

Print ISSN: 1738-3110 / Online ISSN 2093-7717
<http://dx.doi.org/10.15722/jds.13.9.201509.103>

[Field Research]

Factor Analysis for Exploratory Research in the Distribution Science Field

유통과학분야에서 탐색적 연구를 위한 요인분석

Myung-Seong Yim(임명성)*

Received: August 21, 2015. Revised: September 15, 2015. Accepted: September 15, 2015.

Abstract

Purpose – This paper aims to provide a step-by-step approach to factor analytic procedures, such as principal component analysis (PCA) and exploratory factor analysis (EFA), and to offer a guideline for factor analysis. Authors have argued that the results of PCA and EFA are substantially similar. Additionally, they assert that PCA is a more appropriate technique for factor analysis because PCA produces easily interpreted results that are likely to be the basis of better decisions. For these reasons, many researchers have used PCA as a technique instead of EFA. However, these techniques are clearly different. PCA should be used for data reduction. On the other hand, EFA has been tailored to identify any underlying factor structure, a set of measured variables that cause the manifest variables to covary. Thus, it is needed for a guideline and for procedures to use in factor analysis. To date, however, these two techniques have been indiscriminately misused.

Research design, data, and methodology – This research conducted a literature review. For this, we summarized the meaningful and consistent arguments and drew up guidelines and suggested procedures for rigorous EFA.

Results – PCA can be used instead of common factor analysis when all measured variables have high communality. However, common factor analysis is recommended for EFA. First, researchers should evaluate the sample size and check for sampling adequacy before conducting factor analysis. If these conditions are not satisfied, then the next steps cannot be followed. Sample size must be at least 100 with communality above 0.5 and a minimum subject to item ratio of at least 5:1, with a minimum of five items in EFA. Next, Bartlett's sphericity test and the Kaiser-Mayer-Olkin (KMO) measure should be assessed for sampling adequacy. The chi-square value for Bartlett's test should be significant. In addition, a KMO of more than 0.8 is recommended. The next step is to conduct a factor analysis. The analysis is composed of three stages. The first stage determines a rotation technique. Generally, ML or PAF

will suggest to researchers the best results. Selection of one of the two techniques heavily hinges on data normality. ML requires normally distributed data; on the other hand, PAF does not. The second step is associated with determining the number of factors to retain in the EFA. The best way to determine the number of factors to retain is to apply three methods including eigenvalues greater than 1.0, the scree plot test, and the variance extracted. The last step is to select one of two rotation methods: orthogonal or oblique. If the research suggests some variables that are correlated to each other, then the oblique method should be selected for factor rotation because the method assumes all factors are correlated in the research. If not, the orthogonal method is possible for factor rotation.

Conclusions – Recommendations are offered for the best factor analytic practice for empirical research.

Keywords: Principal Component Analysis, Component Analysis, Factor Analysis, Common Factor Analysis, Exploratory Factor Analysis.

JEL Classifications: B41, C00, C01, C10, C12, C18, M19.

1. 서론

1900년대 Charles Spearman에 의해 대중화된 요인 분석은 심리학적 도구를 개발하거나, 관측변수를 정제하거나, 해당 도구의 타당성을 평가할 때, 해당 도구와 관련된 이론을 검정할 때, 새로운 잠재개념(constructs)을 탐색할 때, 개념타당성(construct validity)을 평가하거나 일부의 경우 가설을 검정할 때 등 다양한 상황에서 빈번히 사용되고 있으며, 유통과학뿐만 아니라 다양한 학술분야에서 실증분석을 위해 가장 널리 사용되는 통계적 기법 중 하나이다(Conway & Huffcutt, 2003; Floyd & Widaman, 1995; Tinsley & Tinsley, 1987; Treiblmaier & Filzmoser, 2010). 요인분석(Factor Analysis)이란 큰 수의 상호 연관된 변수들(interrelated variables)을 적은 수의 잠재적 차원(latent or hidden dimensions)으로 축소해주는 분석 기법을 말한다(Tinsley & Tinsley, 1987). 여기서 요인(Factor)이란 용어는 잠재변수(Latent Variable), 가설화된 개념(Hypothetical Construct), 비관측변수(Unobserved Variable), 이론(Theory) 등 다양하게 불려왔다(Floyd & Widaman, 1995; Tinsley & Tinsley, 1987). 요인분석의 목적은 질문을 통해 측정된 관측 변수(observed variables)의 집합 내에 존재하는 상관관계 행

* Assistant Professor, Dept, Business Administration, Sahmyook University, Seoul Korea, Tel: +82-2-3399-1565, E-mail: msyim@syu.ac.kr

렬에서 최대 공통분산(common variance)을 설명해주는 가장 적은 수의 탐색적 개념을 사용하여 간명성을 달성하는 것이다(Beavers et al., 2013; Tinsley & Tinsley, 1987).

요인분석의 유용성에도 불구하고 분석 절차의 복잡성과 유연성은 많은 연구에서 상당한 모호함을 유발하고 있다(Floyd & Widaman, 1995). 많은 경우, 분석 절차의 차이가 연구마다 다른 결과를 산출하고 있다(Floyd & Widaman, 1995).

선행연구를 살펴보면 빈번하게 탐색적 요인 분석과 관련된 선택에서 연구자들이 실수를 범하는 경우가 많다(Conway & Huffcutt, 2003). 대표적으로 공통 요인(common factor) 분석이 더 나옴에도 불구하고 주성분분석(Principal Components Analysis, PCA)을 선택하는 경우, 요인을 추출하기 위한 근거에 대한 근거 부족(일반적으로 고유값(Eigenvalue) 1 이상을 적용함), 사각회전보다는 직각회전의 사용 등 다양한 측면에 대한 고려가 부족한 경우가 여전히 많이 목격된다(Conway & Huffcutt, 2003). 이러한 사례는 유통과학분야에서도 목격되는데 예를 들어, 최근 연구 중에 Asgari & Hosseini(2015)는 연관 데이터를 확보하기 위해 주성분 분석을 사용하였다. 또 다른 예로, Kim(2015)도 요인들의 중복성과 척도의 정제를 위한 요인추출모델로 주성분분석을 사용하였다. 마찬가지로, Chen & Hwang(2014)도 측정도구의 신뢰성 및 타당성을 확보하기 위한 요인 분석 기법으로 주성분분석을 사용했음을 언급하였다. 이처럼 여러 학문분야에서 주성분분석을 요인분석기법으로 사용하는 경우가 많으며 유통과학분야도 예외는 아니다. 따라서 유통과학분야에서도 요인분석의 근본적인 목적을 달성하기 위해서는 요인분석을 위한 절차 및 구체적인 기법의 선택을 위한 가이드라인이 제시되어야 한다.

성분분석(component analysis)과 공통요인분석(common factor analysis)은 이론적으로나 통계적으로 명확하게 차이가 있다(Beavers et al., 2013). 이 두 기법을 명확히 구분하지 못할 경우 이론적으로 견고한 의사결정을 하는데 장애 요인으로 작용할 수 있기에 명확한 구분 하에 해당 기법을 사용하는 것이 무엇보다 중요하다(Beavers et al., 2013).

따라서 본 연구는 측정도구의 개발과 정제를 위해 사용되는 탐색적 요인분석을 위한 가이드라인을 제시하고자 한다. 구체적으로 가장 많이 혼선을 빚고 있는 주성분분석과 요인분석 간의 차이를 제시하고 이후 탐색적 요인 분석을 위한 절차 및 연구자가 충족해야 할 기준을 제시하고자 한다. 각 기준은 기존 연구에서 제시하고 있는 기준을 기반으로 (1) 어떠한 요인 모델(factor model)이 사용되어야 하는지, (2) 몇 개의 요인을 유지할 것인지, (3) 회전 기법은 무엇을 사용할 것인지, (4) 요인 구조의 해(solution)에 대한 해석을 어떻게 할 것인지(Ford et al., 1986)를 중심으로 기술하고자 한다. 이를 통해 앞으로 다양한 실증분석에서 수행되는 요인분석을 위한 명확한 가이드라인을 제시하고자 한다.

2. 요인분석에 대한 이해

다변량 기법(multivariate techniques)은 크게 두 개의 차원으로 구분할 수 있는데, 하나는 다변량 예측(multivariate prediction)과 다른 하나는 다변량 공분산 분석(multivariate covariance analysis)이다(Tinsley & Tinsley, 1987). 다변량 예측에는 회귀분석(regression analysis), 판별함수분석(discriminant function analysis, DFA), 정준 상관 분석(canonical analysis), 다변량 분산분석(multivariate analysis of variance, MANOVA) 등이 해당된다(Tinsley & Tinsley, 1987). 다변량 예측에서 연구자는 다양한 예측변수에 대한 정보를

가지고 응답자의 상태(respondents' statuses)를 예측한다(Tinsley & Tinsley, 1987). 다변량 공분산 분석은 요인분석, 주성분분석, 선형 구조 관계 분석(linear structural relations, LISREL), 군집분석(cluster analysis) 등이 해당된다(Tinsley & Tinsley, 1987).

일 세대 회귀모형(First Generation Regression Models)에서 요인타당성(Factorial Validity)은 주로 탐색적 요인분석에 의해 측정되었다(Gefen & Straub, 2005). 요인 모델은 일반적으로 두 가지 형태가 있는데 하나는 공통요인 분석과 성분 분석이다(Ford et al., 1986). 이중에서 다양한 연구 분야에서 가장 일반적으로 관찰되는 사용기법은 주성분분석이다(Gefen & Straub, 2005).

요인분석은 자료 축소(data reduction)를 위한 수단이 아닌 측정도구의 개념 타당성(construct validity)을 평가하기 위해 사용된다(Floyd & Widaman, 1995). 구체적으로 공변(covary)하는 관측변수(manifest variable)에 영향을 미치는 어떠한 잠재변수(latent variable)를 발견하기 위해 요인분석이 사용된다(Costello & Osborne, 2005). 반면에 주성분분석은 오직 자료 축소(data reduction)를 위한 기법이다(Costello & Osborne, 2005). 따라서 두 기법 간에는 명확한 차이가 존재한다. 하지만 많은 경우 이러한 목적이 무시되고 모두 요인분석을 위해 사용되는 경우가 많다. 그 이유로 세 가지를 제시할 수 있는데 하나는 통계 분석을 위한 도구의 특성 때문이다. 예를 들어 실증분석을 위해 가장 많이 사용되는 IBM사의 SPSS는 요인분석을 위한 도구로 주성분분석을 기본설정(default)으로 제공한다(Costello & Osborne, 2005). 따라서 많은 연구자들은 이를 요인분석을 위한 도구로 사용하기도 한다. 둘째로 주성분 분석이 상대적으로 사용하기 편하다는 장점이다. 총 분산을 사용하는 주성분분석과 비교하여 공통요인만을 사용하는 요인분석의 경우 해석이나 분석 절차가 다소 복잡하다. 따라서 주성분분석 기법이 더욱더 선호되는 경향이 있다. 셋째, 여러 연구자가 제시한 바와 같이 이 두 가지 분석기법의 결과가 유사하다는 점이다. 결국 유사한 결과를 도출해 낸다면 사용과 해석이 용이한 기법이 선호되기 마련이다. 또한 두 접근법 모두 데이터 집합 내에 다른 변수들에 비례하여 주어진 변수에 대한 분산이 얼마나 분포되어 있는가를 평가한다는 공통점이 있다(Ford et al., 1986). 하지만 유사한 결과의 도출은 특정한 상황적 조건이 충족되어야만 가능하다. 모든 경우 동일한 결과를 도출할 수 없다. 따라서 본 연구에서는 명확한 구분을 위해 주성분분석과 (탐색적)요인분석(혹은 공통요인 분석)의 차이점을 구분해 보고자 한다.

2.1. 주성분분석과 (공통)요인분석의 차이

Floyd & Widaman(1995)은 탐색적 요인 분석이 주성분 분석과 공통 요인 분석 등 두 가지 접근법으로 구성된다고 주장하였다. 탐색적 요인 분석 및 주성분 분석 모두 다변량 통계 기법으로 사회과학 연구 및 행동 과학 연구에서 널리 사용되고 있다(Ledesma & Valero-Mora, 2007). 두 분석 절차가 모두 실증적으로 수렴 해(convergent solution)를 제시한다면, 주성분분석과 공통요인분석의 이론적 차이는 중요한 고려대상이 아니다(Snook & Gorshuch, 1989). 하지만 주성분분석과 요인분석의 결과가 유사한 경우가 있기는 하나 두 방법론은 개념적으로 구별된다(Budaev, 2010). 주성분분석은 자료의 축소를 위해서 사용되며, 공통요인분석은 기저의 잠재 변수 측면에서 관측 변수의 집합 간의 관계를 이해하는데 사용된다(Beavers et al., 2013; Floyd & Widaman, 1995). 구체적으로 주성분분석은 차원 축소 기법(dimensionality reduction method)인 반면 요인 분석은 변수간의 상관관계를 설명해주는 잠재변수를 측정하기 위한 기법이다(Budaev, 2010). 따라서 주성분분석은 차원의 수를 줄이는 것을 목적으로 한 경우 이용할 수 있는 방법이

다(Budaev, 2010). 반면에 요인분석은 잠재변수를 결정하고 평가하는 것이 목적일 경우 사용할 수 있다(Budaev, 2010).

이 두 접근법 간의 차이점 다음과 같다. 첫째, 공통성(communality)에 관한 가정이다(Floyd & Widaman, 1995). 공통성이란 관측변수들의 집합의 기저를 형성하는 잠재변수들과 관측변수가 공유하는 분산을 의미한다(Floyd & Widaman, 1995). 주성분분석에서는 관측변수의 상관관계 행렬의 대각선 값들을 1로 지정하는데, 1은 관측변수의 분산의 모든 것을 나타낸다(Floyd & Widaman, 1995; Tinsley & Tinsley, 1987). 이는 주성분분석에서 관계를 검정할 때 세 가지 분산을 구분하지 않고 분산의 전부(공통분산, 고유분산, 오차분산)를 분석에서 사용한다는 것을 의미한다(Beavers et al., 2013; Costello & Osborne, 2005; Tinsley & Tinsley, 1987). 즉, 성분 분석에서는 총 분산이 반영된다(Beavers et al., 2013). 따라서 주성분분석에서 총분산은 관측변수들이 표준화 상태에서 분석되는 관측변수의 수와 일치한다(Alii 2010). 표준화되었다는 것은 평균이 0이고, 표준편차가 1이고 행렬의 대각선 값이 1인 상태를 말한다(Alii 2010). 또한 설명분산은 상관관계 행렬에서 대각선 값의 합과 같다(Alii 2010). 본 분석에서 추출되는 성분의 수는 분석에서 사용되는 관측변수의 수와 동일하다(Alii 2010).

반면에 공통 요인 분석에서는 상관관계 행렬의 대각선에 공통성을 둔다(Floyd & Widaman, 1995). 즉 공통 요인 분석은 각각 변수의 공통 분산만을 반영한다(Costello & Osborne, 2005; Floyd & Widaman, 1995; Tinsley & Tinsley, 1987). 요인분석에서 사용되는 항목분산(item variance, total test variance, or variance of each measured variable)은 하나 이상의 항목(item or test)과 관련된 분산인 공통분산(common variance)과 특정한 항목(specific item or test)에만 해당되는 고유분산(unique variance or specific variance)의 합이다(Floyd & Widaman, 1995; Tinsley & Tinsley, 1987). 공통분산은 한 변수(a variable or a test)의 분산이 다양한 요인(factors)에 나누어져 분포되어 있는 분산이며 잠재변수와 관련된다(Floyd & Widaman, 1995; Tinsley & Tinsley, 1987). 또한 공통분산은 항목의 변동성(variability)을 나타내며, 다른 항목 및 요인과 공유된다(Beavers et al., 2013). 공통분산만 포함시킨다는 것은 더욱더 정확한 해를 도출한다는 것을 의미한다(Beavers et al., 2013). 물론, 일반적으로 여러 연구에서 주성분분석과 공통요인분석이 유사한 결과를 도출한다는 주장도 있으나 이는 신뢰성이 있는 항목들(reliable measures)이 사용되었을 경우에만 해당된다는 전제 조건이 있다(Beavers et al., 2013). 고유분산은 각각의 관측변수와 관련되는 신뢰분산(reliable variance)과 각 변수의 무작위 오차(random error variance)의 조합으로 구성된다(Floyd & Widaman, 1995). 고유분산은 항목의 고유한 특성으로 인해 유발되는 분산을 의미하며, 다른 변수나 요인에 의해 설명될 수 없다(Beavers et al., 2013). 오차 분산(error variance)은 측정과정과 관련되며 비신뢰성(unreliability)을 나타낸다(Beavers et al., 2013). 정리하면, 요인분석은 하나의 항목에 대한 공통분산을 줄여 더 적은 수의 개념적으로 의미가 있는 변수들로 만들기 위해 그리고 각각의 기본 단위(each basic unit(tests or items))가 어떠한 구조를 취하고 있는지를 이해하기 위해 사용된다(Tinsley & Tinsley, 1987). 다양한 연구에서 사용되고 있는 항목(item)은 문항(questionnaire), 측정지표(indicator of latent variable), 검사문항(test), 척도(scale), 측정항목(measure), 측정변수(measured variable), 관측변수(observed variable), 명시변수(manifest variable) 등 다양하게 불리나 동일한 개념이다(Floyd & Widaman, 1995; Costello & Osborne 2005; Tinsley & Tinsley 1987).

실증분석을 통해 규명된 두 방법론의 차이점은 회전기법의 유

형과 관계없이 주성분분석을 통해 도출된 적재값이 공통요인분석보다 높게 나타났다는 것인데, 이는 주성분분석이 공통분산과 고유분산을 모두 포함하기 때문이다(Snook & Gorshuch, 1989).

$$\text{Total variance} = \text{common variance(shared variance)} + \text{specific variance(unique variance)} + \text{error variance(measurement error)}$$

둘째, 주성분분석에서는 자료 내에 존재하는 모든 변동성이 주성분에 의해 설명된다고 가정한다(Budaev, 2010). 주성분분석은 단순한 몇 개의 행렬만을 포함하지만 요인분석은 복잡하고 반복적인 계산이 요구된다(Budaev, 2010).

셋째, 주성분분석에서는 구조적 가정(underlying structural assumptions)이 존재하지 않는다(Beavers et al., 2013). 반면에 공통요인분석에서는 모든 응답 항목이 기반 잠재변수의 산물(product)이라고 가정한다(Beavers et al., 2013).

넷째, Snook & Gorshuch(1989)이 실시한 실증분석에 따르면 공통요인분석보다 주성분분석에서 모집단 유형(population pattern)에 따른 해의 변동과 오류(bias)가 큰 해를 산출한다는 것을 밝혔다(Snook & Gorshuch, 1989). 여기서 정확성(accuracy)은 오류(bias)와 변동(variability)의 합을 말한다(Snook & Gorshuch, 1989). 이러한 경향은 적재값이 낮고 사용된 변수의 수가 적을 경우 더 강하게 나타났다(Snook & Gorshuch, 1989). 즉 공통요인분석을 통해 도출된 해가 주성분분석보다 더 정확하다고 볼 수 있다.

다섯째, 주성분분석에서 초기 추출(initial extraction)은 각각의 선형 결합(linear combination)이 서로 직각(orthogonal)이라고 가정한다(Beavers et al., 2013). 여기서 직각은 서로 독립적 혹은 비상관관계(independent or uncorrelated)에 있다는 것을 의미한다(Beavers et al., 2013). 본 분석에서 선형결합을 만드는 것은 측정항목들 간에 최대 분산을 설명하기 위함이다(Beavers et al., 2013). 따라서 여기서 도출된 성분 중 첫 번째 성분은 항목 내에서 최대분산을 가지며, 다음 성분은 첫 번째 성분에 포함되지 않은 나머지 분산 중 가장 많은 분산을 설명하는 성분이 된다(Beavers et al., 2013).

마지막으로, 성분(component)이란 관측변수의 집합(composite)을 말하며, 각각의 개별 항목 점수는(individual item scores) 성분을 구성·정의(cause or define) 한다(Beavers et al., 2013). 반면에, 요인 분석에서 개별 항목 점수는 기반 요인 혹은 잠재개념의 결과이다(Beavers et al., 2013).

지금까지 제시한 두 기법 간의 차이에도 불구하고 몇몇 학자들은 이 두 기법이 도출해내는 결과값의 차이가 유사하거나 동일하다고 주장하기도 한다. 하지만 여러 실증 분석에서 제시한 바와 같이 이와 같은 결과를 도출하기 위해서는 다음과 같은 여러 가지 상황적 조건이 충족되어야 한다. Gorsuch(1983)는 관측변수의 수가 35개 이상이고, 각각의 관측변수의 공통성이 0.7을 넘는다면 공통성 추정치간의 차이는 무시할 만하며, 요인분석을 위해 주성분 분석을 사용해도 충분하다고 주장하였다. 그는 Monte Carlo 연구를 통해 관측변수의 수가 30개 수준이고 공통성이 0.4이상일 경우에도 연구결과에서 큰 차이가 발견되지는 않았다고 주장하였다(Gorsuch, 1983). 즉, 분석에 사용되는 변수의 수가 충분하고, 각각의 관측변수의 공통성이 높은 경우 두 분석 결과의 차이는 크지 않다는 것이다(Gorsuch, 1983). Floyd & Widaman(1995)도 요인당 관측변수의 수와 공통성이 증가한다면 두 기법간의 차이는 크지 않다고 주장하였다. Gorsuch(1983)는 반대로 관측변수의 수가 상대적으로 적을 경우, 지나치게 높은 요인 적재값과 데이터에 대한 오해를 피하기 위해서 공통요인분석을 사용해야 한다고 주장하였다. 즉, 각각의 차원(요인)에 적재된 관측변수의 수가 적고 상대

적으로 공통성의 값이 작다면 두 분석 간의 차이는 상당히 클 수 있다(Floyd & Widaman, 1995). 따라서 연구자가 적은 수의 잠재 변수를 이용하여 학문적 현상을 이해하고자 할 경우에는 주성분분석보다는 공통요인분석을 사용하는 것이 적절하다(Floyd & Widaman, 1995).

정리하면, 많은 학자들은 실제로 주성분분석이 모수에 대한 유사한 추정치를 산출하는지에 대해 의문을 제기하였다(Floyd & Widaman, 1995). 주성분분석은 연구자의 관심이 자료의 축소일 경우에 사용하는 것인지 요인을 반영하는 모수(parameters)를 확보하는데 사용되지는 않는다(Treblmaier and Filzmoser, 2010). 이처럼 주성분분석은 사회과학에서 주로 탐색적 자료 축소 절차로 사용될 뿐 통계적으로 탐색적 요인 분석과 전혀 다르다(Alili 2010). 따라서 주성분분석은 요인분석을 위한 추출방법이 아니다(Costello & Osborne, 2005).

3. 요인분석을 위한 적합성 검증

연구자는 요인분석을 사용하기 이전에 자신이 수집한 자료가 요인분석을 수행하기에 적합한지를 먼저 평가해야 한다. 다음의 여러 가지 기준들은 자료 행렬(data matrix)이 요인분석에 적합한지(appropriateness)를 평가하는 기준들이며(Tinsley & Tinsley, 1987), 요인분석을 수행하기 위해서는 이 기준들이 충족되어야 한다.

3.1. 표본 크기(sample size)

지금까지 표본(sample), 대상(subject), 관찰대상(observation), 참여자(participant), 개인(individual) 등 다양하게 불린(Alili, 2010; Arrindell & van der Eijde 1985; Floyd & Widaman 1995; Gorsuch 1983; Tinsley & Tinsley 1987) 표본은 수가 클수록 요인 분석에서 더 좋은 결과를 산출한다는 것은 정설로 간주되어 왔다(Alili 2010; Floyd & Widaman, 1995). 그 이유는 첫째, 큰 표본을 사용할 경우 표본 요인 적재값이 모수 요인 적재값을 더욱더 정확하게 추정할 수 있다(Alili 2010; MacCallum et al., 1999). 둘째, 변수 간의 상관관계는 표본마다 변동이 있는데, 큰 수의 표본보다 적은 수의 표본에서 해당 변동이 더욱 심하기 때문에 큰 표본을 사용할 경우 변동가능성이 낮아진다(MacCallum et al., 1999; Tinsley & Tinsley, 1987). 셋째, 표본이 클 경우 표본 오차를 줄여 주기 때문에 반복 분석에서도 매우 유사한 결과를 얻을 수 있어 더욱더 안정적이다(Alili 2010; Beavers et al., 2013; MacCallum et al., 1999). 넷째, 표본 수가 증가함에 따라 무작위 오차는 줄어들 뿐만 아니라 측정항목의 모수도 안정화된다(Tinsley & Tinsley, 1987). 다섯째, 부적절한 표본 수를 사용할 경우 요인 분석 과정에 손상을 줄 수 있으며 이는 곧 신뢰할 수 없는 결과를 도출하게 만든다(Beavers et al., 2013). 여섯째, 표본의 수가 적은 경우 Heywood cases(요인 적재값이 1.0을 초과하는 경우)가 발생할 수 있다(Costello & Osborne, 2005). 일반적으로 Heywood cases는 요인 분석에서 자주 발견되는 것은 아니나 표본의 수가 해를 도출하기에 적합하지 않아서 해의 도출이 실패할 경우에 발생할 수 있다(Costello & Osborne, 2005). 마지막으로, 적은 수의 표본으로 요인을 일반화하기도 어렵기 때문에 요인분석을 위해서는 적절한 수의 표본이 반드시 확보되어야 한다(Alili 2010; Tinsley & Tinsley, 1987).

이와 같이 큰 표본의 장점에도 불구하고 많은 경우(예, 분석이 CEO대상, 기업단위 등) 수집할 수 있는 표본의 수는 제한적이기

에 얼마나 많은 표본이 있어야 올바른 추정이 가능한지에 대한 의문은 꾸준히 제기되어 왔다. 기존의 요인분석에 대한 문헌들에서 표본 수에 대해 많은 관심을 두고 있다는 것은 이와 같은 맥락에서 이해할 수 있다(MacCallum et al., 1999). 하지만 기존 문헌을 검토해보면 표본 크기에 대한 절대적 가이드라인이나 많은 학자들이 자주 차용하는 기준은 존재하지 않아서 여전히 표본 크기에 대한 논란은 진행형이다. 따라서 다음 장에서는 최소 표본에 대한 기준을 문헌 검토를 통해 살펴보았다. 표본에 대한 기준은 절대적 기준과 상대적 기준으로 구분하여 살펴보았다. 절대적 기준은 요인 분석을 위한 최소 표본 수를 명확하게 제시한 경우를 말하며 상대적 기준은 측정항목대비 표본 수의 비율을 반영한 기준을 의미한다.

3.1.1. 정량적 기준: 절대적 기준(a minimum number of cases)

Guadagnoli & Velicer(1988)는 상대적 기준(subject to item ratios)보다는 절대적 최소 표본 기준(absolute minimum sample sizes)이 더 적절하다고 주장하였다. 그들의 주장에서 제시된 최소 표본은 50개에서 400개였다. 선행 연구에 따르면 표본 100개는 취약한 수준(poor), 200개는 일반적인 수준(fair), 300개는 좋은 수준(good), 500개는 매우 좋은 수준(very good), 1,000 이상은 매우 훌륭한 수준(excellent)의 표본 수라고 보았다(Comrey & Lee, 1992; MacCallum et al., 1999; Tinsley & Tinsley, 1987).

다른 학자들이 제시한 최소 표본 기준에 따르면, 100개 이상(Budaev, 2010; Suhr, 2006), 150개 이상(Beavers et al., 2013; Hutcheson & Sofroniou, 1999), 200개 이상(Gorsuch, 1983), 300개 이상(Norušis, 2005) 등 다양하다. 하지만 정량적 표본 기준은 해당 연구에 사용되는 변수의 수를 고려하지 않고 있기 때문에(Tinsley & Tinsley, 1987) 본 기준을 절대적 기준으로 받아들이기에는 한계가 있다는 비판을 받기도 한다. Osborne & Fitzpatrick(2012)은 더욱더 정확하고 안정적인 결과를 도출하는데 있어서 추정되는 모수의 수가 더 중요하게 고려되어야 한다고 주장하였다. Budaev(2010)는 높은 품질의 자료(높은 신뢰성을 가진 관측변수들을 사용하고, 소수의 명확한 잠재변수가 사용되고, 공통성이 높은 경우)가 사용된 경우는 25개의 표본으로도 분석이 가능하다고 주장하기도 하였다. 따라서 절대적 기준에 따른 기준을 반영한 복합적 절대 기준을 살펴보면 다음과 같다.

변수의 수를 고려한 기준: (관측)변수의 수를 고려하여 Lawley & Maxwell(1971)은 변수의 수 대비 51개 많은 표본이 필요하다고 주장하였다.

공통성(communality, h^2)을 고려한 기준: MacCallum et al.(1999)은 지금까지 존재하는 요인 분석을 위한 표본 수에 대한 경험적 규칙이 유용하지도 그리고 타당하지도 않다고 주장하였다. 그들은 지금까지 제시된 기준이 일괄되지도 않고 연구마다 다르며, 변수의 수 그리고 연구의 설계 방법에 따라 상이하다고 주장하였다(MacCallum et al., 1999). 반면 그들은 공통성을 기준으로 표본 수를 결정할 것을 주장하였는데, 그들의 주장에 따르면, 공통성이 높은 경우(0.6 이상) 100개 미만의 표본으로도 적절한 해를 도출할 수 있다고 보았다(MacCallum et al., 1999). 또한 공통성이 0.5 이상에서는 100개에서 200개의 표본으로 있어서 해를 도출할 수 있다고 주장하였다(MacCallum et al., 1999). 반면에 낮은 공통성에서는(0.5미만), 그리고 3~4개의 관측변수를 가진 소수의 요인이 사용된 경우 300개 이상의 표본이 필요하다고 보았다(MacCallum et al., 1999). 또한 공통성이 낮은 반면 요인의 수가 많은 경우 500개 이상의 표본이 필요하다고 주장하였다(MacCallum et al., 1999).

물론 이와 상반된 주장도 제기되었는데, Tinsley & Tinsley(1987)는 표본이 300개를 넘을 경우 항목 대 표본의 수는 특별히 문제되지 않는다고 주장하였다. 즉, 표본의 수가 300개가 넘을 경우 복합적 기준이 사용될 필요가 없다는 것이다.

3.1.2. 정성적 기준: 상대적 기준(a subjects-to-variables ratio)

여러 선행연구들은 절대적 표본 크기보다는 상대적 표본 크기(variable to subject ratio)가 적절한 기준이라고 주장하였다(Budaev, 2010). 그동안 여러 연구에서 절대적 표본 수에 대한 중요성이 강조되어왔지만, 대부분의 연구자들은 상대적 기준(ratio between subjects and variables)을 더 중요하게 여겨왔다(Treblmaier and Filzmoser, 2010). 빈번히 인용되는 최소 표본 기준은 변수의 수 대비 3:1(Budaev, 2010), 4:1(Floyd & Widaman, 1995), 5:1(Gorsuch, 1983; Suhr, 2006), 10:1이다(Bryant & Yarnold, 1995; Nunnally, 1978; Tinsley & Tinsley, 1987; Treblmaier and Filzmoser, 2010). 선행연구를 검토해 본 결과 많은 연구에서 표본의 수가 10:1이하인 경우가 조사 대상 연구의 62.9%였다(Costello & Osborne, 2005). 하지만 이 기준이 지나치게 단순화되었다고 비판받기도 한다(Budaev, 2010; Treblmaier and Filzmoser, 2010). Costello and Osborne(2005)의 연구에 따르면 10:1 기준을 준수한 연구에서 예상되는 요인 구조를 도출한 경우는 60%에 지나지 않았으며, 20:1의 기준을 준수한 연구의 경우 오직 70%의 확률로 견고한 그리고 명확한 요인 구조를 발견할 수 있었다. 따라서 본 기준만을 준용하기보다는 복합적 기준을 준수할 것을 주장하는 경우가 많았다.

Gorsuch(1983)은 탐색적 요인 분석을 수행하기 위해서는 반드시 각 변수 당 표본의 수가 최소 5개 이상이고, 전체 표본의 수가 200개 이상 되어야 한다고 주장하였다. Streiner(1994)는 전체 표본의 수가 100개 이상일 경우 변수 당 표본의 수가 5개 이상 되어야 하며, 전체 표본의 수가 100개 미만일 경우 변수 당 표본의 수는 10개 이상 되어야 한다고 주장하였다. Kass & Tinsley(1979)는 표본 수가 300개 이상이고 각 측정항목 당 표본의 수가 5개에서 10개가 될 경우 표본의 수는 충분하다고 주장하였다. 이러한 복합적 기준에도 불구하고 한 가지 명심할 것은 모든 상황에 적용될 수 있는 절대적인 정성적 기준은 존재하지 않는다는 것이다(Alili 2010). 따라서 요인 분석을 위해서 살펴볼 것은 절대적 기준과 상대적 기준을 모두 만족하는 표본을 사용하는 것이 적절하다.

3.2. 행렬의 유의수준(significance of the matrix)

요인분석을 위한 상관관계 행렬을 위해 수집된 자료 표집 정확성(measures of sampling adequacy), 즉 자료 행렬이 의미가 있는 정보(meaningful information)를 포함하고 있는지 여부를 결정하기 위한 두 가지 방법이 존재한다(Budaev, 2010). 하나는 Bartlett 구형성 검정(sphericity test)이며 다른 하나는 KMO(Kaiser-Meyer-Olkin test of sampling adequacy)이다.

Bartlett(1950)의 chi-square 검정(Bartlett's test of sphericity)은 상관관계 행렬을 평가하기 위해 사용될 수 있다(Budaev, 2010). Bartlett의 chi-square는 변수의 수와 표본의 수간의 관계가 주어진 상황에서 변수들 간의 관계의 강도를 평가함으로써 변수의 요인화 가능성을 제시한다(Beavers et al., 2013; Tinsley & Tinsley, 1987). Bartlett 검정은 항등행렬(identity matrix)로 부터 편차가 존재하지 않는다는 귀무가설(null hypothesis)을 기각하는 오류의 가능성을 나타낸다(Tinsley & Tinsley, 1987). 즉, 귀무가설이 채택될 경우 요인화가 불가능하다는 것을 의미한다(Beavers et al., 2013).

본 검정을 통해 도출된 결과가 통계적으로 유의한 chi-square값을 갖는 경우가 요인 분석을 수행하기 위한 최소 기준이 된다(Tinsley & Tinsley, 1987). 예를 들어 본 값이 $p < 0.001$ 을 나타낼 경우 이는 수집된 데이터가 요인분석을 수행하기에 적합함을 의미한다(Treblmaier and Filzmoser, 2010). 한 가지 주의할 점은 본 검정 방법은 자료의 정규분포(normally distributed)를 요구하기 때문에 정규분포에서 벗어난 자료를 사용할 경우 결과의 편차가 매우 심할 수 있다(Budaev, 2010; Treblmaier and Filzmoser, 2010). 또한, 본 기법은 경우에 따라 상관관계행렬의 질이 낮은 경우에도 귀무가설(null hypothesis)을 기각하는 경우가 있다(Budaev, 2010). 즉, 만약 독립 가설(independence hypothesis)의 기각이 실패했다면, 해당 행렬을 가지고 더 이상의 분석을 수행하기는 어렵다(Budaev, 2010). 반면에 독립 가설의 기각이 Bartlett's test가 행렬이 계량심리학적으로 견고하다는 것을 증명해주는 명확한 지표가 된다는 것을 의미하는 것도 아니다(Budaev, 2010). 따라서 몇몇 학자들은 본 기법을 행렬의 품질이 하한값(lower bound)을 가지는 경우에 사용할 것을 제안하였다(Budaev, 2010).

본 기법의 한계점을 극복할 수 있는 대안으로 Steiger's test가 있다(Budaev, 2010). 본 기법은 정확성이 높다고 평가되며, 표본이 적은 경우에도 사용할 수 있다는 장점이 있기는 하나(Fouladi & Steiger, 1993), 대중적인 통계 프로그램이 아닌 특정 소프트웨어(cortest.mat function in R)만 실행이 가능하다는 한계점이 있어서 아직 널리 사용되고 있지는 못하다(Budaev, 2010).

<Table 1> An Interpretation of KMO Values

KMO Value	Degree of Common Variance
0.90~1.00	Marvelous
0.80~0.89	Meritorious
0.70~0.79	Middling
0.60~0.69	Mediocre
0.50~0.59	Miserable
0.00~0.49	Don't Factor

Source: Beavers et al.(2013)

KMO(Kaiser-Meyer-Olkin test of sampling adequacy)는 측정 항목 내에 공분산을 측정하기 위한 도구이다(Beavers et al., 2013). 즉, 상관관계 행렬의 요인화 가능성을 나타낸다. KMO는 관측 상관관계(observed correlation)와 원천 변수(original variables)간의 편상관관계(partial correlations)를 비교한다(Budaev, 2010). 본 지표에 대한 기준은 다음과 같다. 본 기법의 장점은 수집된 자료의 정규분포를 가정하고 있지 않다는 점이다.

4. 요인분석 수행 절차

연구자는 요인분석을 수행하는데 있어서 다음의 4가지 절차에 대한 의사결정을 수행해야 한다. 의사결정이 필요한 4가지 사항은 공통성 추정치(communality estimate), 요인 추출 기법(method of factor extraction), 요인 회전을 통해 추출된 요인의 수(how many factors to rotate), 회전 절차(rotation procedure) 등이다(Tinsley & Tinsley, 1987). 주지해야 할 점은, 초기 요인 추출 방법에 대한 결정, 유지 요인 수에 대한 결정, 회전 유형에 대한 선택 등은 전적으로 연구의 이론적 목적을 따라야 한다는 것이다.

4.1. 공통성 추정치

하나의 측정항목의 총 분산은 공통 분산, 고유 분산, 오차의 합으로 구성된다(Tinsley & Tinsley, 1987). 공통 분산은 최소 하나 이상의 다른 변수와 공유하는 분산을 말한다. 고유 분산은 신뢰할 수 있는 측정 분산으로 단일 변수에만 해당된다. 오차 분산은 신뢰할 수 없는 측정분산으로 단일 변수에만 해당된다(Tinsley & Tinsley, 1987).

요인 분석은 공통분산에 포함된 잠재적 차원을 식별하는 것이 목적이기 때문에, 총 분산에서 공통분산의 비율(공통성, communality)이 어느 정도인지 파악하는 것이 중요하다. 공통성이란 공통 요인에 의해 설명되는 관측변수의 분산을 나타낸다(Alii 2010). 높은 수준의 공통성은 잠재변수에 의해 강하게 영향을 받고 있다고 볼 수 있다(Alii 2010). 공통성과 고유요인 계수는 다음과 같이 계산된다(Alii 2010).

$D_1^2 = 1 - communality = \%$ 고유 요인(unique factor)에 의해 설명되는 분산

$d_1 = \sqrt{D_1^2} = \sqrt{1 - communality} =$ 고유요인 계수(weight)(모수 추정치, parameter estimate)

일반적 기준에 따르면 본 값이 0.5이상 되어야 한다. 이는 해당 항목의 설명력이 50%이상 된다는 것을 의미한다.

4.2. 요인추출 기법(methods of factor extraction)

주성분분석과 공통요인분석 모두 통계적으로 전체 변수를 작은 수의 성분 혹은 요인으로 줄여준다는 공통점이 있다(Beavers et al., 2013). 하지만 변수에 대한 정확한 해석 및 이해는 선형결합을 추출하는 방법에 따라 달라진다(Beavers et al., 2013).

성분분석에서는 초기 추출에서 총 분산을 포함한다(Beavers et al., 2013). 주성분분석은 성분분석을 위해 가장 널리 사용되는 추출 기법이자, 전체 측정 항목을 대표 성분을 나타내는 적은 수의 항목으로 축소해준다는 가장 적합한 방법이다(Beavers et al., 2013). 반면에 공통요인분석은 요인추출 시 오직 공통분산(공유분산)만을 포함하며, 가장 많이 사용되는 추출기법은 주축요인분석(Principal Axis Factoring, PAF)과 최대우도법(Maximum Likelihood Estimation, ML) 등이 있다(Beavers et al., 2013). 두 추출기법은 연구자에게 변수집합에서 잠재 구조를 이해하는데 도움을 제공함으로써 가장 최적의 결과를 제공한다고 알려져 있다(Costello & Osborne, 2005; Treblmaier & Filzmoser, 2010). 물론 각 기법에 대한 선택은 수집된 데이터가 정규분포를 따르는지 여부에 따라 달라진다. 몇몇 학자들은 PAF, ML 기법 외에 다른 추출 기법이 더 적절하다고 주장하는 하였으나(e.g., Alpha Factoring) 이에 대한 증거는 여전히 매우 부족하다(Costello & Osborne, 2005).

PAF(Principal Axis Factoring, Principal Factor Methods, 주축요인추출법): 본 기법은 분산에 대한 가정이 존재하지 않기 때문에 데이터가 정규분포를 형성하고 있지 않은 경우에도 사용이 가능하다(Beavers et al., 2013). 즉, 수집된 자료의 다변량 정규성이 위배될 경우에 사용할 수 있는 기법이다(Costello & Osborne, 2005).

PAF나 PCA 모두 자료의 정규분포가 선행조건이 아니며 타원형 대칭성(elliptical symmetry)만 요구된다(Treblmaier and Filzmoser, 2010). 단, 주의할 것은 자료의 정규분포가 보장되지 않은 상황에서 사용되기 때문에 비정규분포 자료와 이상치의 발생으로 인해 결과가 부정적 영향을 받을 수 있다는 점은 인지하고

있어야 한다(Treblmaier and Filzmoser, 2010).

ML(Maximum Likelihood Estimation, 최대우도법): 본 기법은 다변량 정규성을 가정하고 있다(Beavers et al., 2013). 즉 수집된 자료가 정규분포 형태를 띠고 있을 경우 본 기법을 선택하는 것이 최적을 결과를 얻기 위한 방법이다(Costello & Osborne, 2005; Treblmaier and Filzmoser, 2010). 본 기법은 요인 간의 상관관계와 신뢰구간뿐만 아니라 각각의 항목(요인 적재값)의 통계적 유의수준과 다양한 유형의 요인 구조(모형)의 적합도(fit)도 산출한다(Beavers et al., 2013; Costello & Osborne, 2005).

이외에도 ULS(Unweighted Least Squares), GLS(Generalized Least Squares), AF(Alpha Factoring), IF(Image Factoring) 등 다양한 기법이 존재하나, 이 추출기법들의 각각의 장점과 약점에 대한 정보는 매우 미약한 수준이다(Costello & Osborne, 2005).

4.3. (유지)요인의 수(the number of factors to retain)

요인 추출의 목적은 통계적으로나 이론적으로나 관련이 없는 요인을 제거하고, 수집된 데이터를 가장 적절히 나타낼 수 있는 충분한 요인을 선택하는 것이다(Beavers et al., 2013). 따라서 요인분석을 통해 도출된 결과는 요인 회전 전에 얼마나 많은 요인을 유지해야 하는지에 달려있다고 해도 과언이 아니다(Ford et al. 1986). 얼마나 많은 요인 혹은 성분을 유지할 것인지를 결정하는 것이 중요한 이유는 다음과 같다. 첫째, 유지 요인의 결정은 추출 기법의 선택, 요인 회전 기법 선택 등 다른 기법보다 탐색적 요인 분석 결과에 더 많은 영향을 미친다(Ledesma & Valero-Mora, 2007). 둘째, 탐색적 요인 분석은 변수의 집단 내에 존재하는 상관관계의 표현 및 축소 간에 균형을 맞추는 것이 중요하기 때문에 중요한 요인과 중요하지 않은 요인을 명확히 구분하는 것이 매우 중요하다(Ledesma & Valero-Mora, 2007). 마지막으로 요인의 수를 결정하는 과정에서 오류가 발생할 경우 해의 도출과 탐색적 요인 분석의 결과의 해석에 상당한 영향을 미칠 수 있다(Ledesma & Valero-Mora, 2007).

기존의 많은 연구에서는 분석에 부정적 영향을 미칠 수 있어서 제거되어야 할 요인임에도 불구하고 해당 요인을 유지하여 최대한 많은 요인을 유지할 것을 주장하기도 하였다(Beavers et al., 2013). 하지만, 인위적으로 너무 많은 요인을 유지할 경우 오류가 있는 해를 도출할 가능성이 높다(Beavers et al., 2013). 반면에, 하나 혹은 두 개의 요인을 가진 해를 사용하는 것은 요인 구조의 정확한 대표성을 제공하지 못하기 때문에 주의해야 한다(Beavers et al., 2013). 따라서 적절한 기법을 통해 합리적 수의 요인을 유지하는 것이 중요하다. 합리적인 수를 위해 사용될 수 있는 기법들은 다음과 같다.

통계적 검정(statistical tests): 통계적 검정은 탐색적 요인분석에서 자주 수행되지는 않는데, 그 이유는 연구자들이 주로 모형의 적합도(model fit)를 위한 통계적 추정치를 산출하지 못하는 주축(principal axes or least squares(최소 자승)) 추정방법을 사용하기 때문이다(Floyd & Widaman, 1995). 하지만 모형 적합도를 산출할 수 있는 다른 통계적 추정법들이 존재하는데, 최대 우도법(maximum likelihood estimation, MLE), 일반최소제곱법(general least squares, GLS), 점근분포무관 추정법(asymptotically distribution free estimation, ADF) 등이 사용될 수 있다(Floyd & Widaman, 1995). 이러한 기법들은 카이제곱(χ^2) 검정을 통하여 특정한 수의 요인이 추출된 이후 관측 변수간의 유의한 잔차 공변값(significant residual covariation)을 표현한다(Floyd & Widaman, 1995). 만약 카이제곱 검정치가 유의하다면 하나 이상의 요인을 가지고 있는 모형에 대한 기각을 결정하기 위한 근거로 사용된다

(Floyd & Widaman, 1995). 통계적 검정법의 가장 큰 문제는 표본의 수에 민감하다는 것인데, 표본의 수가 클 경우(예, $N > 500$) 지나치게 많은 요인을 유지하도록 유도하거나, 상대적으로 낮은 수준의 잔차 공변량으로 인해 적절한 요인을 기각하는 경우가 발생할 수 있으며, 표본의 수가 작을 경우(예, $N < 100$) 너무 적은 요인을 유지하도록 유도할 수 있다는 문제점이 있다(Floyd & Widaman, 1995).

Eigenvalue(고유치) 기준(Latent Root Criterion, Kaiser Criterion, K1): 요인 유지 결정에서 가장 널리 사용되는 기준은 수학적/심리 측정학적 기준이다(Floyd & Widaman, 1995). 주성분 분석에서 가장 많이 사용되는 기준은 고유치(eigenvalue) > 1 혹은 Kaiser-Cuttman 기준(criterion)이다(Floyd & Widaman, 1995). 본 기법은 고유값(eigenvalue)이 1 이상 되는 요인들을 선택하는 방법이다(Gefen & Straub, 2005). 고유치란 성분에 의해 설명되는 분산의 양을 의미하며, 모든 고유치의 합은 주성분 분석에서 관측변수의 수와 일치한다(Floyd & Widaman, 1995). 따라서 고유치 < 1이라는 것은 하나의 성분이 하나의 변수(single variance)에 의해 설명되는 분산보다 더 적은 부분을 설명하고 있다는 것을 의미한다(Floyd & Widaman, 1995). 주성분 분석의 목적은 변수들의 집합을 축소하는 것이기 때문에, 고유치 < 1인 성분은 본 목적에 부합되지 않는다(Floyd & Widaman, 1995). 따라서 고유치 > 1인 성분만 유지된다(Floyd & Widaman, 1995). 본 방법은 표본 대 변의 비율이 높을 때(10:1) 가장 적절하다고 본다(Ford et al., 1986). 하지만 본 기준은 다양한 상황에서 항상 적합한 것은 아니다(Floyd & Widaman, 1995). 수학적 기준의 사용 목적은 자료 축소에 기원을 두고 있다(Floyd & Widaman, 1995). 하지만 자료 축소라는 목적은 Kaiser-Guttman 규칙에 적합하지 않다고 보기도 하는데, 그 이유는 유지해야 할 차원의 수를 지나치게 과대평가하기 때문이다(Floyd & Widaman, 1995). 반면에, 다른 학자는 본 기준이 요인의 수를 과소평가할 수 있다고 주장하였다(Floyd & Widaman, 1995).

Cattell's Scree plot: 위의 두 가지 방법 외에 연구자들의 경험적 기준에 의해 사용되는 요인을 유지하기 위한 방법이 존재한다. 이러한 방법 중에 가장 빈번히 사용되는 방법은 스크리 검정(scree test)이다(Floyd & Widaman, 1995). 대중적인 통계적 소프트웨어에서 기본설정으로 제공되는 기법은 고유치 1 이상 값이다. 하지만 본 기준은 요인에 대한 과다추출(overextraction) 혹은 과소추출(underextraction)의 위험이 존재한다. 선행연구에서 수행한 Monte Carlo 분석 결과를 살펴보면 36%의 표본에서 과다추출(너무 많은 요인이 추출됨) 현상이 발견되었다(Costello & Osborne, 2005). 따라서 많은 학자들은 Scree test를 최선의 선택으로 인정하는 경우가 많다(Costello & Osborne, 2005). 스크리 차트는 요인과 대응 고유치를 그래프로 나타내는 방법이다. X축은 요인(혹은 성분)을 나타내며, Y축은 고유치를 나타낸다(Beavers et al., 2013). 본 차트에서 첫 번째 요인(성분)은 가장 많은 분산을 설명하기 때문에 가장 큰 고유치를 가진다(Beavers et al., 2013). 다음으로 각각의 고유치는 서서히 감소하는 형태를 띠며, 차트의 그래프는 발꿈치(elbow, natural bend, break point) 형태를 보인다(Beavers et al., 2013). 이때 곡선이 급격한 지점(break)의 위에 있는 요인들이 추출된다(Costello & Osborne, 2005). 기존 학자들에 따르면 본 기준이 고유치 1의 기준보다 더 정확한 결과를 산출할 수 있다고 평가하였다(Ledesma & Valero-Mora, 2007). 본 기준의 가장 큰 문제점은 경계값을 선정하는 기준이 상당히 주관적이라는 것이다(Beavers et al., 2013; Ledesma & Valero-Mora, 2007). 즉, 유지 요인과 비유지 요인 간의 명확한 경계를 나누어 주는 객관적 기준이 존재하지 않아 학자마다 다른 기준을 가지고

보유 요인 수를 결정할 수 있다(Ledesma & Valero-Mora, 2007). 이와 같은 문제점으로 인해 종종 요인의 과다추출의 문제를 발생시키기도 한다(Beavers et al., 2013). 따라서 본 방법보다 더 나은 유지 요인 결정방법이 존재한다면 이 방법을 사용하지 말아야 한다(Ledesma & Valero-Mora, 2007).

Velicer's MAP test/criteria(Minimum Average Partial, 편최소평균): Velicer(1976)의해 소개된 본 기법은 주성분분석을 응용한 방법으로 편상관행렬을 사용한다(Ledesma & Valero-Mora, 2007). 본 방법은 얼마나 많은 성분을 추출할 것인지를 결정할 때 탐색적 요인 분석에서 사용되는 공통 요인의 개념을 활용한다(Ledesma & Valero-Mora, 2007). 선행 연구에 따르면 본 기법은 고유치 1이나 스크리 차트 보다 우수한 기법이라고 소개하고 있다(Ledesma & Valero-Mora, 2007). 하지만 본 기법을 사용하기 위해서는 각각의 성분이 높은 요인 적재값을 갖는 최소한 두 개 이상의 측정항목을 가지고 있어야 한다(Ledesma & Valero-Mora, 2007). 또한 너무 낮은 요인 적재값을 갖는 경우 해당 성분은 유지 성분에서 제외된다(Ledesma & Valero-Mora, 2007). 이와 관련된 문제점은 낮은 요인 적재값을 갖는 경우 그리고 성분 당 변수의 수가 적은 경우에는 중요한 성분이라 하더라도 과소추정될 가능성이 높다는 점이다(Ledesma & Valero-Mora, 2007).

Variance Extracted(or Percentage of Variance, 추출 분산): 본 기법은 특정한 양의 추출 분산 비율을 가지고 요인을 추출하는 방법이다(Beavers et al., 2013). 가장 많이 사용되는 기준은 75~90%의 분산 설명력이다(Beavers et al., 2013). 물론 몇몇 연구에서는 50%도 적절하다고 평가하기도 한다(Beavers et al., 2013). 주의해야 할 점은 성분분석의 경우 설명되는 분산의 양이 공통분산과 비교하였을 때 더욱 많다는 것을 기억해야 한다는 것이다(Beavers et al., 2013). 본 기법은 상당히 포괄적인 기준을 사용하기에 단독으로 사용하기 보다는 다른 기준과 함께 사용하는 것이 적절하다(Beavers et al., 2013).

Parallel Analysis(병렬분석): 요인의 수를 결정하는 다른 방법은 병렬 분석이 있다. Horn(1965)에 의해 소개된 병렬분석은 유지 요인 결정을 위해 가장 많이 추천되는 방법 중 하나이다(Ledesma & Valero-Mora, 2007). 또한 본 기법은 탐색적 요인 분석뿐만 아니라 주성분분석에서도 요인이나 성분 결정 시 우수한 기법으로 인정받고 있다(Ledesma & Valero-Mora, 2007). 본 분석은 타당한 요인 구조를 가진 실제 데이터(real data)로 부터 도출된 고유값이 기저요인(underlying factor)이 없는 무작위 데이터(random data)로 부터 추출된 고유값보다 상당히 커야 한다는 것이다(Ford et al., 1986). 따라서 본 분석을 통해 결정되는 요인의 수는 무작위 데이터로부터 도출된 기대되는 고유값보다 실제 고유값이 몇 개나 큰 지로 결정된다(Ford et al., 1986). 본 기법은 유지 요인을 결정할 때 무작위 변수를 생성하여 사용한다(Ledesma & Valero-Mora, 2007). 병렬분석은 상관관계 행렬로부터 추출된 고유치와 비상관된 정규 변수로부터 확보된 고유치를 비교하여 유지 요인을 결정한다(Ledesma & Valero-Mora, 2007). 만약 상관관계 행렬로부터 추출된 고유치가 무작위 비상관 자료로부터 확보된 고유치의 평균보다 높다면 유지해야 하는 요인으로 간주된다(Ledesma & Valero-Mora, 2007). 전연한 상관관계 행렬에서 특징적인 것은 상관관계 행렬의 대각선 값에 다승상관자승치(squared multiple correlation)가 표현된다는 것이다(Ledesma & Valero-Mora, 2007). 하지만 본 기법은 계산하는 방법이 상당히 복잡하여 많은 통계 프로그램에서 제공되지 않기 때문에 연구자들에 의해 많이 사용되지는 않는다(Ledesma & Valero-Mora, 2007). 대중적으로 빈번하게 사용되는 SPSS나 SAS 사용자를 위해 매크로(macro)로 제공되긴 하지만 완전히 독립적인 프로그램으로 제공되는 경우는 드물다.

하지만 최근에 GUI(Graphical User Interface) 기반의 ViSta(Visual Statistics System, 본 프로그램은 <http://forrest.psych.unc.edu/>에서 무료로 다운 받을 수 있음)라는 무료 프로그램이 소개되었는데, 본 프로그램에 플러그 인으로 제공되는 ViSta-PARAN(ViSta를 위한 다양한 플러그 인 프로그램은 <http://www.mdp.edu.ar/psicologia/vista/vista.htm>에서 무료로 다운 받을 수 있음)은 본 분석을 독립적으로 수행하기 적절하여 기존 통계 프로그램이 가지고 있는 문제점을 해결하였다(Ledesma & Valero-Mora, 2007).

위의 여러 가지 요인 유지 기법 중에서 MAP이나 PA는 다른 기법보다 더 정확하다고 평가받고 있기는 하지만 연구자들 사이에서 널리 사용되는 통계적 소프트웨어에서 제공되는 경우가 많지 않아 상대적으로 사용되는 빈도는 높지 않다(Costello & Osborne, 2005).

이미 전언하였듯이 각각의 기준은 장점뿐만 아니라 고유한 단점도 가지고 있다. 따라서 하나의 기준만을 사용하기 보다는 여러 기준을 복합적으로 사용하는 것이 효과적이다.

4.4. 요인 회전 기법(rotation of factors)

탐색적 요인은 두 단계로 구성되는데, 첫 번째 단계는 요인을 추출하는 것이며, 다음 단계는 측정 항목의 기반 요인(underlying factors)들에 대한 더 나은 구조를 제공하기 위해 회전을 하는 것이다(Gefen & Straub, 2005). 회전의 목적은 자료 구조를 단순화(simplify) 그리고 명확하게(clarify) 표현하는 것이다(Costello & Osborne, 2005).

요인 회전을 수행하게 되면 해석이 용이한 해(interpretable solution)를 도출함과 동시에(Gefen & Straub, 2005) 표본을 활용한 개념의 일반화가 가능하다(Tinsley & Tinsley, 1987). 따라서 요인 회전(factor rotation)은 심리학적 의미(meaningfulness)와 신뢰성(reliability), 요인의 재현성(reproducibility)을 향상시켜준다(Ford et al., 1986). 회전을 위한 다양한 축이 존재하기 때문에, 회전 문제(rotation problem)를 해결하기 위한 유일한 해(unique solution)는 존재하지 않는다(Ford et al., 1986). 일반적으로 많이 사용되는 요인 회전 절차는 직각회전(orthogonal)과 사각회전(oblique)이 있다.

4.4.1. 직각회전(orthogonal rotation)

직각 요인 회전 절차는 서로 독립적이지(independent) 비상관(uncorrelated) 요인을 도출한다(Tinsley & Tinsley, 1987). 즉, 직각회전은 통계적으로 상관이 없는(uncorrelated) 요인을 도출한다(Ford et al., 1986; Costello & Osborne, 2005). 따라서 직각회전은 분석의 목적이 요인 점수(factor scores)를 산출하거나 이론적 가설들이 비상관 차원(uncorrelated dimensions)이라고 판단될 때 사용된다(Beavers et al., 2013).

직각회전에 대한 지지자들은 직각회전 기법의 단순성(simplicity), 개념적 명확성(conceptual clarity), 그리고 이후 분석을 위한 유연성(amenability)이 장점이다(Ford et al., 1986). 본 기법의 단점은 직각회전은 많은 데이터 속에 존재하는 관계를 반영하고 있지 않고 있기 때문에 이론적으로 항상 적합한 회전법은 아니라는 점이다(Beavers et al., 2013).

직각회전 기법 중 가장 우수하며 널리 사용되는 기법은 Varimax 기법이다(Beavers et al., 2013; Tinsley & Tinsley, 1987).

Varimax: 요인 회전 기법 중 여러 연구에서 가장 빈번히 사용되는 기법은 Varimax rotation(베리맥스 회전)이다. 본 기법은 직각/직교 회전(orthogonal rotation) 기법 중 하나로 측정 항목의 적재 값이 다른 요인들에 높게 적재되는 것을 최소화해준다(Gefen &

Straub, 2005). 또한 본 방법은 직각의 단순 구조(orthogonal simple structure)가 적절한 경우에 적합한 방법이다(Ford et al., 1986).

이외에도 Quartimax, Equimax 등 다양한 기법들이 존재하나 다른 직각 회전 기법들도 Varimax 회전과 유사한 결과를 도출한다(Beavers et al., 2013; Costello & Osborne, 2005; Ford et al., 1986).

4.4.2. 사각회전(oblique rotation)

기존연구에서는 직각회전을 사용할 것을 추천한 이유는 해석이 매우 용이한 결과를 도출할 수 있다는 장점 때문인데, 이는 잘못된 주장이다(Costello & Osborne, 2005). 사회과학 연구에서 기본 가정으로 요인 간에는 어느 정도의 상관관계가 있다고 본다(Costello & Osborne, 2005). 따라서 이 가정을 무시하고 요인 간의 상관관계를 고려하지 않는다면 연구자가 분석 상에서 의미 있는 정보를 놓칠 가능성이 높아진다(Costello & Osborne, 2005). 만약 실제 요인 간의 상관관계가 없는 경우 직각회전과 사각회전은 거의 동일한 결과를 산출한다(Costello & Osborne, 2005).

요인 회전을 위한 또 다른 방법은 사각회전(Oblique or Nonorthogonal rotation)이다. 사각회전은 서로 상관관계(correlated)를 형성하고 있는 요인들을 생성한다(Ford et al., 1986). 즉, 사각 요인 회전 절차는 회전이 완료된 이후 상관관계가 있는 요인을 도출한다(Tinsley & Tinsley, 1987). 뿐만 아니라 사각회전은 요인 간에 유의한 상관관계가 존재하지 않는 경우에도 사용할 수 있다(Beavers et al., 2013).

사각회전을 선호하는 학자들은 본 기법이 직각회전보다 더 단순한 구조를 도출함과 동시에 심리적 현상을 표현하는데 더 적절하다고 주장한다(Tinsley & Tinsley, 1987). Conway & Huffcutt(2003)에 따르면 사각회전이 직각회전보다 우수함에도 불구하고 많은 연구자들이(이들이 조사에 의하면 80%이상의 연구자들) 사각회전보다는 직각회전을 사용하고 있다고 주장하였다.

사각회전 절차가 직각회전 절차에 비해 이론적 기반에서 우위에 있다고 하더라도 사각회전을 수행하기 위해서는 복잡한 탐색적 가설(complicated exploratory hypotheses)뿐만 아니라 각 요인 속에 내재되어 있는 잠재적 차원과 요인들 간의 상관관계에 내재되어 있는 잠재적 차원에 대한 설명도 필요하다(Tinsley & Tinsley, 1987). 즉, 사각회전은 직각회전에 패턴(pattern matrix) 및 구조 행렬(structure matrix)을 생성함으로 통계적 복잡성을 증가시켰다(Ford et al., 1986). 이러한 복잡성은 연구자로 하여금 정교함과 해석상에 주의를 요구한다(Ford et al., 1986). 물론 이러한 복잡성은 요인 간의 상관관계의 형태를 제공한다는 장점이 있다(Ford et al., 1986). 현실에서 어떠한 개념간의 상관관계가 없는 경우는 매우 드물기 때문에 사각회전이 변수들 간의 복잡성을 더욱더 명확히 표현해준다고 볼 수 있다(Ford et al., 1986). 하지만, 사각회전은 직교성(orthogonality)의 상실로 인해 다중공선성(multicollinearity)이 발생할 가능성을 증가시킨다는 문제점도 있다(Gefen & Straub, 2005).

전언하였듯이 직각회전이 하나의 행렬만을 제시하는 반면 사각회전은 두 개의 행렬을 제시하기에 다소 복잡해 보일 수 있으나 본질적으로 해석하는 방법과 절차는 동일하다(Costello & Osborne, 2005).

사각회전에 적절한 회전 기법은 매우 다양하다. 사각회전을 위해 사용되는 기법들은 Direct Oblimin Method, Promax, Quartimin, Orthoblique, Procrustes, Covarimin, McCammon, Harris-Kaiser Criteria 등이 알려져 있다(Beavers et al., 2013; Costello &

Osborne, 2005; Tinsley & Tinsley, 1987; Treiblmaier & Filzmoser, 2010). 학자마다 선호하는 기법들이 다른데, Tinsley & Tinsley(1987)는 Promax, Harris-Kaiser 기준이 적합하다고 주장하였다. 하지만 위에서 언급한 사각회전을 위한 여러 가지 회전 기법 중에 가장 선호되거나 우수하다고 인정받는 회전기법은 존재하지 않는다(Beavers et al., 2013; Costello & Osborne, 2005). 또한 대부분의 기법이 유사한 결과를 산출한다(Fabrigar et al., 1999). 따라서 회전 기법의 선택은 연구자 혹은 연구에서 사용되는 소프트웨어에 따라 달라질 수 있다.

5. 결론 및 함의

기존의 많은 연구를 살펴보면 실증분석을 위한 기본 단계로 요인분석을 수행한 경우가 많다. 여기서 요인 분석은 주로 탐색적 요인 분석이다. 그만큼 대중적 분석기법인 요인분석은 특정 학문 분야가 아닌 다양한 학문분야에서 사용되어 왔다. 하지만 여러 가지 이유로 인해 요인분석이 성분분석으로 대체되는 경우가 빈번하게 관찰된다. 하지만 명확히 주성분분석과 요인분석은 다른 가정과 다른 접근기법이 요구된다. 그럼에도 불구하고 이 두 분석기법이 혼재되어 사용되는 경우가 많다.

본 연구의 목적은 이러한 상황에서 명확한 요인분석을 수행하기 위한 가이드라인과 절차를 선행연구의 검토를 통해 제시하는 것이다. 이를 통해 요인분석의 오류를 줄이고자 하였다. 본 연구에서 검토한 문헌을 중심으로 요인분석을 위한 판단 기준 및 절차를 제시하면 다음과 같다.

우선 많은 수의 자료를 적은 수의 자료로 축소하는 것이 목적이라면 주성분분석법을 선택하는 것이 적절하다. 반면에 데이터 속에 숨겨진 잠재적 요인 구조를 확인하는 것이 목적이라면 탐색적 요인 분석을 선택해야 한다. 즉 연구자가 이론을 기반으로 요인 구조를 찾고자 한다면 탐색적 요인분석을 수행해야 한다. 주성분 분석을 요인분석을 대신하여 사용하려 하는 경우에는 분석에 사용되는 변수의 수가 많고, 각각의 관측변수의 공통성이 높은 경우(0.7이상)에만 가능하다.

요인분석을 선택하게 되면 다음 단계는 수집된 자료가 요인분석을 수행하기에 적합한지 여부를 판단하기 위해 데이터의 품질을 확인하는 단계이다. 물론 이 단계는 주성분분석에도 적용할 수 있다. 중요한 것은 이 단계의 조건이 충족되지 않으면 다음 단계를 수행하기 어렵다는 것이다. 본 단계는 크게 2단계, 세부적으로 4가지 절차가 존재한다. 하나는 정량적 기준이며 다른 하나는 정성적 기준이다. 정량적 기준은 표본의 수가 요인분석을 수행하기에 적합한지를 평가하는 단계로 절대적 기준과 상대적 기준이 존재한다. 다음으로 정성적 기준은 데이터의 적합성을 평가하는 단계로 Bartlett 구형성 검정(sphericity test)과 KMO(Kaiser-Meyer-Olkin test of sampling adequacy)가 존재한다.

정량적 기준은 요인 분석을 수행하기에 필요한 최소 표본 기준을 충족하였는지를 평가하는 단계이다. 표본이 많으면 많을수록 분석에 더 적합하다는 것은 정설로 받아들여지고 있다. 하지만 연구를 위해 항상 최대의 표본을 수집하는 것이 용이한 것은 아니다. 따라서 많은 연구자들은 요인 분석을 위한 최소 표본에 대해 관심이 두어왔다. 주의할 점은 모든 학자들의 동의하는 완벽한 절대적 표본이나 상대적 표본은 존재하지 않는다는 점이다. 따라서 두 기준을 복합적으로 사용하는 것이 적합하다. 가장 많이 사용되는 절대적 기준은 모든 관측변수의 공통성이 0.5이상인 상황에서 표본 수 100개이다. 그리고 상대적 기준은 한 요인에 최소 측

정항목이 5개 이상인 상황에서 관측변수 대비 5:1 기준이다. 따라서 최소 표본이 100개이며, 관측변수 대비 표본의 수가 5배이면 요인분석을 위한 적합한 표본 수라 볼 수 있다.

다음으로 수집된 데이터의 적합성(sampling adequacy)을 평가하는 정량적 기준을 살펴보면, Bartlett 구형성 검정은 데이터의 정규분포가 가정된 상황에서 사용될 수 있다. 따라서 데이터의 정규성 검사 후 본 방법을 사용하거나 KMO를 부수적인 평가지표로 사용하는 것이 적절하다. 즉 각각의 지표는 데이터의 정규성이 확인된 경우라면 Bartlett 구형성 검정만 사용해도 무방하지만 그렇지 못할 경우 KMO를 평가한 후 Bartlett 구형성 검정을 부수적인 평가지표로 사용하는 것이 적합하다고 판단된다.

이와 같은 과정을 통해 데이터의 적합성이 확보된다면 다음 단계는 요인분석을 수행하는 절차가 남았다. 요인 분석 수행 절차는 다음의 프로세스를 통해 진행된다. 첫째, 요인추출기법의 선택, 둘째, 유지 요인 수 결정 기법의 선택, 셋째, 요인 회전 기법의 선택 등이다.

요인추출기법의 선택은 많은 학자들에게 가장 적절하다고 평가되는 PAF와 ML 중에 선택할 수 있으며, 각각의 기법은 데이터의 특성에 따라 선택 기준이 세부적으로 나뉜다. 만약 수집된 데이터의 정규성이 확보된 상태라면 ML기법을 사용하는 것이 적절하며, 그렇지 못하고 데이터의 정규성이 확보되지 못한 상태라면 PAF기법을 선택하는 것이 적절하다.

다음으로 유지 요인 수의 결정은 여러 가지 기법 중 PA가 가장 선호되는 방법이나 일반적으로 많이 사용되는 통계 분석 프로그램에서 이 기법을 제공하는 경우가 많지 않아 해당 기법 외에 다른 기법을 혼합적으로 사용하는 것이 적절할 것으로 판단된다. PA 외에 다른 기법들은 각자 고유의 문제점을 가지고 있기 때문에 하나의 기법만으로 유지요인을 결정하는 것은 결과에 대한 신뢰성을 낮출 수 있다. 따라서 복합적인 기법을 혼용하는 것이 적절하다. 가장 일반적으로 사용될 수 있으며, 다양한 통계 패키지에서 제공되는 기법들은 Eigenvalue 1초과, Scree Plot, 추출분산 평가 기법 등이며 이를 혼용하면 적절한 유지 요인을 결정하는데 도움이 될 것으로 판단된다.

마지막으로 요인회전 기법은 직각과 사각회전 중에 선택해야 한다. 직각은 모든 요인들이 서로 상관관계가 있다는 것을 가정하지 않고 있다. 하지만 많은 연구에서 특히 인과관계를 검증하는 연구에서는 상관관계 분석을 통해 각 요인 간에 일정수준의 상관관계가 존재한다는 것을 가정하고 이를 수치적으로 제시한다. 따라서 요인간의 상관관계를 무시하고 이를 기본 가정에서 고려하는 사각회전이 요인분석을 위해 적절한 회전기법이다. 사각회전 기법도 세부적으로 다양한 회전 기법들이 존재하나 일반적으로 지금까지 규명된 바에 따르면 대부분의 기법의 유사한 수준의 유용성이 인정받기 때문에 세부적 회전 기법은 연구자의 선호도에 따라 선택되어도 무방할 것으로 판단된다.

본 연구는 그동안 명확한 구분없이 사용되어 왔던 요인 분석 방법인 탐색적 요인분석과 주성분분석을 구분하기 위한 방법과 탐색적 요인 분석을 위한 가이드라인 및 절차를 제시하기 위해 수행되었다. 본 결과를 통해 향후 연구에서는 명확한 절차를 기반으로 더 나은 요인분석이 수행되기를 기대해 본다.

References

- Alii, L. K. (2010). Sample Size and Subject to Item Ratio in Principal Components Analysis and Exploratory Factor

- Analysis. *Journal of Biometrics & Biostatistics*, 1(2), 1-6.
- Arrindell, W. A., & van der Ende, J. (1985). An Empirical Test of the Utility of the Observations-to-Variables Ratio in Factor & Components Analysis. *Applied Psychological Measurement*, 9, 165-178.
- Asgari, O., & Hosseini, S. (2015). Exploring the Antecedents Affecting Attitude, Satisfaction, and Loyalty towards Korean Cosmetic Brands. *Journal of Distribution Science*, 16(6), 45-60.
- Bartlett, M. S. (1950). Tests of Significance in Factor Analysis. *British Journal of Psychology*, 3, 77-85.
- Beavers, A. S., Lounsbury, J. W., Richards, J. K., Huck, S. W., Skolits, G. J., & Esquivel, S. L. (2013). Practical Considerations for Using Exploratory Factor Analysis in Educational Research. *Practical Assessment, Research & Evaluation*, 18(6), 1-13.
- Bryant, F. B., & Yarnold, P. R. (1997). Principal-Components Analysis and Exploratory and Confirmatory Factor Analysis. In L. G. Grimm & P. R. Yarnold (ed.), *Reading and Understanding Multivariate Statistics* (pp. 99-136), Washington, DC.: American Psychological Association.
- Budaev, S. V. (2010). Using Principal Components & Factor Analysis in Animal Behaviour Research: Caveats and Guidelines. *Ethology: International Journal of Behavioural Biology*, 116(5), 472-480.
- Chen, Y., & Hwang, C. S. (2014). Consumer Values, Preference, and Purchase Intention for Luxury Fashion Brands: Post-teen Korean and Chinese Women. *Journal of Distribution Science*, 12(12), 107-118.
- Comrey, A. L., & Lee, H. B. (1992). *A First Course in Factor Analysis*. Hillsdale, NJ: Erlbaum.
- Conway, J. M., & Huffcutt, A. I. (2003). A Review & Evaluation of Exploratory Factor Analysis Practices in Organizational Research. *Organizational Research Methods*, 6(2), 147-168.
- Costello, A. B., & Osborne, J. W. (2005). Best Practices in Exploratory Factor Analysis: Four Recommendations for Getting the Most from Your Analysis. *Practical Assessment, Research & Evaluation*, 10(7), 1-9.
- Fabrigar, L. R., Wegender, D. T., MacCallum, R. C., & Strahan, E. J. (1999). Evaluating Use of Exploratory Factor Analysis in Psychological Research. *Psychological Methods*, 4(3), 272-299.
- Floyd, F. J., & Widaman, K. F. F. (1995). Factor Analysis in the Development & Refinement of Clinical Assessment Instruments. *Psychological Assessment*, 7(3), 286-299.
- Ford, J. K., MacCallum, R. C., & Tait, M. (1986). The Application of Exploratory Factor Analysis in Applied Psychology: A Critical Review & Analysis. *Personnel Psychology*, 39, 291-314.
- Fouladi, R. T., & Steiger, J. H. (1993). Tests of Multivariate Independence: A Critical Analysis of "Monte Carlo Study of Testing the Significance of Correlation Matrices": Comment. *Educational and Psychological Measurement*, 53, 927-932.
- Gefen, D., & Straub, D. A. (2005). Practical Guide to Factorial Validity Using PLS-Graph: Tutorial & Annotated Example. *Communications of the Association for Information Systems*, 16, 91-109.
- Gorsuch, R. L., (1983). *Factor Analysis*, 2nd eds. Hillsdale, NJ: Erlbaum.
- Guadagnoli, E., & Velicer, W. F. (1988). Relation of Sample Size to the Stability of Component Patterns. *Psychological Bulletin*, 103, 265-275.
- Horn, J. L. (1965). A Rationale and Test for the Number of Factors in Factor Analysis. *Psychometrika*, 30, 179-185.
- Hutcheson, G., & Sofroniou, N. (1999). *The Multivariate Social Scientist: Introductory Statistics Using Generalized Linear Models*. Thousand Oaks, CA: Sage Publications.
- Kass, R. A., & Tinsley, H. E. A. (1979). Factor Analysis. *Journal of Research*, 11, 120-138.
- Kim, J. L. (2015). A Study on the Effects of Factor of Service Quality, Service Guarantee and Service Value in General Super Market. *Journal of Distribution Science*, 13(1), 93-103.
- Lawley, D. N., & Maxwell, A. E. (1971). *Factor Analysis in a Statistical Method*. Butterworth, London.
- Ledesma, R. D., & Valero-Mora, P. (2007). Determining the Number of Factors to Retain in EFA: An Easy-to-Use Computer Program for Carrying Out Parallel Analysis. *Practical Assessment, Research & Evaluation*, 12(2), 1-11.
- MacCallum, R. C., & Widaman, K. F. (1999). Sample Size in Factor Analysis. *Psychological Methods*, 4(1), 84-99.
- Norušis, M., (2005). *Advanced Statistical Procedures Companion*. Prentice Hall, Upper Saddle River, NJ.
- Nunnally, J. C. (1978). *Psychometric Theory*, 2nd eds. McGraw Hill, New York.
- Osborne, J. W., & Fitzpatrick, D. C. (2012). Replication Analysis in Exploratory Factor Analysis: What It Is & Why It Makes Your Analysis Better. *Practical Assessment, Research & Evaluation*, 17(15), 1-8.
- Snook, S. C., Gorsuch, R. L. (1989). Component Analysis Versus Common Factor Analysis: A Monte Carlo Study. *Psychological Bulletin*, 106(1), 148-154.
- Streiner, D. L. (1994). Figuring Out Factors: The Use & Misuse of Factor Analysis. *Canadian Journal of Psychiatry*, 39, 135-140.
- Suhr, D. (2006). Exploratory or Confirmatory Factor Analysis. SAS Users Group International Conference, Cary: SAS Institute, Inc., 1-17.
- Tinsley, H. E. A., & Tinsley, D. J. (1987). Uses of Factor Analysis in Counseling Psychology Research. *Journal of Counseling Psychology*, 34(4), 414-424.
- Treiblmaier, H., & Filzmoser, P. (2010). Exploratory Factor Analysis Revisited: How Robust Methods Support the Detection of Hidden Multivariate Data Structures in IS Research. *Information & Management*, 47, 197-207.
- Velicer, W. F. (1976). Determining the Number of Components from the Matrix of Partial Correlations. *Psychometrika*, 41, 321-327.