

Testing the Factorial Validity of Scores From a Measuring Instrument

측정 도구에서 얻은 점수의
요인 타당성 테스트

Second-Order Confirmatory Factor Analysis Model

2차 확인적 요인 분석 모델

chapter 5

교육과정학 전공 김지혜, 조윤영
교육사회학 전공 장승민

INDEX

- ◆ Introduction
- ◆ The Hypothesized Model
- ◆ Analysis of Categorical Data
- ◆ Categorical Variables Analyzed as Continuous Variables
- ◆ Categorical Variables Analyzed as Categorical Variables

- ◆ Mplus Input File Specification and Output File Results: Input File
- ◆ Mplus Input File Specification and Output File Results: Output File

- ◆ Article Review: 그릿의 요인 구조의 안정성에 관한 연구

Introduction

- 이 장에서 설명할 응용 프로그램은 두 가지 점에서 이전 장의 응용 프로그램과는 다름
 - ✓ 첫째, 2차 요인 구조로 구성된 확인적 요인분석(CFA) 모델에 중점을 둠
 - ✓ 둘째, 3, 4장의 분석은 연속형 데이터를 기반으로 한 반면, 본 장에서는 범주형 데이터를 기반으로 함
 - 청소년 커뮤니티 표본과 관련된 Beck Depression Inventory-II의 중국어 버전(C-BDI-II)과 관련된 가정된 요인 구조(hypothesized factorial structure)를 테스트할 것임
- * C-BDI-II:**
- 우울증의 인지적, 행동적, 정서적, 신체적 구성요소와 관련된 증상을 측정하는 21개 항목 척도, 그러나 21번 항목은 부적절한 것으로 간주되어 삭제되고 20개 항목만이 사용됨
 - 응답자들은 각 항목에 대해 0에서 3까지로 평가되는 4개의 진술 중 자신의 감정을 가장 정확하게 설명하는 진술을 선택함. 점수가 높을수록 우울의 정도가 심각
 - 본 검사의 확인적 요인분석 절차에 대한 타당성은 Tanaka, Huba(1984)가 제공한 증거와 Byrne 등의 관련 연구(1993~1998)에서 BDI 점수 데이터가 계층적 계승 구조로 가장 적절하게 표현된다는 증거를 기반으로 함

The Hypothesized Model

- 이 장에서 테스트된 CFA 모델은 다음 사항들을 선형적으로 가정함
 - (a) C-BDI-II에 대한 반응은 3가지 1차 요인(부정적 태도, 수행의 어려움, 신체적 요소)과 1개의 2차 요인(일반 우울증)으로 설명할 수 있음
 - (b) 각 항목은 측정하도록 설계된 1차 요인에 대해 0이 아닌 로딩을 가지며, 다른 두 개의 1차 요인에 대한 로딩은 0임
 - (c) 각 항목과 관련된 잔차는 상관 관계가 없음
 - (d) 세 가지 1차 요인 간의 공변량은 2차 요인에 대한 회귀로 완전히 설명됨
- 또한, 3, 4장에서 검토한 CFA 모델과 달리 식별 및 잠재 변수 척도화 목적으로 1.0으로 고정되어 있음
 - => 요인 1의 경우 CBDI2_3, 요인 2의 경우 CBDI_12, 요인 3의 경우 CBDI2_16에 지정됨

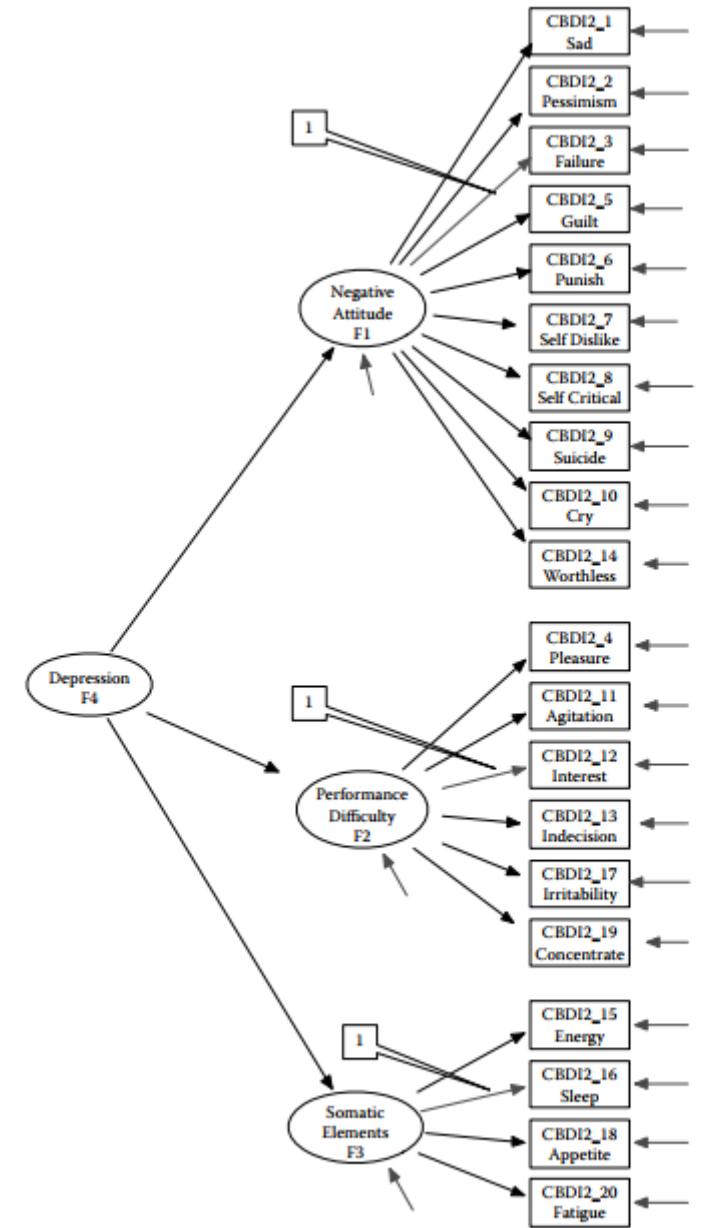


Figure 5.1. Hypothesized second-order model of factorial structure for the Chinese version of the Beck Depression Inventory II.

Analysis of Categorical Data

범주형 데이터 분석

- 지금까지는 최대 우도(ML) 추정(3장)과 강력한 ML(MLM) 추정(4장)을 기반으로 분석을 함. 이 두 추정 절차의 중요한 가정은 관찰된 변수의 규모가 연속적이라는 것임. 이 장에서 사용하는 C-BDI-II 데이터는 4개의 척도점을 갖는 항목으로 구성되며, 데이터는 이러한 범주적 특성을 고려하는 분석이 적합함
- 수년 동안 연구자들은 범주형 데이터를 마치 연속적인 것처럼 취급함. 구조방정식에서 데이터의 범주적 특성을 다루기 위해 잘 개발된 전략이 없었기 때문임. Mplus의 전신인 LISCOMP가 개발되고, 이를 고려하기 시작한 것은 약 2000년대 초반부터였음. SEM 컴퓨터 소프트웨어 패키지에 통합된 방법론적 발전이 데이터의 범주적 특성을 다루는 것에 대해 연구자의 인식을 높이고, 적절한 분석 절차를 적용하도록 하는 데에 중요한 역할을 함
- 그럼에도 불구하고, 범주형 데이터를 연속적인 것처럼 처리하면 어떻게 되는지, 데이터가 범주형으로 처리되는지 연속형으로 처리되는지에 따른 차이가 거의 없는 경우는 없는지 질문을 제기할 수 있음. 이러한 문제를 다룬 문헌에 대해 간략히 검토해보고자 함

Categorical Variables Analyzed as Continuous Variables

연속형 변수로 분석된 범주형 변수

- 이 문제를 다룬 Monte Carlo 연구 검토 결과(West, Finch & Curran, 1995)
 - ✓ 첫째, Pearson 상관 계수는 정렬된 범주형 척도로 재구성된 동일한 두 변수 간 계산 시보다, 두 개의 연속형 변수 간 계산 시 더 높은 것으로 나타남(가장 큰 감쇠(attenuation)는 범주가 5개 미만인 변수와 높은 왜곡도를 나타내는 변수에서 발생. 후자의 조건은 반대 방향으로 기울어진 변수에 의해 더욱 악화됨)
 - ✓ 둘째, 범주형 변수가 정규 분포에 근접하는 경우:
 1. 범주 수는 모형 적합도의 χ^2 우도비 검정에 거의 영향을 미치지 않음. 그러나 왜도, 특히 미분 왜도(변수가 반대 방향으로 치우침)가 증가하면 χ^2 값이 점점 더 부풀려짐
 2. 요인 로딩과 요인 상관 관계는 약간 과소평가됨. 그러나 범주 수가 3개 미만이고, 왜도가 1.0보다 크며 변수 전반에 걸쳐 차등 왜도가 발생하는 경우 과소평가가 더욱 중요해짐
 3. 잔차 분산 추정치는 다른 매개변수보다 항목 2에 언급된 범주형 및 왜도 문제에 가장 민감함
 4. 모든 매개변수에 대한 표준 오차 추정치는 너무 낮은 경향이 있으며, 분포가 크게 차등적으로 치우쳐 있을수록 더 낮음
- 즉, 범주 수가 많고 데이터가 정규 분포에 근접할 때 데이터의 순서성을 다루지 못하는 경우 무시할 수 있음
 - Bentler & Chou(1987), "변수가 4개 이상의 범주를 가질 때 연속형 변수로 분석해도 무방하다."
 - 이후에도 이러한 초기 주장을 뒷받침하는 연구들

Categorical Variables Analyzed as Categorical Variables

범주형 변수로 분석된 범주형 변수

● *The Theory*

- 관찰된 변수의 범주형 특성을 다룰 때, 연구자는 각 변수에 기본 연속 척도가 있다고 자동적으로 가정함. 따라서 범주는 실제로 연속 척도를 가지며 각 임계값 쌍(또는 초기 척도 점)이 연속 척도의 일부를 나타내는 관찰되지 않은 변수에 대한 대략적인 측정값으로 간주될 수 있음
 - 이러한 측정의 조잡함은 구성의 연속 규모를 고정된 수의 정렬된 범주로 분할하는 데서 발생함. 이는 Likert 척도 데이터 분석에서 다음 두 가지 오류를 가져올 수 있음(O'Brien, 1985):
 - (a) 연속 척도를 범주 척도로 분할함으로써 발생하는 분류 오류
 - (b) 너비가 다른 범주로 인해 발생하는 변환 오류
 - 이 장에서 사용하고 있는 4점 척도로 구성되어 다각적 범주형 변수를 나타내는 측정 도구에서, y 는 범주형 변수에서, y^* 는 관찰되지 않은 기본 연속 변수일 때, 임계값(τ)은 다음과 같이 개념화됨:
 - If $y^* \leq \tau_1$, y is scored 1.
 - If $\tau_1 < y^* \leq \tau_2$, y is scored 2.
 - If $\tau_2 < y^* \leq \tau_3$, y is scored 3.
 - If $\tau_3 < y^*$, y is scored 4.
- * 임계값 수는 항상 범주 수보다 1이 적다.

Categorical Variables Analyzed as Categorical Variables

범주형 변수로 분석된 범주형 변수

● *The Theory*

- ✓ 범주형 데이터로 구조방정식 모델을 테스트할 때, 분석은 연속형 데이터의 경우처럼 더 이상 표본 분산-공분산 행렬(S)을 기반으로 하지 않음. 오히려 올바른 상관 행렬을 기반으로 해야 함
- ✓ 상관 변수가 모두 순서 척도 (ordinal scale)인 경우, 결과 행렬은 다항성 상관관계((polychoric correlation))를 손상시킴
- ✓ 한 변수는 순서 척도이고 다른 변수는 연속 척도(continuous scale)인 경우, 결과 행렬은 다계열 상관관계 (polyserial correlations)로 구성됨
- ✓ 두 변수가 이분형(dichotomous)인 경우 다항성 상관의 이 특별한 경우를 사중 상관(tetrachoric correlation)이라고 함
- ✓ 다계열 상관관계가 순서형 변수(ordinal variable)가 아닌 이분형(dichotomous) 변수를 포함하는 경우 다계열 상관관계를 이차 상관관계(biserial correlation)라고도 함

Categorical Variables Analyzed as Categorical Variables

범주형 변수로 분석된 범주형 변수

● *Underlying Assumptions*

- 범주형 데이터의 사용과 관련된 적용은 세 가지 매우 중요한 가정을 기반으로 함
 - (a) 각 범주형 관찰 변수의 기본은 관찰된 잠재 대응물(observed latent counterpart)이며, 그 척도는 연속적이고 정규 분포임
 - (b) 표본 크기는 관련 상관 행렬의 신뢰성 있는 추정을 가능하게 할 만큼 충분히 큼
 - (c) 관찰된 변수의 수는 최소로 유지됨
- 한편, 이러한 일련의 가정은 이 방법론의 주요 약점이 됨(Bentler, 2005)
 - 각 범주형 변수가 연속적이고 정규 분포 척도를 갖는다는 것은 충족하기 어려운 기준임(비현실적)
 - 위의 가정 (b), (c)의 기초가 되는 이론적 근거는 범주형 변수를 사용할 때 분석이 임계값 수에 관측 변수 수를 곱하여 구성된 도수분포표에서 상관 행렬 추정까지 진행되어야 한다는 점인데, 사례가 0이거나 0에 가까운 셀이 발생하는 경우 추정이 어려워짐. 이는 응답 범주 수에 비해 표본 크기가 작거나, 변수 수가 지나치게 크거나, 임계값의 수가 많은 경우 발생함. 즉, 관찰된 변수의 수 및 변수에 대한 임계값 수가 많을수록, 표본 크기가 작을수록 셀이 0에 가까운 사례를 0으로 구성할 가능성이 더 커져 추정이 어려움

Categorical Variables Analyzed as Categorical Variables

범주형 변수로 분석된 범주형 변수

● *General Analytic Strategies*

- 2000년대 초반까지만 해도 범주형 데이터 분석에 대한 두 가지 주요 접근 방식이 이 연구 분야를 지배함. 그러나, 이는 대부분의 실무자에게 비현실적이고 충족하기 어려운, 극도로 제한적인 가정으로 인해 상쇄됨. 특히 안정적 추정치 산출을 위해서는 예외적으로 큰 표본 크기가 필요함
- 이러한 어려움을 해결하는 시도로 범주형 데이터의 모델링 및 테스트에 대한 여러 다른 접근 방식이 개발되었지만, 기본 추정량(estimator)은 세 가지 뿐임
 - ✓ 첫째, 추정 평균 및/또는 평균에 대한 비가중 최소 제곱
 - ✓ 둘째, ULS 및 DWLS 추정치(WLSM)만을 기반으로 한 분산
 - ✓ DWLS 추정치(WLSMV)의 평균 및 분산에 대한 수정(물론 WLS와 SWLS 라벨링을 엮는 것은 쓸데없이 복잡함)

=> WLSMV 추정량이 범주형 데이터의 CFA 모델링에서 가장 잘 수행됨. Mplus 는 현재 적어도 하나의 이진 또는 순서 범주형 지표 변수로 구성된 데이터에 사용할 수 있는 7개의 추정기를 제공함(Muthen & Muthen, 2007-2010). WLSMV 추정기가 기본값

Mplus Input File Specification and Output File Results: Input File

● 4가지 특징

- 1) CFA 모델은 계층적으로 구조화되어있음.
- 2) 분석에 기반한 데이터: 서열 척도로 측정됨.
- 3) 식별 및 척도를 위해 1.0으로 고정된 Factor-Loading path는 Mplus 기본 경로와 다름.
- 4) 고차 구조와 관련된 하나의 동등 제약이 지정되어있음.

Mplus Input File Specification and Output File Results: Input File

Figure 5.2

TITLE: CFA of Chinese BDI2 for Hong Kong Adolescents

DATA:

FILE IS "C:/Mplus/Files/bdihk2c2.dat";

VARIABLE:

NAMES ARE LINKVAR GENDER AGE CBDI2_1 – CBDI2_20;
CATEGORICAL ARE CBDI2_1 – CBDI2_20;
USEVARIABLES ARE CBDI2_1 – CBDI2_20;

MODEL:

F1 by CBDI2_1 – CBDI2_3 CBDI2_5 - CBDI2_10 CBDI2_14;
F2 by CBDI2_4 CBDI2_11 - CBDI2_13 CBDI2_17 CBDI2_19;
F3 by CBDI2_15 CBDI2_16 CBDI2_18 CBDI2_20;
F4 by F1-F3;

OUTPUT: MODINDICES STDYX;

Figure 5.3

TITLE: CFA of Chinese BDI2 for Hong Kong Adolescents(CBDI.FREE.EAUATE)

DATA:

FILE IS "C:/Mplus/Files/bdihk2c2.dat";

VARIABLE:

NAMES ARE LINKVAR GENDER AGE CBDI2_1 – CBDI2_20;
CATEGORICAL ARE CBDI2_1 – CBDI2_20;
USEVARIABLES ARE CBDI2_1 – CBDI2_20;

ANALYSIS:

ESTIMATOR IS WLSMV;

MODEL:

F1 by CBDI2_1* CBDI2_2 CBDI2_3@1 CBDI2_5 - CBDI2_10 CBDI2_14;
F2 by CBDI2_4* CBDI2_11 CBDI2_12@1 CBDI2_13 CBDI2_17 CBDI2_19;
F3 by CBDI2_15* CBDI2_16@1 CBDI2_18 CBDI2_20;

F4 by F1* F2 F3;

F4@1;

F1 F3 (1);

OUTPUT: MODINDICES STDYX;

Mplus Input File Specification and Output File Results: Input File

Figure 5.2 Mplus input file based on usual program defaults

TITLE: CFA of Chinese BDI2 for Hong Kong Adolescents

DATA:

FILE IS "C:/Mplus/Files/bdihk2c2.dat";

VARIABLE:

NAMES ARE LINKVAR GENDER AGE CBDI2_1 – CBDI2_20;

CATEGORICAL ARE CBDI2_1 – CBDI2_20;

USEVARIABLES ARE CBDI2_1 – CBDI2_20;

MODEL:

F1 by CBDI2_1 – CBDI2_3 CBDI2_5 - CBDI2_10 CBDI2_14;

F2 by CBDI2_4 CBDI2_11 - CBDI2_13 CBDI2_17 CBDI2_19;

F3 by CBDI2_15 CBDI2_16 CBDI2_18 CBDI2_20;

F4 by F1-F3;

OUTPUT: MODINDICES STDYX;

- WLSMV(Weighted least square mean and variance: 가중 최소 제곱 평균 및 분산)은 Mplus의 범주형 변수 분석을 위한 기본 추정량이므로 지정할 필요가 없음.
- 첫번째 로딩 경로는 자동으로 1.0로 제한되어 지정할 필요가 없음.
- 2차 수준에서는 동등 제약 조건이 적용되지 않음.

Mplus Input File Specification and Output File Results: Input File

Figure 5.3 Mplus input file with alternative specification and no program defaults

TITLE: CFA of Chinese BDI2 for Hong Kong Adolescents(CBDI.FREE.EAUATE)

DATA:

FILE IS "C:/Mplus/Files/bdihk2c2.dat";

VARIABLE:

NAMES ARE LINKVAR GENDER AGE CBDI2_1 – CBDI2_20;

CATEGORICAL ARE CBDI2_1 – CBDI2_20;

USEVARIABLES ARE CBDI2_1 – CBDI2_20;

ANALYSIS:

ESTIMATOR IS WLSMV;

MODEL:

F1 by CBDI2_1* CBDI2_2 CBDI2_3@1 CBDI2_5 - CBDI2_10 CBDI2_14;

F2 by CBDI2_4* CBDI2_11 CBDI2_12@1 CBDI2_13 CBDI2_17 CBDI2_19;

F3 by CBDI2_15* CBDI2_16@1 CBDI2_18 CBDI2_20;

F4 by F1* F2 F3;

F4@1;

F1 F3 (1);

OUTPUT: MODINDICES STDYX;

- 명확히 하기 위해 **ANALYSIS 명령**을 포함하고 WLSMV 추정기를 지정하였으나 필수는 아님.
- **MODEL 명령**에서 CBDI2_1, CBDI2_4, CBDI2_15에는 **별표가 할당**되어있음: 일반적인 기본 제약 조건인 1.0과 대조적으로 자유롭게 추정할 수 있음을 드러냄.
- 대신 CBDI2_3, CBDI2_12, CBDI2_16의 값이 1.0으로 제한되므로 모델의 참조 지표 변수로 사용됨. -> **CBDI2_3@1** 등으로 적절하게 고정된 매개변수로 지정된 것을 볼 수 있음.

Mplus Input File Specification and Output File Results: Input File

Figure 5.3 Mplus input file with alternative specification and no program defaults

TITLE: CFA of Chinese BDI2 for Hong Kong Adolescents(CBDI.FREE.EAUATE)

DATA:

FILE IS "C:/Mplus/Files/bdihk2c2.dat";

VARIABLE:

NAMES ARE LINKVAR GENDER AGE CBDI2_1 – CBDI2_20;

CATEGORICAL ARE CBDI2_1 – CBDI2_20;

USEVARIABLES ARE CBDI2_1 – CBDI2_20;

ANALYSIS:

ESTIMATOR IS WLSMV;

MODEL:

F1 by CBDI2_1* CBDI2_2 CBDI2_3@1 CBDI2_5 - CBDI2_10 CBDI2_14;

F2 by CBDI2_4* CBDI2_11 CBDI2_12@1 CBDI2_13 CBDI2_17 CBDI2_19;

F3 by CBDI2_15* CBDI2_16@1 CBDI2_18 CBDI2_20;

F4 by F1* F2 F3;

F4@1;

F1 F3 (1);

OUTPUT: MODINDICES STDYX;

- "F4 by F1* F2 F3" 와 "F4@1"은 참조 지표 변수의 수정을 나타냄.
- 2차 요인 로딩 경로(F4 -> F1)를 1.0으로 자동으로 제한함.
- ✓ 연구자들은 고차 모델을 구성할 때 특히 고차 요인 로딩을 추정하는 데 관심이 있음. 이는 모델의 복잡성과 다양한 변수 간 관계의 깊이를 더 잘 이해하기 위함임. 따라서 parameter를 자유롭게 추정하는 것을 어느 정도 선호함.
- ✓ SEM의 중요한 결과는 회귀 경로나 분산 중 하나를 추정할 수 있지만, 둘 다 추정할 수 없다는 것임. 이는 모델 식별 문제와 관련이 있으며, 모델이 통계적으로 명확하고 유의미한 결과를 제공하기 위해 필요한 조건임.

Mplus Input File Specification and Output File Results: Input File

Figure 5.3 Mplus input file with alternative specification and no program defaults

TITLE: CFA of Chinese BDI2 for Hong Kong Adolescents(CBDI.FREE.EAUATE)

DATA:

FILE IS "C:/Mplus/Files/bdihk2c2.dat";

VARIABLE:

NAMES ARE LINKVAR GENDER AGE CBDI2_1 – CBDI2_20;

CATEGORICAL ARE CBDI2_1 – CBDI2_20;

USEVARIABLES ARE CBDI2_1 – CBDI2_20;

ANALYSIS:

ESTIMATOR IS WLSMV;

MODEL:

F1 by CBDI2_1* CBDI2_2 CBDI2_3@1 CBDI2_5 - CBDI2_10 CBDI2_14;

F2 by CBDI2_4* CBDI2_11 CBDI2_12@1 CBDI2_13 CBDI2_17 CBDI2_19;

F3 by CBDI2_15* CBDI2_16@1 CBDI2_18 CBDI2_20;

F4 by F1* F2 F3;

F4@1;

F1 F3 (1);

OUTPUT: MODINDICES STDYX;

- "F4 by F1* F2 F3" 와 "F4@1"은 참조 지표 변수의 수정을 나타냄.
- 2차 요인 로딩 경로(F4 -> F1)를 1.0으로 자동으로 제한함.
- ✓ 따라서 모든 2차 요인 로딩을 자유롭게 추정하려면 2차 요인(F4, 우울증)의 분산을 1.0 값으로 제한해야 함. 그렇지 않으면 기본적으로 자동으로 자유롭게 추정됨.
- ✓ 이 추론을 따르지 않았을 때의 결과는 다음의 장과 같음.

Mplus Input File Specification and Output File Results: Input File

ESTIMATION TERMINATED NORMALLY

THE STANDARD ERRORS OF THE MODEL PARAMETER ESTIMATES COULD NOT BE COMPUTED. THE MODEL MAY NOT BE IDENTIFIED. CHECK YOUR MODEL. PROBLEM INVOLVING **PARAMETER 84.**

THE CONDITION NUMBER IS -0.806D-16.

Figure 5.4. Mplus error message regarding higher order factor specifications.

	PSI			
	F1	F2	F3	F4
F1	81			
F2	0	82		
F3	0	0	83	
F4	0	0	0	84

Figure 5.5. Mplus TECH1 output for factor covariance matrix (PSI).

- F4의 분산을 1.00으로 고정하지 않고 파일을 실행한 결과, 그림5.4 처럼 오류 메시지가 뜸.
- 그림5.4에는 모든 2차 요인 로딩을 자유롭게 추정한 결과 84가 문제의 원인으로 식별함.
- 그림5.5와 같이 출력 파일을 확인할 때, 84가 결국 F4임을 알 수 있음. 따라서 이를 식별하려면 1.0으로 제한되어야 함을 알 수 있음.
- 따라서 그림5.3에서 "F4@1"이 할당되어 있음.

Mplus Input File Specification and Output File Results: Input File

Figure 5.3 Mplus input file with alternative specification and no program defaults

TITLE: CFA of Chinese BDI2 for Hong Kong Adolescents(CBDI.FREE.EAUATE)

DATA:

FILE IS "C:/Mplus/Files/bdihk2c2.dat";

VARIABLE:

NAMES ARE LINKVAR GENDER AGE CBDI2_1 – CBDI2_20;
CATEGORICAL ARE CBDI2_1 – CBDI2_20;
USEVARIABLES ARE CBDI2_1 – CBDI2_20;

ANALYSIS:

ESTIMATOR IS WLSMV;

MODEL:

F1 by CBDI2_1* CBDI2_2 CBDI2_3@1 CBDI2_5 - CBDI2_10 CBDI2_14;
F2 by CBDI2_4* CBDI2_11 CBDI2_12@1 CBDI2_13 CBDI2_17 CBDI2_19;
F3 by CBDI2_15* CBDI2_16@1 CBDI2_18 CBDI2_20;

F4 by F1* F2 F3;
F4@1;
F1 F3 (1);

OUTPUT: MODINDICES STDYX;

- “F1 F3 (1)”은 동등성 제약을 의미함.
 - F1과 F3과 관련된 잔차 분산이 동일하게 제한되어야 함을 의미함.
- 괄호안의 1은 제약 조건이 하나만 지정되었음을 의미하며, F1과 F3 모두 해당됨.
- Mplus에서는 지정된 각 제약 조건이 별도의 줄에 표시되어야 함.
 - 고차 모델이 현재는 3개이지만, 6개의 요인으로 구성되어 있을 경우, 두 개의 추가 잔차 분산을 동일시하려는 경우 제약 조건은 “F1 F3 F5 F6 (1)”과 표시됨.
 - 반면, F5와 F6의 추정치가 동일하지만 완전히 다른 제약 조건이라면 “F5 F6 (2)”과 같이 표시됨.

Mplus Input File Specification and Output File Results: Input File

- 2장과 3장에서 지속적으로 “가설화된 모델과 관련된 자유도를 계산하는 것”이 중요하다고 강조하였음.
 - 통계적 식별 상태를 확인하기 위함임.
- 계층적 모델의 경우, 모델의 상위 부분의 식별 상태를 확인하는 것도 중요하다고 강조함.
 - 예를 들어, 2차 요인의 경우, 1차 요인이 관측 변수를 대체하여 2차원 모델 식별의 기반으로 사용됨.
 - 현재 사례에서 오직 3개의 1차 요인만 지정되어 있기 때문에, 모델 상위 수준에서 적어도 하나의 매개변수 제약이 없다면, 상위 구조는 정확히 식별됨.
 - 3개 1차 요인으로 6개($4 \times 3 / 2$)의 정보(3 요인 부하, 3 잔차 분산)가 있으며, 이로 인해 모델이 정확히 식별됨.
 - 가설 모델 초기 테스트 결과 F1과 F3과 관련된 잔차 분산의 추정치가 매우 가까웠다고 나와 이 두 매개변수는 동일하다고 제약되었고, 모델의 상위 수준에서도 자유도를 하나 제공하게 되었음.

Mplus Input File Specification and Output File Results: Output File

Table 5.1 Mplus Output for Hypothesized Model: Selected Summary Information

Summary of Analysis					
Number of groups					
1					
Number of observations					
486					
Number of dependent variables					
20					
Number of independent variables					
0					
Number of continuous latent variables					
4					
Observed Dependent Variables					
Binary and Ordered Categorical (Ordinal)					
CBDI2_1	CBDI2_2	CBDI2_3	CBDI2_4	CBDI2_5	CBDI2_6
CBDI2_7	CBDI2_8	CBDI2_9	CBDI2_10	CBDI2_11	CBDI2_12
CBDI2_13	CBDI2_14	CBDI2_15	CBDI2_16	CBDI2_17	CBDI2_18
CBDI2_19	CBDI2_20				
Continuous Latent Variables					
F1	F2	F3	F4		
Estimator					
WLSMV					

- 표5.1은 Mplus 출력 파일에 처음 나타나는 요약 정보를 나타냄.
- Sample 크기: 486
- 종속변수의 수: 20개(20개의 관찰된 지표 변수)
- 독립변수의 수: 0개
- 연속 잠재변수: 4개(3개의 1차요인과 1개의 2차 요인을 나타냄)
- 분석은 WLSMV의 추정량을 기반으로 함을 의미함.

Mplus Input File Specification and Output File Results: Output File

Table 5.2 Mplus Output for Hypothesized Model:
Summary of Categorical Proportions

CBDI2_1		CBDI2_8		CBDI2_15	
Category 1 0.508		Category 1 0.630		Category 1 0.430	
Category 2 0.247		Category 2 0.278		Category 2 0.397	
Category 3 0.222		Category 3 0.053		Category 3 0.142	
Category 4 0.023		Category 4 0.039		Category 4 0.031	
CBDI2_2		CBDI2_9		CBDI2_16	
Category 1 0.607		Category 1 0.790		Category 1 0.276	
Category 2 0.313		Category 2 0.167		Category 2 0.481	
Category 3 0.060		Category 3 0.039		Category 3 0.218	
Category 4 0.021		Category 4 0.004		Category 4 0.025	
CBDI2_3		CBDI2_10		CBDI2_17	
Category 1 0.465		Category 1 0.739		Category 1 0.516	
Category 2 0.261		Category 2 0.140		Category 2 0.342	
Category 3 0.255		Category 3 0.049		Category 3 0.130	
Category 4 0.019		Category 4 0.072		Category 4 0.012	
CBDI2_4		CBDI2_11		CBDI2_18	
Category 1 0.603		Category 1 0.352		Category 1 0.558	
Category 2 0.323		Category 2 0.438		Category 2 0.311	
Category 3 0.051		Category 3 0.169		Category 3 0.088	
Category 4 0.023		Category 4 0.041		Category 4 0.043	
CBDI2_5		CBDI2_12		CBDI2_19	
Category 1 0.644		Category 1 0.582		Category 1 0.346	
Category 2 0.267		Category 2 0.321		Category 2 0.424	
Category 3 0.051		Category 3 0.078		Category 3 0.189	
Category 4 0.037		Category 4 0.019		Category 4 0.041	
CBDI2_6		CBDI2_13		CBDI2_20	
Category 1 0.632		Category 1 0.527		Category 1 0.267	
Category 2 0.210		Category 2 0.377		Category 2 0.556	
Category 3 0.068		Category 3 0.076		Category 3 0.154	
Category 4 0.091		Category 4 0.021		Category 4 0.023	
CBDI2_7		CBDI2_14			
Category 1 0.628		Category 1 0.630			
Category 2 0.247		Category 2 0.235			
Category 3 0.086		Category 3 0.117			
Category 4 0.039		Category 4 0.019			

- 표5.2에 제시된 것은 4가지 범주 각각을 지지한 표본 응답자 비율임.
- CBDI2_1은 "나는 슬프다" 라는 진술에 대한 응답으로, 51%(0.508)가 슬픈 감정이 없음을 나타냄.
- 20개 항목 중 16개 항목에 대해 대부분의 응답자가 범주 1를 선택하여 우울증 증상의 증거가 없는 것으로 나타남.
- Leptokurtic 패턴을 반영하므로 비정규 분포를 따름을 알 수 있음.

Mplus Input File Specification and Output File Results: Output File

Table 5.3 Mplus Output for Hypothesized Model: Selected Goodness-of-Fit Statistics

Tests of Model Fit	
Chi-Square Test of Model Fit	
Value	200.504*
Degrees of freedom	82**
p-value	0.0000
CFI/TLI	
CFI	0.958
TLI	0.987
Number of free parameters	82
Root Mean Square Error of Approximation (RMSEA)	
Estimate	0.055
Weighted Root Mean Square Residual (WRMR)	
Value	0.947

* The chi-square value for MLM, MLMV, MLR, ULSMV, WLSM, and WLSMV cannot be used for chi-square difference tests. MLM, MLR, and WLSM chi-square difference testing is described in the Mplus "Technical Appendices" at the Mplus website, <http://www.statmodel.com>. See "chi-square difference testing" in the index of the Mplus User's Guide (Muthén & Muthén, 2007–2010).

** The degrees of freedom for MLMV, ULSMV, and WLSMV are estimated according to a formula given in the Mplus "Technical Appendices" at <http://www.statmodel.com>. See "degrees of freedom" in the index of the Mplus User's Guide (Muthén & Muthén, 2007–2010).

- 표5.3은 **적합도 통계**를 의미함.
- WLSMV robust χ^2 : 데이터의 범주형 특성과 비정규 특성을 모두 고려한 일반적인 카이제곱 통계의 척도화된 버전을 나타냄.

* WLSMV 추정량과 관련된 직접적인 카이제곱 차이 테스트는 "정규분포를 따르지 않기 때문에" 허용되지 않고 있음을 경고함.

** WLSMV 추정기를 사용할 때, 자유도도 정규 분포 연속 데이터에서 사용되는 CFA 절차와 필연적으로 다르다는 것을 경고함.

Mplus Input File Specification and Output File Results: Output File

Table 5.3 Mplus Output for Hypothesized Model: Selected Goodness-of-Fit Statistics

Tests of Model Fit	
Chi-Square Test of Model Fit	
Value	200.504*
Degrees of freedom	82**
p-value	0.0000
CFI/TLI	
CFI	0.958
TLI	0.987
Number of free parameters	82
Root Mean Square Error of Approximation (RMSEA)	
Estimate	0.055
Weighted Root Mean Square Residual (WRMR)	
Value	0.947

* The chi-square value for MLM, MLMV, MLR, ULSMV, WLSM, and WLSMV cannot be used for chi-square difference tests. MLM, MLR, and WLSM chi-square difference testing is described in the Mplus "Technical Appendices" at the Mplus website, <http://www.statmodel.com>. See "chi-square difference testing" in the index of the Mplus User's Guide (Muthén & Muthén, 2007–2010).

** The degrees of freedom for MLMV, ULSMV, and WLSMV are estimated according to a formula given in the Mplus "Technical Appendices" at <http://www.statmodel.com>. See "degrees of freedom" in the index of the Mplus User's Guide (Muthén & Muthén, 2007–2010).

- 아래의 모델 적합 통계를 살펴보면 가설 모델이 데이터에 잘 맞는 것으로 나타남.
 - CFA: 0.958
 - TLI: 0.987
 - RMSEA: 0.055
- WRMR이 범주형 데이터(Yu, 2022)를 사용했을 때 SRMR보다 더 나은 성능을 발휘하는 것으로 밝혀졌음. 0.95의 컷오프 기준이 연속형 데이터와 범주형 데이터에 적합한 모델 적합성을 나타내는 것으로 간주됨.

Mplus Input File Specification and Output File Results: Output File

Table 5.4 Mplus Output for Hypothesized Model: Parameter Estimates

Model Results				
	Estimate	Standard Error (SE)	Estimate/SE	Two-Tailed p-Value
F1 BY				
CBDI2_1	1.173	0.063	18.607	0.000
CBDI2_2	1.178	0.064	18.438	0.000
CBDI2_3	1.000	0.000	999.000	999.000
CBDI2_5	0.976	0.070	13.928	0.000
CBDI2_6	0.874	0.068	12.825	0.000
CBDI2_7	1.233	0.068	18.188	0.000
CBDI2_8	1.034	0.064	16.208	0.000
CBDI2_9	0.959	0.075	12.750	0.000
CBDI2_10	0.831	0.079	10.519	0.000
CBDI2_14	1.233	0.068	18.115	0.000
F2 BY				
CBDI2_4	0.884	0.042	20.828	0.000
CBDI2_11	0.939	0.041	23.129	0.000
CBDI2_12	1.000	0.000	999.000	999.000
CBDI2_13	0.931	0.043	21.882	0.000
CBDI2_17	0.936	0.044	21.104	0.000
CBDI2_19	0.807	0.045	18.094	0.000

- 표5.4는 표준화되지 않은 매개변수 추정값을 의미함.
- CBDI2_3, CBDI2_12, CBDI2_16은 Factor Loading에 할당된 참조 변수이므로 값은 1.000임.
- 1차 요인과 2차 요인 로딩 모두 통계적으로 유의미한 것으로 나타남.

Mplus Input File Specification and Output File Results: Output File

Table 5.4 Mplus Output for Hypothesized Model: Parameter Estimates (continued)

Model Results				
	Estimate	Standard Error (SE)	Estimate/SE	Two-Tailed p-Value
F3 BY				
CBDI2_15	1.481	0.107	13.784	0.000
CBDI2_16	1.000	0.000	999.000	999.000
CBDI2_18	0.799	0.103	7.791	0.000
CBDI2_20	1.388	0.108	12.873	0.000
F4 BY				
F1	0.599	0.033	18.276	0.000
F2	0.784	0.028	27.693	0.000
F3	0.502	0.039	12.783	0.000
Thresholds				
CBDI2_1\$1	0.021	0.057	0.363	0.717
CBDI2_1\$2	0.691	0.062	11.130	0.000
CBDI2_1\$3	2.002	0.126	15.953	0.000
CBDI2_2\$1	0.271	0.058	4.712	0.000
CBDI2_2\$2	1.403	0.083	16.970	0.000
CBDI2_2\$3	2.042	0.130	15.728	0.000
•				
•				
•				
•				
CBDI2_20\$1	-0.620	0.061	-10.169	0.000
CBDI2_20\$2	0.927	0.067	13.902	0.000
CBDI2_20\$3	2.002	0.126	15.953	0.000
Variances				
F4	1.000	0.000	999.000	999.000
Residual Variances				
F1	0.077	0.012	6.128	0.000
F2	0.033	0.021	1.565	0.118
F3	0.077	0.012	6.128	0.000

- Thresholds의 값은 C-BDI-II 항목의 특정한 임계값을 의미함.
- 20개 항목 각각에 대해 3개의 임계값을 표시될 것으로 예상됨.
- Mplus에서는 매개변수 이름 뒤에 달러 기호(\$)로 할당하여 표시함. 따라서 CBDI2_1\$1, CBDI2_1\$2, CBDI2_1\$3으로 표기됨.
- 고차 요인(F4)에 대한 분산이며, 이는 1.000으로 고정되어 있음.
- 각 1차 요인 로딩과 관련된 잔차 분산을 의미함. 1차 요인은 Endogenous Latent Variables이므로 분산을 추정할 수 없으므로 잔차의 분산이 대신 계산되며, Mplus에서 기본값임.
- F1과 F3이 동일하게 제한되었기 때문에, 0.077값을 가지며, 이는 통계적으로 유의함. 그러나 요인 2에 대한 잔차 분산은 유의미하지 않는 것으로 나타남.

Mplus Input File Specification and Output File Results: Output File

Table 5.4 Mplus Output for Hypothesized Model: Parameter Estimates (continued)

Model Results				
	Estimate	Standard Error (SE)	Estimate/SE	Two-Tailed p-Value
F3 BY				
CBDI2_15	1.481	0.107	13.784	0.000
CBDI2_16	1.000	0.000	999.000	999.000
CBDI2_18	0.799	0.103	7.791	0.000
CBDI2_20	1.388	0.108	12.873	0.000
F4 BY				
F1	0.599	0.033	18.276	0.000
F2	0.784	0.028	27.693	0.000
F3	0.502	0.039	12.783	0.000
Thresholds				
CBDI2_1\$1	0.021	0.057	0.363	0.717
CBDI2_1\$2	0.691	0.062	11.130	0.000
CBDI2_1\$3	2.002	0.126	15.953	0.000
CBDI2_2\$1	0.271	0.058	4.712	0.000
CBDI2_2\$2	1.403	0.083	16.970	0.000
CBDI2_2\$3	2.042	0.130	15.728	0.000
•				
•				
•				
•				
•				
CBDI2_20\$1	-0.620	0.061	-10.169	0.000
CBDI2_20\$2	0.927	0.067	13.902	0.000
CBDI2_20\$3	2.002	0.126	15.953	0.000
Variances				
F4	1.000	0.000	999.000	999.000
Residual Variances				
F1	0.077	0.012	6.128	0.000
F2	0.033	0.021	1.565	0.118
F3	0.077	0.012	6.128	0.000

- 관찰된 범주형 변수에 대한 잔차 분산이 누락됨.
- SEM 모델에서 연속형 변수가 아닌 범주형 변수가 포함된 경우 분석은 '표본 공분산 행렬'이 아닌 '표본 상관 행렬(S)'에 기반해야 함.
- 이 상관 행렬은 기본 연속 변수 y^* 를 나타내기 때문에 범주형 변수의 잔차 분산이 식별되지 않으므로 추정되지 않음.

Mplus Input File Specification and Output File Results: Output File

Table 5.5 Mplus Output for Hypothesized Model: Standardized Parameter Estimates				
Standardized Model Results: STDYX Standardization				
	Estimate	Standard Error (SE)	Estimate/SE	Two-Tailed p-Value
F1 BY				
CBDI2_1	0.774	0.027	28.514	0.000
CBDI2_2	0.778	0.029	27.046	0.000
CBDI2_3	0.660	0.031	21.145	0.000
CBDI2_5	0.644	0.037	17.281	0.000
CBDI2_6	0.577	0.042	13.878	0.000
CBDI2_7	0.814	0.024	33.656	0.000
CBDI2_8	0.682	0.034	20.071	0.000
CBDI2_9	0.633	0.043	14.817	0.000
CBDI2_10	0.548	0.045	12.061	0.000
CBDI2_14	0.813	0.024	34.219	0.000
F2 BY				
CBDI2_4	0.712	0.032	22.522	0.000
CBDI2_11	0.756	0.026	29.583	0.000
CBDI2_12	0.805	0.025	31.949	0.000
CBDI2_13	0.749	0.029	25.905	0.000
CBDI2_17	0.753	0.028	26.917	0.000
CBDI2_19	0.649	0.033	19.818	0.000
F3 BY				
CBDI2_15	0.849	0.026	32.972	0.000
CBDI2_16	0.573	0.038	15.241	0.000
CBDI2_18	0.458	0.050	9.150	0.000
CBDI2_20	0.796	0.028	28.331	0.000
F4 BY				
F1	0.908	0.016	57.837	0.000
F2	0.974	0.017	57.942	0.000
F3	0.876	0.021	41.206	0.000
Thresholds				
CBDI2_1\$1	0.021	0.057	0.363	0.717
CBDI2_1\$2	0.691	0.062	11.130	0.000
CBDI2_1\$3	2.002	0.126	15.953	0.000
CBDI2_2\$1	0.271	0.058	4.712	0.000

Table 5.5 Mplus Output for Hypothesized Model: Standardized Parameter Estimates (continued)				
Standardized Model Results: STDYX Standardization				
	Estimate	Standard Error (SE)	Estimate/SE	Two-Tailed p-Value
CBDI2_2\$2	1.403	0.083	16.970	0.000
CBDI2_2\$3	2.042	0.130	15.728	0.000
•				
•				
•				
•				
CBDI2_20\$1	-0.620	0.061	-10.169	0.000
CBDI2_20\$2	0.927	0.067	13.902	0.000
CBDI2_20\$3	2.002	0.126	15.953	0.000
Variances				
F4	1.000	0.000	999.000	999.000
Residual Variances				
F1	0.176	0.029	6.169	0.000
F2	0.051	0.033	1.569	0.117
F3	0.233	0.037	6.256	0.000

Table 5.6 Mplus Output for Hypothesized Model: Reliability Estimates and Modification Indices					
R ²					
Observed Variable	Estimate	Standard Error (SE)	Estimate/SE	Two-Tailed p-Value	Residual Variance
CBDI2_1	0.599	0.042	14.257	0.000	0.401
CBDI2_2	0.605	0.045	13.523	0.000	0.395
CBDI2_3	0.436	0.041	10.573	0.000	0.564
CBDI2_4	0.506	0.045	11.261	0.000	0.494
CBDI2_5	0.415	0.048	8.640	0.000	0.585
CBDI2_6	0.332	0.048	6.939	0.000	0.668
CBDI2_7	0.662	0.039	16.828	0.000	0.338
CBDI2_8	0.466	0.046	10.036	0.000	0.534
CBDI2_9	0.400	0.054	7.408	0.000	0.600
CBDI2_10	0.301	0.050	6.031	0.000	0.699
CBDI2_11	0.571	0.039	14.791	0.000	0.429
CBDI2_12	0.647	0.041	15.975	0.000	0.353
CBDI2_13	0.561	0.043	12.953	0.000	0.439
CBDI2_14	0.662	0.039	17.109	0.000	0.338
CBDI2_15	0.721	0.044	16.486	0.000	0.279
CBDI2_16	0.329	0.043	7.621	0.000	0.671
CBDI2_17	0.568	0.042	13.458	0.000	0.432
CBDI2_18	0.210	0.046	4.575	0.000	0.790
CBDI2_19	0.421	0.043	9.909	0.000	0.579
CBDI2_20	0.634	0.045	14.166	0.000	0.366
Latent Variable					
F1	0.824	0.029	28.919	0.000	
F2	0.949	0.033	28.971	0.000	
F3	0.767	0.037	20.603	0.000	
Model Modification Indices (MIs)*					
Modification Index	Expected parameter change (EPC)	Standard EPC	StdYX EPC	StdYX EPC	

* Modification indices for direct effects of observed dependent variables regressed on covariates and residual covariances among observed dependent variables may not be included. To include these, request MODINDICES (ALL). Minimum MI value for printing the modification index: 10.000. No MIs above the minimum value.

Mplus Input File Specification and Output File Results: Output File

- 표5.5와 표5.6의 추정치가 통계적으로 연결되어 있음. 따라서 결합된 해석에 중요한 영향을 미침.
- 표5.5는 표5.4에 나열된 각 매개변수에 대한 표준화된 추정치를 제시함.
- 표5.6은 3가지 1차 잠재요인 외에 각 범주형 변수에 대한 신뢰도 추정치를 보고함.
 - 관찰된 변수가 범주형인 CFA 모델 분석은 y^* 분산이 1.0으로 표준화됨.
 - 표준화되지 않은 추정치보다는 표준화된 추정치에 초점을 맞춰야 함.
 - 예를 들어, 연속형 변수에 대한 요인 로딩 추정치는 기본 요인에 의해 설명되는 관측 변수의 분산 비율로 해석, 범주형 변수에 대한 요인 로딩 추정치는 제공된 표준화된 요인 로딩을 기반으로 함.
 - CBDI2_1의 표준화된 로딩은 0.774임. 이 로딩을 제공하면 0.599라는 값이 나오며, 표5.6에 보고된 R^2 추정치를 나타냄.
 - 가설 모델의 요인 1에 설명될 수 있는 CBDI2_1의 기본 연속 및 잠재 측면(y)의 분산 비율을 나타내는 것으로 해석
 - 즉, CBDI2_1에 대한 결과는 해당 항목의 분산(범주 항목의 잠재 연속 측면)의 60%가 연결된 부정적 태도의 구성으로 설명될 수 있음.
 - 잔차 분산은 연관된 잠재 요인으로 설명되지 않는 y^* 분산의 비율을 나타냄.
 - 따라서 CBDI2_1에 초점을 맞춰 1.00-표준화된 하중(0.599)의 제공을 빼면 0.401의 값을 얻음. 이 값은 CBDI2_1에 대한 잔차 분산으로 보고됨.

Mplus Input File Specification and Output File Results: Output File

- MIs(Model Modification Indices)는 구조방정식 모델링에서 모델을 개선할 수 있는 가능성을 가리키는 통계적 지표임. 즉, MIs는 특정 경로나 parameter에 제약을 추가하거나 제거함으로써 모델의 적합도를 어떻게 개선할 수 있는지를 나타내는 값임. 즉, 모델의 특정 부분이 데이터에 얼마나 잘 맞지 않는지, 그리고 이를 수정함으로써 모델 적합도를 어떻게 향상시킬 수 있는지를 수치로 보여줌.
- MIs가 없다고 보고되므로, 데이터에 매우 잘 맞는 것으로 판단되며, 2차 요인 구조 가설이 확실하다고 할 수 있음. 즉 그림5.1이 홍콩 청소년에게 사용하기에 적합함을 의미함.

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

단기 종단자료를 이용한 그릿 요인구조의 안정성 연구*

임효진(서울교육대학교, 부교수)*

< 요약 >

본 연구에서는 모형에 근거한 그릿의 요인구조를 비교하고, 이에 근거하여 시간에 따른 그릿의 안정성을 살펴보고자 하였다. 그릿은 장기적인 목표를 위한 열정과 인내로 정의되며, 하위 요인은 흥미유지와 노력지속으로 구분된다. 연구를 위해 137명의 대학생을 대상으로 약 11개월의 간격을 두고 설문조사를 실시하였다. 분석 방법으로는 종단적 확인적 요인분석을 사용하였고 각종 적합도 지수를 통해 모형들을 비교하였다. 연구 결과 첫째, 그릿의 개념적 정의를 반영한 요인 모형으로 단일 요인 모형, 2요인 모형, 2차 요인 모형, 2중 요인 모형을 두 시점에서 비교한 결과 두 시점 모두 2요인 모형 또는 이와 동치 모형인 2차 요인 모형이 자료를 잘 설명하고 있었다. 둘째, 종단적 안정성을 살펴보기에 앞서 두 시점에 걸쳐 동일한 요인을 측정하고 있는지의 여부를 각 모형에 대해 알아본 결과 2요인 모형과 2차 요인 모형에서는 흥미유지의 부분측정동일성, 노력지속의 완전측정동일성이 확인되었다. 이어 그릿의 안정성 계수를 살펴본 결과 전체적인 그릿 요인들의 안정성이 확인되었으며, 이는 이전 시점의 그릿 수준이 높을수록 이후 시점의 그릿이 높음을 의미한다. 단, 흥미유지, 노력지속을 구분한 2요인 모형의 종단적 안정성에 있어서는 흥미유지보다 노력지속의 안정성 계수가 높아, 하위 요인에 따른 차이점을 확인할 수 있었다.

주제어 : 그릿, 흥미유지, 노력지속, 요인구조, 종단적 안정성

1. 그릿의 개념과 요인 구조

2. 연구방법 및 모형

3. 연구결과

* 본 논문은 2016년 정부(교육부)의 재원으로 한국연구재단의 지원(NRF-2016S1A3A2925401)을 받아 수행된 연구임.

† 서울시 서초구 서초중앙로 96 서울교육대학교, hyolim@snu.ac.kr

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

1. 그릿의 개념 및 요인 구조

그릿은 장기적인 목표를 위한 열정과 인내로 정의되며, 열정은 '흥미유지', 인내는 '노력지속'의 특성을 각각 반영하고 있다 (Duckworth et al., 2007). 흥미 유지는 비교적 장기간에 걸쳐 목표나 관심을 꾸준히 유지하는 수준을 나타내고, 노력 지속은 목표달성을 위해 어려움이나 장애물을 극복하는 수준을 나타낸다. 또한 그릿은 성격적 특성의 하나로, 유사 구인으로는 성실성, 자기통제, 자기조절, 노력조절, 과제지속 등이 대표적이다. 그러나 그릿과 유사 구인에 대한 차이를 검증하기 위한 심도 있는 연구는 현재 부족한 상태다.

Duckworth와 Quinn(2009)는 아동부터 성인까지의 표본을 대상으로 단일 요인 모형과 '흥미유지'와 '노력지속'을 1차 요인 (first-order factor)으로, 그릿을 2차 요인(second-order factor)으로 가정한 2차 요인 모형을 설정, 비교하였다. 연구의 결과 성인집단의 경우 2차 요인 모형이 적합 하였고, 군인 집단에서는 2차 요인 모형이 적합 하였으나, 아동과 대학생 집단에서는 적합도가 좋지 않았다. Wolters와 Hussain(2015)은 대학생 집단을 대상으로 단일 요인, 2요인, 2차요인 모형을 비교하였는데, 2요인 모형만 적합도가 좋았으며, Muenks 등(2017)은 고등학생 집단에서는 2요인 모형, 대학생 집단에는 2중 요인 모형이 자료를 잘 설명한다고 보고하였다.

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

2. 연구방법 및 모형

1) 연구대상

: 본 연구에서는 서울시에 소재한 4년제 대학에 재학 중인 137명의 자료를 분석에 사용하였다. 2016년 기준으로, 1학년은 50명(36.5%), 2학년은 46명(33.6%), 3학년은 40명(29.2%), 4학년은 1명(0.7%)이다.

2) 연구도구

: 그릿 척도는 Duckworth와 Quinn(2009)의 성인용 GIRT-S(GRIT-SHORT)의 문항을 사용하였다. 이 척도는 12문항으로 된 GRIT-O(GRIT-ORIGIANL, Duckworth et al., 2007)의 간편형 척도이며, Duckworth와 Quinn이 성인과 아동 등 4개 표본을 사용하여 원척도와의 상관계수가 낮은 4개의 문항 (흥미유지 2개, 노력지속 2개)를 삭제하고 최종 8개 문항으로 타당화 한 척도다.

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

2. 연구방법 및 모형

변수	문항	문항내용	문항내용과 신뢰도	
			신뢰도(Cronbach's alpha)	
			1차	2차
흥미 유지	C1	나는 목표를 세우지만 나중에 그것과는 다른 일을 하곤 한다.	.84	.86
	C2	때때로 새로운 생각이나 일 때문에 원래 하고 있는 생각이나 일이 방해를 받는다.		
	C3	나는 어떤 생각이나 일에 잠깐 집중하다가 곧 흥미를 잃은 적이 있다.		
	C4	나는 2-3개월 넘게 걸리는 일에 계속해서 집중하기 어렵다.		
노력 지속	P1	나는 시작한 것은 뭐든지 끝장을 본다.	.87	.84
	P2	어려움은 나를 꺾지 못한다.		
	P3	나는 부지런하다.		
	P4	나는 끊임없이 노력한다.		

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

2. 연구방법 및 모형

3) 분석방법

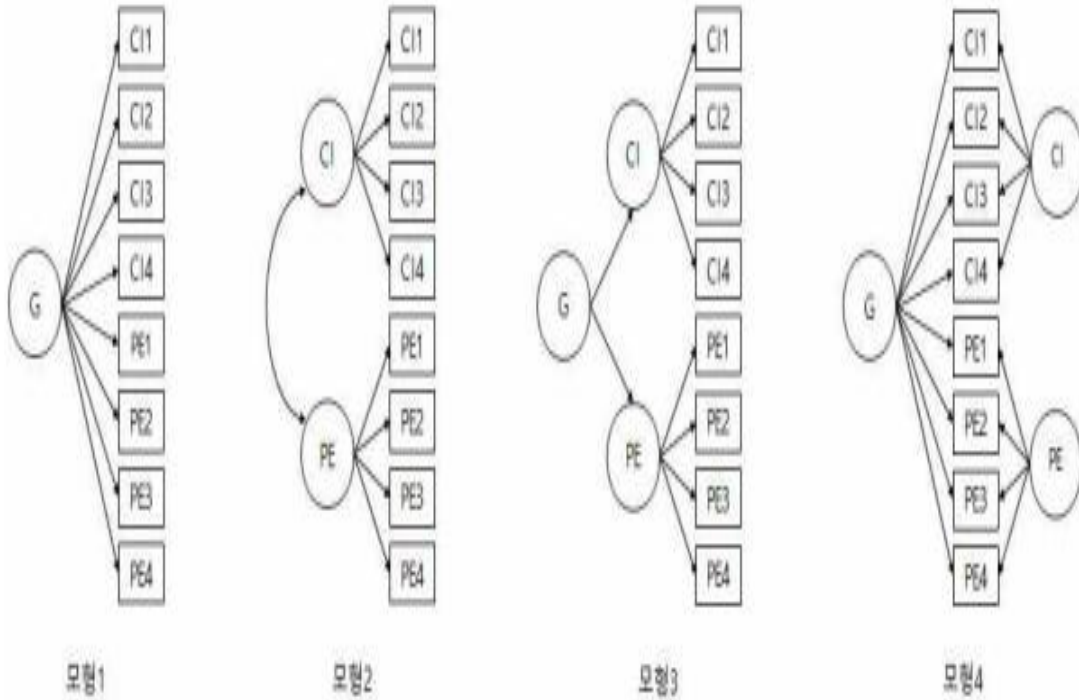
: 분석 방법으로는 종단적 확인적 요인분석(longitudinal confirmatory factor analysis)을 사용하였다. 이는 이론적 가정에 따라 설정된 모형을 비교할 수 있고, 특히 위계적 요인구조(예: 고차 요인 모형)를 가진 모형을 검증하기에 적합하다.

추정 방식은 완전 정보 최대우도법(Full information Maximum Likelihood Estimation Method: FIML)을 사용하였으며, 각종 적합도 지수들로 모형이 자료에 부합하는지를 검토하였다.

사용한 적합도 지수는 CFI(Comparative FitIndex), TLI(Tucker-Lewis' Index), RMSEA(Root Mean Square Error of Approximation)이며 일반적으로 CFI와 TLI는 그 값이 클수록, RMSEA는 작을수록 적절한데, CFI, TLI는 .90이상일때 양호한 것으로 평가되며, RMSEA는 .08이하면 양호한 적합도, .05이하면 매우 좋은 적합도로 간주된다(Browne & Cudeck, 1993; Hu & Bentler, 1999).

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

2. 연구방법 및 모형



모형1) Duckworth 등(2007, 2011)이 제안한 단일 요인 모형

모형2) 임효진(2017, 2018), Muenks(2017)이 검증한 2요인 모형

모형3) Duckworth와 Quinn(2009)이 제안한 2차 요인 모형

모형4) Muenks 등 (2017), 임효진(2017)의 연구에 근거한 2중 요인 모형

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

3. 연구결과

	C21	C22	C23	C24	P21	P22	P23	P24	C2 ^{주2}	P2 ^{주2}	M	SD	왜도	첨도
C11	-	.690***	.491***	.621***	.024	.042	.172*	.203*	.785***	.101	3.00	1.07	-0.08	-0.83
C12	.605***	-	.518***	.704***	.084	.071	.064	.078	.825***	.062	2.89	1.17	0.16	-0.71
C13	.565***	.648***	-	.560***	.111	.159	.127	.101	.718***	.129	2.64	1.05	0.27	-0.62
C14	.547***	.493***	.563***	-	.190*	.153	.199*	.168*	.688***	.216*	3.35	1.20	-0.36	-0.81
P11	.192*	.264**	.261**	.390***	-	.592***	.518***	.515***	-.007	.779***	3.82	0.85	-0.43	-0.81
P12	.296***	.332***	.359***	.480***	.615***	-	.578***	.514***	-.012	.765***	3.95	0.82	-0.57	-0.27
P13	.223**	.206*	.276**	.459***	.549***	.612***	-	.764***	.056	.800***	4.28	0.72	-0.92	0.34
P14	.241**	.247**	.316***	.457***	.624***	.595***	.845***	-	.058	.783***	4.22	0.75	-0.88	0.74
C1 ^{주1}	.783***	.779***	.802***	.689***	.194*	.287**	.199*	.196*			2.97	0.72	0.08	0.10
P1 ^{주1}	.260**	.298***	.306***	.503***	.801***	.809***	.817***	.845***			4.04	0.70	-0.54	-0.15
M	3.23	3.03	2.87	3.58	3.84	3.91	4.20	4.12	2.66	4.10				
SD	1.06	1.07	1.10	1.06	0.98	0.94	0.84	0.85	0.77	0.59				
왜도	-0.15	-0.03	0.34	-0.42	-0.32	-0.32	-0.73	-0.60	-0.06	-0.22				
첨도	-0.90	-1.06	-0.68	-0.78	-0.49	-0.56	0.16	-0.25	0.03	-0.51				

* $p < .05$, ** $p < .01$, *** $p < .001$, C-흥미유지, P-노력지속(대각선 위: 2차, 대각선 아래: 1차)

주1. C1, P1의 상관계수는 1차년도 각 문항과 해당 하위요인의 수치임. $r_{C1,P1}=.215$, $p < .05$.

주2. C2, P2의 상관계수는 2차년도 각 문항과 해당 하위요인의 수치임. $r_{C2,P2}=-.043$, ns.

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

3. 연구결과

<표 3> 요인모형들의 적합도 지수

	1차					2차				
	χ^2	df	CFI	TLI	RMSEA	χ^2	df	CFI	TLI	RMSEA
모형 1	202.36	20	.682	.555	.258	263.625	20	.513	.318	.298
수정모형	52.382	16	.937	.889	.129	69.284	16	.894	.814	.156
모형 2	56.039	19	.941	.913	.114	39.236	19	.960	.940	.088
수정모형	32.260	18	.975	.961	.076	23.572	18	.989	.983	.048
모형 3	모형 2와 동치모형					모형 2와 동치모형				
모형 4	NC [*]					NC				

NC= non-convergent.

<표 5> 2요인 모형의 측정불변성 검증(적합도 지수)

	χ^2	df	CFI	TLI	RMSEA
기저모형	179.439	99	.931	.916	.077
완전측정동일성(흥미유지 측정변수의 요인 계수를 두 시점 간 동일하게 제약을 가한 모형)	189.307	102	.924	.910	.079
부분측정동일성(흥미유지 측정변수 중 한 문항의 요인 계수를 두 시점 간 자유롭게 추정하도록 한 모형)	179.631	101	.931	.918	.075
완전측정동일성(위 모형에 더하여 노력지속 측정변수의 요인 계수를 두 시점 간 동일하게 제약을 가한 모형)	183.900	104	.930	.919	.075

<표 6> 2차 요인 모형의 측정불변성 검증(적합도 지수)

	χ^2	df	CFI	TLI	RMSEA
기저모형	184.113	100	.927	.912	.078
완전측정동일성(흥미유지 측정변수의 요인 계수를 두 시점 간 동일하게 제약을 가한 모형)	195.208	103	.919	.906	.081
부분측정동일성(흥미유지 측정변수 중 한 문항의 요인 계수를 두 시점 간 자유롭게 추정하도록 한 모형)	185.146	102	.927	.915	.077
완전측정동일성(위 모형에 더하여 노력지속 측정변수의 요인 계수를 두 시점 간 동일하게 제약을 가한 모형)	192.323	105	.924	.913	.078

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

3. 연구결과

그릿의 요인 구조를 단일 요인, 2요인, 2차 요인, 2중 요인으로 구분하여 각 모형의 적합도를 알아보았다. 모형 비교에 의하면 그릿은 단일 요인이라기 보다는 '흥미유지', '노력지속'으로 서로 관련 있으나 구분되는 두 요인으로 나뉘는 것을 알 수 있다.

최근의 그릿 연구들에서는 흥미유지, 노력지속을 구분한 2요인 모형을 사용한 연구들이 지지되고 있는 추세이다. 단일 요인 모형은 Duckworth와 Quinn(2009), Wolter와 Hussain(2015), 임효진(2017a, 2017b)에서도 지속적으로 모형 적합도가 나쁜 것으로 나타나고 있다. 이에 비해 2요인 모형은 초등학생(임효진, 2017b), 고등학생(Muenks et al., 2017), 대학생(임효진, 2017a), 성인(Hwang et al., 2018) 집단에서 일관되게 자료를 잘 설명하는 것으로 나타나 그릿에 있어서 두 요인이 독립적인 역할을 하고 있음을 알 수 있다.

반면 2차 요인에 대해서는 원래의 개념적 정의를 반영하고 있음에도 불구하고 이를 사용한 연구들이 거의 없는 편이다. 이는 그릿에 대한 개념적 검토와 이에 대한 경험적 검증의 선행 연구들이 부족했기 때문에, 단일 요인의 작용에 관심을 가지거나(예: 류재준, 임효진, 2018) 혹은 흥미유지, 노력지속의 차별적 작용(예: Hwang et al., 2018)을 확인한 연구들이 주를 이루었기 때문인 것으로 파악된다. 또한 위계가 존재하는 다수의 요인들을 가정하는 모형(고차요인 모형)은 특수한 경우에만 존재하며, 이러한 모형을 검증하는 경우 자체가 드물기 때문에 현재까지 관련 연구들의 수가 적은 것으로 보인다.

Article Review : 그릿의 요인 구조의 안정성에 관한 연구

*Appendix

본 연구의 척도에서 '흥미유지' 요인의 문항들은 모두 부정진술문으로, '노력지속' 요인의 문항들은 모두 긍정진술문으로 구성되어 있다는 점에 주목하였다. 자기보고식 검사에서는 응답자의 경향성 혹은 문항의 내용이 아니라 문항의 형식 때문에 나타나는 방법효과가 반응 편파성(bias)을 일으킬 수 있다 (Quilty, Oakman, & Risko, 2006). 즉 긍정, 부정진술문이 혼합된 심리적 평정 척도의 경우 전체 신뢰도가 낮아지거나(Schriesheim & Hill, 1981), 측정오차가 크게 증가할 가능성이 있다(Tepper & Tepper, 1993). 이러한 척도의 요인구조를 검토할 때에는 방법효과(method effect)를 고려해야 하며 (Horan, DiStefano, & Motl, 2013), 이는 문항형태로 인한 방법효과를 통제하면 검사가 측정하고자 하는 요인구조가 보다 정확히 나타날 수 있기 때문이다 (홍세희, 노언경, 정송, 2011).

부정문항의 방법효과를 고려하기 위해 4개의 문항에 오차 간 공분산을 설정한 모형과, 방법효과가 2차적 요인으로 존재한다고 가정하는 모형. 같은 방법으로 긍정문항의 방법효과를 고려하기 위해 4개의 문항에 오차간 공분산을 허용한 모형. 긍정문항의 방법효과를 2차적 요인으로 설정한 모형 등을 생각해 볼 수 있다.