법은 구조방정식모형에서 가장 대표적으로 사용하는 추정방법으로, 분석할 자료의 다변량 정규성이 만족되어야 한다는 강한 가정이 존재한다.

## 2. 점근분포무관법에 의한 검정 통계량

점근분포무관법(Asymptotic Distribution Free, ADF)은 가중최소제곱법(weighted least square)으로도 불리며 Browne(1984)에 의해 발전되었다. 그리고 실제로 데이터가 정규성을 만족하지 않을 때 가장 많이 사용 되는 방법 중 하나이다. 점근분포무관법에 대한 적합함수는 다음과 같다.

$$F_{ADF} = (s - \sigma(\theta))' W^{-1}(s - \sigma(\theta))$$

여기에서  $s$ 는 표본공분산행렬의 중복되지 않은 요소들을 통해 만들어지며 차수가  $p(p+1)/2 \times 1$ 인 벡터이고,  $\sigma(\theta)$ 는 모형에 의해 얻어진 공분산행렬의 중복되지 않은 요소들을 통해 만들어지며 차수가  $p(p+1)/2 \times 1$ 인 벡터이다. 그리고  $W$ 는 가중치행렬로서 차수가  $p(p+1)/2 \times p(p+1)/2$ 인 행렬이다. 가중치행렬은 표본 분산 및 공분산의 점근적 공분산행렬로서 4차 적률과 2차 적률의 추정치가 조합된 요소로 이루어져 있고 이를 다음과 같이 나타낼 수 있다.

$$[W]_{ij,gh} = s_{ijgh} - s_{ij}s_{gh}, \quad i \geq j, g \geq h$$

여기서,

$$s_{ijgh} = \frac{\sum_{n=1}^N (X_{ni} - \bar{X}_i)(X_{nj} - \bar{X}_j)(X_{ng} - \bar{X}_g)(X_{nh} - \bar{X}_h)}{N}$$

이다. 여기서  $N$ 은 표본의 수이다. 4차 적률은 첨도를 말하므로 가중치행렬은 첨도와 연관이 있음을 알 수 있다. 하지만 점근분포무관법은 가중치행렬의 계산이 상당히 복잡하다. 또한 관찰변수가 많으면 사용하기가 어렵고, 표본의 수가 커야 한다는 단점이 있다(Muthen & Kaplan 1992).

### 3. Satorra-Bentler scaled 검정 통계량

Satorra-Bentler(S-B) scaled 통계량을 이용하는 방법은 자료의 정규성이 만족되지 않는 경우에 사용하는 대체적인 방법이다. 정규성이 만족되지 않을 때 ML 방법으로 모수를 추정하면 결과가 정규성이 가정된 경우와 차이가 없지만, 그에 비해 표준오차는 25~50%까지 크게 감소하게 된다. 이는 표준오차가 과소 추정이 된다는 말이며 이로 인하여 유의한 모수 추정치가 더 많이 발생하게 된다. 또한 모형의  $\chi^2$  검정통계량은 그 값이 크게 증가한다. 이는 검정통계량이 과대추정이 되어 이 또한 모형이 과하게 기각이 된다. 따라서 정규성을 만족하지 않는 자료에 대한 ML 방법은 편향적인 결과가 나타나게 된다.

이를 해결하기 위하여 Satorra-Bentler scaled 카이제곱 검정 통계량은 다음과 같은 관계식으로 표현할 수 있다.

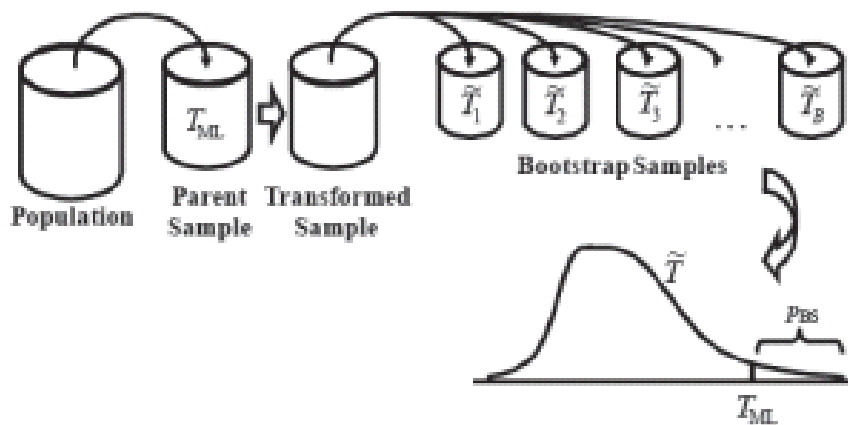
$$\text{S-B scaled } \chi^2 = d^{-1} (ML \chi^2)$$

S-B  $\chi^2$  통계량은 ML  $\chi^2$  통계량에 영향을 받는다. 여기에서  $d$ 는 변수들의 다변량 첨도의 크기에 부분적으로 관련된 scaling factor이다(Satorra & Bentler 1994).

만일 첨도가 없다면 ML의  $\chi^2$  통계량과 S-B  $\chi^2$  통계량은 같은 값을 가지게 되지만, 첨도가 커지면 두 통계량은 차이가 발생하고 이는 모형의 적합도에 영향을 미친다(Hancock & Mueller 2013). 다른 적합도 검정 통계량들도 ML의  $\chi^2$  통계량 대신에 S-B  $\chi^2$  통계량에 근거하여 계산된다. S-B scaled 통계량은 비정규성에 로버스트한 것으로 알려져 있다.

#### 4. Bollen-Stine 부트스트랩에 의한 유의확률

마지막으로 부트스트랩에 대해서 소개하고자 한다. naive한 부트스트랩도 있지만 여기에서는 Bollen-Stine(B-S)의 부트스트랩에 대해서 소개한다. Bollen과 Stine(1992)은 naive한 부트스트랩과는 다르게 모집단으로부터 1회 리샘플링한 표본을 변환하고, 이 변환된 표본을 부트스트랩 샘플링을 실시하는 방법을 소개하였다. 리샘플링한 표본을 parent sample이라고 표현한다. <그림 1>은 Bollen-Stine의 부트스트랩을 쉽게 이해할 수 있도록 첨부하였다(Hoyle 2012).



<그림 1> Bollen-Stine 부트스트랩 절차

여기에서 재표집한 표본을 변환하는 과정은  $Y$ 를 원래 데이터의 행렬이라고 하고,  $S$ 는  $Y'Y/(N-1)$ 이고  $Y$ 의 표본 공분산행렬이다.  $\hat{\Sigma}$ 는 추정된 모형 공분산행렬이라고 한다면, 변환된 데이터인  $Z$ 는  $YS^{-1/2}\hat{\Sigma}^{-1/2}$ 로 표현되며  $Z'Z/(N-1)$

$= \hat{\Sigma}$ 로 설명된다. 이제 원 데이터인  $Y$ 대신에 변환된  $Z$ 를 부트스트랩을 실시한다. 검정통계량은 공분산행렬인  $Z$ 와  $\hat{\Sigma}$ 가 같기 때문에 0이 된다.

이렇게 부트스트랩을 실시한 후 계산된 유의확률을 통해 모형의 적합성 여부를 판단한다. 따라서 이 방법을 적용하면 추정량과 검정통계량은 최대우도법에 의한 것과 동일하며, 다만 유의확률의 차이만 있게 된다.

### Ⅲ. 청소년들의 안녕감 자료 분석

#### 1. 분석 자료 및 분석 프로그램

분석에 사용된 자료는 청소년의 기질, 성격, 사회적 지지가 안녕감에 어떠한 영향을 미치는지에 대해 알아보기 위한 자료로서 광주·전남 지역 소재의 중학생 649명을 대상으로 하였다. 먼저 기질(temperament)을 나타내기 위한 변수로써 사회적 민감성(RD), 인내력(P)이 있고, 성격(character)을 나타내기 위한 변수에는 자율성(SD), 연대감(C), 자기초월(ST)이 있다. 사회적 지지(social support)에는 가족의 지지(FAMILY), 친구의 지지(FRIEND), 교사의 지지(TEACHER)가 있다. 마지막으로 안녕감(well-being)을 나타내는 변수로는 정서적 안녕감(EMOTIONAL), 심리적 안녕감(MENTAL), 사회적 안녕감(SOCIAL)이 있다.

구조방정식모형 분석을 할 수 있는 프로그램은 다양하다. SAS는 PROC CALIS를 이용하여 구조방정식 분석이 가능하며, R 소프트웨어는 sem 패키지와 lavaan 패키지를 이용하여 분석이 가능하다. SAS와 R 소프트웨어는 모형을 직접 입력해야 한다는 불편함이 있지만 AMOS는 GUI(graphic user interface)을 기반으로 한 프로그램으로서, 분석을 위한 구조방정식 모형의 작성을 쉽게 할 수 있으며 SPSS와 연계된다는 장점이 있다. 하지만 AMOS는 유료프로그램이고, 이에 반해 R 소프트웨어는 오픈소스로서 많은 통계분석을 할 수 있다. 하지만 R 소프트웨어의 lavaan 패키지는 아직까지도 개발 중이며 개선을 해야 할 점이 많이 남아 있다. SAS는 유료 프로그램이지만 통계분석에 가장 강력한 프로그램이며 구조방정식모형 분석을 하기 위한 개발이 모두 이루어진 상태이다.

SAS에서는 구조방정식모형에 대해서 Bollen-Stine 부트스트랩이 불가능하며, AMOS에서는 Satorra-Bentler scaled 통계량을 제공하지 않는다. R 소프트웨어는 두 방법 모두 분석이 가능하다(Narayanan 2012). 본 연구에서는 AMOS 24를 이용하여 구조방정식모형 분석을 하였고, Satorra-Bentler scaled 검정 통계량은 SAS 9.4를 이용하여 계산하였다.

표본의 크기에 따른 여러 방법들의 수행 정도를 비교하기 위하여, 649명의 전체 표본에서 100, 200, 300명을 단순 랜덤추출법에 의해 재추출하여 분석을 실시하였다.

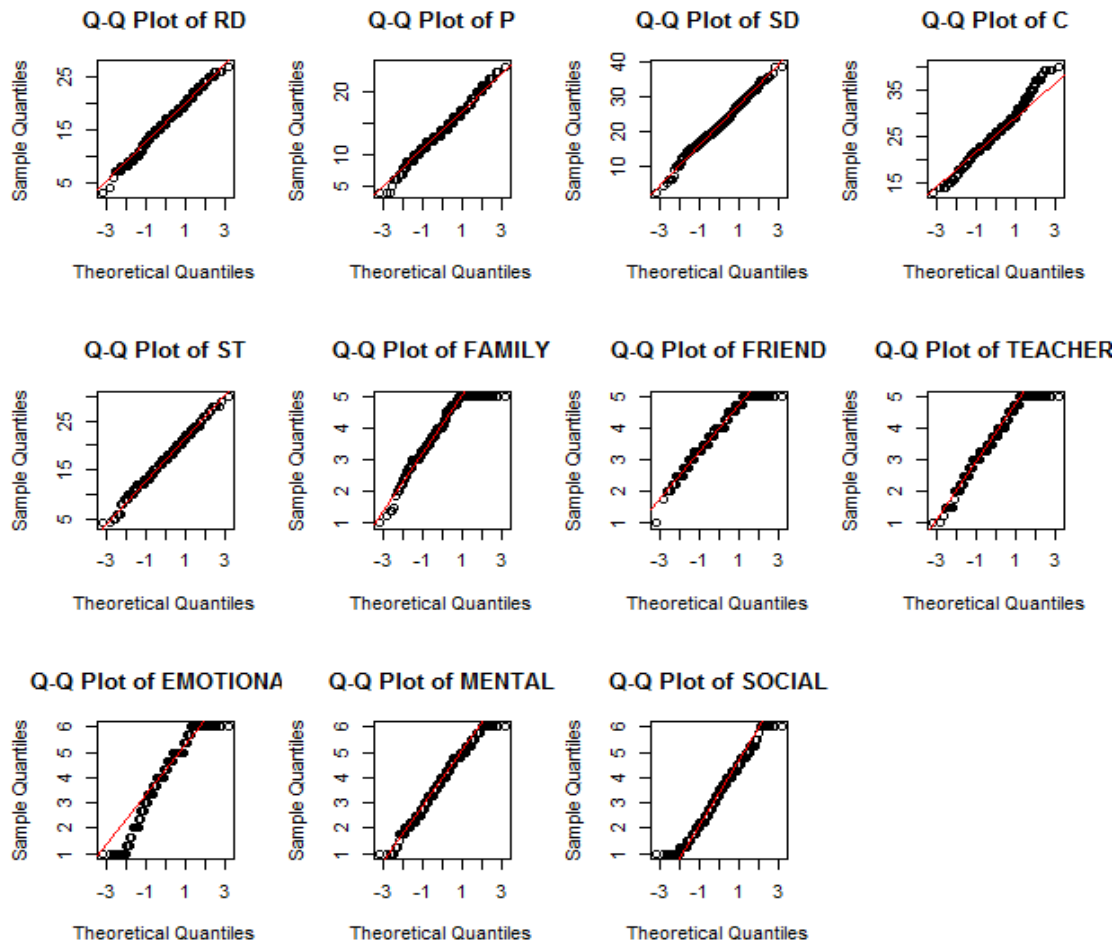
## 2. 기술통계 및 정규성 검정

구조방정식모형을 적합하기에 전 각 변수의 기술통계량을 <표 1>에서 확인하였다. 각 변수의 평균, 표준편차, 왜도, 첨도를 제시하였다. 왜도와 첨도의 값이 영에 가까우면 정규성을 만족하는 것으로 생각할 수 있는데, 본 연구의 자료는 그 기준에서 어긋나지 않는 것을 보여준다.

<표 1> 청소년 안녕감 자료의 기술통계량

| 변수      | 평균    | 표준편차 | 왜도    | 첨도    |
|---------|-------|------|-------|-------|
| 사회적 민감성 | 16.32 | 3.83 | -0.12 | 0.15  |
| 인내력     | 13.87 | 3.12 | 0.11  | 0.42  |
| 자율성     | 21.94 | 5.60 | 0.03  | 0.36  |
| 연대감     | 25.78 | 4.45 | 0.31  | 0.59  |
| 자기초월    | 17.29 | 4.25 | 0.02  | 0.10  |
| 가족의 지지  | 4.10  | 0.79 | -0.70 | 0.18  |
| 친구의 지지  | 4.00  | 0.69 | -0.46 | 0.10  |
| 교사의 지지  | 3.79  | 0.82 | -0.43 | -0.01 |
| 정서적 안녕감 | 4.21  | 1.21 | -0.61 | -0.02 |
| 심리적 안녕감 | 3.92  | 1.06 | -0.25 | -0.38 |
| 사회적 안녕감 | 3.35  | 1.16 | 0.02  | -0.56 |

<그림 2>는 각 변수의 Q-Q plot이다. <그림 2>의 RD, P, SD, ST (사회적 민감성, 인내력, 자율성, 자기초월)의 경우는 관측자료가 직선 위에 있어 정규성을 만족하는 것으로 보이지만 나머지 변수들은 그렇지 못하다고 할 수 있다.



<그림 2> 청소년 안녕감 자료의 Q-Q plot

S Gao, P Mokhtarian, & R Johnston의 2008년 연구보고서에 의하면 Curran, West, & Finch(1996)가 제시한 왜도와 첨도의 범위 안에 있는 자료라고 하더라도 정규성이 위배되는 경우가 종종 있다고 한다. 대안으로 왜도와 첨도를 각각의 표준편차로 나누어 표준화한 값을 Critical Ratio(CR)라고 하며 이를 사용하는 것을 추천하고 있다. 개별 변수에 대해서 비정규성을 판단하는 기준으로 왜도와 첨도의 CR이 점근적으로 표준정규분포에 따른다는 가정 하에서  $\pm 1.96$ 을 벗어나

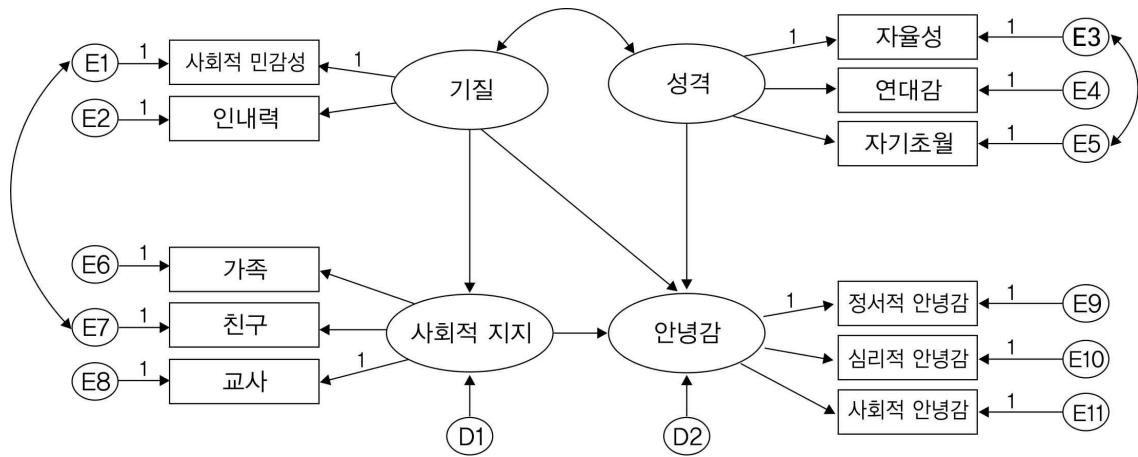
는지 확인하였다. <표 2>에 결과를 제시하였으며, 왜도의 경우는 기준을 벗어나는 변수가 6개이고, 첨도는 벗어나는 변수가 4개로 나타났다. CR에 의한 비정규성 판단이 Q-Q plot에 의한 판단과 거의 일치하는 것을 보여준다. 또한 다변량 첨도의 CR이 17.163으로 다변량 정규성 가정을 만족시키지 못하는 것으로 나타났다. 따라서 다변량 정규성 가정이 필요한 ML 방법으로 구조방정식 모형을 분석하기에는 적절하지 않다. 위에서 언급한 비정규성 데이터의 분석방법이 적절해 보인다.

**<표 2> 왜도와 첨도의 CR에 의한 단변량/다변량 정규성 검정**

| 변수      | 왜도의 CR | 첨도의 CR |
|---------|--------|--------|
| 사회적 민감성 | -1.291 | 0.742  |
| 인내력     | 1.138  | 2.135  |
| 자율성     | 0.307  | 1.834  |
| 연대감     | 3.201  | 3.022  |
| 자기초월    | 0.242  | 0.457  |
| 가족의 지지  | -7.336 | 0.871  |
| 친구의 지지  | -4.813 | 0.466  |
| 교사의 지지  | -4.457 | -0.049 |
| 정서적 안녕감 | -6.345 | -0.103 |
| 심리적 안녕감 | -2.631 | -1.977 |
| 사회적 안녕감 | 0.232  | -2.924 |
| 다변량     |        | 17.163 |

### 3. 구조방정식모형의 적합

해당 자료의 분석 목적은 중학생들의 기질과 성격이 사회적 지지를 매개로 하여 안녕감의 변화에 미치는 영향을 확인하기 위한 것이다. <그림 3>은 잠재변수들과 관찰변수들의 관계를 경로도로 나타낸 것이다.



〈그림 3〉 구조방정식모형 경로도

잠재변수인 기질과 성격은 서로 공분산의 관계가 있고, 관측변수들 간의 공분산 관계는 모형의 수정지수(modification indices)를 통해 카이제곱의 변화에 큰 영향을 미치는 관계를 포함시켰다.

#### 4. 구조방정식모형 적합도 검정 결과 분석

구조방정식모형의 적합도 검정은 표본의 크기에 영향을 받는다. 본 연구에서 분석에 사용한 검정 통계량 중에서는 SRMR과 RMSEA이 대체로 표본크기에 무관한 통계량으로 알려져 있다.(Hu & Bentler 1999) 총 표본인 649명에 대한 분석을 먼저 실시하여 적합도 결과와 경로계수들을 살펴본 후, 단순무작위추출을 통해 표본의 수를 각각 100, 200, 300개씩 10회 반복 추출하여 표본크기가 미치는 영향을 고려한 비교·분석을 실시하였다. 그리고 10회의 분석으로 나온 적합도 지수들에 대해 평균을 구하고, 카이제곱 통계량은 모형의 적합성에 대한 기각율을 구하여 분석방법별로 기각 여부에 차이가 있는지를 확인하였다.

다변량 정규성을 가정한 최대우도 추정에 의한 통계량, Satorra-Bentler scaled 검정 통계량, Bollen-Stine 부트스트랩에 의한 통계량과 점근분포무관 추정에 의한 통계량을 비교하였다. 부트스트랩은 반복횟수를 1,000회 실시하였다.

&lt;표 3&gt; 청소년 안녕감의 전체표본에 대한 구조방정식 모형의 적합도 통계량

| 방법            | 카이제곱 통계량   | GFI   | AGFI  | SRMR  | RMSEA | NFI   | AIC     |
|---------------|------------|-------|-------|-------|-------|-------|---------|
| ML            | 104.396*** | 0.971 | 0.948 | 0.036 | 0.053 | 0.961 | 162.396 |
| ADF           | 85.533***  | 0.963 | 0.935 | 0.050 | 0.045 | 0.874 | 143.553 |
| S-B scaled    | 91.725***  | 0.971 | 0.948 | 0.037 | 0.048 | 0.961 | 149.725 |
| B-S Bootstrap | 104.396*** | 0.971 | 0.948 | 0.036 | 0.053 | 0.961 | 162.396 |

\*  $p < 0.05$ , \*\*  $p < 0.01$ , \*\*\*  $p < 0.001$ 

&lt;표 4&gt; 청소년 안녕감의 전체표본에 대한 구조방정식 모형의 경로계수 추정량

| 경로          | ML       |       | ADF      |       | S-B scaled |       | B-S bootstrap |       |
|-------------|----------|-------|----------|-------|------------|-------|---------------|-------|
|             | $\beta$  | S.E   | $\beta$  | S.E   | $\beta$    | S.E   | $\beta$       | S.E   |
| 사회적 민감성←기질  | 1        |       | 1        |       | 1          |       | 1             |       |
| 인내력←기질      | 1.218*** | 0.109 | 1.190*** | 0.095 | 1.218***   | 0.108 | 1.218***      | 0.108 |
| 자율성←성격      | 1        |       | 1        |       | 1          |       | 1             |       |
| 연대감←성격      | 0.644*** | 0.046 | 0.703*** | 0.048 | 0.644***   | 0.052 | 0.644***      | 0.052 |
| 자기초월←성격     | 0.351*** | 0.046 | 0.342*** | 0.047 | 0.351***   | 0.050 | 0.351***      | 0.050 |
| 가족←사회적 지지   | 1        |       | 1        |       | 1          |       | 1             |       |
| 친구←사회적 지지   | 0.819*** | 0.049 | 0.874*** | 0.045 | 0.819***   | 0.048 | 0.819***      | 0.048 |
| 교사←사회적 지지   | 0.988*** | 0.059 | 1.053*** | 0.056 | 0.988***   | 0.060 | 0.988***      | 0.060 |
| 정서적 안녕감←안녕감 | 1        |       | 1        |       | 1          |       | 1             |       |
| 심리적 안녕감←안녕감 | 1.040*** | 0.053 | 1.103*** | 0.051 | 1.040***   | 0.054 | 1.040***      | 0.054 |
| 사회적 안녕감←안녕감 | 1.072*** | 0.057 | 1.112*** | 0.054 | 1.072***   | 0.056 | 1.072***      | 0.056 |
| 사회적 지지←기질   | 0.199*** | 0.021 | 0.195*** | 0.018 | 0.199***   | 0.021 | 0.199***      | 0.021 |
| 안녕감←기질      | 0.168**  | 0.062 | 0.130*   | 0.056 | 0.168**    | 0.063 | 0.168**       | 0.063 |
| 안녕감←성격      | 0.052    | 0.030 | 0.081**  | 0.028 | 0.052      | 0.031 | 0.052         | 0.031 |
| 안녕감←사회적 지지  | 0.329*** | 0.078 | 0.223**  | 0.077 | 0.329***   | 0.085 | 0.329***      | 0.085 |

\*  $p < 0.05$ , \*\*  $p < 0.01$ , \*\*\*  $p < 0.001$ 

<표 3>에서 검정 방법 별로 확인을 해보면 먼저 카이제곱 통계량은 모두 유의 확률이 0.05보다 작았다. 이는 귀무가설을 기각하여 모형이 적절하지 않다는 결

과를 의미하지만 위에서 언급했다시피 데이터와 모형의 관계, 그리고 표본의 수에 큰 영향을 받는다. 최대우도법에 대한 카이제곱 통계량은 104이고 점근분포무관법에 의한 카이제곱 통계량은 85이다. Satorra-Bentler scaled 검정통계량은 91이고, 부트스트랩의 카이제곱 통계량은 104이다. 부트스트랩은 분석 자체가 최대우도법을 기반으로 분석을 하기 때문에 카이제곱 포함 모든 적합도 지수는 같다. 점근분포무관법과 Satorra-Bentler scaled 검정 통계량을 보았을 때, 정규성을 가정한 최대우도법에 비해 카이제곱 통계량이 10이상 작지만 유의확률 측면에서는 큰 차이는 없다. 다른 적합도 지수들을 비교해도 최대우도법과 큰 차이가 없음을 알 수 있다.

추정방법별로 구한 모수의 추정량은 <표 4>와 같다.

Satorra-Bentler scaled 방법과 Bollen-Stine 부트스트랩을 이용한 방법은 모두 최대우도법을 기반으로 하고 있기 때문에 세 방법의 추정량은 동일하다. 다만 점근분포무관법은 다소 차이가 있었고 이로 인해 잠재변수들 간의 관계 중에서 성격이 안녕감에 영향을 미치는 관계에 대해서만 유의확률의 변화가 있었다. 하지만 이 결과로는 어떠한 방법이 비정규 데이터를 분석하는 데 가장 효과적인 방법인지를 알 수는 없다. 따라서 표본의 수를 달리하여 여러 적합도 검정 통계량을 분석한 결과를 바탕으로 비교하고자 한다.

<표 5>는 표본의 수를 100, 200, 300개로 증가하면서 각각 10회 무작위 반복 추출하여 평균한 적합도 지수들에 대한 결과이다. 표본의 크기가 커지면서 검정력이 커질 것으로 기대하는 측면에서 평가한다면 <표 5>에서 굵게 표시된 부분은 그와 반대되는 양상을 보여주고 있는 결과이다.

최대우도법에 의하면, 표본 수가 커짐에 따라 RMSEA가 일관되지 않은 양상을 보이지만, 카이제곱 통계량을 비롯한 다른 통계량은 기대되는 결과를 보인다.

점근분포무관법은 표본의 수에 따른 양상이 더욱 일관되지 않은 경향을 여러 적합도 지수에서 확인할 수 있다. 이는 모형이 복잡할수록, 표본의 크기가 1,000보다 작은 경우에 영향을 받는 특성을 갖기 때문임을 확인할 수 있다.

Satorra-Bentler scaled ML 추정에 의한 방법은 표본의 크기에 따른 모든 적합도 통계량이 일관된 모습을 보인다.

Bollen-Stine의 부트스트랩은 모형에 최대우도법을 사용하는 것은 같지만, 유

의확률을 계산할 때 부트스트랩을 이용하는 것으로 모형의 적합도 지수는 최대우도법의 적합도지수와 이론적으로 같다. <표 5>의 차이는 무작위 추출된 표본의 차이에서 기인한다.

<표 5> 청소년 안녕감의 부분표본에 대한 구조방정식 모형의 적합도 통계량

| 방법            | 표본 크기 | 카이제곱 통계량      | 적합모형 기각율*   | GFI          | AGFI         | SRMR  | RMSEA        | NFI   | AIC            |
|---------------|-------|---------------|-------------|--------------|--------------|-------|--------------|-------|----------------|
| ML            | 100   | 49.580        | 4/10        | 0.919        | 0.856        | 0.059 | 0.055        | 0.898 | 107.580        |
|               | 200   | 63.766        | 8/10        | 0.941        | 0.896        | 0.050 | <b>0.062</b> | 0.922 | 125.193        |
|               | 300   | 67.508        | 10/10       | 0.961        | 0.930        | 0.042 | 0.052        | 0.947 | 125.508        |
| ADF           | 100   | 63.627        | 7/10        | 0.946        | 0.903        | 0.103 | 0.081        | 0.865 | 121.627        |
|               | 200   | <b>60.988</b> | 9/10        | <b>0.946</b> | <b>0.903</b> | 0.077 | 0.056        | 0.870 | <b>118.988</b> |
|               | 300   | 66.163        | <b>9/10</b> | <b>0.916</b> | 0.913        | 0.062 | 0.050        | 0.879 | 124.163        |
| S-B scaled    | 100   | 50.938        | 5/10        | 0.912        | 0.844        | 0.064 | 0.055        | 0.875 | 108.938        |
|               | 200   | 53.736        | 6/10        | 0.949        | 0.909        | 0.048 | 0.044        | 0.932 | 111.736        |
|               | 300   | 58.763        | 7/10        | 0.962        | 0.931        | 0.040 | 0.043        | 0.947 | 116.763        |
| B-S bootstrap | 100   | 49.580        | 1/10        | 0.919        | 0.856        | 0.059 | 0.055        | 0.898 | 107.580        |
|               | 200   | 63.766        | 5/10        | 0.945        | 0.901        | 0.049 | <b>0.059</b> | 0.928 | 121.766        |
|               | 300   | 67.508        | 8/10        | 0.961        | 0.930        | 0.042 | 0.052        | 0.947 | 125.508        |

\* 적합모형 기각률은 카이제곱 검정통계량에 의해 계산됨.

표본크기에 따른 값의 변동이 일관성이 결여된 결과를 보이는 ADF 방법을 제외하면, 척도가 다르므로 단언할 수는 없지만, 비슷한 값을 갖는 AGFI, SRMR, NFI에 비해서 SRMR과 RMSEA는 표본의 크기에 따른 값의 변동폭이 더 작은 결과를 보인다. SRMR과 RMSEA는 상대적으로 표본크기에 무관한 것으로 알려져 있는데 본 연구에서도 이 사실을 확인할 수 있었다.

#### IV. 결론

비정규성 자료에 대한 구조방정식모형 분석을 실제 관측된 사례에 적용하여 적합도 지수를 비교했을 때 전체 표본에 대한 결과는, 최대우도법에 의한 카이제

곱통계량, Satorra-Bentler scaled 검정통계량, Bollen-Stine 부트스트랩의 카이제곱통계량과 점근분포무관법에 의한 검정통계량은 차이가 있었으나 유의확률에 변화가 있을 정도의 차이는 아니었다.

단순무작위추출한 비정규 데이터를 분석한 결과 앞의 네 방법 간에 카이제곱통계량에 의한 기각율에 차이가 있었다. 또한 점근분포무관법은 표본의 크기가 커짐에 따른 적합도 지수가 개선되는 양상이 일관성이 없는 경우가 많았다. 이는 점근분포무관법은 표본크기가 1,000을 넘는 경우 사용할 것을 추천하는 Muthen & Kaplan (1992)의 연구 결과와 맥락을 같이 한다. Satorra-Bentler scaled 추정에 의한 검정통계량은 표본의 수나 변수의 수 등에 민감하게 영향을 받지 않는다. 또한 표본의 수가 커질수록 적합도 지수는 더욱 개선되었다. 반면 카이제곱 통계량은 표본의 수에 민감하게 영향을 받기 때문에 표본의 수가 커질수록 증가하는 경향을 보였다. 이러한 결과를 종합적으로 보았을 때 비정규 자료는 최대우도법과 점근분포무관법보다 Satorra-Bentler scaled 검정통계량을 이용한 방법이나 Bollen-Stine 부트스트랩을 이용하여 분석을 진행하는 것이 좀 더 로버스트한 결과를 얻을 수 있다.

본 논문에서는 김규성(2016)의 여러 통계적 쟁점 중 정규분포를 따르지 않는 데이터에 대한 문제에 대해서 연구를 하였다. 구조방정식모형이 계속해서 발전하고 있고 많은 분야에서 사용하기 때문에, 이러한 쟁점들을 놓치고 분석의 오류를 범하는 사례도 늘고 있다. 구조방정식모형에 관한 여러 가지 다른 통계적 쟁점들에 대해서도 추후 살펴보려고 한다.

## 참고문헌

- 김규성. 2016. “구조방정식모형의 통계적 쟁점.” 《조사연구》 17(1): 31-53.
- 김정희·신수진·박진화. 2015. “국내 간호학 학회지에 출판된 구조방정식모형 연구의 방법론적 질 평가.” 《대한간호학회지》 45(2): 159-168.
- 박일혁·엄한주·이기봉. 2014. “체육학연구에서 확인적 요인분석 이용의 문제점과 올바른 적용.” 《한국체육측정평가학회지》 16(1): 1-22.

- 조형대. 2011. “교육학 분야에서 구조방정식 활용의 문제점과 대안제시.” 고려대학교 석사학위논문.
- Browne, M.W. 1984. “Asymptotic Distribution-free Methods in the Analysis of Covariance Structure.” *British Journal of Mathematical and Statistical Psychology* 37: 62-83.
- Gao, S., P. Mokhtarian, and R. Johnston. 2008. “Non-normality of Data in Structural Equation Models.” *Research Report*: UCD-ITS-RR-08-47.
- Gerbing, D.W. and J.C. Anderson. 1992. “Monte Carlo Evaluations of Goodness of Fit Indices for Structural Equation Models.” *Sociological Methods and Research* 21(2): 132-160.
- Hancock, G.R. and R.O. Mueller. 2013. *Structural Equation Modeling a Second Course*(2nd ed.). Charlotte: Information Age Publishing Inc.
- Hoyle, R.H., 2012. *Handbook of Structural Equation Modeling*. New York: The Guilford Press.
- Hu, L. and P.M. Bentler. 1999. “Cutoff Criteria for Fit Indexes in Covariance Structure Analysis: Conventional Criteria Versus New Alternatives.” *Structural Equation Modeling* 6(1): 1 - 55.
- Joreskog, K.G. 1970. “A General Method for Analysis of Covariance Structures.” *Biometrika* 57(2): 239-251.
- Joreskog, K.G. and D. Sorbom. 1989. *LISREL 7. A Guide to the Program and Applications*. Chicago: SPSS.
- Kline, R.B. 2015. *Principles and Practice of Structural Equation Modeling*(3rd ed.). New York: Guilford Press.
- Marsh, H.W., K.-T. Hau, and Z. Wen. 2004. “In Search of Golden Rules: Comment on Hypothesis-testing: Approaches to Setting Cutoff Values for Fit Indexes and Dangers in Overgeneralizing Hu and Bentler’s (1999) Findings.” *Structural Equation Modeling* 11(3): 320-341.
- McDonald, R.P. 1980. “A Simple Comprehensive Model for the Analysis of Covariance Structures: Some Remarks on Applications.” *British Journal of Mathematical and Statistical Psychology* 33(2): 161-183.
- Muthen, B. and D. Kaplan. 1992. “A Comparison of Some Methodologies for

- the Factor Analysis of Non-normal Likert Variables: A Note on Size of the Model.” *British Journal of Mathematical and Statistical Psychology* 45: 19-30.
- Narayanan, A. 2012. “A Review of Eight Software Packages for Structural Equation Modeling.” *The American Statistician* 66(2): 129-138.
- Satorra, A. and P.M. Bentler. 1988. “Scaling Corrections for Chi-square Statistics in Covariance Structure Analysis.” *Proceedings of the Business and Economic Statistics Section of the American Statistical Association* 308- 313.
- Satorra, A. and P.M. Bentler. 1994. “Corrections to Test Statistics and Standard Error in Covariance Structure Analysis.” in A. Von Eye & C.C. Clogg (eds.). *Latent Variables Analysis: Applications for Developmental Research* (pp. 399-419). Thousand Oaks, CA: Sage.
- Tenko, R. and A.M. George. 2012. *A First Course in Structural Equation Modeling*. Mahwah, NJ: Lawrence Erlbaum Assoc Inc.

<접수 2017/08/23, 수정 2017/09/28, 게재확정 2017/10/18>