

**ANKARA ÜNİVERSİTESİ
EĞİTİM BİLİMLERİ ENSTİTÜSÜ
ÖLÇME VE DEĞERLENDİRME ANABİLİM DALI**

**KAYIP VERİLERİN VARLIĞINDA
İKİ KATEGORİLİ PUANLANAN MADDELERDEN OLUŞAN
TESTLERİN PSİKOMETRİK ÖZELLİKLERİNİN İNCELENMESİ**

DOKTORA TEZİ

Ergül DEMİR

**Tez Danışmanı
Prof. Dr. Nizamettin KOÇ**

**ANKARA
Şubat, 2013**

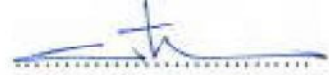
Eđitim Bilimleri Enstitüsü M¼d¼rl¼đ¼'ne,

Bu alıřma j¼rimiz tarafından lme ve Deęerlendirme Anabilim Dalında
Doktora Tezi olarak kabul edilmiřtir.

Başkan Prof. Dr. Nizamettin KO



¼ye Prof. Dr. Ezel TAVřANIL



¼ye Prof. Dr. řener B¼Y¼KZT¼RK



¼ye Do. Dr. Duygu ANIL



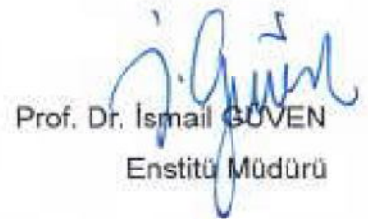
¼ye Yrd. Do. Dr. mer KUTLU



Onay

Yukarıdaki imzaların, adı geen đretim ¼yelerine ait olduęunu onaylarım.

13.10/2013



Prof. Dr. İsmail G¼VEN
Enstit¼ M¼d¼r¼

TEZ BİLDİRİMİ

Tez içindeki bütün bilgilerin etik davranış ve kurallar çerçevesinde elde edilerek sunulduğunu, ayrıca tez yazım kurallarına uygun olarak hazırlanan bu çalışmada bana ait olmayan her türlü ifade ve bilginin kaynağına eksiksiz atıf yapıldığını bildiririm.



Ergül DEMİR

ÖZET
KAYIP VERİLERİN VARLIĞINDA
İKİ KATEGORİLİ PUANLANAN MADDELERDEN OLUŞAN TESTLERİN
PSİKOMETRİK ÖZELLİKLERİNİN İNCELENMESİ

DEMİR, Ergül

Doktora Tezi, Eğitimde Ölçme ve Değerlendirme Anabilim Dalı

Tez Danışmanı: Prof. Dr. Nizamettin KOÇ

Şubat 2013, 130 sayfa

Bu çalışmada, kayıp verilerin varlığında iki kategorili puanlanan maddelerden oluşan testlerin psikometrik özellikleri, farklı kayıp veri yöntemleri kullanılarak incelenmektedir. Bu tür testlerde psikometrik özelliklerin kestiriminde uygun olan ve uygun olmayan kayıp veri yöntemlerinin belirlenebilmesi amaçlanmaktadır.

Temel araştırma türünde, ayrıntı ve ilişki saptamaya yönelik bir çalışmadır. Testin psikometrik özelliklerine yönelik incelemeler ve kayıp veri analizlerinde, SBS 2011 Matematik Testi A Kitapçığı verileri kullanılmıştır. Bu veriler 527517 yanıtlayıcıya yönelik yanıt örüntüleri ile cinsiyet, okul türü ve il değişkenlerini içermektedir.

Veri analizlerinde 12 farklı kayıp veri yöntemi kullanılmıştır. Kayıp veri yöntemleri olarak silmeye dayalı yöntemlerden ‘dizin silme yöntemi’, basit atama yöntemlerinden ‘0 atama’, ‘seri ortalamaları ataması’, ‘gözlem birimi ortalaması ataması’, ‘yakın noktalar ortalama ataması’, ‘yakın noktalar medyan ataması’, ‘doğrusal interpolasyon’ ve ‘dorusal eğilim noktası ataması’ yöntemleri, en çok olabilirlik yöntemlerinden ‘regresyon atama’, ‘beklenti-maksimizasyon algoritması’ ve ‘veri çoğaltma’ yöntemleri, çoklu veri atama yöntemlerinden ise ‘MarkovChain Monte Carlo’ yöntemi kullanılmıştır.

Elde edilen bulgular, ihmal edilebilir olmayan kayıp verilerin varlığında, iki kategorili puanlanan maddelerden oluşan testlerin psikometrik özelliklerine yönelik inceleme ve analizlerde, uygun bir kayıp veri yönteminin kullanılmasının gerekli olduğunu göstermektedir. Silmeye dayalı yöntemler ve 0 Atama yöntemi, bu tür veriler için uygun yöntemler değildir. Basit atama yöntemlerinin ise yanlı kestirimler üretme olasılığı yüksektir. En çok olabilirlik ve çoklu veri atama yöntemleri, bu tür verilerde kullanılması en uygun kayıp veri yöntemleri olarak değerlendirilmektedir.

Anahtar sözcükler: kayıp veri, yanıtlanmama, ihmal edilebilirlik, SBS

SUMMARY

PSYCHOMETRIC PROPERTIES OF TESTS COMPOSED OF DICHOTOMOUS ITEMS IN THE PRESENCE OF MISSING DATA

DEMİR, Ergül

Phd, Educational Measurement and Evaluation Department

Thesis Advisor: Prof. Dr. Nizamettin KOÇ

February 2013, 130 pages

In this study, psychometric properties of tests composed of dichotomous items are examined by using different missing value methods. It's aimed to determine available and unavailable missing value methods for estimating the psychometric properties in this kind of tests.

The type of this study is a basic research to determine the details and relationships. SBS 2011 Math Test A Booklet was taken into account for estimating psychometric properties of the tests and it is used for missing value analyses. Data set contains the response patterns of 527517 participants and some other information with gender and school type and province as independent variables.

Twelve different missing value methods were used to obtain the estimations. This methods are 'listwise deletion' and '0 imputation' and 'serial mean imputation' and 'unit's mean imputation' and 'mean of nearby points imputation' and 'median of nearby imputation' and 'linear interpolation' and 'linear trend at point' and 'regression imputation' and 'expectation-maximization algorithms' and 'data augmentation' and 'Markov Chain Monte Carlo'.

Results show that available missing value methods should be used for estimating the psychometric properties of tests composed of dichotomous items in the presence of missing values. Methods based on deletion and 0 imputation are not available for this kind of data. Simple imputation methods are likely to produce biased estimates. Methods based on maximum likelihood and multiple imputation are considered to be most available methods for this kind of data.

Keywords: missing value, nonresponse, ignorability, SBS

ÖNSÖZ

Bilimsel bilgi üretiminde araştırmacının temel kaygısı, gerçek değerlere en az hata ile ulaşabilmektir. Çoğu zaman araştırmacının elinde ancak gözlemler ve bu gözlemlerden elde edilen veriler bulunmaktadır. Gözlenen verilerin niteliği ve karakteristiği, ölçmeye karışan hata miktarının olasılığı ve dolayısıyla gerçek değerlere ne düzeyde yaklaşılabildiği hakkında bilgi sağlamaktadır. Bu nedenle bir araştırmada, öngörülen analizler öncesinde, kullanılan verilerin karakteristiğinin belirlenmesi ve olabildiğince ayrıntılı bir şekilde incelenmesi, araştırmacının önemli sorumlulukları arasında yer almaktadır.

Kayıp veriler, amaçlanan bir gözlemin yapılamadığını ya da olmadığını tanımlamaktadır. Birçok değişken dikkate alınarak gerçekleştirilen bir gözlemlerde, değişkenlerin birinde ya da bir kaçında ihmal edilebilir olmayan düzeyde kayıp veri bulunması, amaçlanan gözlemlerin yeterli düzeyde yapılamadığını göstermektedir. Bu durumda gerçek değerlere yönelik bir kestirim elde edilebilmesi de mümkün değildir.

Kayıp verilerin kestirimlerde yol açacağı olası yanlışlığın, belirli sınırlılıklar dâhilinde, istatistiksel olarak kontrol altına alınması mümkündür. Bu amaçla geliştirilmiş ve geliştirilmekte olan kayıp veri yöntemleri bulunmaktadır. Elbette ki bu yöntemlerin kullanımı, öncelikle gerekli farkındalık ve bilgi düzeyi ile ilişkilidir.

Kayıp verilerin göz ardı edilmesi ve veri analizlerine kayıp veri yokmuş gibi devam edilmesinin, yanlışlık içermesi olasılığı yüksek kestirimler elde edilmesine ve bu kestirimlere bağlı olarak ‘yanlış’ kararlara varılmasına yol açacağı, öngörülebilir ve beklenen bir durumdur. Diğer taraftan örneğin seçme ve yerleştirme amacıyla kullanılan testlerde, kayıp verilerin bilimsel bir gerekçe ve kanıt olmaksızın göz ardı edilmesinin ortaya çıkarabileceği olumsuz sonuçların, gerek testi alan bireyler gerekse ilgili diğer paydaşlar açısından telafi edilebilmesi ne derece mümkün olabilir?

Bu tür sorgulamalar ve akıl yürütmeler, kayıp veriler ve kayıp veri yöntemleri üzerine çalışma ihtiyacını beraberinde getirmektedir. Bu ihtiyaç, bu çalışmayı hazırlayanın yüksek lisans ve doktora eğitimleri sürecinin bir ürünü ve sonucudur.

TEŞEKKÜR

Bu çalışmaya sağladıkları katkıların yanı sıra doktora eğitimi sürecimin hemen her aşamasında yanımda olan, gösterdikleri insanî ve akademik duruş itibarıyla büyük saygı duyduğum değerli öğretmenlerim Sayın Prof. Dr. Nizamettin KOÇ'a, Sayın Prof. Dr. Ezel TAVŞANCIL'a, Sayın Prof. Dr. Şener BÜYÜKÖZTÜRK'e, Sayın Doç. Dr. Duygu ANIL'a ve Sayın Yrd. Doç. Dr. Ömer KUTLU'ya teşekkürlerimi sunarım.

Doktora eğitimlerim süresince, yapmış oldukları önemli çalışmaları izleme ve paylaşma imkânı bulduğum ve kendilerinden çok şey öğrendiğim değerli öğretmenlerim Sayın Prof. Dr. Nûkhet DEMİRTAŞLI'ya ve Sayın Doç. Dr. ÖmayÇOKLUK'a ayrıca teşekkürlerimi sunarım.

Sürecin en başında ölçme ve değerlendirme alanında çalışmam yönündeki teşviklerinden ve süreçteki önemli katkılarından dolayı değerli öğretmenim Sayın Prof. Dr. Selahattin GELBAL'a teşekkürlerimi sunmayı bir borç bilirim.

İÇİNDEKİLER

	Sayfa
ONAY.....	i
BİLDİRİM.....	ii
ÖZET.....	iii
SUMMARY.....	iv
ÖNSÖZ.....	v
TEŞEKKÜR.....	vi
İÇİNDEKİLER.....	vii
ÇİZELGELER DİZİNİ.....	viii
ŞEKİLLER DİZİNİ.....	ix
KISALTMALAR.....	x
BÖLÜM 1.....	1
GİRİŞ.....	1
Problem.....	1
Amaç.....	16
Önem.....	17
Sınırlılıklar.....	18
BÖLÜM 2.....	20
KAVRAMSAL ÇERÇEVE.....	20
Kayıp Veri Örüntüsü.....	20
Kayıp Veri Mekanizmaları.....	22
Kayıp Veri Yöntemleri.....	27
İlgili Araştırmalar.....	41
BÖLÜM 3.....	47
YÖNTEM.....	47
Araştırmanın Modeli.....	47
Evren.....	47
Verilerin Toplanması.....	47
Verilerin Analizi.....	49
BÖLÜM 4.....	56
BULGULAR VE YORUMLAR.....	56
Kayıp Veri Örüntüsü.....	56
Kayıp Veri Mekanizması.....	63
Yapı Geçerliği ve Tek Boyutluluk Özelliği.....	75
Madde Parametreleri.....	83
Test Parametreleri.....	90
BÖLÜM 5.....	98
SONUÇ VE ÖNERİLER.....	98
Sonuç.....	98
Öneriler.....	105
KAYNAKÇA.....	107
EKLER.....	111
EK A. SBS 2011 Matematik Testi A Kitapçığı.....	112
EK B. Veri Talep Dilekçesi.....	117
EK C. Özgeçmiş.....	118

ÇİZELGELER DİZİNİ

	Sayfa
Çizelge 3.1. Gözlenen ve Kayıp Verilerin Dağılımı.....	56
Çizelge 3.2. Doğru Yanıtlanan, Yanlış Yanıtlanan ve Kayıp Veri İçeren Madde Sayılarının Dağılımı.....	57
Çizelge 3.3. Doğru Yanıtlanan, Yanlış Yanıtlanan ve Kayıp Veri İçeren Maddelerin Dağılımlarına Yönelik Parametreler.....	59
Çizelge 3.4. Maddeler Düzeyinde Tepkilerin Dağılımı.....	60
Çizelge 3.5. Gözlem Birimleri Düzeyinde En Sık Görülen Kayıp Veri Örüntülerinin Dağılımı.....	62
Çizelge 3.6. Gözlem Birimleri Düzeyinde En Az Sıklıkta Görülen Kayıp Veri Örüntüleri.....	62
Çizelge 3.7. Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Cinsiyetlerine Göre Dağılımı.....	63
Çizelge 3.8. Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Okul Türlerine Göre Dağılımı.....	65
Çizelge 3.9. Gözlenen Verilerde Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Yaşadıkları Bölgelere Göre Dağılımı.....	68
Çizelge 3.10. Madde Düzeyinde Kayıp İçeren Yanıt Örüntülerinde Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Yaşadıkları Bölgelere Göre Dağılımı.....	69
Çizelge 3.11. Yanıt Örüntüleri Birim Düzeyinde Kayıp Verilerden Oluşan Yanıtlayıcıların Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Bölgelere Göre Dağılımı.....	70
Çizelge 3.12. Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Yaşadıkları Bölgelere Göre Dağılımı.....	71
Çizelge 3.13. Kayıp Veri Yöntemlerine Göre Yapı Geçerliği Parametreleri(Açımlayıcı Faktör Analizi - 1).....	76
Çizelge 3.14. Kayıp Veri Yöntemlerine Göre Yapı Geçerliği Parametreleri(Açımlayıcı Faktör Analizi - 2).....	77
Çizelge 3.15. Kayıp Veri Yöntemlerine Göre DFA Model Uyum İyiliği Parametreleri.....	81
Çizelge 3.16. Kayıp Veri Yöntemlerine Göre Madde Güçlük İndeksleri.....	84
Çizelge 3.17. Kayıp Veri Yöntemlerine Göre Madde Ayırıcılık İndeksleri.....	86
Çizelge 3.18. Kayıp Veri Yöntemlerine Göre Madde (Nokta Çift Serili) Korelasyon Katsayıları.....	87
Çizelge 3.19. Kayıp Veri Yöntemlerine Göre Madde (Çift Serili) Korelasyon Katsayıları.....	88
Çizelge 3.20. Kayıp Veri Yöntemlerine Göre Madde Güvenirlik Katsayıları..	89
Çizelge 3.21. Kayıp Veri Yöntemlerine Göre Toplam Puanların Dağılımı.....	90
Çizelge 3.22. Kayıp Veri Yöntemlerine Göre Test Parametreleri.....	94
Çizelge 3.23. Kayıp Veri Yöntemlerine Göre Madde Parametre Aralıkları ve Test Parametreleri.....	96

ŞEKİLLER DİZİNİ

	Sayfa
Şekil 1.1. Tipik Kayıp Veri Örüntüleri.....	21
Şekil 1.2. Monotonik Yapıda Kayıp Veri İçeren Değişkenler.....	34
Şekil 3.1. SBS 2011 Matematik Testi A Kitapçığı Verilerine Yönelik Ölçme Modeli (ConceptionalDiagram).....	81

KISALTMALAR

AFA	Açımlayıcı faktör analizi (explatoryfactoranalysis)
BA	Basit atama (simpleimputation)
BM	Beklenti-maksimizasyon (expectation-maximization) algoritması
BO	Gözlem birimi ortalaması ataması (unit'smeanimputation)
ÇKK	Çift serili korelasyon katsayısı
ÇVA	Çoklu veri atama (multipleimputation)
DEÇO	Doğrudan en çok olabilirlik (directmaximumlikelihood)
DFA	Doğrulamalı faktör analizi (confirmatoryfactoranalysis)
Dİ	Doğrusal interpolasyon (linearinterpolation)
DE	Doğrusal nokta eğilimi (lineartrend at point)
DS	Dizin silme (listwisedeletion)
EÇO	En çok olabilirlik (maximumlikelihood)
EHA	Erişilebilir hücre analizi (availablecaseanalysis)
HEÇO	Ham en çok olabilirlik (raw ML)
HS	Hücre silme (casewisedeletion)
KGY	Kayıp gösterge yöntemi (missingindicatormethod)
KDD	Kukla değişken düzeltmesi (Dummyvariableadjustment)
KOA	Koşullu ortalama atama (conditionalmeanimputation)
KS	Kısmi silme (pairwisedeletion)
MCMC	MarkovChain Monte Carlo
MTK	Madde tepki kuramı (itemresponsetheory)
NÇKK	Nokta çift serili korelasyon katsayısı
OA	Ortalama atama (meansubstitution)
RA	Regresyon atama (regressionestimation)
SA	Seçkisiz atama (randomimputation)
SBS	Seviye Belirleme Sınavı
SG	Seçkisiz gözlenen (observed at random)
SK	Seçkisiz kayıp (missing at random)
SO	Seri ortalaması ataması (serialmeanimputation)
SOK	Seçkisiz olmayan kayıp (missing not at random)
TA	Tanımlayıcı atama (deterministicimputation)
TEÇO	Tam bilgi en çok olabilirlik (fullinformation ML)
THA	Tam hücre analizi (completecaseanalysis)
TSK	Tam seçkisiz kayıp (missingcompletely at radom)
YNM	Yakın noktalar medyan ataması (median of nearbypoints)
YNO	Yakın noktalar ortalaması ataması (mean of nearbypoints)
VÇ	Veri çoğaltma (dataaugmentation)

BÖLÜM 1

GİRİŞ

Bu bölümde çalışmanın problemi, amacı, önemi ve sınırlılıkları yer almaktadır. Dikkate alınan araştırma soruları, amaç başlığı altında verilmektedir. Konuya ilişkin kavramsal tartışmalar ve ilgili araştırmalar, hacmi dikkate alınarak ayrı bir bölüm halinde düzenlenmiştir.

Problem

Tipik bir veri setinde yer alan bazı değişkenlere ilişkin bilgilerin ya da gözlemlerin bulunmaması araştırmalarda sıklıkla karşılaşılabilen bir durumdur. Örneğin gelir düzeyinin bir değişken olarak alındığı hane halkı araştırmalarında yanıtlayıcıların, gelirlerini beyan etmekten kaçınmaları ve bu yöndeki soruları yanıtlamamaları olasıdır. Endüstriyel deneysel çalışmalarda, deneysel çalışma süreci ile ilgisi olmaksızın, mekanik aksaklıklardan dolayı bazı kayıplar söz konusu olabilir. Benzer şekilde yüz yüze görüşmede görüşülen kişinin bazı soruları yanıtlamayı reddetmesi de karşılaşılabilecek bir durumdur. Örneğin siyasi tercihlerin sorgulandığı bir araştırmada bazı bireylerin, bir adayı diğerine göre tercih etme gerekçelerini açıklamaları mümkün olmayabilir. Kâğıt-kalem testleri, anket, ölçek gibi yanıtlayıcının bizzat yanıtlaması beklenen veri toplama araçlarının kullanımında, veri toplama aracının uzunluğu, madde sayısı, ifadelerin anlaşılabilirliği gibi özelliklere bağlı olarak, yanıtlayıcıların bazı soruları atlamaları ya da yanıtlamayı unutmaları söz konusu olabilmektedir. Bazı durumlarda yanıtlayıcı, yöneltilen soru hakkında bilgi sahibi değildir ve sadece “herhangi bir bilgim yok” demeyi tercih etmektedir. Araştırma konusu hakkında yeterli düzeyde bilgi ve deneyim sahibi olmayan bireylerin yanıtlayıcı olarak kabul edildiği araştırmalarda, bu bireylere yöneltilen pek çok soru ‘uygulanamaz’ olarak değerlendirilir. Örneğin evli olmayan bir bireye yöneltilen “kaç yıldır evlisiniz?” ya da “evliliğinizin sosyal hayatınıza ne düzeyde katkı sağladığını düşünüyorsunuz?” gibi sorular, uygulanabilir değildir. Bu durumda sorular

ya yanıtlanamaz ya da verilen yanıtlar anlamlı olmayacağı için dikkate alınmaz. Boylamsal çalışmalarda, ilk veri toplama aşamasında yer alan bazı yanıtlayıcılar, ölüm, yer değişikliği gibi nedenlerden dolayı ikinci veya diğer aşamalarda yer almayabilmektedir. Çoklu alan uygulamalarında, bazı uygulamalara yönelik verilerin kaybolması da karşılaşılabilecek durumlardan bir değeridir (Little ve Rubin, 1987; Allison, 2002).

Gelir durumunun beyan edilmemesi, mekanik aksaklıklardan kaynaklanan kayıplar ve benzeri durumlarda, gözlenemeyen verilerin kayıp veri olarak kabul edilmesi davranışı doğaldır. Bu tür kayıp verilerin oluşmaması için araştırmanın daha titiz bir şekilde yapılandırılması ya da endüstriyel araç-gereçlerin bakımının daha dikkatli bir şekilde yapılması gibi öneriler getirilebilir. Böylelikle gerçek değerlerin temelini oluşturan gözlemlerin eksiksiz bir şekilde elde edilme olasılığı artacaktır. Diğer taraftan siyasi tercihlerin sorgulandığı araştırma örneğinde ya da nüfus araştırmaları örneğinde olduğu gibi bazı araştırmalarda elde edilen sonuçlar, yanıtlanmama durumuna bağlı olarak maskelenmiş ve yanlı olabilmektedir. Bu tür durumlarda gözlenemeyen verilerin kayıp veri olarak kabul edilmesi davranışı, önceki örneklere göre daha az doğaldır (Little ve Rubin, 1987).

Rubin (1976), veri kayıplarının önemli ve yaygın bir sorun olduğunu gerçek bir örnek üzerinden hareketle açıklamaktadır. Pek çok sosyoekonomik değişkenin dikkate alındığı bir geniş ölçekli tarama araştırmasında, irtibat kurulan ailelerden 1967 ve 1970 yıllarında bilgiler toplanmıştır. Boylamsal bir çalışma olarak tasarlanan araştırmada, 1970 yılı verilerinin 1967 yılına göre oldukça büyük miktarda kayıp veri içerdiği belirlenmiştir. Ailelerin birçoğu yer değiştirmiş ve 1967'deki adreslerinde bulunamamıştır. Buna rağmen söz konusu veri kayıplarının ihmal edilebilir olup olmadığına bakılmaksızın, eldeki eksik veri seti ile analizlere devam edilmiştir. Rubin, söz konusu analizlerden elde edilen kestirimlerin manidar düzeyde yanlılık içerdiğini belirlemiştir. Rubin'in kuramsal ve teknik ispatları, bu tür durumlarda parametre kestirimi yapılabilmesi için veri kayıplarının 'ihmal edilebilir (ignorable)' düzeyde olması gerektiğini ortaya koymaktadır.

Temel bir yaklaşım olarak bir yanıtın eksikliği, ölçülen değişkene yönelik örnek uzayda ek bir nokta ile temsil edilir. Bu nokta evrenin bir 'yanıtlanmama' ya da 'kayıp veri' alt tabakasını oluşturmaktadır. Bir takım

özel kodlamalar aracılığı ile bir değişkene yönelik yanıtlanmama alt tabakasının tanımlanması mümkündür. Bu tür bir tanımlamada gerektiğinde birden fazla kodlama kullanılabilmektedir. Örneğin 'bilgisi olmama', 'yanıtlamayı reddetme' ya da 'yasal sınırların dışında olma' gibi farklı kodlamalar yapılabilmektedir. Yanıtlanmama durumu, veri matrisine 'gözlenemeyen veri' olarak girilir. Böylelikle yanıtlanmama durumuna bağlı olarak kayıp veriler içeren bir veri seti ortaya çıkar (Little ve Rubin, 1987).

İstatistiksel analizlerde kayıp veri sorunu, önemli bir tartışma alanıdır. Graham (2009)'a göre bu tartışmaların çıkış noktası, araştırmacılar tarafından kullanılan ve çoğu yirminci yüzyılın başlarında geliştirilmiş olan analitik süreçlerin, eksiksiz veri setleri üzerinden yapılandırılmış olmasıdır. Bu analitik yaklaşımların neredeyse tamamı, kayıp veriler içeren yanıtlara yönelik herhangi bir çözüm mekanizması içermemektedir. İlk önemli çalışmaların yapılmaya başlandığı 1970'lerin sonlarına kadar kayıp verilerin, istatistiksel analizlerde bir sorun olarak görülmediği anlaşılmaktadır.

Afifi, Elashoff, Hartley, Hocking, Orchard, Woodbury, Rubin, Dempster, Laird ve Heckman gibi araştırmacılar tarafından 1970'lerde yürütülen çalışmalar, kayıp veriler üzerine ilk önemli örnekler olarak gösterilmektedir (Little ve Rubin, 1987). Bununla birlikte asıl kırılmanın, Little ve Rubin tarafından hazırlanan 'Analysis with Missing Data' ve yine Rubin tarafından hazırlanan 'Multiple Imputation for Nonresponse in Surveys' adlı kitapların 1987 yılında yayımlanmasıyla yaşandığı belirtilmektedir. Bu gelişmelere bağlı olarak istatistiksel analizlerde kayıp verilerin yol açabileceği sorunlar üzerinde durulmuş ve kayıp veri sorunu ile başa çıkabilmek için istatistiksel yöntemler geliştirilmiştir (Graham, 2009).

Kayıp verilerin istatistik alanında bile çoğunlukla göz ardı edilen bir konu olduğunu ifade eden Allison (2002), istatistik kitaplarının büyük çoğunluğunda kayıp veriler ve kayıp verilerle başa çıkma yollarına yönelik hiçbir bilgi bulunmadığını belirlemektedir. Allison'a benzer şekilde Graham (2009)'a göre de bu durum, standart istatistiksel yöntemlerin doğasından kaynaklanmaktadır. Standart istatistiksel yöntemlerin neredeyse tamamı, tam bilgi gerektirmektedir. Analizde yer alan değişkenlere yönelik her bir hücre dolu olmalıdır.

Standart istatistiksel yöntemler, dikdörtgensel veri setlerinin analizine yönelik olarak geliştirilmiştir. Bir matris şeklinde hazırlanan bu veri setlerinde satırlar gözlemleri, sütunlar ise değişkenleri temsil etmektedir. Matrisin 'hücre' olarak adlandırılan her bir birimine gerçek sayı değerleri girilir. Gerçek sayılar yaş, gelir gibi sürekli değişkenleri, eğitim düzeyi gibi sıralı kategorik değişkenleri ya da cinsiyet, uyruk gibi sınıflamalı kategorik değişkenleri temsil eder. Bir değişkene yönelik bir gözlemin bulunmaması durumunda, söz konusu gözlemi temsil eden hücre boş kalır (Little ve Rubin, 1987).

Psikolojik ölçmelerin temeli, gözlenen verilerden hareketle gözlenemeyen yani gizil değişkenlere yönelik çıkarımların yapılmasıdır. Yanıtlayıcılar tarafından, kendilerine yöneltilen maddelerin boş bırakılması ya da atlanması, bu tür bir çıkarımın yapılabilmesi önündeki en önemli engeli oluşturmaktadır. Bir maddeye verilen tepki gözlenemediğinde, örtük ya da gizil özellik hakkında bir çıkarım yapılabilmesi de doğal olarak mümkün değildir (Hohensinn ve Kubinger, 2011).

'Yanıtlanmama (nonresponse)', 'boş bırakma (omitted)' ve 'kayıp veri (missing value)' tanımlamaları, literatürde oldukça yakın anlamda ve sıklıkla birbirini yerine kullanılmaktadır. Esas itibarıyla boş bırakma, yanıtlanmama durumuna, yanıtlanmama durumu ise kayıp verilerin oluşmasına yol açmaktadır. Örneğin testi alan birey, kazara ya da unutarak bir maddeyi yanıtlamayabilir ya da örneğin bazı sağlık sorunlarından dolayı testi tamamlayana kadar bekleyemez ve uygulamadan erken ayrılır (Culbertson, 2011). Araştırmalarda irtibat kurulan bazı bireyler, kendilerine yöneltilen soruların tamamını ya da bir kısmını yanıtlamayı reddedebilir. Bu ve benzeri örnekler, bireyler, hane halkı, çalışanlar gibi gözlem birimleri içeren araştırmalarda oldukça yaygındır. Bu tür yanıtlanmama durumları sadece örneklem üzerinde yürütülen araştırmalarda ve özellikle tarama araştırmalarında değil aynı zamanda tamsayım araştırmalarında da söz konusu olabilmektedir. Tamsayım araştırmalarında, evreni oluşturan her bir birimden bilgi toplanması esastır. Örneğin ülke genelinde yürütülen bir nüfus sayımında, ülkedeki her bir bireyle irtibat kurularak yaş, cinsiyet, öğrenim durumu gibi özelliklere yönelik bilgiler kaydedilir. Örneklem üzerinde yürütülen araştırmalarda ise bu tür bilgiler ancak örneklemi oluşturan birimlerden elde edilir. Bu nedenle örneklem üzerinde yürütülen iyi

yapılandırılmış bir araştırmada, evrene yönelik güvenilir ve doğru çıkarımlar yapılabilmesi, öncelikle örneklemin dikkatli bir şekilde seçilmesine bağlıdır (Rubin, 1987).

Yanıtlanmama durumuna bağlı olarak ortaya çıkan kayıp veri sorunu, esas itibarıyla veri toplanmasında benimsenen yöntemin ve sürecin bir sonucudur (Rubin, 1987). De Ayala, Plake ve Impara (2001), yanıtlanmamaya yol açan olası üç durum üzerinde durmaktadır:

- i. Yanıtlanmama durumu, testin yapısına ve tasarımına bağlı olarak, test geliştirme sürecinde planlanmış olabilmekte ve doğal bir şekilde ortaya çıkabilmektedir. Örneğin 'bireye uyarlanmış (adaptive)' testlerde, matris yapıda testlerde ya da eşdeğer formlar içeren testlerde bireyler, soruların tamamına muhatap olmamaktadır. Bu tür durumlarda kayıp veri olarak ortaya çıkan yanıtlanmama durumları, 'uygulanmamış (not administered) madde' olarak kodlanır ve bu kayıplar önceden planlandığı şekliyle parametre kestirimlerinde bir soruna yol açmaz.
- ii. Hız testlerinde, yanıtlanmama durumu söz konusu olabilmektedir. Bu tür testlerde, test uygulaması, belirli bir zaman aralığında gerçekleştirilmektedir ve bazı yanıtlayıcılar öngörülen sürede tüm maddelere erişemeyebilmektedir. Erişememeye bağlı yanıtlanmama durumu, yanıtlayıcının kararı olmadığı gibi genellikle testin geliştirilme sürecinde planlanan ve öngörülen bir durum da değildir. Bununla birlikte bu tür yanıtlanmama durumunun, test ile ölçülmek istenen özellik açısından bireyler arasındaki farklılıklarla ilişkili olduğu ileri sürülebilir.
- iii. Yanıtlayıcılar, zamanları yeterli olsa bile bilinçli bir şekilde bir maddeyi boş bırakmayı tercih edebilmektedir. Bu durum yanıtı bilmemeleri ya da bilgilerine güvenmemeleri gibi birçok nedene bağlı olabilmektedir.

Lord (1983), araştırmalarda yanıtlanmamaya bağlı olarak kayıp verilerin oluşmasını, 'erişilememe (not reached)' ve 'boş bırakma (omitted)' olmak üzere iki durumda ele almaktadır. Bir testin hız testi olup olmasına bağlı olarak, yanıtlayıcıların bazılarının, testte yer alan maddelerden bazısına ve özellikle testin son maddelerine yanıt verememesi durumu, erişilememe durumu olarak tanımlanmaktadır. Yanıtlayıcının doğru yanıtı

bilmemesi ya da bilgisine güvenmemesi ve riske girmek istememesine bağlı olarak maddeyi yanıtlamaktan kaçınması durumu ise boş bırakma olarak tanımlanmaktadır. Buna göre yanıtlanmamaya bağlı olarak kayıp verilerin oluşması, hem yanıtlayıcının yeterli düzeyi hem de testin yapısı ile ilişkili görülmektedir.

Rubin (1987), yanıtlanmama durumunu 'birim' ve 'madde' düzeyinde olmak üzere iki grupta ele almaktadır. Bir araştırmada gözlem birimi tarafından, kendisine yöneltilen tüm maddelerin yanıtlanmaması durumu 'birim düzeyinde yanıtlanmama (unit nonresponse)', bazı maddelerin yanıtlanmaması durumu ise 'madde düzeyinde yanıtlanmama (item nonresponse)' olarak tanımlanmaktadır. Birim düzeyinde yanıtlanmama, örnekleme hatası şüphesini, madde düzeyinde yanıtlanmama ise madde yanlılığı şüphesini ortaya çıkarabilmektedir.

Puma, Olsen, Bell ve Price (2009)'a göre verilerde kayıpların ortaya çıkmasına neden olan en belirgin yol, bazı yanıtlayıcılardan ya da gözlemlerden hiçbir bilgi elde edilememesidir. Bu durum, araştırma literatüründe 'birim düzeyinde yanıtlanmama (unit nonresponse)' olarak tanımlanmaktadır. Örneğin eğitim alanında öğrencilerin bir kısmı, ailelerin onayı olmadığı için ya da başka bir okula nakil gittikleri için test uygulamasına katılmayabilmektedir. Benzer şekilde öğretmenlerden veri toplanan bir araştırmada, öğretmenlerin bir kısmı test almayı reddedebilmekte ya da uygulamanın yapıldığı zaman okulda bulunamayabilmektedir. Bir veri toplama aracının tamamına yanıt verilmemesi durumuna ek olarak sıklıkla karşılaşılan bir diğer sorun, yanıtlayıcının bazı maddeleri yanıtlamayı reddetmesi durumunda ortaya çıkmaktadır. Yanıtlayıcı, özellikle zihinsel davranışları ölçen testlerde ya da tipik performans testlerinde, yanıtı bilmeme, madde ile ölçülmek istenen özellik açısından yeterli olmama, maddeyi kazara atlama gibi nedenlerle bazı maddeleri yanıtlamayabilmektedir. Bu durum ise 'madde düzeyinde yanıtlanmama (item nonresponse)' olarak tanımlanmaktadır.

Yanıtlanmama ya da boş bırakmaya bağlı olarak ortaya çıkan temel sorun, amaçlanan gözlemlerde kayıp verilerin oluşmasıdır. Kayıp veriler, veri setinin daralması ve buna bağlı olarak yapılacak kestirimlerin gücünün azalması anlamına gelmektedir. Kayıp veriler içeren bir veri seti üzerinde,

eksiksiz veri setlerine göre yapılandırılmış standart analiz yöntemlerinin kullanılması da mümkün değildir. Dahası, yanıtlayanlar ile yanıtlanmayanlar arasındaki, sıklıkla sistematik olan farklılıklardan dolayı, olası bir yanlılık söz konusudur. Yanıtlanmayanların ve yanıtlanmama gerekçesinin genellikle bilinmiyor olması da, bu tür bir yanlılığın giderilmesini zorlaştırmaktadır (Rubin, 1987).

Araştırmalarda yanıtlanmama durumuna bağlı olarak ‘yanıtlanmama yanlılığı (nonresponse bias)’ ya da ‘boş bırakma/atlama yanlılığı (omitted response bias)’ oluşabilmektedir. Bir araştırmada yüksek ‘yanıtlanmama oranı (nonresponse rate)’, manidar düzeyde bir sistematik hata şüphesi ortaya çıkarmaktadır. Bu durumda, yanıtlayıcılardan elde edilen veriler üzerinde, yanlılık incelemelerinin yapılması gerekli görülmektedir. Potansiyel bir yanıtlanmama yanlılığının incelenmesi, genellikle, bazı temel değişkenler açısından yanıtlayanlar ve yanıtlanmayanlar arasında manidar bir farklılık bulunup bulunmadığının test edilmesi yaklaşımına dayanmaktadır (Whiple ve Muffo, 1982; Culbertson, 2011). Örneğin Kanuk ve Berenson (1975), e-mail yolu ile veri toplamada yanıtlanma oranlarını inceledikleri araştırmalarında, eğitim düzeyinin önemli bir yordayıcı olduğu sonucuna ulaşmışlardır. Buna göre eğitim düzeyi yükseldikçe, yanıtlanma oranı da artmaktadır. Diğer bir deyişle, yanıtlanmama durumu, eğitim düzeyi açısından sistematik hata oluşturma potansiyeline sahiptir. Diğer taraftan yüksek ‘yanıtlanmama oranı (nonresponse rate)’, düşük ‘yanıtlanma oranı (response rate)’ ortaya çıkarmaktadır. Düşük yanıtlanma oranı, örneklemin temsil edebilirlik gücünü azaltmakta ve buna bağlı olarak kestirimlerin güvenilirliğini olumsuz etkilemektedir. Dolayısıyla bir araştırmada, özellikle dış geçerlik ve güvenilirlik açısından, yanıtlanma oranının olabildiğince yüksek olması istenir. Bu nedenle bir araştırmada gerekli bilgilerin elde edilebilmesi için gözlem birimlerinin teşvik edilmesinin yanı sıra zamanlama ve kullanılan yöntem açısından titiz olunması da gerekli görülmektedir (Whiple ve Muffo, 1982).

Brick (2001), Amerikan Ulusal Eğitim İstatistikleri Merkezi (U.S.NCES) tarafından anaokuluna başlayan çocukların hazırbulunuşluk düzeylerine yönelik olarak yürütülen bir araştırmanın verilerini kullanarak, yanıtlanmama yanlılığını araştırmıştır. Çocuklara yönelik cinsiyet, okul türü, anne-baba eğitim ve gelir düzeyi gibi değişkenler açısından yanıtlayanlar ve

yanıtlamayanlar arasındaki manidar farkları test etmiştir. Buna bağlı olarak yanıtlanmama yanlılığının kaynağını belirlemeye çalışmıştır. Okul türü değişkeninin, yanıtlama yapan ve yapmayanlar arasındaki en iyi açıklayıcı değişken olduğunu belirlemiştir. Okul türü değişkeni ile birlikte diğer olası yanlılık kaynaklarına yönelik, ortalama düzeltme yöntemi ile kayıp veri ataması yapmış ve analizleri tekrarlanmıştır. Araştırmada, yapılan düzeltmelerin yanıtlanmama yanlılığını ve manidar farkları yeterince gideremediği sonucuna ulaşılmıştır. Yanıtlanmama yanlılığına yönelik anlamlı bir düzeltmenin yapılabilmesi için yanıtlama yapmayan bireylere yönelik daha ayrıntılı bilgilere ihtiyaç duyulduğu ifade edilmektedir.

Rubin (1987)'e göre bir araştırmada yanıtlanmama durumunun bir sorun haline gelmesi, yanıtlanmama oranına bağlıdır. Groves (2006)'a göre ise yanıtlanmama oranı ile yanıtlanmama yanlılığı arasında doğrudan bir ilişki bulunmamaktadır. Bununla birlikte, yanıtlanmama oranının düşük olması, söz konusu sistematik hatanın manidar olmadığına yönelik bir kanıt olarak görülebilmektedir.

Yukarıdaki açıklama ve tartışmalar kayıp verilerin yol açabileceği sorunları ortaya koymaktadır. Peng, Harwell, Liou ve Ehman (2007), kayıp verilere bağlı olarak ortaya çıkabilecek dört önemli problemi, ilgili literatür bağlamında özetlemektedir:

- i. İstatistiksel bir modelden hareketle yapılacak kestirimlerde kayıp verilerin yol açacağı yanlılık, en önemli sorun olarak ifade edilmektedir. Örneğin yanıtlamayı reddedenlerin profilinin yanıtlayanların profilinden farklı olması, bu tür bir yanlılığın doğal nedeni olabilmektedir. Bu tür bir örnekte yanıtlayanlardan oluşan örneklem, 'seçkisizlik (randomization)' özelliğini kaybetmiştir. Örneklemin evreni temsil edebilirlik düzeyi düşüktür. Araştırma sonuçları, ancak yanıtlayanlarla sınırlıdır. Dolayısıyla verilen kararlar yanlılık içermektedir.
- ii. Kayıp veriler, bilgi eksikliğine ve buna bağlı olarak istatistiksel analizin gücünün azalmasına yol açabilmektedir. Bazı değişkenlerin analiz dışı bırakılması, örneğin t testi gibi istatistiksel testlerde serbestlik derecesinin hata miktarını artırmaktadır. Bu artış istatistiksel gücün azalmasına ve standart hatanın artmasına yol açmaktadır.

- iii. Yaygın olarak kullanılan istatistiksel yöntemlerin, kayıp veriler içeren veri setlerinde kullanımı zordur. Örneğin kayıp veriler, faktör analizinde dengesizliğe yol açmaktadır. Özel istatistik yazılımlarının kullanımını gerektiren çok değişkenli istatistiksel yöntemler de ancak eksiksiz veri setlerinde uygulanabilir durumdadır.
- iv. Kayıp veriler, değerlendirilebilir kaynakların boşa harcanmasına yol açabilmektedir. Boylamsal çalışmalar, geniş ölçekli testler, tamsayım ve benzeri çalışmalarda verilerin toplanması, önemli bir zaman ve maliyet içermektedir. Bu tür araştırmalarda araştırmacının yüksek yanıtlanma oranları ve yanıtlayıcılara yönelik tam bir profil elde edebilmesi, ekstra çaba, zaman ve maliyet anlamına gelmektedir.

Kayıp Verilerin İhmal Edilebilirliği

Allison (2002)'a göre araştırmacılar genellikle, kayıp verilerin gözlenen ölçmelerde manidar bir farka yol açmadığını kanıtlama eğilimindedir. Örneğin, gelirlerini beyan eden ve etmeyen yanıtlayıcılar arasında, dikkate alınan değişkenler açısından manidar bir fark oluşmadığına yönelik kanıt sağlama çabası oldukça yaygındır. Daha genel bir durum olarak araştırmacılar, ne anlama geldiğini açıkça anlamamalarına rağmen sıklıkla verilerinin 'seçkisiz kayıp (missing at random)' içerdiğini varsaymakta ve bu varsayıma bağlı olarak kayıp verileri analiz dışı bırakmaktadırlar.

Demir ve Parlak (2012), eğitim araştırmalarında kayıp veri sorununu ele aldıkları çalışmalarında araştırmacıların, genellikle eksiksiz veri seti üzerinde çalışma eğiliminde olduklarını belirlemiştir. Kayıp verilerin karakteristiğinin incelendiği araştırma sayısı oldukça azdır. Kayıp veriler, sıklıkla herhangi bir istatistiksel kanıt olmaksızın ihmal edilmekte ve analiz dışı bırakılmaktadır. Benzer bir araştırmada Groves (2006) araştırmacıların, herhangi bir istatistiksel kontrol kullanmaksızın ve sıklıkla herhangi bir rasyonel gerekçe sunmaksızın, değişkenler düzeyinde %15 ile %50 arasında kayıp veriyi ihmal edebildiğini belirlemiştir.

Rubin (1976)'e göre çoğu araştırmada kayıp verilerin analiz dışı bırakılmasının temelinde bu verilerin, bir takım varsayımlara bağlı olarak 'ihmal edilebilir (ignorable)' görülmesi yatmaktadır. Örneğin kayıp verilerin seçkisiz olarak oluştuğu, bu nedenle veri dağılımında bir sapma ya da farklılık oluşmayacağı, çok değişkenli normal dağılım varsayımının sağlanması durumunda her bir değişkene yönelik kayıp veri oluşma olasılığının eşit olacağı ya da bağımlı değişkenin, gözlenen değişkenle ilgisi olmayacak şekilde kayıp veri içerebileceği gibi varsayımlara bağlı olarak kayıp verilerin analiz dışı bırakıldığı araştırmalara sıklıkla rastlanmaktadır. Allison (2002)'a göre esas itibariyle ihmal edilebilirlik, kestirim sürecinin bir parçası olarak kayıp veri mekanizmasının modellenmesine gerek olmadığı anlamına gelmektedir.

Kayıp verilerin analiz dışı bırakılabilmesi için öncelikle kayıp verilerin ihmal edilebilir olduğunun kanıtlanması gerekmektedir. Bu tür bir kanıt, kayıp değerlerle, seçkisiz olarak belirlenen gözlenen verilerin örtüşeceği varsayımını destekler. Dolayısıyla tam ya da eksiksiz veri seti ile kayıp veri içeren veri setlerinin dağılımları ve bu veriler üzerinden yapılacak çıkarımlar arasında manidar bir fark olmadığı varsayımına ulaşılır (Rubin, 1976).

Kayıp verilerin ihmal edilebilir olup olmaması, kayıp verilerin karakteristiği ve kayıp verilerin oluşmasına yol açan süreçlerle ilişkilidir. Kayıp veri örüntüsü ve mekanizması, uygun analiz yöntemlerinin belirlenmesinde ve sonuçların yorumlanmasında anahtar rol oynamaktadır. Bazı durumlarda kayıp veri örüntüsü ve mekanizması, istatistikçinin kontrolü altındadır. Örneğin örneklem üzerinde yürütülen bir araştırmada kayıp veri örüntüsü ve kayıp veriye yol açan mekanizma, basitçe örnekleme sürecinden kaynaklanıyor olabilir. Örnekleme sürecinin olasılıklı örneklemeyle dayalı olarak yapılmış olması durumunda, kayıp veri örüntüsü ve mekanizmasının kontrol altında ve ihmal edilebilir olduğu ileri sürülebilir. Örnekleme çatısı eksikse ya da bazı birimler yanıt vermemişse, gözlenen verilerin oluşmasına yol açan mekanizma iyi bir şekilde anlaşılamaz. Dolayısıyla veri analizi, kayıp veri örüntüsüne ve bu örüntüyü açlayan kayıp veri mekanizmasına yönelik varsayımlara bağlıdır (Little ve Rubin, 1987).

Kayıp veri mekanizmasının ihmal edilebilir olması, kayıp verilerin analiz dışı bırakılabilmesi ile araştırmacının işini oldukça kolaylaştıracaktır. Fakat aksi durumda kayıp veri mekanizmasının modellenmesi gerekir. İhmal edilemez kayıp veriler içeren veri seti, genellikle hangi modellerin daha uygun olacağı konusunda yeterli bilgi sağlayamaz. Araştırma sonuçlarının seçilen modele oldukça duyarlı olacağı da açıktır. Dolayısıyla ihmal edilebilir olmayan kayıp verilerle etkili bir kestirim yapılabilmesi için, kayıp veri sürecinin doğasına yönelik oldukça iyi ön bilgilere ihtiyaç vardır (Schaffer, 1997; Allison, 2002).

İki Kategorili Puanlanan Maddelerden Oluşan Testlerde Kayıp Veri Sorunu

Madde puanlamada benimsenen iki temel biçim, 'ikili puanlama (dichotomously)' ve 'çoklu puanlama (polytomously)' olarak tanımlanmaktadır. Bu tanımlamalara bağlı olarak bir ölçme aracında kullanılabilecek maddeler, puanlama biçimine göre ikiye ayrılmaktadır: (i) iki kategorili (dichotomous) puanlanan maddeler ve (ii) çoklu puanlanan (polytomous) maddeler. İki kategorili puanlanan maddelerde, esas olarak ilgilenilen özelliğin varlığı ya da yokluğu belirlenmeye çalışılmaktadır. Bu belirleme, doğru-yanlış, var-yok, evet-hayır gibi sunulan seçeneklere göre yapılan yanıtlamalara dayanmaktadır. Başarı testlerinde yaygın bir şekilde kullanılan 'çoktan seçmeli' maddeler de puanlama biçimi açısından ele alındığında, iki kategorili puanlanan maddeler arasında değerlendirilebilmektedir (Erkuş, 2012).

Bir ölçme aracında yer alan maddelerin, ölçek üzerinde yer alan her bir noktada ayırıcılık özelliğine sahip olması beklenmektedir. Tavşancıl (2010), maddelerin bu 'ayırım gözetme işlevinin' özellikle tutum ölçeklerinde, tutum maddelerinin karşılaması gereken önemli ölçütlerden biri olduğunu ifade etmektedir. Ölçek üzerindeki her bir noktada ayırıcılık özelliği, madde puanlarının süreklilik gösterdiği ideal durumda sağlanabilmektedir. Sürekli verilerin elde edilebilmesi için maddelerin en az eşit aralık ölçeğinde olması gerekir. İki kategorili puanlanan maddeler ise sınıflama ölçeğinde

değişkenlerdir ve 1-0 yapay süreksiz veri üretebilmektedir. Aiken (1995), iki kategorili puanlanan maddelerden oluşan testlerde, 20 maddeden az madde bulunması durumunda, testle ölçülmek istenen özelliğin süreklilik gösterme olasılığının düştüğünü belirtmektedir. Bu bağlamda Erkuş (2012), iki kategorili puanlanan maddelerin, ancak gerçek süreksiz yapıdaki özelliklerin ölçülmesinde kullanılmasını önermektedir.

İki kategorili puanlanan maddelerden oluşan testlerde veri seti, kesikli verilerden oluşmaktadır. Bu tür veri setlerine yönelik analizlerin, sürekli veriler içeren veri setlerine yönelik analizlerden daha zor ve karmaşık olduğu bilinmektedir.

Pratikte, bir veri setinde kayıp verilerin oluşmasını engellemek oldukça zordur. İki kategorili puanlanan maddelerde, yanıtlanmamaya bağlı olarak oluşan kayıp verilerde en yaygın yaklaşım, bu kayıp verilerin ‘yanlış’ olarak kabul edilmesidir. Bu yaklaşımın yaygın bir alternatifi ise kayıp verilerin ‘uygulanmamış (not administered)’ olarak kodlanmasıdır. Söz konusu bu her iki yaklaşımın da yanlış ve hatalı kestirimlere yol açtığı bilinmektedir. Bu nedenle bu tür kayıp verilere yönelik olarak geliştirilmiş daha karmaşık kayıp veri yöntemlerinin kullanımı gerekli görülmektedir (Hohensinn ve Kubinger, 2011; Culbertson, 2011).

Lord (1973, 1974), yetenek ve parametre kestirimlerinde, atlanan ya da boş bırakılan maddelerin ‘yanlış’ olarak kodlanmasının doğru bir yaklaşım olmadığını göstermiştir. Özellikle hız testlerinde, testin bazı maddelerine ve genellikle son maddelerine erişemeyen yanıtlayıcıların boş bıraktığı maddelerin, ‘yanlış’ ya da ‘uygulanmamış’ olarak kodlanması, kestirim sürecinde önemli bir sorun olarak ifade edilmektedir. Bu yaklaşım, yaygın bir şekilde kullanılmakla birlikte kayıp verilere yönelik optimal bir yöntem olarak görülmemektedir. Ancak boş bırakılan maddelerin tamamen seçkisiz oluşması koşulunun sağlanması durumunda, yanıtlama yapan ve yapmayan bireylere yönelik kestirimlerin karşılaştırılması mümkün olabilmektedir.

Literatürde özellikle iki kategorili puanlanan maddeler içeren çok boyutlu testlerde parametre kestirimlerine yönelik istatistiksel tekniklerin istenen düzeyde geliştirilememiş olduğu belirtilmektedir. Bununla birlikte ‘madde tepki kuramı-MTK (item response theory)’ temelinde geliştirilmiş modellerin, iki kategorili puanlanan maddelere yönelik veri setlerinde sıklıkla

kullanıldığı görülmektedir (McKinley ve Reckase, 1983; Little ve Rubin, 1987).

MTK, gözlenemeyen yetenek parametresi θ verildiğinde, bireyin bir maddeyi yanıtlama olasılığı ile madde karakteristiklerini 'koşullu bağımsız (conditional independent)' olarak ilişkilendirmekte ve modellemektedir. Klasik test kuramından farklı olarak yetenek kestirimlerinin dikkate alınması, farklı yetenek düzeylerindeki bireylere farklı testlerin uygulanabilmesini olanaklı kılmaktadır. Dolayısıyla MTK temelinde geliştirilen bir testte bireylerin, testte yer alan bazı sorulara yanıt vermeleri, yetenek kestirimleri için yeterli olabilmektedir. MTK'nın sağladığı bu avantaj, bir veri setinde, erişilemeyen, boş bırakılan ya da atlanan maddelere bağlı olarak oluşan kayıp verilere yönelik, alternatif çözümler üretilebilmesini sağlamıştır (Lord, 1980).

Finch (2008), kayıp verilerin bulunması durumunda madde tepki kuramı temelinde, iki kategorili puanlanan veriler üzerinde, parametre ve yetenek kestirimlerine yönelik yöntemleri incelemiştir. 500 ve 1000 veriden oluşan iki ayrı simülatif veri tabanı üzerinde, 'beklenti-maksimizasyon algoritması (BM),' 'çoklu veri atama (ÇVA)' yöntemleri gibi farklı kayıp veri yöntemlerini kullanarak, madde ayıricılığı, madde güçlüğü ve olasılık kestirimleri gerçekleştirmiştir. Ayrıca kayıp verilere yönelik olarak 'yanlış', 'uygulanmamış' gibi yaygın kodlama yaklaşımlarını da ele almıştır. Sonrasında veri eksiltme ile bu veri tabanlarını 'tam seçkisiz kayıp (TSK)', 'seçkisiz kayıp (SK)' ve 'seçkisiz olmayan kayıp (SOK)' varsayımlarına uygun hale getirerek kestirimleri tekrarlamıştır. Araştırmada, iki kategorili puanlanan maddelerde kayıp verilerin 'yanlış' ya da 'uygulanmamış' olarak kodlanmasının, TSK ya da SK varsayımları karşılanırsa bile, kestirimlerde yanlışlığa yol açtığı sonucuna ulaşılmıştır. Ayrıca titiz parametre kestirimlerinin sağlanabilmesi için en uygun yöntemin, uygun veri atama yöntemleri ve özellikle ÇVA yöntemleri olduğu belirlenmiştir.

Schaffer (1997), kayıp veriler yerine uygun yöntemlerle ve özellikle ÇVA yöntemleri ile atama yapılmış olan kategorik verilerden oluşan veri setlerinde, normallik varsayımına dayalı modellerin kullanılabileceğini belirtmektedir. Bununla birlikte Allison (2006), iki kategorili puanlanan verilerde ÇVA yöntemlerinin kullanılabilirliğini incelediği araştırmasında, ortalama kestirimlerinin yanlışlık içerebileceği bulgusuna ulaşmıştır.

İki kategorili puanlanan maddelere yönelik en bilindik modellerden birisi madde tepki kuramı temelinde geliştirilmiş Rasch modelidir. Rasch modelinde yetenek kestirimi, madde güçlük indeksine ve sabit değerli ayırıcılık gücü indeksine bağlı olarak oluşturulan bir olasılık-yoğunluk fonksiyonu ile ifade edilmektedir. Rasch modelinde, 'koşulsuz en çok olabilirlik kestirimi (unconditional maximum likelihood estimation)', 'koşullu en çok olabilirlik kestirimi (conditional maximum likelihood estimation)', 'katılmalı en çok olabilirlik kestirimi (joint maximum likelihood estimation)', 'marjinal en çok olabilirlik kestirimi (marginal maximum likelihood estimation)', 'kısmi koşullu kestirim (pairwise conditional estimation)' gibi olasılıklı ve ötelemeli kestirim yöntemlerine yönelik algoritmalar kullanılabilmektedir. Bu algoritmalar, genellikle kayıp verilere yönelik bir düzeltme içermemektedir ve bu nedenle eksiksiz veri seti gerektirmektedir. Eksiksiz veri seti, 'dizin silme (DS)' ve 'kısmî silme (KS)' gibi silmeye dayalı yöntemler ya da 'basit atama (BA)' yöntemleri ile elde edilebilmektedir. Eksiksiz veri setinin ÇVA yöntemleri ile elde edilmesi de mümkündür. Dolayısıyla kayıp verilerin varlığında Rasch modeli ve benzeri modellerin kullanılabilmesi için öncelikle uygun kayıp veri yöntemleri ile eksiksiz veri seti elde edilmesi ve ileri analizlerin bu veri seti üzerinden gerçekleştirilmesi gerekmektedir (Linacre, 2004).

Hohensinn ve Kubinger (2011), 1 ve 0 şeklinde iki kategorili puanlanan veri örüntülerine yönelik Rasch model üzerinde, kayıp veriler ile madde uyumu ve model geçerliği arasındaki ilişkileri incelemişlerdir. Simülatif veriler üzerinden gerçekleştirilen analizlerde üç ayrı veri tabanı oluşturulmuştur. İlk veri tabanı, kayıp veri içermemektedir. İkinci veri tabanında, seçkisiz şekilde veri eksiltilerek kayıp veriler oluşturulmuş ve bu kayıp veriler 'yanlış' olarak değerlendirilerek 0 kodlanmıştır. Üçüncü veri tabanında ise yine seçkisiz şekilde veri eksiltme ile kayıp veriler oluşturulmuş ve bu kayıp veriler 'uygulanmamış (not administered)' olarak kodlanarak analiz dışı bırakılmıştır. Kayıp verilerin 'yanlış' ya da 'uygulanmamış' olarak kodlanmasının, madde ve yetenek parametreleri kestirimlerinde manidar farklılıklar ortaya çıkardığı sonucuna ulaşılmıştır.

Bir testin son maddelerine erişilememesi, genellikle hız testlerinde karşılaşılan bir durum olarak, kayıp verilerin ortaya çıkmasının nedenlerinden biridir. Mislevy ve Wu (1996), zaman sınırlılığına bağlı olarak bir testin son

maddelerine ulaşılamaması durumunda, TSK ve SK varsayımlarının ihlâl edildiğini ispatlamışlardır. Bu durumda yetenek kestirimlerine yönelik olarak doğrudan olabilirlik ya da Bayesian temelli yöntemlerin kullanılabileceği ifade edilmiştir.

DeMars (2002), iki kategorili puanlanan maddelerden oluşan bir testin son maddelerine erişilememesi durumunda madde tepki kuramı temelinde parametre kestirimlerini araştırmıştır. Simülatif veriler üzerinden gerçekleştirilen analizlerde, 1 parametrelili lojistik modelde yetenek kestirimleri gerçekleştirilmiş ve veri seti, düşük ve yüksek yetenek düzeyi olmak üzere iki gruba ayrılmıştır. Madde güçlük kestirimleri bu yetenek düzeylerinde ayrı ayrı gerçekleştirilmiştir. TSK ve SK varsayımları altında, ‘katılımlı EÇO (joint MLE)’ ve ‘marjinal EÇO (marginal MLE)’ yöntemlerine göre madde güçlük kestirimleri elde edilmiştir. Araştırmada, kayıp veriler içeren maddelere yönelik kestirimlerin manidar düzeyde yanlışlık içerdiği sonucuna ulaşılmıştır. Lord (1983) ise DeMars (2002)’in aksine, ihmal edilebilir olmayan kayıp verilerin varlığında ‘katılımlı EÇO’ gibi olasılık temelli bazı yöntemlerin kullanılabildiğini belirtmektedir.

Özellikle iki kategorili puanlanan maddelerden oluşan testlerde, yanıltıcının bilinçli bir şekilde boş bıraktığı ya da atladığı maddelere bağlı olarak oluşan kayıp verilerin ihmal edilebilir olmadığı, üzerinde uzlaşılan bir belirlemedir. Bu tür kayıp verilerin varlığında, elde edilen kestirimlerin manidar düzeyde yanlışlık içerme olasılığı yüksektir. Yanlışlık olasılığını düşürmede en yüksek performans gösteren kayıp veri yöntemlerinin ise EÇO ve ÇVA yöntemleri olduğu sıklıkla belirtilmektedir.

Kayıp veri sorununa yönelik olarak geliştirilmiş pek çok çözüm yönteminden bazılarının, diğerlerine göre daha ‘iyi’ olduğu ileri sürülebilir. Fakat bu yöntemlerden hiç birisi gerçek anlamda ‘iyi’ olarak tanımlanamaz. Gerçek çözüm, kayıp veri olmaması ya da kayıp veri miktarının ihmal edilebilir düzeyde olmasıdır. Bu nedenle bir araştırmada kayıp veri miktarının en aza indirgenmesi çabası önemlidir. Özensiz ve dikkatsiz bir şekilde yürütülen bir araştırmada herhangi bir istatistiksel düzeltmenin bir anlamı da yoktur (Allison, 2002).

Araştırmalarda kayıp veri sorununun ancak 1970 sonrasında tartışılmaya başlanması ve kayıp veri yöntemlerinin ancak 1990 sonrasında istatistiksel yazılımlara yansması, geliştirilmiş birçok kayıp veri yönteminin henüz beklentileri karşılayacak düzeyde olmamasını bir ölçüde açıklamaktadır. Normallik varsayımının sağlanamadığı modellerde, ihmal edilebilir olmayan kayıp verilerin varlığında, kategorik ve iki kategorili puanlanan değişkenlerde, çok değişkenli ve çok düzeyli analizlerin büyük çoğunluğunda kayıp veriler, hâlâ önemli bir sorun olarak tartışılmaktadır.

Çok kategorili puanlanan maddelerden oluşan testlerde ya da sürekli değişkenlerde kayıp verilerle başa çıkabilme, üzerinde daha fazla çalışılmış bir alandır ve bu tür verilere yönelik olarak görece daha fazla yöntem geliştirildiği görülmektedir. İki kategorili puanlanan maddelerden oluşan testlerde ise kayıp verilerle başa çıkabilme, madde karakteristiklerinin doğal bir sonucu olarak oldukça zordur. Maddelerin karakteristiği, kayıp veri yöntemleri açısından önemli bir sınırlılık oluşturmaktadır. Bu nedenle kayıp verilerin doğasına yönelik olabildiğince ayrıntılı ön bilgilere ihtiyaç vardır. İleri analizler öncesinde kayıp veri mekanizması ve örüntüsü ayrıntılı bir şekilde incelenmelidir. Toplam puanların anlamlı olduğunun belirlenebilmesi için yapı geçerliği ve tek boyutluluk özellikleri sağlanmalıdır. Bu ön bilgilere bağlı olarak kullanılabilecek uygun kayıp veri yönteminin ve bu yöntemin testin psikometrik özelliklerine yönelik kestirimlerde yeterli performans gösterip göstermediğinin belirlenmesi gereklidir.

Amaç

Bu çalışmanın amacı; kayıp veriler içeren ve iki kategorili puanlanan maddelerden oluşan testlerin psikometrik özelliklerine yönelik istatistiksel kestirimlerde, kullanılması uygun olan ve uygun olmayan kayıp veri yöntemlerinin belirlenmesidir. Bu amaçla, gerçek bir veri seti üzerinde, farklı kayıp veri yöntemlerine göre testin psikometrik özelliklerine yönelik parametre kestirimleri gerçekleştirilmiştir.

Belirlenen amaç doğrultusunda veri setine yönelik olarak aşağıdaki sorulara yanıt aranmaktadır:

1. Kayıp veriler ihmal edilebilir midir?
 - a. Kayıp veri örüntüsü nasıldır?
 - b. Kayıp veri mekanizması nasıldır?
2. Kayıp veri yöntemlerine göre test, yapı geçerliği ve tek boyutluluk özelliğine sahip midir? Yapı geçerliği ve tek boyutluluk özelliğine yönelik olarak kestirilen parametreler, kullanılan kayıp veri yöntemlerine göre nasıl bir değişim göstermektedir?
3. Kayıp veri yöntemlerine göre kestirilen
 - a. Madde parametreleri nelerdir?
 - b. Test parametreleri nelerdir?

Bu parametreler, kullanılan kayıp veri yöntemlerine göre nasıl bir değişim göstermektedir?

Önem

Analiz süreçlerinde kayıp verilerin dikkate alınmasının bir gereklilik olduğu ve kayıp verilerin dikkate alındığı araştırmalarda, daha az yanlılık ve sistematik hata içeren sonuçlara ulaşabilme olasılığının yüksek olduğu açıktır. İki kategorili puanlanan maddelere yönelik kayıp veri yöntemleri, sürekli verilere yönelik yöntemlere göre daha karmaşıktır. Dahası bu yöntemlerin, titiz kestirimler sağlayabildiklerine yönelik henüz yeterli kanıtların elde edilemediği ileri sürülebilir. Bu çalışma, iki kategorili puanlanan verilerde kullanılabilecek uygun olan ve uygun olmayan kayıp veri yöntemlerine yönelik kanıt sağlaması açısından önemli görülmektedir.

Kayıp veri yöntemlerine yönelik araştırmalarda sıklıkla simülatif veriler üzerinde çalışıldığı görülmektedir. Gerçek veriler üzerinde gerçekleştirilen çalışmaların azlığı, kullanılan yöntemlerin performansının değerlendirilmesinde önemli bir sınırlılık oluşturmaktadır. Bu çalışmada, analizlerin gerçek veri seti üzerinde gerçekleştirilmesi, bir takım sorgulamaları beraberinde getirmekle birlikte, söz konusu kayıp veri

yöntemlerinin performansının değerlendirilmesinde bir avantaj olarak değerlendirilmektedir.

Kayıp veri yöntemlerine yönelik araştırmalar, literatürde belli bir birikime ulaşmıştır. Bununla birlikte kayıp verilerin varlığında iki kategorili puanlanan maddelerden oluşan testlere yönelik kestirimlerde kullanılabilecek yöntemler, üzerinde daha az çalışılmış bir alandır. Bu çalışma, ilgili literatüre katkı sağlaması ve özellikle iki kategorili puanlanan maddelere yönelik kayıp veri sorununa dikkat çekmesi beklentisi ile önemli görülmektedir.

Sınırlılıklar

Bu çalışma;

1. Testin incelenen psikometrik özellikleri,
2. Parametre kestirimlerinde dikkate alınan test kuramı,
3. Kayıp veri örüntüsü ve mekanizmalarına yönelik incelemelerde dikkate alınan değişkenler ve
4. Kullanılan kayıp veri yöntemleri

kapsamında sınırlandırılmıştır.

Bu çalışmada SBS 2011 Matematik Testi A Kitapçığının incelenen psikometrik özellikleri, (i) yanıt örüntülerinin yapısı ve özellikleri, (ii) yapı geçerliği ve tek boyutluluk özelliği, (iii) madde ve test puanlarının karakteristikleri olarak belirlenmiştir.

Bu çalışmada, madde ve test parametrelerinin kestirimleri, ‘klasik test kuramı (KTK)’ temelinde gerçekleştirilmiştir. SBS, ‘madde tepki kuramı (MTK)’ temelinde geliştirilen bir test olmadığı ve değişmezlik gibi MTK’nın temel varsayımlarını karşıladığına yönelik istatistiksel kanıtlar bulunmadığı için bu çalışmada KTK yaklaşımı tercih edilmiştir.

Bu çalışmada, kayıp veri örüntüsü ve mekanizmasına yönelik inceleme ve analizler, SBS 2011 Matematik Testi A Kitapçığı verilerinde yer alan değişkenler ve bu veriler kullanılarak üretilebilen değişkenlerle sınırlandırılmıştır. Bu kapsamda söz konusu inceleme ve analizlerde, yanıtlayıcılara yönelik 1-0 yanıt örüntüleri, cinsiyet, okul türü ve bölge, değişkenleri, toplam doğru ve yanlış yanıt sayıları, kayıp veri içeren madde

sayıları ve yanıt örüntülerinde kayıp veri bulunma durumunu tanımlayan bir 'kukla değişken' dikkate alınmıştır.

Bu çalışmada kullanılan kayıp veri yöntemleri, kayıp veri örüntüsü, kayıp veri mekanizması ve ilgili istatistik yazılımlarının özellikleri de dikkate alınarak 12 yöntemle sınırlandırılmıştır. Amaca uygun olarak belirlenen bu kayıp veri yöntemleri; 'dizin silme (DS)', '0 atama', 'seri ortalamaları ataması (SO)', 'gözlem birimleri ortalaması ataması (BO)', 'yakın değerle ortalama atama (YNO)', 'yakın değerler medyan atama (YNM)', 'doğrusal interpolasyon (Dİ)', 'doğrusal nokta eğilimi (DE)', 'regresyon atama (RA)', 'beklenti-maksimizasyon algoritması (BM)', 'veri çoğaltma (VÇ)' ve 'Markov Chain Monte Carlo (MCMC)' yöntemleridir.

BÖLÜM 2

KAVRAMSAL ÇERÇEVE

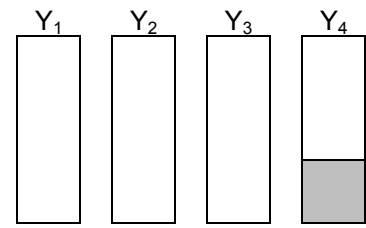
Bu bölümde öncelikle kayıp verilerin ihmal edilebilir olup olmadığı kararına gerekçe oluşturan kayıp veri örüntüsü ve kayıp veri mekanizmaları açıklanmaktadır. Sonrasında kayıp verilerin ihmal edilebilir olup olmaması durumlarında kullanılabilecek kayıp veri yöntemleri özet olarak tanıtılmaktadır. Son olarak ilgili araştırmalar değerlendirilmektedir.

Kayıp Veri Örüntüsü

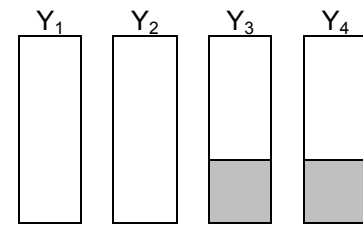
'Kayıp veri örüntüsü (missing data pattern)' ve 'kayıp veri mekanizması (missing data mechanisms)' ifadeleri, zaman zaman birbirini yerine kullanılmakla birlikte, esasında farklı anlamlara sahiptir. Kayıp veri örüntüsü, gözlenen ve kayıp verilerin bir veri setindeki konumlanma biçimlerini, kayıp veri mekanizması ise ölçülen değişkenler ile kayıp veri oluşma olasılığı arasındaki olası ilişkileri tanımlamaktadır. Basitçe kayıp veri örüntüsü, bir verideki kaybın konumunu ifade eder ve bu kaybın neden oluştuğunu açıklamaz. Kayıp veri mekanizması ise nedensel bir ilişkinin kanıtı olmamakla birlikte gözlenen ve kayıp veriler arasındaki çoğunlukla matematiksel ilişkileri ortaya koymaktadır (Enders, 2010).

Kayıp veri örüntülerini, kayıp verilerin veri setindeki konumlanma biçimlerine bağlı olarak, Şekil 1.1'de verildiği gibi sınıflandırmak mümkündür. Veri setinde yer alan değişkenlerden sadece birinde kayıp veri bulunması (univariate pattern), daha çok deneysel desenlerde karşılaşılan bir durumdur. Birim düzeyinde kayıp veri oluşması (unit nonresponse pattern), tarama araştırmalarında, örneklem yapısına ve özelliklerine bağlı olarak sıklıkla ortaya çıkabilmektedir. Monoton kayıp veriler (monotone pattern), özellikle boylamsal çalışmalarda, yanıtlayıcıların bazılarının adres değişikliği, ölüm gibi nedenlerle araştırma dışı kalmasına bağlı olarak oluşabilmektedir. Kayıp verilere yönelik en yaygın örüntü, genel örüntü (general pattern) olarak ifade edilmektedir. Genel örüntüde kayıp veriler, farklı değişkenlerde düzensiz bir

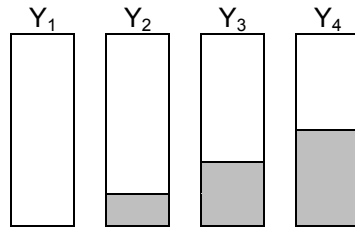
şekilde konumlanmış durumdadır ve bu durum ilk bakışta kayıp verilerin seçkisiz olduğu yönünde bir öngörü oluşturabilir. Fakat bu tür bir belirleme, ancak kayıp veri mekanizmasının incelenmesinden sonra yapılabilir. Çok fazla madde içeren veri toplama araçlarının kullanımı söz konusu olduğunda yanıtlayıcıların yanıtlama yükünün azaltılması amacıyla kullanımı işlevsel olan 'üç formlu anket (three-form questionnaire)' yöntemi gibi yöntemlerde kayıp veriler, planlı ve sistematik şekilde (planned pattern) oluşabilmektedir. Yapısal eşitlik modelleri gibi gizil değişkenler içeren analizlerde, gizil değişkene yönelik herhangi bir gözlem söz konusu değildir. Bu nedenle bu tür modellerde gizil değişken başlı başına bir kayıp veri örüntüsü (latent variable pattern) oluşturabilmektedir (Enders, 2010).



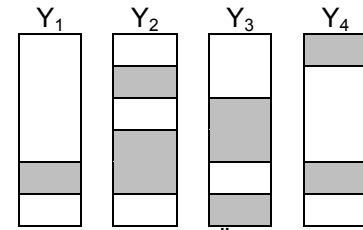
(A) Tek Değişkende Kayıp Veri
(Univariate Pattern)



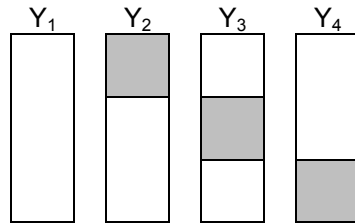
(B) Birim Düzeyinde Yanıtlanmama
(Unit Nonresponse Pattern)



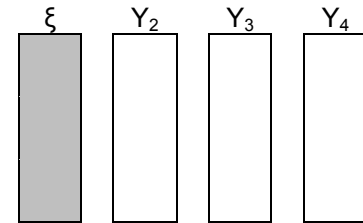
(C) Monoton Kayıp Veri
(Monotone Pattern)



(D) Genel Örüntü
(General Pattern)



(E) Planlı Örüntü
(Planned Missing Pattern)



(F) Gizil Değişken
(Latent Variable Pattern)

Kaynak: Enders, C.K. (2010), *Applied Missing Data Analysis*. New York: The Guilford Press, sy.4

Şekil 1.1. Tipik Kayıp Veri Örüntüleri

Kayıp Veri Mekanizmaları

Kayıp veri mekanizması, kayıp verilerin ve kayıp veri içeren değişkenlerin karakteristiğini belirlemede, bu belirlemeye bağlı olarak kullanılacak kayıp veri yöntem ve analizlerinin tercihinde anahtar rol oynamaktadır. Kayıp veri mekanizmaları üzerine, temelde benzer yaklaşımlarla, farklı araştırmacılar tarafından farklı sınıflandırmalar ve kavramlaştırmalar yapılmıştır. Bunlardan üçü aşağıda açıklanmaktadır.

Rubin'in Kayıp Veri Mekanizmalarına Yönelik Sınıflaması

Kayıp veri mekanizmaları üzerine bilinen ilk kavramlaştırmalar, Rubin (1976) tarafından yapılmıştır. Kayıp verilerin oluşmasına yol açan süreçlere yönelik üç durum tanımlanmıştır. Bu tanımlamalarda iki parametreye atıf yapılmaktadır. θ , veriye yönelik parametreyi, Φ ise kayıp veri sürecine yönelik parametreyi temsil etmektedir. Örneğin Φ , veri setinde yer alan kayıp veri göstergesinin koşullu dağılımına yönelik bir parametre olabilir. Buna göre;

1. Kayıp veri, 'seçkisiz kayıp-SK (missing at random)' olabilir. Bu durumda Φ parametresinin olası her bir değeri için, kayıp ve gözlenen veriler verildiğinde kayıp verilerin gözlenen örüntüsünün koşullu olasılığı, kayıp verilerin olası tüm değerleri için aynıdır.
2. Gözlenen veri, 'seçkisiz gözlenen-SG (observed at random)' olabilir. Bu durumda kayıp verilerin ve Φ parametresinin olası her bir değeri için, kayıp ve gözlenen veriler verildiğinde kayıp verilerin gözlenen örüntüsünün koşullu olasılığı, gözlenen verilerin olası tüm değerleri için aynıdır.
3. Φ parametresi, θ parametresinden farklı olabilir. Bu durumda Φ ve θ arasında, parametre uzayının sınırlılıkları ya da olasılık dağılımları aracılığıyla oluşan hiçbir bağ yoktur (Rubin, 1976, s.582).

θ 'ya yönelik örneklem dağılımı çıkarımları yapıldığında, kayıp veriye yol açan süreçlerin ihmal edilebilir olması, ancak kayıp veri SK koşulunu ve gözlenen veri SG koşulunu sağladığında mümkündür. Fakat çıkarımların sonuçlandırılması, kayıp verilerin gözlenen örüntülerine yönelik olarak genellikle koşulsaldır.

Allison'ın Kayıp Veri Mekanizmalarına Yönelik Sınıflaması

Allison (2002), Rubin'in sınıflamasının oldukça teknik olduğunu belirterek daha genel düzeyde bu mekanizmaları iki varsayımla ilişkilendirmektedir: (i) Tam Seçkisiz Kayıp – TSK (Missing Completely at Random) ve (ii) Seçkisiz Kayıp – SK (Missing at Random).

TSK oldukça güçlü bir varsayımdır. Bununla birlikte TSK varsayımının sağlanması ve gerekçelendirilebilmesi zaman ve çaba gerektirmektedir. TSK varsayımı, bir değişkendeki kayıp verilerin olasılığının, bu değişkenin kendi değeriyle ya da veri setindeki diğer herhangi bir değişkenin değeriyle ilişkili olmadığını ifade etmektedir. Bu varsayımın veri setindeki tüm değişkenler için anlamlı olması durumunda, kayıp veriler dışında kalan tam veri seti, orijinal gözlemler kümesinin basit seçkisiz örnekleme olarak kabul edilebilir.

TSK, aynı zamanda bir değişkende kayıp veri oluşma olasılığı ile diğer bir değişkende kayıp veri oluşma olasılığının ilişkilendirilebilmesine de olanak sağlamaktadır. Örneğin bir araştırma yaşları ile ilgili bilgi vermeyi reddeden yanıtlayıcıların gelirleri ile ilgili bilgi vermeme olasılıkları, gelirini beyan edenlere göre daha yüksektir. Bu durumda TSK varsayımı halen geçerlidir. TSK varsayımı, diğer yanıtlayıcılar gelirlerini beyan ederken ortalamanın altında bir yaşta olan yanıtlayıcıların gelirlerini beyan etmemeleri durumunda ihlal edilmiş olur. Bu tür bir ihlalin bulunup bulunmadığını belirlemek kolaydır. Örnekleme, gelirlerini beyan edenler ve etmeyenler olmak üzere ikiye ayrılır. Yaş ortalamaları arasında manidar bir fark oluşup oluşmadığı test edilir. Gelirini beyan eden ve etmeyenler arasında yaş değişkenine göre anlamlı fark bulunmaması, kayıp veri içeren veri seti ile kayıp veri içermeyen veri seti arasında, gözlenen değişkenler açısından herhangi bir sistematik farklılığın olmadığını gösterir. Bu durumda verilerin 'seçkisiz gözlenen (observed at random)' olduğu ileri sürülür. Diğer taraftan bu tür bir teste bağlı olarak TSK varsayımının karşılandığı ileri sürülemez. Herhangi bir değişkende kayıp veri oluşma olasılığı ile bu değişkenin değerleri arasında herhangi bir ilişki olmadığının da kanıtlanması gerekmektedir.

SK, TSK'ye göre daha zayıf bir varsayımdır. SK varsayımı, bir değişkendeki kayıp verilerin olasılığının, analizdeki diğer değişkenler kontrol altına alındığında, bu değişkenin kendi değeri ile ilişkisiz olduğunu ifade

etmektedir. X, daima gözlenen değişken ve Y, kayıp veri içeren değişken olmak üzere SK varsayımı, daha formel bir gösterimle ifade edilebilir:

$$P(Y_{\text{kayıp}}/Y, X) = P(Y_{\text{kayıp}}/X)$$

Bu gösterim, Y ve X'in verilmesi durumunda Y'deki kayıp veri koşullu olasılığı ile sadece X'in verilmesi durumunda Y'deki kayıp veri koşullu olasılığının eşit olduğu anlamına gelmektedir. Örneğin, gelire yönelik kayıp veri olasılığının evlilik durumuna bağlı olması ve aynı zamanda evlilik durumu değişkeninin her bir kategorisinde, gelir kaybı olasılığının gelirle ilişkisi olmaması, SK varsayımının anlamlı olduğunu gösterir. Olumsuz bir örnek olarak, diğer gözlenen değişkenler kontrol altına alındığında, özel bir değişkene yönelik kayıp veri oluşmasına neden olan bireylerin, bu değişkene yönelik verilerin elde edilebildiği bireylere göre, bu değişkene yönelik daha düşük ya da daha yüksek düzeyde veri değerlerine sahip olması durumunda SK varsayımı ihlal edilmiş olur.

SK varsayımının manidarlığının test edilmesinin mümkün olmadığı açıktır. Doğal olarak kayıp verilerin değeri bilinmemektedir. Bu nedenle herhangi bir sistematik farklılık bulunup bulunmadığını görmek amacıyla kayıp veri içeren ve içermeyen değerler karşılaştırılamamaktadır.

Enders'in Kayıp Veri Mekanizmalarına Yönelik Sınıflaması

Enders (2010), kayıp veri mekanizmalarını, Rubin (1976)'ın temel yaklaşımını esas almakla birlikte Allison (2002)'dan farklı olarak üçüncü bir varsayımla da ilişkilendirmektedir. Buna göre kayıp veri mekanizmalarını tanımlayan üç varsayım bulunmaktadır: (i) Seçkisiz Kayıp – SK (Missing at Random), (ii) Tam Seçkisiz Kayıp – TSK (Missing Completely at Random) ve (iii) Seçkisiz Olmayan Kayıp – SOK (Not Missing at Random).

Bir Y değişkenindeki kayıp veri oluşma olasılığının, bu değişkenin kendisi dışında analizde yer alan diğer bazı değişkenler ile ilişkili olması durumunda, söz konusu kayıp verilerin seçkisiz olduğu ileri sürülebilir. Diğer bir deyişle kayıp verilerin seçkisizliği, diğer değişkenler kontrol altına

alındığında, bir Y değişkeninde kayıp veri oluşma olasılığının bu değişkenin değerleri ile ilişkili olmaması durumudur. Esas itibarıyla SK varsayımı, ölçülen bazı değişkenlerle kayıp veri oluşma olasılığı arasındaki sistematik ilişkiyi ifade etmektedir. Örneğin performans puanlarında, belli bir zekâ puanının altında kayıp verilerin oluşması, bu tür bir sistematik ilişkinin varlığını ortaya koyar. Bu durumda performans değişkenindeki kayıp verilerin, SK varsayımını karşıladığı ileri sürülebilir.

SK mekanizmasının temel sınırlılığı, bir Y değişkenindeki kayıp veri oluşma olasılığının Y değişkeninin kendisi dışında diğer bazı değişkenlerin bir fonksiyonu olduğunun doğrulanabilir olmamasıdır. Bu tür bir doğrulama, kayıp veri içeren bir değişkende bu kayıpların oluşma olasılığının, değişkenin kendi değerleri ile ilişkisiz olmasını gerektirir. Kayıp değerler bilinmediğinden dolayı, söz konusu bu ilişkisizliğin kanıtlanması da mümkün değildir. Örneğin okuma becerilerine yönelik olarak uluslararası ölçekte yapılan bir testte, A ülkesine yönelik okuma becerisi puanlarının diğer ülkelere göre daha fazla miktarda kayıp veri içermesi, kayıp verilerin SK varsayımını karşıladığı yönünde bir öngörü oluşturur. Ülke değişkeni ile okuma becerisi puanlarında oluşan kayıplar arasında sistematik bir ilişkinin varlığı ileri sürülebilir. Sonrasında okuma becerisi puanlarında kayıp verilerin oluşma olasılığının, bu puanların bizzat kendisi ile ilişkili olmadığının kanıtlanması gerekir. Kayıp verilerin gerçek değerleri bilinmeksizin, bu tür bir belirlemenin mümkün olmadığı da açıktır.

TSK varsayımının formel tanımı, bir Y değişkeninde kayıp veri oluşma olasılığının, Y değişkenin kendi değeri ve diğer değişkenlerin değerleri arasında manidar bir ilişkinin bulunmaması durumunu ifade etmektedir. Diğer bir deyişle TSK, bir değişkendeki kayıp verilerin seçkisiz oluşması durumudur. TSK varsayımının karşılanması durumunda, bir değişkendeki gözlenen verilerin gerçek değerlerin seçkisiz bir örnekleme olduğu ileri sürülebilir. Örneğin bir işyerinde çalışanların iş performanslarının her ay ölçülmesi ve altı ay sonunda elde edilen verilerin değerlendirilmesi sürecinde, bazı çalışanlara yönelik bazı ölçmelerin, hastalık, işten ayrılma gibi nedenlerle yapılamaması ve buna bağlı olarak kayıp verilerin oluşması olasıdır. Bu durumda söz konusu kayıp verilerin seçkisiz oluştuğu dolayısıyla kayıp veri mekanizmasının TSK varsayımını karşıladığı ileri sürülebilir.

SK varsayımının aksine TSK varsayımının, bir değişkende kayıp veri bulunup bulunmamasına bağlı olarak tanımlanacak bir gösterge değişken yardımıyla istatistiksel olarak test edilmesi mümkündür. TSK varsayımının sağlanması, gözlem yapılamayan gözlem birimleri ile gözlem yapılan gözlem birimlerinin aynı ya da benzer karakteristiğe sahip olmasını gerektirir. Bu tür istatistiksel bir inceleme, t-testi, ANOVA ya da basit regresyonel desenler kullanılarak yapılabilmektedir. Bu amaçla geliştirilmiş 'Little's MCAR test' gibi bazı yöntemler de bulunmaktadır.

Bir Y değişkenindeki kayıp veri oluşma olasılığı, bu değişkenin kendi değerleri ile ilişkili ise kayıp veri mekanizması seçkisizlik özelliğini sağlamaz. Bu durumda kayıp veri mekanizması SOK varsayımına karşılık gelir. Örneğin bir başarı testinde, test puanı belli bir değerin altında olan bireylerin puanlarının silinerek ihmal edilmesi, kayıp verilerin sistematik bir şekilde oluşmasına yol açar. Test puanlarında kayıp verilerin oluşma olasılığı, bireyin test performansı ile ilişkilidir. Bu durum seçkisizlik özelliğinin ihlaline işaret eder ve kayıp veri mekanizması SOK olarak tanımlanabilir.

SOK varsayımının istatistiksel olarak test edilebilmesi için kayıp verilerin gerçek değerlerinin bilinmesi gerektiği açıktır. Bu nedenle SK varsayımında olduğu gibi SOK varsayımının da istatistiksel açıdan doğrulanması mümkün değildir.

Rubin (1987), Allison (2002) ve Enders (2010), kayıp veri mekanizmasının 'ihmal edilebilir (ignorable)' ya da 'ihmal edilebilir olmayan (nonignorable)' olduğu kararının verilebilmesine yönelik bazı ilişkiler ve bu ilişkilere bağlı işevuruk varsayımlar tanımlamışlardır. Buna göre kayıp veriler, TSK varsayımının karşılanması durumunda ihmal edilebilirdir. SK ya da SG varsayımının karşılanması, kayıp verilerin ihmal edilebilir olduğuna yönelik yeterli bir kanıt sağlamamaktadır. Ayrıca kayıp veri sürecine yönelik parametrelerin kestirilen parametrelerle ilişkisiz olması da gerekmektedir. TSK, SK ya da SG varsayımlarının karşılanmaması durumunda ise kayıp veri mekanizması, doğrudan kayıp verinin kendi değeriyle ilişkilidir ve ihmal edilebilir değildir.

Kayıp Veri Yöntemleri

Kayıp verilerin ihmal edilebilir olup olmaması durumlarına göre kullanılabilecek kayıp veri yöntemlerini, genel olarak iki grupta sınıflandırmak mümkündür: i) Silmeye ve basit atamaya dayalı yöntemler, ii) Olasılıklı ve ötelemeli veri atama yöntemleri. Olasılıklı ve ötelemeli yöntemler de kendi içerisinde iki grupta ele alınmaktadır: i) En çok olabilirlik yaklaşımına dayalı yöntemler ve ii) çoklu atama yaklaşımına dayalı yöntemler. Bu sınıflama aynı zamanda kayıp veri yönteminin tarihsel gelişim aşamalarını da temsil etmektedir.

Silme ve Basit Atamaya Dayalı Kayıp Veri Yöntemleri

Kayıp verilere yönelik olarak ilk akla gelen ve alışıldık yöntemler, genel olarak kayıp verilerin analiz dışı bırakılması ya da kayıp veriler yerine basit veri atama yöntemlerini içermektedir. Kullanılacak iyi bir yöntemin, (i) kayıp verilerden kaynaklanabilecek yanlılığı minimize etmesi, (ii) elverişli bilgilerin kullanımını maksimize etmesi ve (iii) belirsizliklere yönelik standart hata, güven aralığı ve olasılık değerlerine yönelik titiz kestirimler sağlaması beklenir (Allison, 2009).

Kayıp verilere yönelik silme ya da basit atamaya dayalı bilindik yöntemler, istatistik yazılımlarında yaygın bir şekilde kullanılmaktadır. Bununla birlikte bu yöntemlerin özellikle kayıp verilerin sınırlı miktarda olduğu durumlar dışında kullanılmaması önerilmektedir. Ayrıca daha önceki bölümlerde ele alındığı şekliyle herhangi bir kayıp veri yöntemi, kayıp verilere yol açan mekanizmalarla yakından ilişkilidir. Silme ve basit atamaya dayalı kayıp veri yöntemlerinin ise genellikle TSK varsayımı altında kullanılabilir olduğu vurgulanmaktadır (Little ve Rubin, 1987).

Aşağıda, kayıp verilerle başa çıkmada silme ve basit atamaya dayalı yöntemlerden ‘dizin silme’, ‘kısmi silme’, ‘kukla değişken düzeltmesi’ ve yaygın olarak kullanılan ‘basit atama’ yöntemleri, Rubin (1987), Little ve Rubin (1987), Allison (2002, 2009) ve Enders (2010)’ın yaklaşım ve açıklamaları dikkate alınarak özetlenmektedir.

Dizin Silme (DS) Yöntemi

Kayıp veriler için kullanılan en bilindik çözüm; herhangi bir değişkene yönelik kayıp verinin analizden çıkarılmasıdır. Böylece kayıp veri içermeyen, eksiksiz bir veri seti elde edilir ve artık bilindik istatistiksel analizlerin herhangi biri kolaylıkla uygulanabilir. Kayıp verilere yönelik bu yaygın yaklaşım 'dizin silme-DS (listwise deletion)', 'hücre silme-HS (casewise deletion)' ya da 'tam hücre analizi-THA (complete case analysis-)' olarak tanımlanmaktadır.

Kullanımı oldukça kolay olan DS, ilk bakışta pratik ve basit bir çözüm gibi durmaktadır. Bununla birlikte DS, temelde bazı verilerin ihmal edilmesine yönelik bir yaklaşımdır. Dolayısıyla bir veri kaybı söz konusudur. İhmal edilen verilerin niteliğinin ve miktarının, ulaşılan bulguların niteliğini etkileyeceği açıktır. Veri kaybı, standart hatanın artmasına, güven aralığının genişlemesine ve hipotez testinin gücünün azalmasına yol açmaktadır.

DS gibi kayıp verilerin ihmaline yönelik yöntemler, hedef evrenin sınırlı bir örnekleme yönelik çıkarımlar için daha kullanılabilir olmakla birlikte özellikle hedef evrene yönelik çıkarımlar için uygun değildir. Allison (2002), DS'nin söz konusu bu dezavantajını, 20 değişken içeren bir çoklu regresyon modeli kestirimi örneği ile açıklamaktadır. Örnekte veriler 1000 kişilik bir gruptan elde edilmiştir. Değişkenlerin her biri %5 kayıp veri içermektedir. Bir değişkendeki kayıp veri oluşma olasılığı, diğer bir değişkendeki kayıp veri oluşma olasılığından bağımsızdır. DS yönteminin kullanımı durumunda eksiksiz veri seti, her bir değişkene yönelik ancak 360 veriden oluşmaktadır, geriye kalan 640 veri ise analiz dışı bırakılmaktadır.

DS yönteminin uygulanabilirliği, kayıp veri mekanizmasının temel varsayımlarına göre değişmektedir. TSK, temel olarak seçkisiz örneklemin varlığını ortaya koymaktadır. Basit seçkisiz örneklemin yanlılığa yol açmayacağı bilinmektedir. Dolayısıyla verilerin TSK varsayımını karşılaması durumunda DS, parametre kestirimine yönelik herhangi bir yanlılığa yol açmamaktadır. Diğer taraftan TSK varsayımının karşılanmaması fakat SK varsayımının karşılanması durumunda ise DS yönteminin yanlılık üretmesi olasıdır. Örneğin, bir araştırmada kadınların %85'inin ve erkeklerin %60'ının gelirlerini beyan etmesi, TSK varsayımının ihlal edildiğini gösterir. Aynı zamanda kadın ve erkeklere yönelik verilerde oluşan gelir kayıplarının gelir

değişkenine bağlı olmaması, SK varsayımının karşılandığını gösterir. Bu tür bir örnekte DS yönteminin uygulanması durumunda, ortalama gelirlerde yanlılık oluşacağı açıktır.

DS yönteminin söz konusu dezavantajlarına rağmen avantajları da bulunmaktadır. Örneğin DS'nin TSK ve SK varsayımlara karşı oldukça kararlı ve sağlam olduğu ifade edilmektedir. Bu durum özellikle regresyon analizlerinde önemli bir üstünlük sağlamaktadır. Bir diğer örnek, DS yönteminin, gerçek standart hata kestirimindeki titizliğidir. Bu ve benzeri olumlu yanları nedeniyle DS, geleneksel kayıp veri yöntemleri arasında 'dürüst (honest)' bir yöntem olarak değerlendirilmektedir.

Kısmi Silme (KS) Yöntemi

DS, özel bir değişkenin değerlerinin, diğer bir değişendeki kayıp veriye bağlı olarak analiz dışı bırakılmasından dolayı, ortalamaların ve marjinal dağılım fonksiyonlarının kestirimini gerektiren tek değişkenli analizlerde kullanışlı görülmemektedir. Bu durumda 'kısmi silme-KS (pairwise deletion)' yöntemi, alternatif bir yöntem olarak ortaya çıkmaktadır. 'Erişilebilir hücre analizi-EHA (available case analysis)' olarak da bilinen KS, alternatif bir yöntem olarak özellikle doğrusal regresyon analizi, faktör analizi ve birçok yapısal eşitlik modelinde kullanılabilir. Bu ve benzeri doğrusal modellemeye dayalı istatistiksel analizler, esas olarak örneklem ortalaması, standart sapma, korelasyon ve kovaryans matrisleri gibi istatistiklerin kestirimini içermektedir. Bu özet istatistiklerin kestirimi, ilgilenilen parametrelerin kestiriminde öncü rol oynamaktadır. KS yöntemi, kestirim sürecinde, erişilebilir olan tüm verilerin kullanımına dayanmaktadır.

TSK varsayımının karşılanması durumunda KS yönteminin kararlı ve özellikle geniş örneklerde yansız parametre kestirimleri üretebileceği ifade edilmektedir. TSK varsayımı altında KS yöntemi, DS yöntemine göre daha düşük düzeyde örneklem çeşitliliği içerir. Dolayısıyla KS'nin DS'ye göre daha iyi kestirimler üretmesi beklenebilir. Diğer taraftan verilerin SK varsayımını karşılayıp KS ya da TSK varsayımlarının karşılamaması durumunda ise kestirimler ciddi yanlılık içerebilmektedir.

Kestirim sürecinde daha fazla veriyi dikkate almasına bağlı olarak KS'nin, DS'ye göre daha iyi bir yöntem olduğu ileri sürülebilir. Bununla birlikte KS'nin en önemli problemi, örneklem büyüklüğüne fazlasıyla duyarlı olmasıdır. Kayıp veri örüntüsüne bağlı olarak erişilebilir örneklem büyüklüğünün yeterli düzeyde olmaması, kovaryans matrisinin pozitif tanımlı olmasını engellemektedir. Diğer bir deyişle kovaryans ve korelasyon matrisleri hesaplanamayabilmektedir. Bu sorun istatistiksel analiz yazılımlarında KS'nin kullanılması durumunda da kendisini göstermektedir.

Kukla Değişken Düzeltmesi (KDD)

'Kukla değişken düzeltmesi-KKD (dummy variable adjustment)' ya da 'kayıp gösterge yöntemi-KGY (missing indicator method)', sıklıkla regresyon analizine yönelik özel bir kayıp veri yöntemi olarak Cohen ve Cohen'in 1985 yılındaki çalışmalarına dayandırılmaktadır. KDD, basitçe kayıp veri içeren yordayıcı değişkenlerin tamamında, verinin kayıp veri olup olmamasına bağlı olarak bir kukla değişkeni oluşturulması ve bu değişkenin de bir yordayıcı olarak analize dâhil edilmesi fikrine dayanmaktadır. Buna göre bir X değişkenine yönelik D kukla değişkeninde, X değişkenindeki kayıp veriler 1, diğer veriler 0 olarak kodlanır. Ayrıca bir X* değişkeni tanımlanır;

$$X^* = \begin{cases} X & ; \text{veri kayıp veri değil ise} \\ c & ; \text{veri kayıp veri ise} \end{cases}$$

X* değişkeninde c, genellikle erişilen veri değerlerinin ortalaması kullanılmakla birlikte, herhangi bir katsayı olabilir. c katsayısı sadece D değişkeni ile ilişkilidir. X*, herhangi bir c katsayısı için değişmez kalır. Dolayısıyla c katsayısının seçiminde herhangi bir sınırlılık söz konusu değildir. Yapılan tanımlamalar ile bir Y bağımlı değişkeni üzerinde X* ve D değişkenlerini içeren model üzerinde regresyon analizi süreci yürütülür. Benzer şekilde tanımlanacak diğer bağımsız değişkenlerin de modele dâhil edilmesi mümkündür.

KDD yönteminin en önemli üstünlüğü, kayıp verilere yönelik erişilebilir tüm bilgileri kullanmasıdır. Bununla birlikte KDD'nin, TSK varsayımı karşılanırsa bile regresyon katsayılarının kestiriminde yanlılığa yol açtığı da ifade edilmektedir.

KDD'ye benzer bir yöntem, yordayıcı değişkenin kategorik olması durumunda regresyon analizine yönelik olarak geliştirilmiştir. Bu yöntemde de kategorik yordayıcı değişkeninde yer alan kayıp verilere yönelik bir kukla değişkeni tanımlanarak analiz süreci yürütülmektedir. Fakat bu yöntemin, titiz bir standart hata kestirimi sağlamakla birlikte beklenmedik şekilde yanlılık ürettiği de belirtilmektedir.

Basit Atama (BA) Yöntemleri

'Basit atama-BA (simple imputation)', her bir kayıp veriye kabul edilebilir bir tahmini değerin atanması ve veri setinde kayıp veri yokmuş gibi analizlere devam edilmesine dayalı bir takım yöntemleri ifade etmektedir.

En basit ve bilindik veri atama yöntemi, 'ortalama atama-OA (mean substitution)' yöntemidir. Ortalama atama, 'seri ortalaması atanması-SO (serial mean imputation)' şeklinde dikey olarak ya da 'gözlem birimleri ortalaması-BO (unit's mean imputation)' şeklinde yatay olarak yapılabilmektedir. Ayrıca kayıp veriler yerine değer atamada, 'yakın noktalar ortalama ataması-YNÖ (mean of nearby points)' ve 'yakın noktalar medyan ataması-YNM (median of nearby points)' gibi bazı değerlerin ortalamalarını dikkate alan yöntemler de bulunmaktadır. Ancak bu yöntemin varyans-kovaryans kestiriminde yanlılık ürettiği, artık iyi bilinen bir durumdur.

Kayıp veriler yerine veri atanmasında çoklu regresyon kullanımının ortalama atamaya dayalı yöntemlere göre daha iyi bir yöntem olduğu belirtilmektedir. 'Koşullu ortalama atama-KOA (conditional mean imputation)' ya da 'Buck yöntemi' olarak da bilinen bu yöntemde, regresyon modelinde yer alan bağımsız değişkenlerden biri kayıp veriler içermektedir. Bu değişken, diğer bağımsız değişkenlerle birlikte regresyon analizine tabi tutulur. Kestirilen denklem kullanılarak kayıp veriler yerine atama yapılır. Bu yöntem, kayıp veri içeren bağımsız değişken sayısı arttıkça daha

zorlaşmaktadır. Veri atama sadece bağımsız değişkenlere bağlı olmak üzere, TSK varsayımı karşılanıyorsa, en küçük kareler katsayıları tutarlı ve kararlı olmaktadır. Bu durumda özellikle geniş örneklemelerde yansız kestirimler elde edilebilmektedir. Diğer taraftan neredeyse tüm veri atama yöntemlerinin temel bir problemi olarak, bazı parametrelerin kestiriminde yanlılık oluşabilmektedir. Veri atamasından sonra kayıp veri yokmuş gibi analizlere devam edilmesi, standart hatanın daha düşük kestirimine ve buna bağlı olarak test istatistiklerinin ise daha yüksek kestirimine yol açmaktadır.

Basit atama yöntemleri, veri atama sürecinde kayıp verilere yönelik belirsizlik olması durumunda, herhangi bir düzeltme sağlayamamaktadır. Bu nedenle basit veri atama yöntemlerinin 'dürüst olmayan (non honest)' yöntemler olduğu belirtilmektedir.

Geleneksel yöntemlerin eksiklikleri ve dezavantajları, kayıp verilere yönelik yeni yaklaşımlar üzerinde çalışılması ihtiyacını da beraberinde getirmiştir. Alternatif olarak geliştirilen pek çok yöntem arasında, 1990 başlarında teorik alt yapısı şekillendirilen ve 1990 sonlarında uygulama boyutuyla olgunlaştırılan 'en çok olabilirlik-EÇO (maximum likelihood)' ve 'çoklu veri atama-ÇVA (multiple imputation)' yaklaşımlarının öne çıktığı ve giderek daha yaygın bir şekilde kullanıldığı ifade edilmektedir.

En Çok Olabilirlik Yöntemleri

Aşağıda, kayıp verilerle başa çıkmada en çok olabilirlik yaklaşımı, bu yaklaşıma dayalı yöntemlerden 'beklenti-maksimizasyon algoritması' ve 'doğrudan en çok olabilirlik' yöntemleri, Dempster, Laird ve Rubin (1977), Rubin (1987), Little ve Rubin (1987), Allison (2002, 2009)'ın çalışmaları dikkate alınarak özetlenmektedir.

'En çok olabilirlik-EÇO (maximum likelihood)', bir yöntem olmaktan çok parametre kestiriminde olasılık temelli bir yaklaşımdır. Kayıp veri yöntemi olarak, bu yaklaşım temelinde geliştirilmiş 'doğrudan maksimum olabilirlik-DEÇO (direct maximum likelihood)' ya da 'beklenti-maksimizasyon algoritması-BM (expectation-maximization algorithm)' gibi bazı yöntemler bulunmaktadır.

EÇO yaklaşımının temel ilkesi, kestirim için, gözlemlerin olasılığını maksimum yapacak verilerin seçilmesidir. EÇO yaklaşımında verilerin olasılığı, veriler ve bilinmeyen parametrelere bağlı bir formülle ifade edilir. θ parametresi ve Y değişkeni için $f(Y/\theta)$, olasılık ya da 'olasılık yoğunluk (probability density)' fonksiyonudur. Bu fonksiyon kestirilmek istenen θ parametresinin bazı değerleri verildiğinde Y 'nin tek bir değerine yönelik gözlemin olasılığını vermektedir. Genel bir varsayım olarak gözlemlerin bağımsız olması durumunda, örnekleme yönelik olabilirlik, bireysel gözlemlere yönelik tüm olabilirliklerin ürünüdür. Buna göre n sayıda gözlem içeren bir örnekleme yönelik olabilirlik formülü, gözlemlerin olasılık yoğunluk fonksiyonlarının çoklu çarpımı ile ifade edilir (Little ve Rubin, 1987, s.79-80; Allison 2002, s.13):

$$L(\theta) = \prod_{i=1}^n f(y_i / \theta)$$

Bu formülde $f(y/\theta)$ ile ifade edilen olasılık yoğunluk fonksiyonu, Y değişkeninin ve θ parametresinin özelliğine göre değişmektedir. Örneğin Y değişkeninin 1-0 kodlaması ile verilen iki kategorili bir değişken (dichotomous) olması ve θ parametresinin, $Y=1$ olasılığını ifade etmesi durumunda olabilirlik formülü aşağıdaki şekilde verilmektedir:

$$L(\theta) = \prod_{i=1}^n \theta^{y_i} (1 - \theta)^{1-y_i}$$

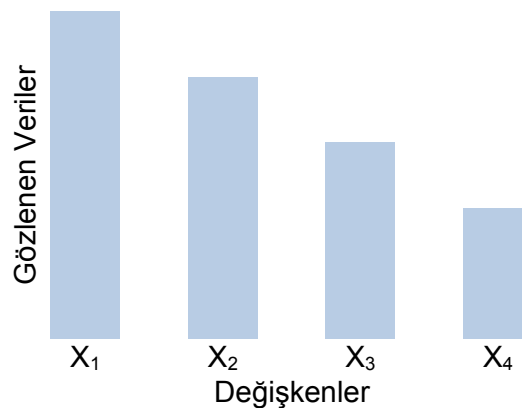
Olabilirlik fonksiyonu $L(\theta)$ oluşturulduktan sonra, θ parametresinin kestiriminde, olabilirliği artırmaya dayalı birçok farklı tekniğin kullanımı mümkündür.

EÇO'ya dayalı yöntemlerde, kestiricilerde bulunması beklenen bazı özellikler, varsayım olarak dikkate alınmaktadır. Buna göre EÇO'nun üç temel varsayımı söz konusudur: (i) Kararlılık (consistency), (ii) asimptotik etkililik (asymptotic efficiency) ve (iii) asimptotik normallik (asymptotic normality). Kararlılık, geniş örneklemlerde kestirimlerin, yaklaşık olarak yansız bir şekilde yapılabilirliğini ifade etmektedir. Etkililik, gerçek standart hatanın, en azından diğer kararlı kestiricilere yönelik standart hata kadar

küçük olmasının bir göstergesidir. Asimptotik etkililik ise etkililik özeliğinin yaklaşık olarak sağlanmasının yeterli olacağını ifade etmektedir. Asimptotik normallik ise tekrarlı örneklemelere bağlı olarak kestirimlerin ‘yaklaşık normal dağılıma’ sahip olması anlamına gelmektedir. EÇO yöntemlerinin temel varsayımlarının örneklem büyüklüğüne duyarlı olduğu görülmektedir. Dolayısıyla EÇO yöntemlerinin geniş örneklemelerde, dar örneklemelere göre daha iyi kestirimler üreteceği açıktır.

Temel varsayımların sağlanması durumunda EÇO, kayıp verilere yönelik oldukça iyi bir yöntem olarak görülmektedir. EÇO yöntemlerinin özellikle SK varsayımının sağlandığı fakat TSK varsayımının sağlanmadığı durumlarda daha iyi kestirimler ürettiği belirtilmektedir. Kayıp veri mekanizmasının ihmal edilebilir olması yani SK varsayımının sağlanması durumunda, bir olabilirlik fonksiyonu elde etmek mümkündür. Önceki açıklamalara bağlı olarak bu fonksiyon, kayıp verilerin her bir olası değerine yönelik olabilirliklerin bütününden oluşacaktır.

EÇO yöntemlerinin özellikle monotonik kayıp verilerde kullanımı oldukça kolaydır. Monotonik kayıp veriler özellikle boylamsal araştırmalarda ve panel araştırmalarında sıklıkla ortaya çıkabilmektedir. Sadece bir değişken kayıp veri içeriyorsa, kayıp veri örüntüsü doğal olarak monotonik olmaktadır. Diğer taraftan birden fazla değişkende kayıp veri olması durumunda, monotonik bir yapının oluşması, bu değişkenlerdeki kayıp verilerin belirli bir örüntüye sahip olmasını gerektirmektedir. Bu durumda değişkenlerdeki kayıp veri örüntüsünün Şekil 1.2’deki gibi olması gerekir.



Şekil 1.2. Monotonik Yapıda Kayıp Veri İçeren Değişkenler

Monotonik yapıda kayıp veri içeren değişkenlerin sayısına bağlı olarak olabilirlik fonksiyonu, farklı şekillerde oluşturulmaktadır. Olabilirlik fonksiyonu, sadece bir değişkenin kayıp veri içermesi durumunda, birden fazla değişkende monotonik yapıda kayıp veri bulunması durumuna göre daha sadedir ve daha kolay üretilebilir. Örneğin X ve Y gibi iki değişkene yönelik veri toplamak amacıyla n sayıda gözlem yapıldığı varsayalım. İlk m sayıda gözlemden, hem X hem de Y değişkenine yönelik veriler toplanmış olsun. Sonraki n-m sayıda gözlem ise sadece Y değişkenine yönelik olsun. Bu durumda sadece X değişkeninde kayıp veriler oluşur. X ve Y değişkenlerinin dağılımı verildiğinde, bilinmeyen θ parametreleri için tek bir gözleme yönelik tüm verileri içeren olasılık yoğunluk fonksiyonu, $f(x, y/\theta)$ şeklinde ifade edilir. Kayıp veriler içerdiğinden dolayı X değişkenindeki herhangi bir kayıp verinin olabilirliği, ancak Y değişkeninin ‘marjinal’ dağılımının bir ifadesi olur:

$$g(y/\theta) = \sum_x f(x, y/\theta)$$

X değişkeni sürekli ise marjinal dağılım fonksiyonunda toplam simgesi yerine integral simgesi getirilir. Bu durumda olabilirlik fonksiyonu, olasılık yoğunluk ve marjinal dağılım fonksiyonlarının bütününe ifade edecek şekilde oluşur:

$$L(\theta) = \prod_{i=1}^m f(x_i, y_i/\theta) \prod_{i=m+1}^n g(y_i/\theta)$$

Bir diğer örnek olarak Y değişkeninin yanı sıra monotonik yapıda kayıp veri içeren X_1, X_2, X_3 ve X_4 değişkenleri varsayalım. Bu durumda Y’nin marjinal dağılım fonksiyonunun yanı sıra, Y verildiğinde X’in koşullu olasılık dağılımını ifade eden $h(x/y)$ fonksiyonunun da dikkate alınması gerekmektedir. Dolayısıyla oluşturulacak olabilirlik fonksiyonu λ ve ϕ gibi iki farklı parametre içerecektir:

$$L(\lambda, \phi) = \prod_{i=1}^m h(x_i, y_i/\lambda) \prod_{i=1}^n g(y_i/\phi)$$

Verilen örneklerden de görülebileceği gibi ML yaklaşımının kayıp verilerde uygulanabilmesi için ilgili tüm değişkenlerin ‘katılımlı dağılımına (joint distribution)’ ya da ‘marginal dağılımına (marginal distribution)’ yönelik bir modele ve olabilirliği maksimize edecek bir sayısal yöntem ihtiyacı duyulmaktadır. Örneğin tüm değişkenlerin kategorik olması durumunda, ‘log-linear’ ya da ‘multinomial’ modeller kullanılabilir. Diğer taraftan tüm değişkenlerin sürekli olması durumunda, tipik olarak ‘çok değişkenli normallik (multivariate normality)’ varsayımı altında bir modelin kullanılması gerekir. Her bir değişkenin normal dağılıma sahip olmasını gerektiren bir durum olarak çok değişkenli normal dağılım, ortalamaları 0 olan hatalar içeren diğer değişkenlerin doğrusal bir fonksiyonla ifade edilebileceğini de gösteren, oldukça güçlü bir varsayımdır. Çok değişkenli normal dağılım altındaki bir modelde, BM ve DEÇO gibi yöntemlerle olabilirliğin maksimize edilmesi mümkün olabilmektedir.

Beklenti-Maksimizasyon Algoritması (BM)

‘Beklenti-maksimizasyon algoritması-BM (expectation-maximization algorithm)’, bazı verilerin kayıp olması durumunda EÇO kestirimleri elde edilmesinde kullanılan oldukça genel bir yöntem olarak tanımlanmaktadır. BM yöntemi iki aşamadan oluşan ‘ötleme (iterative)’ bir yöntemdir. ‘Beklenti (expectation)’ aşamasında, beklentinin hesaplanmasına yönelik parametre kestirimlerinin mevcut değerleri kullanılarak, kayıp veri içeren değişkenlerden, beklenen logaritmik olabilirlik değerleri elde edilir. ‘Maksimizasyon (maximization)’ aşamasında ise bir önceki aşamada elde edilen logaritmik olabilirlik değerleri, parametre kestirimlerine yönelik yeni değerler elde edilecek şekilde maksimize edilir. BM yönteminde bu iki aşama ötleme bir süreç içerisinde benzer kestirimler elde edilene kadar defalarca tekrar edilir.

Çok değişkenli normal dağılım altındaki bir modelde BM algoritması aracılığıyla kestirilen parametreler; ortalama, varyans ve kovaryanstır. Bu nedenle BM yöntemi, ‘ötleme doğrusal regresyon ataması (iterative linear

regresyon multiple)' olarak da ifade edilebilmektedir. BM yönteminin uygulanması, genel olarak beş aşama içermektedir:

1. Kayıp veriler dışında kalan verilerle ortalama, varyans ve kovaryanslar hesaplanır. Bunun için DS ya da KS gibi geleneksel yöntemlerden biri kullanılabilir.
2. Her bir kayıp veri örüntüsü için, gözlenen değişkenlere dayalı kayıp verilerin kestirimine yönelik regresyon denklemi yapılandırılır. Regresyon parametreleri, bir önceki aşamada kestirilen ortalama, varyans ve kovaryanslar üzerinden hesaplanır.
3. Regresyon denklemleri, değişkenler ve kayıp veri içeren hücrelerin değerlerinin kestiriminde kullanılır.
4. Gerçek veriler ve yerine atama yapılan verilerin tamamı, ortalama, varyans ve kovaryansların tekrar hesaplanmasında kullanılır.
5. İkinci aşamaya geri dönülür ve benzerlik elde edilene kadar işlemlere devam edilir (Allison, 2002, s.19-20; Allison, 2009, s.78).

BM algoritmalarının temel çıktısı, ortalama, varyans ve kovaryansların EÇO kestirimleridir. Dolayısıyla kestirilen bu parametre değerleri, aynı zamanda kestirim sürecinin bir parçası olarak kullanılmaktadır. Bu parametreler dışındaki diğer parametrelerin kestirim sürecine dâhil edilmesi durumunda, yanlış kestirimler elde edileceği belirtilmektedir. Kayıp veri yöntemi olarak BM algoritmalarının kullanımında en önemli dezavantajın ise BM algoritmasının standart hata kestirimi üretmemesi olduğu ifade edilmektedir.

Doğrudan En Çok Olabilirlik (DEÇO)

BM yönteminin bir alternatifi olarak 'doğrudan en çok olabilirlik-DEÇO (direct maximum likelihood)' yöntemi, standart hata kestirimine de imkân sağlamaktadır. DEÇO yönteminde temel girdi olarak kovaryans matrisi yerine ham veriler kullanılmaktadır. Bu nedenle bu yöntem 'ham en çok olabilirlik-HEÇO (raw ML)' ya da 'tam bilgi en çok olabilirlik-TEÇO (full information ML)' yöntemi olarak da bilinmektedir. DEÇO yönteminde, ilgilenilen doğrusal model özelleştirilmiştir. Bu modelde ortalamalar ve kovaryans matrisi, modeldeki parametrelerin bir fonksiyonu olarak ifade edilir. DEÇO yönteminde olabilirlik fonksiyonu, modeldeki parametrelere göre doğrudan maksimize edilir. Standart hataların, bilgi matrisinin negatif tersinin

hesaplanması gibi geleneksel EÇO yöntemleri ile belirlenmesi mümkündür. Bununla birlikte çok değişkenli normal dağılım altındaki bir modelin gerektirdiği temel varsayımlar DEÇO yönteminde de aynen geçerlidir.

EÇO yaklaşımı temelinde ortaya konulan yöntemler, geleneksel yöntemlere göre oldukça avantajlıdır. Bununla birlikte varsayımların sağlanmasındaki zorluk ve karmaşık hesaplama süreçleri bu yeni yöntemlerin kullanımında önemli sınırlılıklar oluşturmaktadır. Hâlihazırda EÇO yaklaşımı ancak doğrusal ve logaritmik doğrusal modellere uygulanabilir durumdadır. Örneğin lojistik regresyon, Poisson regresyon ya da Cox regresyon gibi çok değişkenli modellerde EÇO yaklaşımının uygulanabilirliğine yönelik teorik bir alt yapı henüz kurulamamıştır. Bir başka sınırlılık olarak EÇO yöntemlerinin oldukça karmaşık hesaplamalar içermesi, özel istatistiksel yazılımların kullanımını gerektirmektedir. Bu amaçla geliştirilmiş LEM ya da AMOS gibi özel yazılımların en önemli sorunu ise EÇO kestiriminde genellikle birbirinden farklı sonuçlar üretmeleridir.

Çoklu Veri Atama Yöntemleri

Aşağıda, kayıp verilerle başa çıkmada çoklu veri atama yaklaşımına dayalı yöntemler, Rubin (1987), Schaffer (1997) ve Allison (2002, 2009)'ın çalışmaları dikkate alınarak özetlenmektedir.

'Çoklu veri atama-ÇVA (multiple imputation)', her bir kayıp ya da eksik veri yerine, olasılıkların bir dağılımını yansıtan, kabul edilebilir iki ya da daha fazla verinin atanmasını öngören bir yaklaşımdır. ÇVA yaklaşımı temelinde geliştirilmiş 'tanımlayıcı atama-TA (deterministic imputation)', 'seçkisiz atama-SA (random imputation)' yöntemleri ya da bu ve benzeri yöntemleri kullanan 'Markov Chain Monte Carlo (MCMC)' yöntemi gibi daha özel yöntemler bulunmaktadır.

Geleneksel bir kayıp veri yöntemi olarak BA yönteminde, her bir kayıp veri yerine 1 tane veri atanmaktadır. ÇVA yöntemlerinde ise kayıp veri yerine m sayıda veri ataması yapılmakta ve m sayıda tamamlanmış veri seti elde edilmektedir. Her bir veri seti, atanan veriler gerçek verilermiş gibi kabul edilerek, standart eksiksiz veri süreçlerine göre analiz edilmektedir. ÇVA

yöntemlerinde 2 ile 10 arası tamamlanmış veri setinin kullanımının mümkün olduğu belirtilmektedir. Ne çok ne az yani ılımlı düzeyde kayıp veri bulunması durumunda, SAS istatistiksel yazılımında öngörüldüğü gibi 5 veri setinin, etkili bir parametre kestirimi için yeterli olabileceği de ifade edilmektedir.

ÇVA yöntemlerinde verilerin atanmasında TA kullanılabileceği gibi SA da kullanılabilmektedir. SA, TA'nın ortaya çıkardığı yanlışlık sorunu çözmektedir. Fakat hala bir sorun bulunmaktadır. Atanan verilerin gerçek verilerymiş gibi kabul edilmesi, standart hata kestirimlerinin sonuçlarının çok düşük ve buna bağlı olarak test istatistiklerinin çok yüksek çıkmasına yol açmaktadır. Bu sorun SA kullanıldığında daha azdır. SA'ya bağlı olarak ilgilenilen parametrelerin kestirimi, her bir tamamlanmış veri setinde belirgin şekilde farklı olmaktadır. Bu çeşitlilik, standart hata düzeltmesinde kullanılmaktadır.

ÇVA yöntemlerinin temel varsayımları, EÇO yöntemleri ile aynıdır. Kararlılık, asimptotik normallik ve asimptotik etkililik, ÇVA yöntemlerinde de aranan özelliklerdir. EÇO yöntemlerinde olduğu gibi ÇVA yöntemlerinde de bu özellikler ancak SK varsayımının karşılanması ya da kayıp veri mekanizmasına yönelik doğru bir modelin kurulabilmesi durumunda söz konusu olabilmektedir.

Temel varsayımların yanı sıra ÇVA yöntemlerinin kullanılabilmesi için, yapılan atamaları bir araya getirecek bir model gereklidir. Karmaşık yapıların karmaşık modeller gerektireceği de açıktır. ÇVA yöntemleri, EÇO yöntemlerine göre model seçiminde daha az duyarlıdır. Çünkü basitçe, ÇVA yöntemlerinde model, sadece kayıp verilere atama yapılmasında kullanılmakta, EÇO yöntemlerinde olduğu gibi diğer parametrelerin kestiriminde kullanılmamaktadır.

EÇO kestirimlerinde olduğu gibi ÇVA yöntemleri için de en kullanışlı model, çok değişkenli normal modeldir. Bazı değişkenler normal dağılım varsayımını karşılamasa da çok değişkenli modelde ÇVA yöntemlerinin oldukça başarılı olduğu ifade edilmektedir. Çok değişkenli normal modelde ÇVA yönteminin kullanımında takip edilen süreçler BM yöntemine oldukça benzerdir. Aradaki fark ÇVA yöntemlerinde, kestirilen kayıp veri değerlerine, standart hataların da eklenmesidir. Bu ekleme, genellikle tanımlayıcı atama yönteminin kullanılması durumunda ortaya çıkan varyans kestirimi yanlışlığını

gidermeye yönelik bir düzeltme sağlamaktadır. Öncelikle kayıp veri içeren her bir değişken için bu değişkenin diğer değişkenler üzerindeki doğrusal regresyon kestirimi yapılır. Kestirilen regresyon denklemleri, kayıp veriler içeren hücrelere atanacak verilerin kestiriminde kullanılır. Son olarak, her bir kestirilen veriye, bu değişkenin artık normal dağılımından çıkarılan seçkisiz bir değer eklenir. Bu son aşama oldukça zordur ve genellikle özel istatistik yazılımlarının kullanımını gerektirir.

EÇO yöntemleri ile karşılaştırıldığında ÇVA yöntemlerinin iki önemli üstünlüğü bulunmaktadır. İlk olarak ÇVA yöntemleri her çeşit veri ve modele uygulanabilmektedir. İkinci olarak ÇVA yöntemlerinde, LES ve AMOS gibi özel yazılımlar yerine geleneksel istatistik yazılımlarının kullanılabilmesidir. Fakat EÇO yöntemlerinin bir sınırlılığı olarak belirtildiği gibi bu tür istatistik yazılımlarının kullanımındaki önemli bir sorun, bu yazılımların her birinin farklı sonuçlar üretmesidir. Bu farklılığı temel sebebi, ÇVA yöntemlerinde veri atamasının genellikle seçkisiz olarak yapılmasıdır.

ÇVA yöntemleri, BA yöntemi ile ortak iki temel üstünlüğe sahiptir: (i) Tam veri yöntemlerini kullanabilme yeteneği ve (ii) veri toplayıcılardan elde edilen tüm bilgileri bir araya getirebilme yeteneği. ÇVA, veri toplayıcılara, hangi değerlere atama yapılacağına yönelik belirsizliği ortaya çıkarmaları için kendi bilgilerini kullanma fırsatı vermektedir. Örneğin yanıtlanmama durumu, bu tür bir belirsizliğin kaynağı olabilmektedir. Örneklem dağılımına ya da bireylerin yanıtlanmama gerekçelerine dayalı olarak yanıtlanmama durumu ortaya çıkabilmekte bu durum ise bir belirsizlik durumu oluşturmaktadır. ÇVA yöntemlerinde bu tür belirsizliğe yönelik olarak, örneklem dağılımındaki çeşitliliği ya da yanıtlanmama gerekçelerine bağlı olarak ortaya çıkabilecek belirsizlikleri yansıtabilecek modeller üretilmektedir.

ÇVA yöntemlerinin BA'ya göre bazı üstünlükleri bulunmaktadır. İlk olarak, veri dağılımını yansıtacak şekilde, seçkisiz veri ataması yapılması, kestirimlerin etkililiğini artırmaktadır. İkinci olarak, birden fazla tamamlanmış veri setinin kullanımı, daha geçerli çıkarımlara olanak sağlamaktadır. Üçüncü olarak, tamamlanmış veri setlerinin tekrarlı kullanımı ile birden fazla model altında tekrarlı seçkisiz veri ataması, çeşitli modellere yönelik çıkarımların duyarlılığını sağlamaktadır.

ÇVA yöntemlerinin BA'ya göre bazı dezavantajları da bulunmaktadır. ÇVA yöntemlerinin kullanılabilmesi için daha fazla çalışmaya ihtiyaç vardır. ÇVA yöntemleri ile atanan veriler, daha fazla alan kaplamaktadır. ÇVA yöntemlerinde analizler daha zor ve karmaşıktır. Bu ve benzeri eksiklikler, çok önemli sorunlar olarak görülmemektedir. Bununla birlikte kayıp verilere yönelik bilgilerin çok dağınık olması durumunda ÇVA, istenen düzeyde anlamlı olmayacaktır. Bu durumda BA yönteminin kullanımı tercih edilmesi önerilmektedir.

İlgili Araştırmalar

Önceki bölümlerde verilen, Lord, Rubin, Little, Schaffer, Allison, Finch, DeMars, Linacre gibi alanın önde gelen uzmanlarının çalışmalarının yanı sıra, kayıp verilerle ilgili araştırmaların yaygınlaşmakta olduğu ve literatürde belirli bir birikime ulaştığı görülmektedir. Bu çalışmada, ilgili diğer araştırmaların, konu bütünlüğünü bozmamak amacıyla ayrı bir başlık altında bu bölümde özet olarak verilmesi uygun görülmüştür.

Yurtdışında Yapılan Araştırmalar

Elashoff ve Elashoff (1971), kayıp veri içeren olasılıklı modellerde oransal karşılaştırmaların yapılabilmesi için kullanılabilecek yöntemleri incelemişlerdir. Kuramsal düzeyde yürüttükleri araştırmalarında simülatif verilerden yararlanmışlardır. Bazı testlerin ve kestirimlerin, kayıp verilere fazlasıyla duyarlı olduklarını, bu tür durumlarda standart yöntemlerin yeterli olmadığını belirlemişlerdir. Araştırmada, kayıp verilere yol açan mekanizmanın karmaşık olması ve özellikle kayıp verilerin ihmal edilebilir olamaması durumunda, standart yöntemlerin yanlışlık ürettiği sonucuna varılmıştır.

Brown (1983), örneklem büyüklüğüne bağlı olarak, çok değişkenli analizlerde kayıp veri sorununu ele almıştır. Bu amaçla, DS, KS ve kısmi EÇO yöntemleri kullanılarak, tek faktörlü modelde faktör yükleri ve kovaryans

matrisi kestirimleri gerçekleştirilmiştir. Kullanılan yöntemler, faktör yapısı altında parametre kestirimi yeterliliği ölçüt alınarak karşılaştırılmıştır. Kuramsal düzeyde yürütülen araştırmada, konu oldukça teknik bir yaklaşımla ele alınmış ve gerçek ya da simülatif veriler kullanılmamıştır. Araştırma sonucunda, kullanılan kayıp veri yönteminin etkililiğinin kayıp veri örüntüsüne bağlı olduğu görülmüştür. Buna göre kayıp verilerden dolayı modelde bir dengesizliğin bulunduğu durumlarda, DS ve KS yöntemlerinin her ikisi de iyi performans göstermektedir. Dengelenmiş modelde ise DS yöntemi, KS yöntemine göre daha iyi işlemektedir. Kovaryans matrisleri ve örüntülerin ranjinin geniş olduğu durumlarda, basit regresyonel kestirimlerde, EÇO yöntemlerinin yüksek düzeyde etkili olduğu sonucuna ulaşılmıştır. EÇO dışında kalan diğer regresyon temelli kestirim süreçlerinin ise EÇO yöntemlerine göre daha az etkili olacağı ifade edilmiştir.

Raghunathan, Lepkowski, Hoewyk ve Solenberger (2001), kayıp veri atamasına yönelik çok değişkenli bir teknik geliştirmişlerdir. Bu amaçla, SK varsayımı altında kayıp veri içeren kompleks veri yapısında, bir veri atama süreci tanımlanmıştır. Önerilen teknik, kayıp veri atamada, bir regresyon modeli dizisinin kullanımını içermektedir. Bu amaçla doğrusal, lojistik, Poisson, logit ya da karma regresyon modellerinin kullanımı mümkündür. Söz konusu tekniğin geliştirilmesi sürecinde, simülatif veri setleri kullanılarak benzetim çalışması yapılmıştır. Oldukça teknik düzeyde yürütülen araştırma sonucunda, regresyon serisine dayalı çok değişkenli kayıp veri atama tekniği, ÇVA analizlerinde alternatif bir yöntem olarak sunulmuştur.

Acock (2005), birçok kayıp veri yönteminin güçlü ve zayıf yönlerini özetledikten sonra bu yöntemlerin bazılarını karşılaştırmak amacıyla gerçek veriler üzerinden farklı yöntemlere göre analizler gerçekleştirmiştir. Veri seti olarak, 818 eksiksiz gözlem içeren ve 2002 yılında gerçekleştirilen Genel Sosyal Araştırma verilerini kullanmıştır. Analizlerde çocuk sayısı, mutluluk düzeyi, gelir, yaş ve eğitim düzeyi değişkenlerini dikkate almıştır. Dikkate alınan bu değişkenler üzerinden, TSK varsayımı altında %50 kayıp veri içeren bir veri seti ve %40 seçkisiz kayıp veri içeren ikinci bir veri seti oluşturmuştur. Eksiksiz veri seti ile birlikte, diğer iki veri setini de kullanarak DS, KS, BA, OA, BM, TEÇO ve ÇVA yöntemlerine göre regresyon kestirimleri üretmiştir. Bu kestirimlerde SPSS, NORM, STATA, HLM ve Mplus

istatistiksel yazılımlarını kullanmıştır. Araştırmada; Mplus ve HLM yazılımları ile TEÇO yönteminin kullanımının, ÇVA yöntemleri kadar iyi sonuçlar ürettiği sonucuna ulaşmıştır. Bu yöntemlerin, geleneksel yöntemlere ve diğer yöntemlere göre daha avantajlı olduğu da ifade edilmektedir.

Peng ve arkadaşları (2007), eğitim araştırmalarında kayıp veri sorununun nasıl ele alındığını anlamak amacıyla bir araştırma yürütmüştür. Bu amaçla eğitim araştırmalarının temel ilgi alanlarını ve araştırma yöntemlerini yansıttığı varsayılan 11 eğitim dergisinde 1998-2002 yılları arasında yayımlanan nicel araştırma makalelerini incelenmiştir. Bu kapsamda 918 makale içerisinde yer alan toplam 1087 çalışma tespit edilmiştir. Her bir çalışma, kullanılan yöntemler, bulgular, tartışma bölümleri ve özet tabloları dikkate alınarak iki araştırmacı tarafından incelenmiştir. Bu incelemede toplam örneklem büyüklüğü, raporlanan test istatistiklerinin serbestlik derecesi ve kayıp verilere yönelik araştırmacının yaklaşımına yönelik bilgiler belirlenmeye çalışılmıştır. Araştırma sonucunda 1087 çalışmadan 305 (%28)'inde herhangi bir kayıp veri sorunundan bahsedilmediği, 587 (%54)'sinde kayıp veri sorununun çözümüne yönelik bazı yöntemlerin kullanıldığı ve 195 (%18)'inde ise kayıp verilerden bahsedilmekle birlikte kullanılan yöntem ya da kayıp veri sorunun bulunduğuna yönelik serbestlik derecesi, örneklem büyüklüğü gibi herhangi bir bilgiye yer verilmediği görülmüştür. Kayıp verilere yönelik kanıtlar içeren 597 çalışmadan 569 (%97)'unda kayıp veri sorununun nasıl ele alındığı ayrıntılı bir şekilde raporlanmıştır. Bu 569 çalışmadan 509 (%89.5)'unda DS, 43 (%7.6)'ünde KS ve 17 (%2.9)'sinde ise OA, RA ve BM yöntemleri kullanılmıştır. Araştırma sonuçları eğitim araştırmacılarının, özellikle deneysel desenlerde, temel kayıp veri yöntemlerini yeterli düzeyde uygulamadıklarını göstermektedir. Diğer taraftan söz konusu dergilerin de araştırmacıları, DS ve KS dışındaki yöntemleri kullanmaları yönünde teşvik etmediği de ifade edilmektedir. Peng ve arkadaşları tarafından yürütülen araştırma, eğitim araştırmalarında kayıp veri sorununun önemini ve boyutlarını ortaya koymaktadır. Temel sorun olarak, araştırmacıların geleneksel yöntemler dışına çıkamamaları, vurgu yapılan noktalardan biridir.

Eberle ve Toutenburg (1999), kayıp veri yöntemleri içeren bazı istatistik yazılımlarını, karşılaştırmalı olarak incelemiştir. Araştırmada 'MINITAB', 'SPSS', 'SYSTAT', 'SAS', 'STATISTICA', 'StatXact', 'Stata', 'LogXact', 'S-PLUS' ve 'JMP' yazılımları, belirlenen bazı ölçütler çerçevesinde, kuramsal olarak değerlendirilmiştir. Değerlendirme ölçütleri olarak, (i) kayıp verilerin kodlanma şekli ve kodlama esnekliği, (ii) analiz çıktılarında kayıp verilerin sunumu ve gösterimi, (iii) TSK varsayımının test edilebilir olup olmaması, (iv) kayıp verilere yönelik betimleyici istatistiklerin kestirilmesinde sunulan opsiyonlar, (v) kovaryans ve korelasyonların hesaplanmasında sunulan opsiyonlar, (vi) güven aralıklarının belirlenmesinde sunulan opsiyonlar, (vii) regresyon analizi, varyans analizi, kümeleme analizi, diskriminant analizi ve zaman serilerinde kayıp veri yöntemlerinin kullanımına yönelik olarak sunulan opsiyonlar, (viii) makroların programlanabilme olasılığı ve (ix) kılavuz ve çevrimiçi yardım hizmetlerinin niteliği dikkate alınmıştır. Araştırmada, kayıp veri sorununa yönelik olarak geliştirilen yöntemlerin kullanımının, yazılımdan yazılıma farklılık gösterdiği görülmüştür. İncelenen yazılımların hiçbiri, olası kayıp veri yöntemlerinin tamamını içermemektedir. En yaygın ve bilindik yazılımlar bile genellikle sınırlı sayıda yöntemin kullanımına olanak sağlamaktadır. Dolayısıyla kayıp veri yöntemlerinin kullanımına yönelik mükemmel bir yazılım bulunmamaktadır.

Yurtiçinde Yapılan Araştırmalar

Bal (2003) tarafından yürütülen bir araştırmada, çok gruplu veri setlerinde eksik gözlem sorunu ele alınmıştır. Öncelikle bazı kayıp veri yöntemlerinin karşılaştırılması amacıyla bir benzetim çalışması yürütülmüştür. Simülatif 50, 100, 200, 300, 400 ve 500 birimlik veri setleri, TSK varsayımı altında %5, %10, %15 ve %20 eksiltilerek kayıp veri içeren veri setleri elde edilmiştir. Bu veri setleri üzerinde farklı kayıp veri yöntemleri uygulanmış ve yapılan kestirimler tam veri setleri üzerinden yapılan kestirimlerle karşılaştırılmıştır. Benzetim çalışmasında DS, KS, OA, RA ve BM yöntemlerinin, düşük hacimli örneklerde tutarsız sonuçlar ürettiği görülmüştür. Genel olarak, tam veri setine en yakın kestirimlerin ise BM

algoritmasının kullanılması durumunda gerçekleştiği belirlenmiştir. BM algoritmasının özellikle 200 ve üzeri birime sahip örneklerde, %5 ile %20 arası kayıp veri oranı olması durumunda, diğer yöntemlerden daha iyi işlediği belirtilmektedir. Araştırmanın ikinci aşamasında, gerçek veriler üzerinden bir uygulama yapılmıştır. Bu uygulamada TSK yapıda kayıp veri içerdiği belirlenen bir veri seti üzerinde, herhangi bir yöntem kullanmaksızın yapılan kestirimler ile bir kayıp veri yöntemi kullanılarak elde edilen kestirimler karşılaştırılmıştır. Kayıp veri yöntemi olarak, benzetim çalışması sonuçları dikkate alınmış ve BM algoritması kullanılmıştır. Gerçek uygulama sonucunda kayıp veri içeren üç değişkenden ikisinde, kestirimler arasında manidar fark belirlenmiştir. Buna bağlı olarak, özellikle çok gruplu veri setlerinde, uygun bir kayıp veri yönteminin kullanılmasının önemi ve gerekliliği vurgulanmıştır.

Dural (2010) tarafından yürütülen bir araştırmada, kayıp veri yöntemlerinin çok göstergeli örtük gelişme modelleri üzerindeki etkisi, örneklem büyüklüğü ve kayıp veri türüne bağlı olarak incelenmiştir. Bir benzetim çalışması olarak simülatif veriler üzerinde yürütülen araştırmada, farklı sayıda veri içeren üç bağımsız değişken dikkate alınmıştır. Bu değişkenler kullanılarak, çok göstergeli bir örtük gelişme modeli tanımlanmıştır. Değişkenlerde yer alan veriler TSK ve SK varsayımları altında eksiltilerek, kayıp veri içeren model elde edilmiştir. Model testleri, DS, KS, RA, BM ve ÇVA yöntemleri kullanılarak gerçekleştirilmiştir. Bulgular, çok göstergeli örtük gelişme modellerinde KS yönteminin, yakınsama oranı, model uyumu ve yanlılık ölçütleri açısından, diğer yöntemlere göre daha iyi bir yöntem olduğunu göstermiştir. Buna karşılık standart hata yanlılığı ve testlerin gücünün düşmesi açısından, KS yönteminin uygun olmadığı görülmüştür. Dikkate alınan tüm değerlendirme ölçütleri açısından, çok göstergeli örtük gelişme modelleri için en uygun yöntemlerin, BM yöntemi ve ÇVA yöntemleri olduğu sonucuna ulaşılmıştır.

Çokluk ve Kayrı (2011), kayıp veri atama yöntemlerinin performansını, ölçek güvenilirliği ve geçerliği açısından incelemiştir. Seri ortalamaları ataması-SO (serial mean imputation), 'yakın noktalar ortalama ataması-YNÖ (mean of nearby points)', 'yakın noktalar medyan ataması-YNM (median of nearby points)', 'doğrusal interpolasyon-Dİ (linear interpolation)' ve 'doğrusal

eğilim noktası ataması-DE (linear trend at point)' yöntemlerinin dikkate alındığı araştırmada, 'kadercilik (fatalism) ölçeği' ile elde edilen veriler kullanılmıştır. Söz konusu ölçeğin uygulanması ile 200 yanıtlayıcıya yönelik bir tam veri seti elde edilmiştir. Veri setinde seçkisiz yöntemle %15-%20 arası ve %0-%50 arası eksiltmelerle, kayıp veri içeren veri setleri elde edilmiştir. Veri setleri üzerinde yürütülen analizlerin çıktılarına göre faktör yapıları, düzeltilmiş test-madde korelasyonları ve iç tutarlılık göstergesi olarak Cronbach α katsayılarına yönelik kestirimler karşılaştırılmıştır. Kayıp veri yöntemlerinin, orijinal veri setinde olduğu gibi, tek faktörlü yapıyı belirlediği görülmüştür. Fakat bu yöntemler, açıklanan varyans yüzdelerinde bir düşüşe yol açmıştır. Benzer sonuçlar faktör yüklerine yönelik öz değerlerde ve Cronbach α katsayılarında da kendini göstermiştir.

Kayıp veri sorununa ve yöntemlerine yönelik olarak ilgili literatürde yer alan, bu bölümde ya da önceki bölümlerde ele alınan araştırmaların büyük çoğunluğu, oldukça teknik çalışmalardır. Genellikle kayıp veri yöntemleri, özet olarak tartışılmakta ve örnek bir uygulama üzerinden karşılaştırılmaktadır. Kayıp veri yöntemlerinin kullanımına yönelik örnek uygulamaların genellikle simülatif veriler üzerinden yapıldığı görülmektedir. Beklenen bir sonuç olarak EÇO ve ÇVA yaklaşımları temelinde geliştirilen yöntemlerin, kayıp veri sorunu ile başa çıkmada daha işlevsel ve uygun olduğu belirtilmektedir. İki kategorili puanlanan maddelerden oluşan testlere yönelik araştırmaların ise daha seyrek olduğu görülmektedir.

Üzerinde durulan önemli bir diğer araştırma konusu, kayıp veri analizlerinde kullanılabilen istatistiksel yazılımların performansıdır. Kayıp veri analizleri içeren yazılımlar, temel yaklaşım ve yöntemlerde benzerlik göstermekle birlikte, genellikle farklı algoritmalar kullanmaktadır. Bu nedenle elde edilen kayıp veri kestirimleri, kullanılan yazılıma göre değişebilmektedir. Bu durumda kullanılan istatistiksel yazılımın performansı, üzerinde durulması gereken bir konu olarak ortaya çıkmaktadır.

BÖLÜM 3

YÖNTEM

Bu bölümde araştırmanın modeli, evreni, verilerin toplanması ve analizine yönelik bilgiler ve açıklamalara yer verilmektedir.

Araştırmanın Modeli

Bu çalışma, temel araştırma türünde, betimleyici bir araştırma olarak tasarlanmıştır. Kayıp veriler içeren bir testin psikometrik özellikleri ile kayıp veriler ve kayıp veri oluşma olasılığı arasındaki ilişkiler, hazır bir veri seti üzerinde yürütülen ikincil düzey analizlerle incelenmiştir.

Evren

Bu çalışmanın evreni, SBS 2011 Matematik Testi A Kitapçığını alan 527517 yanıtlayıcıdan oluşmaktadır. Kayıp verilere yönelik inceleme ve analizler, evrenden elde edilen veri seti üzerinde yürütülmüştür.

Verilerin Toplanması

Bu çalışmada veri toplama aracı olarak, SBS 2011 Matematik Testi A Kitapçığı kullanılmıştır. SBS Matematik Testi, matematik başarısına yönelik çoktan seçmeli ve dört seçenekli 20 maddeden oluşmaktadır. Test sonuçlarının elde edilmesinde yanıtlayıcı tepkileri, 1-0 şeklinde yapay süreksiz olarak puanlanmaktadır.

SBS (Seviye Belirleme Sınavı), Türkiye genelinde, Millî Eğitim Bakanlığı (MEB) tarafından, ilköğretim kademesinin 8. sınıfından mezun olma aşamasında olan öğrencilere yönelik olarak gerçekleştirilen ve her yıl düzenlenen bir başarı testidir. Test sonucunda, yanıtlayıcıların başarı

puanları, okul başarıları da dikkate alınarak bir üst eğitim kademesine geçişte, seçme ve yerleştirme amacıyla kullanılmaktadır.

SBS, Türkçe Testi, Matematik Testi, Fen Bilimleri Testi, Sosyal Bilimler Testi ve Yabancı Dil Testi olmak üzere beş alt testten oluşmaktadır. Bu testlere yönelik maddelerin yazımında, ilgili ders programında yer alan konu ve kazanımların dikkate alındığı belirtilmektedir. Her ne kadar test geliştirme süreçlerine uygunluğu açısından eleştirilebilecek olsa da SBS testlerinin geliştirilmesi ve uygulanması süreçlerinin kendi içerisinde bir mantığı olduğu anlaşılmaktadır¹.

SBS 2011 Matematik Testi A Kitapçığına yönelik veri seti Millî Eğitim Bakanlığının ilgili birimleri ile yapılan resmi yazışmalar sonucunda elde edilmiştir. Söz konusu ham veri seti, testi alan 527517 yanıtlayıcıya yönelik A-B-C-D şeklinde kodlanan yanıt örüntüleri ile cinsiyet, okul türü ve il değişkenlerini içermektedir. Cinsiyet değişkeni, 'erkek' ve 'kız' olmak üzere iki kategorilidir. Okul türü değişkeni, 'devlet' ve 'özel okul' olmak üzere iki kategorilidir. İl değişkeni ise Türkiye'nin 81 iline göre kategorik bir değişken olarak tanımlıdır.

SBS 2011 Matematik Testi A Kitapçığına yönelik ham veri setinde, bu çalışmada öngörülen analizler dikkate alınarak gerekli düzenlemeler yapılmıştır. Öncelikle testin cevap anahtarı dikkate alınarak, yanıt örüntüleri 1-0 örüntüsüne dönüştürülmüştür. Bu dönüştürmede doğru yanıtlar 1, yanlış yanıtlar 0 olarak puanlanmıştır. Kayıp veriler ise özel kodlama ile tanımlanmıştır. İkinci olarak ham veri setinde yer alan il değişkeni, Türkiye 'İstatistikî Bölge Birimleri Sınıflaması (İBBS) Düzey 1' dikkate alınarak, 12 bölge esasına göre yeniden kodlanmış ve 'bölge' değişkeni olarak veri setine yerleştirilmiştir. Üçüncü olarak 'toplam doğru yanıt sayıları', 'toplam yanlış yanıt sayıları' ve 'kayıp veri içeren madde sayıları' her bir gözlem birimi düzeyinde hesaplanmış, değişken olarak veri setine yerleştirilmiştir. Son olarak gözlem birimleri düzeyinde kayıp veri bulunup bulunmaması durumu, bir 'kukla değişken' olarak tanımlanmıştır. Böylelikle ham veri seti, 1-0 yanıt örüntülerinin yanı sıra 7 bağımsız değişken içerecek şekilde ileri analizlere uygun hale getirilmiştir.

¹ SBS ve Ortaöğretime Geçiş Sistemi hakkında ayrıntılı bilgi için Millî Eğitim Bakanlığının ilgili internet sitesine; http://oges.meb.gov.tr/sbs_istat.htm adresinden erişilebilmektedir.

Verilerin Analizi

SBS 2011 Matematik Testi A Kitapçığı verilerine yönelik analizlerde ilk olarak, kayıp veri örüntüsü incelenmiştir. Kayıp verilerin veri setindeki konumlanma biçiminin betimlenmesi ve tanımlanması amaçlanmaktadır. Bu incelemede, doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının dağılımları dikkate alınmıştır. Her bir dağılıma yönelik olarak sıklık, ortalama, standart sapma, varyans, basıklık ve çarpıklık katsayıları kestirimleri gerçekleştirilmiştir. Dağılımların biçimi ve normal dağılımdan sapma gösterip göstermediği, betimsel olarak değerlendirilmiştir. Ayrıca en sık görülen ve en az sıklıkta görülen yanıt örüntüleri betimsel olarak incelenmiştir.

İkinci aşamada, kayıp veri mekanizmasına yönelik inceleme ve analizler gerçekleştirilmiştir. Kayıp veri mekanizmasının, TSK, SK ve SOK olarak önceki bölümlerde açıklanan varsayımlardan hangisini karşıladığının belirlenmesi amaçlanmaktadır. Bu mekanizmalar, özellikle iyi tanımlanmış değişkenler olmaksızın, doğrudan test edilebilir olmamakla birlikte bir takım istatistiksel kanıtlardan hareketle değerlendirilebilir görülmektedir. Kayıp veri mekanizmasının tanımlanmasına yönelik istatistiksel kanıtlar genellikle kayıp veri mekanizması ile ilişkili olmaması gereken cinsiyet, sosyoekonomik düzey gibi değişkenlerin kategorilerine göre yanıt örüntülerinde manidar bir fark oluşup oluşmadığının incelenmesi yaklaşım ve yöntemlerini içermektedir. Sıklıkla başvuru alan diğer bir yöntem ise genel yanıtlayıcı kitlesi ile yanıt örüntülerinde kayıp veri bulunmayan yanıtlayıcı kitlesinin profillerinin karşılaştırılmasıdır.

Bu çalışmada, kayıp veri mekanizması, cinsiyet, okul türü ve bölge değişkenleri dikkate alınarak gerçekleştirilen analizlerle tanımlanmaya çalışılmıştır. Cinsiyet, okul türü ve bölge değişkenlerinin kategorileri arasında, yanıtlayıcıların, doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının manidar düzeyde farklılaşıp farklılaşmadığı test edilmiştir. Bu incelemede, yanıtlayıcıların tamamının yanı sıra yanıt örüntüleri gözlenen verilerden oluşan, madde düzeyinde kayıp veri içeren ve birim düzeyinde kayıp veri içeren yanıtlayıcılar da ayrı ayrı gruplar olarak dikkate alınmıştır.

Tüm yanıtlayıcılar ile yanıt örüntülerinde kayıp veri bulunmayan yanıtlayıcıların profilleri karşılaştırılmıştır.

Doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının dağılımları, kayıp veri örüntüsüne yönelik inceleme ve analizlerde değerlendirilmiştir. Bu değerlendirmeler, söz konusu dağılımların normal dağılımdan manidar düzeyde farklılaştığını göstermektedir. Ayrıca söz konusu dağılımların cinsiyet, okul türü ve bölge değişkenlerinin kategorileri arasında varyansların homojenliği varsayımının karşılanıp karşılanmadığı, 'Levene Varyansların Homojenliği Testi' kestirimlerine bağlı olarak test edilmiş, üretilen olasılık değerleri $p < 0.05$ olarak elde edilmiştir. Bu kestirimler, cinsiyet, okul türü ve bölge değişkenlerinin kategorileri arasında varyansların manidar düzeyde farklılaştığını ve grupların varyans dağılımları açısından homojen olmadığını göstermektedir. Normallik ve varyansların homojenliği varsayımlarının karşılanmadığı, ayrıca doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının yapay süreksiz değişkenler olan iki kategorili puanlanan maddelerden üretildiği de dikkate alınarak, verilerin parametrik özellik göstermediği kararına varılmıştır. Bu karara bağlı olarak, ileri analizlerde parametrik olmayan testlerin kullanımı uygun görülmüştür.

Cinsiyet ve okul türü değişkenlerinin kategorileri arasında kayıp veri oluşma olasılığı açısından manidar fark bulunup bulunmadığı, parametrik olmayan testlerden 'Mann-Whitney U Testi' ile test edilmiştir. Bölge değişkeni ise 12 kategori içermektedir. Kayıp veri oluşma olasılığı açısından bu kategoriler arasındaki farkın manidarlığının test edilmesinde, parametrik olmayan testlerden 'Kruskal-Wallis Testi' kullanılmıştır. Manidar fark belirlenmesi durumunda çoklu karşılaştırma testleri ile farkın kaynağı belirlenmeye çalışılmıştır. Bu amaçla parametrik olmayan verilere daha uygun olduğu düşünülen çoklu karşılaştırma testlerinden 'Tamhane Testi' kullanılmıştır. Elde edilen kestirimler, tüm yanıtlayıcılar dikkate alınarak değerlendirildiği gibi yanıt örüntüleri gözlenen verilerden oluşan, madde ve birim düzeyinde kayıp veri içeren yanıtlayıcılar düzeyinde de ayrı ayrı değerlendirilmiştir. Tüm yanıtlayıcılar ile yanıt örüntülerinde kayıp veri bulunmayan yanıtlayıcılara yönelik kestirimler karşılaştırılmıştır.

Üçüncü aşamada, bu çalışmada kullanılabilecek kayıp veri yöntemleri belirlenmiş, bu yöntemlere göre silme, atama ya da veri çoğaltma ile eksiksiz veri setleri elde edilmiştir. Bu belirlemede kayıp veri örüntüsü ve kayıp veri mekanizmasına yönelik inceleme ve analizlerin yanı sıra istatistiksel yazılımların özellikleri ve erişilebilirlikleri de dikkate alınmıştır. İki kategorili puanlanan maddelerden oluşan testlere yönelik olarak kullanılabilecek 12 kayıp veri yöntemi belirlenmiştir. Bu yöntemlerden DS yöntemi silmeye dayalı, 0 Atama, SO, BO, YNO, YNM, Dİ ve DE yöntemleri basit atamaya dayalı, RA, BM ve VÇ yöntemleri en çok olabilirlik yaklaşımına dayalı, MCMC ise çoklu atama yaklaşımına dayalı yöntemlerdir.

Önceki bölümlerde ele alındığı şekliyle, ilgili literatürde, 0 Atama yönteminin dezavantajları üzerinde bir uzlaşma bulunmaktadır. 0 Atama yöntemi, kestirimlerde manidar düzeyde yanlılığa yol açmaktadır ve bu yöntem ile elde edilen bulgular anlamlı görülmemektedir. Bununla birlikte, 0 Atama yöntemi, kayıp verilere yönelik olarak, bilinçli ya da bilinçsiz, oldukça yaygın bir şekilde kullanılmaktadır. Bu nedenle bu çalışmada, diğer yöntemlerin yanı sıra 0 Atama yönteminin de dikkate alınması uygun bulunmuştur.

Kayıp veri yöntemi olarak ilgili literatürde sıklıkla kullanılan 'kısmi silme (KS)' yöntemi, ikili karşılaştırmalara dayalı bir yöntemdir. KS yöntemi, maddeler düzeyinde ortalama, standart hata, standart sapma, varyans ve kovaryans kestirimlerinde sadece gözlenen değerleri dikkate almaktadır. Gözlenen değerler üzerinden elde edilen bu kestirimler tek bir değişkene yöneliktir. Diğer taraftan KS yönteminde toplam puanların belirlenmesinde, çoklu karşılaştırma yapılmakta ve kayıp veri içeren tüm birimler analiz dışı bırakılmaktadır. Bu durumda elde edilen toplam puanlar, DS yöntemi ile elde edilen puanlarla örtüşmektedir. Bu kapsamda bu çalışmada silmeye dayalı yöntemlerden DS yönteminin kullanılması daha uygun bulunmuştur.

Bu çalışmada kullanılan kayıp veri yöntemlerinden RA yöntemi, EÇÖ yöntemleri içerisinde değerlendirilmiştir. Sürekli verilerde RA, en küçük kareler ortalaması hesaplamasına dayalı olarak kestirim üretmektedir. Kategorik verilerde ise en küçük kareler hesaplaması anlamlı değildir. Bu nedenle kategorik verilerde RA, en çok olabilirlik hesaplaması yerine en çok olabilirlik yaklaşımı temelinde olasılıklı hesaplama yapmaktadır.

Dördüncü aşamada, farklı kayıp veri yöntemlerine göre elde edilmiş eksiksiz veri setleri üzerinde, 'açımlayıcı faktör analizi (AFA)' ve 'doğrulayıcı faktör analizi (DFA)' kullanılarak, testin yapı geçerliği ve tek boyutluluk özellikleri incelenmiştir. Analizlerde AFA kullanılmasının amacı, maddeleri tanımlanmış boyutlar altında sınıflandırarak testi yapılandırmak değildir. Test maddelerinin, ağırlıklı olarak baskın tek bir boyut altında toplanıp toplanmadığının belirlenmesi amaçlanmaktadır. Benzer şekilde DFA, tek bir gizil değişken içeren ölçme modeline göre model-veri uyumunun sağlanıp sağlanmadığının test edilmesi amacıyla kullanılmıştır.

Açımlayıcı faktör analizi uygulamalarında, kullanılan veri setinin analize uygunluğunun incelenmesinde sıklıkla başvurulanan testlerden ikisi 'Kaiser-Meyer-Olkin Örneklem Büyüklüğünün Uygunluğu Testi (KMO)' ve 'Bartlett'in Küresellik Testi'dir. Bu çalışmada, her bir kayıp veri yöntemi için KMO katsayıları .94 ve üzerinde elde edilmiştir. Bu kestirimler veri seti büyüklüğünün faktör analizi için uygun olduğunu göstermektedir. Temelde bir χ^2 testi olan Bartlett Testi kestirimleri ise analizlerde kullanılan veri setinin, her bir kayıp veri yöntemi için çok değişkenli normal dağılım ve doğrusallık varsayımlarını karşıladığına yönelik kanıt sağlamaktadır. Bartlett Testi kestirimlerine yönelik olasılık değerleri $p < 0.05$ olarak elde edilmiştir. KMO ve Bartlett testlerine bağlı olarak belirlenen her bir kayıp veri yöntemi için, verilerin faktör analizine uygun olduğu anlaşılmaktadır.

Bu çalışmada kullanılan veri seti iki kategorili puanlanan maddelerden oluştuğu için AFA'da 'Temel Eksenler Faktörlendirmesi (Principal Axis Factoring - TEF)' ve 'En Çok Olabilirlik (Maximum Likelihood - EÇO)' desenleri kullanılmıştır. EÇO deseni, olasılıklı ve ötelemeli bir yöntem olup yanıt örüntülerine bağlı olarak 3 olasılıklı öteleme ile kestirimler gerçekleştirilmiştir. Özellikle EÇO deseninin normallik varsayımına karşı duyarlı olmasından dolayı veriler üzerinde 'oblique' döndürme ve 'Kaiser normalleştirme' ile düzeltme yapılmıştır.

SBS 2011 Matematik Testi A Kitapçığına yönelik olarak farklı kayıp veri yöntemleri ile elde edilmiş eksiksiz veri setleri üzerinde iki AFA uygulaması gerçekleştirilmiştir. Faktör sayısı ilk uygulamada sınırlandırılmamıştır. İkinci uygulamada ise faktör sayısı 1 olarak belirlenmiş ve maddeler tek faktör altında toplanmaya zorlanmıştır. Faktör sayılarının

belirlenmesinde özdeğerler yaklaşımı ile özdeğeri 1.00 ve üzerinde olan faktörler dikkate alınmıştır.

AFA'da veri setinin karakteristiğine uygun olarak kullanılan EÇO deseni, olasılıklı bir model kurmakta ve bu modelin uyum iyiliğine yönelik χ^2 kestirimleri üretmektedir. Model uyumunda χ^2/sd oranı dikkate alınmaktadır. Bununla birlikte χ^2 , veri setinin büyüklüğüne karşı fazlaca duyarlı olduğundan, AFA çıktılarına bağlı olarak model uyumunun değerlendirilmesi yeterli görülmemektedir. Bu nedenle bu çalışmada, veri-model uyumu, AFA çıktılarının yanı sıra DFA uyum iyiliği indeksleri de dikkate alınarak değerlendirilmiştir.

DFA, sürekli verilerde, kovaryans ya da korelasyon matrisleri üzerinden hareketle kestirimler üretmektedir. İki kategorili puanlanan veri setlerinde ise DFA, asimetrik kovaryans matrisleri üzerinden hesaplama yapmaktadır. Bu nedenle bu çalışmada DFA, asimetrik kovaryans matrisleri tanımlanarak uygulanmıştır.

Beşinci aşamada, farklı kayıp veri yöntemleri ile elde edilen eksiksiz veri setleri üzerinde, madde parametre kestirimleri gerçekleştirilmiştir. Bu kestirimlerde, 'klasik test kuramı (KTK)' temel alınmıştır. KTK temelinde, iki kategorili puanlanan maddelerden oluşan testlere yönelik madde parametreleri, madde güçlük indeksi, madde varyans ve standart sapması, madde ayırıcılık indeksi, madde güvenilirlik katsayısı gibi temel parametreleri ifade etmektedir.

Kayıp veri yöntemleri, öncelikle ortalama, standart sapma ve varyans-kovaryans matrisi kestirimleri üretmektedir. Kayıp veri yöntemlerinin kullanımına yönelik istatistiksel yazılımların neredeyse tamamı, kayıp veri içeren veri setlerinde sadece bu istatistiklerin kestirimlerini çıktı olarak vermektedir. İleri analizler için bu istatistikler temel oluşturmaktadır, fakat veri setinde ayrıca düzenlemeler yapılması gerekmektedir.

İki kategorili puanlanan maddeler, sınıflama ölçeğinde ve yapay süreksiz değişkenlerdir. Bu değişkenlerde ortalama, standart sapma ve varyans kestirimleri anlamlı değildir. Bu nedenle bu çalışmada, madde parametre kestirimlerinde, ortalama, standart sapma ve varyanslar dikkate alınmamıştır.

Altıncı aşamada, farklı kayıp veri yöntemleri ile elde edilen eksiksiz veri setleri üzerinde test parametre kestirimleri gerçekleştirilmiştir. KTK temelinde gerçekleştirilen kestirimlerde test parametreleri olarak, minimum ve maksimum puanlar, test ortalaması, test ortalama güclüğü, test varyans ve standart sapması, test basıklık ve çarpıklık katsayıları dikkate alınmıştır.

Kestirilen test parametreleri öncelikle farklı kayıp veri yöntemleri ile elde edilen toplam puanların dağılımlarının incelenmesinde kullanılmıştır. Farklı kayıp veri yöntemlerine göre elde edilen toplam puanların normal dağılımdan gelip gelmediği ve bu puanlar arasındaki ilişkiler, betimsel olarak incelenmiş ve değerlendirilmiştir.

Yukarıda açıklanan analizlerin çoğunluğunda, özel amaçlı istatistiksel yazılımlardan yararlanılmıştır. Kullanılan kayıp veri yöntemlerinin tamamını içeren bir istatistiksel yazılım bulunmamaktadır. Bu nedenle, verilerin analizinde, belirlenen kayıp veri yöntemlerinin kullanımına uygun olarak, SPSS 15.0, SPSS 20.0, LISREL 8.7, NORM ve IMPUTE yazılımları kullanılmıştır². Madde ve test parametrelerinin kestirimleri ise ITEMANW yazılımı ile gerçekleştirilmiştir.

SPSS yazılımı, 17 ve üzeri sürümlerinde, kayıp verilerle başa çıkmada silmeye ve basit atamaya dayalı yöntemlerin yanı sıra ÇVA tabanlı yöntemler içermektedir. ÇVA yöntemlerinin kullanımında SAS yazılımı kadar zengin içerikte olmasa da SPSS, bazı önemli yaklaşımlar içermektedir. Bunlardan en belirgini SPSS'te basit veri atama yöntemi olarak VÇ yerine RA yöntemlerinin kullanılmasıdır. Bu durum ölçek türlerine göre farklı atama yöntemlerinin kullanılabilmesine olanak sağlamaktadır. Sıralama ve sınıflama ölçeğindeki değişkenlere yönelik olarak lojistik regresyon atama, sürekli değişkenlere yönelik olarak ise doğrusal regresyon atama yöntemi kullanılabilir. SPSS, ÇVA tabanlı yöntemlere yönelik analiz sürelerinde, her bir öteleme işlemine yönelik ortalama ve standart hata kestirimlerini saklamakta ve kullanıcıya çıktı olarak sunmaktadır. Bu yaklaşım, özellikle zaman serilerinin yapılandırılmasında ve otokorelasyon kestirimlerinde önemli bir işlevsellik sağlamaktadır (Enders, 2010).

² Kayıp veri analizlerinde kullanılabilen AMOS, EQS, Mplus, SAS, mle ve IVEware gibi başkaca yazılımlar bulunmaktadır. Bu yazılımlarla ilgili ayrıntılı bilgi için sırasıyla Arbuckle (2007), Bentler ve Wu (2002), Muthen ve Muthen (2009), Enders (2010), Holman (1991-2003), Raghunathan, Solenberger ve Van Hoewyk (2002) kaynaklarına başvurulabilir.

LISREL, kayıp verilerle başa çıkmada EÇO tabanlı algoritmalar kullanan bir diğer istatistiksel yazılımdır. Bu algoritmalar genel olarak normal dağılım altında yapılandırılmıştır. Bununla birlikte normal dağılıma sahip olmayan ve kayıp veri içeren değişkenlere yönelik olarak 'sağlamlık (robust)' standart hatası ve 'yeniden ölçeklenmiş (rescaled)' test istatistikleri kestirimleri yapılabilmektedir. LISREL, standart hata kestirimlerinde, gözlenen değerler yerine beklenen değerlerden elde edilen bilgi matrisini kullanmaktadır. EÇO tabanlı yöntemler kullanmakla birlikte, 'veri çoğaltmaya-VÇ (data augmentation)' dayalı çoklu veri atama, benzer değer atama ve regresyon atama yöntemleri ile birden fazla veri setini bir araya getirme olanağı sunmaktadır. EÇO tabanlı algoritmalar kullanan diğer yazılımların aksine LISREL'de kayıp verilere yönelik regresyon atama yöntemi, BM temelinde yapılandırılmıştır. LISREL, atanan regresyon denklemlerini bir araya getirmede BM ortalama vektörü ve kovaryans matrisi kestirimlerinden yararlanmaktadır. Fakat bu yaklaşımın da standart regresyon atama yöntemine benzer şekilde yanlışlık içerdiği bilinmektedir (Jöreskog ve Sörbom, 2006).

NORM, veri çoğaltma yöntemleri içeren, tamamen ücretsiz ve açık erişim sağlayan bir yazılımdır. Program veri çoğaltmaya yönelik olarak KDD, yuvarlama, grafiksel benzerlik tanılaması gibi birçok veri çoğaltma yönteminin kullanımına olanak sağlamaktadır. NORM esas itibarıyla bir analiz programı değildir. NORM, tamamlanmış veri setleri üretmekte, istatistiksel kestirimler üretmemektedir. Bununla birlikte BM yönteminin algoritmaları aracılığı ile EÇO temelinde, ortalama vektörü ve kovaryans matrisi kestirimlerine olanak sağlamaktadır (Schaffer, 1999).

BÖLÜM 4

BULGULAR VE YORUMLAR

Belirlenen araştırma soruları çerçevesinde, yürütülen analizler, analizlerden elde edilen çıktılar ve çıktılara yönelik yorumlar bu bölümde sunulmaktadır.

Kayıp Veri Örüntüsü

SBS 2011 Matematik Testi A Kitapçığına yönelik veri setinde, kayıp veri örüntüsüne ilişkin incelemelerde, yanıt örüntülerinin kayıp veri içermesi durumu, toplam doğru ve yanlış yanıt sayıları, kayıp veri içeren madde sayıları, en sık ve en az sıklıkta görülen yanıt örüntüleri dikkate alınmıştır. Testi alan yanıtlayıcılara yönelik gözlenen ve kayıp verilerin dağılımı Çizelge 3.1’de verilmektedir.

Çizelge 3.1. Gözlenen ve Kayıp Verilerin Dağılımı

	N	%
Gözlenen Veri	144553	27.40
Kayıp Veri (Madde Düzeyinde)	382937	72.59
Kayıp Veri (Birim Düzeyinde)	27	.01
Toplam	527517	100.00

Çizelge 3.1’de görüldüğü gibi testi alan gözlem birimi yani yanıtlayıcı sayısı 527517’dir. Yanıtlayıcıların %27.40’ı matematik testinde yer alan maddelerin her birinde, doğru ya da yanlış, yanıtlama yapmıştır. Bu yanıtlayıcılara yönelik yanıt örüntülerinde kayıp veri bulunmamaktadır. Bu durumda kayıp veri içermeyen yani eksiksiz alt veri seti 144553 gözlem biriminden oluşmaktadır. Yanıtlayıcıların 382964 (%72.60)’ü testte yer alan maddelerden en az birinde yanıtlama yapmamıştır. Bu yanıtlayıcıların neredeyse tamamının (%79.59) yanıt örüntülerinde madde düzeyinde kayıplar bulunmaktadır. Birim düzeyindeki kayıplar ise tüm gözlem

birimlerinin %0.01'ini oluşturmaktadır. Bu yanıtlayıcılar, testte yer alan maddelerin tamamında yanıtlama yapmamıştır ya da birden fazla seçenek işaretleme, belirgin olmayan şekilde işaretleme gibi gerekçelerden dolayı yanıtları geçersiz sayılmıştır.

Rubin (1987), Puma ve diğerleri (2009), birim düzeyinde kayıp verilerin örnekleme hatasına, madde düzeyinde kayıp verilerin ise madde yanlışlığına işaret edebileceğini belirtmektedir. SBS 2011 Matematik Testi A Kitapçığına yönelik verilerde birim düzeyinde kayıplar, oldukça azdır. Bu durumda, SBS'nin Türkiye genelinde uygulanan bir test olduğu da dikkate alındığında, kayıp verilerin oluşmasının örnekleme hatasından kaynaklanma olasılığı düşüktür. Diğer taraftan madde düzeyinde kayıp veri miktarının fazlalığı, madde yanlışlığı olasılığının yüksek olduğunu göstermektedir.

SBS 2011 Matematik Testi A Kitapçığını alan yanıtlayıcıların toplam doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının dağılımı Çizelge 3.2'de verilmektedir.

Çizelge 3.2. Doğru Yanıtlanan, Yanlış Yanıtlanan ve Kayıp Veri İçeren Madde Sayılarının Dağılımı

Madde Sayısı	Doğru Yanıt		Yanlış Yanıt		Kayıp Veri	
	N	%	N	%	N	%
20	9025	1.7	507	.1	27	.0
19	10213	1.9	3231	.6	76	.0
18	10105	1.9	9591	1.8	222	.0
17	9677	1.8	19562	3.7	491	.1
16	9149	1.7	29693	5.6	1031	.2
15	9179	1.7	37069	7.0	1877	.4
14	9356	1.8	40584	7.7	3052	.6
13	9702	1.8	42001	8.0	4469	.8
12	10231	1.9	41506	7.9	6357	1.2
11	11486	2.2	39916	7.6	8546	1.6
10	13183	2.5	37608	7.1	12088	2.3
9	16516	3.1	34980	6.6	15406	2.9
8	22760	4.3	32561	6.2	19669	3.7
7	33237	6.3	29204	5.5	25035	4.7
6	48085	9.1	26473	5.0	31296	5.9
5	63439	12.0	23821	4.5	38790	7.4
4	73962	14.0	20948	4.0	46107	8.7
3	70583	13.4	18597	3.5	52919	10.0
2	52775	10.0	16045	3.0	57006	10.8
1	27241	5.2	13531	2.6	58500	11.1
0	7613	1.4	10089	1.9	144553	27.4
Toplam	527517	100.0	527517	100.0	527517	100.0

Çizelge 3.2’de görüldüğü gibi SBS 2011 Matematik Testi A Kitapçığında çoktan seçmeli 20 madde yer almaktadır. İki kategorili puanlanan maddelerden oluşan testin genelinde doğru yanıtlanan, yanlış yanıtlanan ve kayıp veri içeren madde sayıları, 0-20 aralığında değişmektedir.

Testte yer alan maddelerin tamamına doğru yanıt veren yanıtlayıcı sayısı 9025 (%1.7)’tir. Hiçbir doğru yanıt bulunmayan yanıtlayıcı sayısı 7613 (%1.4)’tür. Bu yanıtlayıcıların 27’si, birim düzeyindeki kayıpları oluşturmaktadır ve testin genelinde herhangi bir yanıtlama yapmamıştır. Geriye kalan 7586 yanıtlayıcının ise yanıt örüntülerinde en az bir yanlış yanıt bulunmaktadır. En fazla doğru yanıtlama yapılan madde sayılarının 2 ile 6 arası madde civarında toplandığı görülmektedir. Yanıtlayıcılar arasında 4 maddeye doğru yanıt verenlerin (%14.0) sıklığı daha fazladır.

Hiçbir yanlış yanıtlama yapmayan yanıtlayıcı sayısı 10089 (%1,9)’dur. Bu yanıtlayıcıların 27’si, hiçbir yanıtlama yapmamış olup veri setinin birim düzeyinde kayıp olarak tanımlanan kısmını oluşturmaktadır. Geriye kalan yanıtlayıcılardan 9025’inin tüm yanıtları doğrudur. Diğer 1037 yanıtlayıcı ise tüm maddelerde değil fakat en az bir maddede doğru yanıtlama yapmış olup bu yanıtlayıcıların hiçbir yanlış yanıtı bulunmamaktadır. En fazla yanlış yanıtlama yapılan madde sayılarının 8 ile 15 aralığında yoğunlaştığı görülmektedir. Yanıtlayıcılar arasında 13 maddede yanlış yanıtlama yapanların (%8.0) sıklığı daha fazladır.

Yanıtlayıcıların çoğunluğu oluşturan kısmının (%27.4) yanıt örüntülerinde herhangi bir kayıp bulunmamaktadır. Yanıt örüntülerinde 2 veya daha az kayıp veri içeren yanıtlayıcılar (%49.3), tüm yanıtlayıcıların yarıya yakın kısmını oluşturmaktadır. Yanıt örüntülerinde 15 veya daha fazla kayıp veri içeren yanıtlayıcılar (%0.7) ise yüzdelerik bir dilim oluşturmayacak kadar azdır.

SBS 2011 Matematik Testi A Kitapçığını alan yanıtlayıcıların doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının dağılımına ilişkin kestirimler Çizelge 3.3’te verilmektedir.

Çizelge 3.3. Doğru Yanıtlanan, Yanlış Yanıtlanan ve Kayıp Veri İçeren Maddelerin Dağılımlarına Yönelik Parametreler

	μ	SH (μ)	σ^2	σ	Basıklık	Çarpıklık
Doğru Yanıt	6.54	.007	23.943	4.893	1.212	.604
Yanlış Yanıt	10.01	.006	21.531	4.640	-0.292	-0.789
Kayıp Veri	3.45	.005	12.329	3.511	1.098	.771

Çizelge 3.3'te görüldüğü gibi gözlenen veriler dikkate alındığında, testi alan yanıtlayıcıların ortalama doğru yanıt sayısı 6.54, ortalama yanlış yanıt sayısı 10.01 ve kayıp veri içeren ortalama madde sayısı ise 3.45'tir. Bu kestirimlere yönelik standart hatalar (.007, .006 ve .005), evren büyüklüğüne bağlı olarak oldukça düşüktür. Standart sapma değerleri 5'in altındadır.

Ortalama, basıklık ve çarpıklık katsayılarına göre doğru yanıt sayıları ve kayıp veri içeren madde sayıları, sağa çarpık bir dağılım gösterirken, yanlış yanıt sayıları sola çarpık bir dağılım göstermektedir. Doğru yanıt sayıları ve kayıp veri içeren madde sayılarındaki çarpıklıklar oldukça belirgindir ve bu dağılımların normal dağılımdan saptığı iddia edilebilir. Diğer taraftan yanlış yanıt sayılarındaki sapma, normal dağılıma kıyasla kabul edilebilir ve makul sınırlar içerisindedir.

SBS 2011 Matematik Testi A Kitapçığında yer alan her bir maddede doğru yanıtlama, yanlış yanıtlama sayıları ve kayıp veri sayılarının dağılımları Çizelge 3.4'te verilmektedir.

Çizelge 3.4'te görüldüğü gibi gözlenen veriler dikkate alındığında, test maddeleri düzeyinde doğru yanıtlama yüzdeleri %16.1 ile %61.3 arasında değişmektedir. En fazla doğru yanıtlanan madde 17. maddedir. Bunu sırasıyla 8, 5, 7 ve 15. maddeler izlemektedir. En az doğru yanıtlanan madde ise 13. maddedir. Bunu sırasıyla 19, 6, 1 ve 2. maddeler izlemektedir. Doğru yanıtlama yüzdelerinde düzenli bir artış ya da azalış görülmemektedir.

Test maddeleri düzeyinde yanlış yanıtlama yüzdeleri %30.0 ile %71.5 arasında değişmektedir. En az yanlış yanıtlanan madde 17. maddedir. Bunu sırasıyla 5, 4, 7 ve 11. maddeler izlemektedir. En fazla yanlış yanıtlanan madde 13. maddedir. Bunu sırasıyla 6, 1, 20 ve 2. maddeler izlemektedir. Yanlış yanıtlama yüzdelerinde, düzenli bir artış ya da azalış görülmemektedir.

Çizelge 3.4. Maddeler Düzeyinde Tepkilerin Dağılımı

Madde	Doğru		Yanlış		Kayıp		Toplam	
	N	%	N	%	N	%	N	%
1	131528	24.9	319443	60.6	76546	14.5	527517	100.0
2	131145	24.9	313670	59.5	82702	15.7	527517	100.0
3	147484	28.0	215742	40.9	164291	31.1	527517	100.0
4	160423	30.4	185470	35.2	181624	34.4	527517	100.0
5	228436	43.3	174604	33.1	124477	23.6	527517	100.0
6	130361	24.7	355921	67.5	41235	7.8	527517	100.0
7	213151	40.4	194275	36.8	120091	22.8	527517	100.0
8	235407	44.6	210916	40.0	81194	15.4	527517	100.0
9	155951	29.6	311763	59.1	59803	11.3	527517	100.0
10	184155	34.9	255874	48.5	87488	16.6	527517	100.0
11	190659	36.1	208363	39.5	128495	24.4	527517	100.0
12	167622	31.8	271045	51.4	88850	16.8	527517	100.0
13	84776	16.1	377078	71.5	65663	12.4	527517	100.0
14	156593	29.7	280101	53.1	90823	17.2	527517	100.0
15	211270	40.0	268379	50.9	47868	9.1	527517	100.0
16	173559	32.9	303382	57.5	50576	9.6	527517	100.0
17	323179	61.3	158305	30.0	46033	8.7	527517	100.0
18	138864	26.3	241590	45.8	147063	27.9	527517	100.0
19	89559	17.0	313571	59.4	124387	23.6	527517	100.0
20	198427	37.6	318652	60.4	10438	2.0	527517	100.0

Çizelge 3.4'te görüldüğü gibi test maddeleri düzeyinde kayıp veri yüzdeleri %2.0 ile %34.4 arasında değişmektedir. %5 ve altında kayıp veri içeren bir tane madde (20. madde) bulunmaktadır. %10 ve altında kayıp veri içeren madde sayısı ise beştir. En az kayıp veri içeren madde 20. maddedir. Bunu sırasıyla 6, 17, 15 ve 16. maddeler izlemektedir. En fazla kayıp veri içeren madde ise 4. maddedir. Bunu sırasıyla 3, 18, 11, 5 ve 19. maddeler izlemektedir. Kayıp veri yüzdelerinde, düzenli bir artış ya da azalış görülmemektedir.

Testte yer alan 20 maddenin her biri, kayıp veri içermektedir. Bu maddelerden 19'unda %5'in üzerinde, 14'ünde %10'un üzerinde, 7'sinde ise %20'nin üzerinde kayıp veri bulunmaktadır. Allison (2002), çok değişkenli ve çoklu analizlerde ihmal edilebilir kayıp veri miktarının %5 civarında olması gerektiğini belirtmektedir. Acock (2005)'a göre bu oran %20'ye kadar çıkabilmektedir. Groves (2006)'a göre ihmal edilebilir kayıp veri miktarı, araştırmanın türü ve modeli ile kullanılan verilerin karakteristiğine bağlı olarak değişebilmektedir. Bununla birlikte istatistiksel analiz süreçleri içeren

araştırmalarda kullanılan veri setinin, değişkenler düzeyinde %15'in üzerinde kayıp veri içermemesi önerilmektedir.

SBS 2011 Matematik Testi A Kitapçığına yönelik veri setinde, kayıp verilere bağlı olarak, 28226 farklı yanıt örüntüsünün oluştuğu belirlenmiştir. Bu örüntüler içerisinde en sık görülen 20 örüntü Çizelge 3.5'te, en az sıklıkta görülen 20 örüntüsü ise Çizelge 3.6'da birlikte verilmektedir. Çizelge 3.6'da verilen örüntüler, 5'ten az yanıtlayıcıda görülmektedir.

Çizelge 3.5'te görüldüğü gibi kayıp verilere bağlı yanıt örüntüleri arasında en sık görülen 20 örüntü, 207288 (%40) yanıtlayıcının yanıt örüntülerini temsil etmektedir. En sık görülen örüntü (%27.40), maddelerin tamamında, doğru ya da yanlış, yanıtlama yapılmış olan yani eksiksiz yanıt örüntüsüdür. İkinci sırada sadece 4. maddede, üçüncü sırada sadece 19. maddede, dördüncü sırada sadece 3. maddede, beşinci sırada sadece 18. maddede ve altıncı sırada ise sadece 5. maddede kayıp veri bulunan yanıt örüntüleri gelmektedir. Bu maddeler aynı zamanda en fazla kayıp veri içeren maddelerdir. Çizelge 3.6'ya göre en az sıklıkta görülen yanıt örüntülerinde ise kayıp veriler, belirli maddelerde yığılmamakta, her maddede görülebilmektedir.

Kayıp veri örüntüsüne ilişkin olarak yapılan inceleme ve analizler, her bir maddede kayıp veri bulunduğunu, bir madde dışında kayıp veri miktarlarının %5'in üzerinde olduğunu göstermektedir. Kayıp veriler gerek maddeler düzeyinde gerekse yanıt örüntüleri düzeyinde düzenli, planlı ya da monoton bir dağılım göstermemektedir. Bu durumda kayıp veri örüntüsünün, önceki bölümlerde tanımlandığı şekliyle, 'genel örüntü' yapısına daha uygun olduğu anlaşılmaktadır. Enders (2010), kayıp verilere yönelik en sık rastlanan örüntü yapısının 'genel örüntü' olduğunu belirtmektedir. Genel örüntü, ilk bakışta, kayıp verilerin seçkisiz oluştuğu yönünde bir öngörü oluşturabilmektedir. Bununla birlikte seçkisizlik kararının kayıp veri mekanizmasına bağlı olarak verilmesi daha uygun görülmektedir.

Çizelge 3.5. Gözlem Birimleri Düzeyinde En Sık Görülen Kayıp Veri Örüntülerinin Dağılımı

Örüntü	Madde																				F	%
	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20		
1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	144553	27.40
2	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	7555	1.43
3	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	7205	1.37
4	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	6660	1.26
5	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	6415	1.22
6	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	4215	.80
7	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	3820	.72
8	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	3645	.69
9	0	0	1	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	2395	.45
10	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	2315	.44
11	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	2255	.43
12	0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	2205	.42
13	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	1	0	2150	.41
14	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	0	1845	.35
15	0	0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	1785	.34
16	0	0	1	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	0	1755	.33
17	0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	1720	.33
18	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0	0	1615	.31
19	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1600	.30
20	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1530	.29
Toplam																					207288	39.29

*1, kayıp verinin konumunu belirlemektedir.

Çizelge 3.6. Gözlem Birimleri Düzeyinde En Az Sıklıkta Görülen Kayıp Veri Örüntüleri

Örüntü	Madde																					
	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20		
1	0	0	0	1	0	0	0	0	0	0	0	0	0	1	1	0	0	1	0	0	0	
2	0	0	1	1	1	1	1	0	0	1	1	0	1	0	1	0	0	1	1	0		
3	1	0	1	1	0	0	1	0	1	0	0	0	1	0	0	0	0	0	0	0		
4	1	1	1	0	0	0	0	1	0	1	1	1	0	0	0	0	0	0	1	0		
5	0	0	0	1	0	0	1	1	0	0	1	0	0	1	1	1	0	0	0	0		
6	0	0	1	0	1	0	0	0	0	0	1	0	1	0	0	0	0	1	1	0		
7	0	0	0	0	0	1	1	1	0	0	0	1	0	0	0	0	0	1	1	0		
8	1	1	1	1	1	0	1	0	0	0	1	0	0	0	1	1	0	1	1	0		
9	0	0	1	1	1	0	1	1	1	1	1	1	0	1	0	1	0	0	1	0		
10	1	0	0	0	0	0	1	0	1	1	0	0	0	0	0	0	0	0	0	0		
11	1	0	0	1	1	0	1	1	0	1	0	1	0	1	1	1	0	1	1	0		
12	1	1	1	1	1	0	1	1	1	0	1	0	0	0	0	0	0	1	0	0		
13	0	0	1	0	1	0	0	1	0	0	0	0	1	0	0	0	0	1	0	0		
14	1	0	0	0	1	0	0	0	0	1	0	0	0	0	1	0	0	0	0	0		
15	1	1	1	1	1	0	0	1	1	0	0	1	1	1	0	0	0	1	1	0		
16	0	1	0	1	0	0	0	0	0	1	1	0	0	1	0	0	0	1	1	0		
17	0	0	1	0	1	0	0	0	0	0	1	1	1	0	0	1	0	0	0	0		
18	0	1	1	1	0	0	1	0	0	1	1	0	0	0	0	1	0	0	1	0		
19	0	1	1	0	1	0	1	1	1	0	1	0	1	1	0	0	0	0	1	0		
20	0	0	0	1	1	0	1	1	0	0	1	0	0	1	0	0	1	1	1	0		

*1, kayıp verinin konumunu belirlemektedir.

Kayıp Veri Mekanizması

Bu çalışmada, SBS 2011 Matematik Testi A Kitapçığına yönelik yanıt örüntülerinde kayıp verilerin oluşma olasılığı ile ilişkili olabilecek değişkenlerin belirlenmesi amacıyla kayıp veri mekanizması tanımlanmaya çalışılmıştır. Kayıp veri mekanizmasının tanımlanmasına yönelik incelemeler, cinsiyet, okul türü ve bölge bağımsız değişkenleri dikkate alınarak gerçekleştirilmiştir.

Testi alan yanıtlayıcıların doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının cinsiyetlerine göre manidar düzeyde fark içerip içermediğine yönelik analiz çıktıları Çizelge 3.7’de verilmektedir. Gözlenen yanıt örüntülerinde kayıp veri, birim düzeyinde kayıp veri içeren yanıt örüntülerinde ise doğru ve yanlış yanıtlar bulunmadığı için bu örüntülere yönelik parametrelerden, ortalama, standart sapma ve Mann-Whitney U testi değerleri elde edilememiştir.

Çizelge 3.7. Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Cinsiyetlerine Göre Dağılımı

		Cinsiyet	N	μ	SH (μ)	σ	Mann-Whitney U İstatistiği*
Gözlenen Veri	Doğru	E	83122	8.54	.02	6.04	8.03
		K	61431	9.38	.03	6.47	
	Yanlış	E	83122	11.46	.02	6.04	8.03
		K	61431	10.62	.03	6.47	
Kayıp Veri (Madde Düzeyinde)	Doğru	E	187429	5.51	.01	3.82	9.60
		K	195508	5.80	.01	4.01	
	Yanlış	E	187429	9.55	.01	3.74	2.95
		K	195508	9.63	.01	3.82	
	Kayıp	E	187429	4.94	.01	3.43	12.69
		K	195508	4.57	.01	3.13	
	Kayıp	E	16	20.00	.00	.00	-
		K	11	20.00	.00	.00	
Toplam	Doğru	E	270567	6.44	.01	4.83	4.50
		K	256950	6.65	.01	4.96	
	Yanlış	E	270567	10.14	.01	4.66	7.27
		K	256950	9.87	.01	4.62	
	Kayıp	E	270567	3.42	.01	3.65	8.51
		K	256950	3.48	.01	3.36	

*p<0.05

Çizelge 3.7’de görüldüğü gibi testi alan yanıtlayıcıların yarıdan fazlası (%51.3) erkeklerden oluşmaktadır. Benzer şekilde yanıt örüntüleri gözlenen verilerden oluşan yani yanıt örüntülerinde kayıp veri bulunmayan yanıtlayıcıların yarıdan fazlası da erkeklerden oluşmaktadır. Yanıt örüntülerinde madde düzeyinde kayıp veri bulunan yanıtlayıcıların ise yarıdan fazlası kızdır. Yanıt örüntülerinde birim düzeyinde kayıp bulunan yani hiçbir maddede yanıtlama yapmayan yanıtlayıcı sayısı evren büyüklüğüne göre çok düşük olmakla birlikte, bu yanıtlayıcıların yarıdan fazlası erkeklerden oluşmaktadır.

Kayıp veri içermeyen yanıt örüntüleri dikkate alındığında, doğru yanıt sayıları kızlar lehine, yanlış yanıt sayıları ise erkekler lehine manidar düzeyde farklılaşmaktadır. Kızlar, erkeklere göre daha fazla sayıda doğru yanıtlama ve daha az sayıda yanlış yanıtlama yapmıştır.

Yanıt örüntülerinde madde düzeyinde kayıp veri bulunan yanıtlayıcılar dikkate alındığında, doğru ve yanlış yanıt sayılarında kızlar lehine, kayıp veri içeren madde sayılarında ise erkekler lehine manidar bir farklılık bulunduğu görülmektedir. Kızlar, erkeklere göre daha fazla sayıda doğru ve yanlış yanıtlama yapmıştır. Kızların yanıt örüntülerinde, erkeklerin yanıt örüntülerinde olduğundan daha fazla sayıda kayıp veri bulunduğu anlaşılmaktadır.

Test geneli dikkate alındığında doğru yanıt ve kayıp veri içeren madde sayılarında kızlar lehine, yanlış yanıt sayılarında ise erkekler lehine manidar fark bulunduğu görülmektedir. Testi alan kızlar, test genelinde, erkeklere göre daha fazla sayıda doğru yanıtlama ve daha az sayıda yanlış yanıtlama yapmıştır. Kayıp veri içeren madde sayısı ise kızların yanıt örüntülerinde daha fazladır.

Tüm yanıtlayıcılar ile yanıt örüntüleri kayıp veri içermeyen yanıtlayıcıların profilleri, cinsiyet açısından benzerlik göstermektedir. Her iki grupta da kızların doğru yanıt sayıları erkeklerden fazladır. Erkeklerin yanlış yanıt sayıları ise kizlardan fazladır.

Çizelge 3.7’de verilen çıktılar ve bu çıktılara bağlı değerlendirmeler, testi alan yanıtlayıcılar arasında kızların erkeklere göre doğru olduğuna emin olmadıkları sürece yanıtlama yapmamaya daha eğilimli olduklarını

göstermektedir. Bu durum cinsiyet değişkeninin kayıp veri oluşma olasılığı ile ilişkili olabileceğini göstermektedir.

Kayıp veri mekanizmasını belirleyebilmek amacı ile SBS 2011 Matematik Testi A Kitapçığını alan yanıtlayıcıların yanıt örüntülerinde yapılan bir diğer incelemede 'okul türü' değişkeni dikkate alınmıştır. Yanıtlayıcıların doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının, öğrenim gördükleri okul türüne göre manidar düzeyde fark içerip içermediğine yönelik analiz çıktıları Çizelge 3.8'de verilmektedir. Gözlenen örüntülerde kayıp veri, birim düzeyinde kayıp veri içeren yanıt örüntülerinde ise doğru ve yanlış yanıtlar bulunmadığı için bu örüntülere yönelik parametrelerden ortalama, standart sapma ve Mann-Whitney U Testi değerleri elde edilememiştir.

Çizelge 3.8. Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Okul Türlerine Göre Dağılım

		Okul Türü	N	μ	SH (μ)	σ	Mann-Whitney U Testi*
Gözlenen Veri	Doğru	Devlet	135477	8.30	.01	5.94	54.38
		Özel	9076	17.78	.03	3.22	
	Yanlış	Devlet	135477	11.70	.02	5.94	54.38
		Özel	9076	2.22	.03	3.22	
Kayıp Veri (Madde Düzeyinde)	Doğru	Devlet	374760	5.57	.01	3.86	33.88
		Özel	8177	9.55	.05	4.86	
	Yanlış	Devlet	374760	9.66	.00	3.76	30.38
		Özel	8177	6.65	.04	3.62	
	Kayıp	Devlet	374760	4.77	.01	3.29	14.41
		Özel	8177	3.80	.03	3.05	
Kayıp Veri (Birim Düzeyinde)	Kayıp	Devlet	27	20.00	.00	.00	-
		Özel	0	-	-	-	
Toplam	Doğru	Devlet	510264	6.30	.01	4.66	47.55
		Özel	17253	13.88	.04	5.79	
	Yanlış	Devlet	510264	10.20	.01	4.54	46.73
		Özel	17253	4.32	.03	4.09	
	Kayıp	Devlet	510264	3.51	.01	3.52	23.72
		Özel	17253	1.80	.02	2.83	

* p<0.05

Çizelge 3.8'de görüldüğü gibi testi alan yanıtlayıcıların tamamına yakını (%96.7) devlet okullarında öğrenim görmektedir. Benzer durum gözlenen ve madde düzeyinde kayıp veri içeren yanıt örüntülerinde de bulunmaktadır. Birim düzeyinde kayıp veri içeren yanıt örüntülerine sahip yanıtlayıcı sayısı ise 27 olup bu yanıtlayıcıların tamamı devlet okullarında öğrenim görmektedir.

Gözlenen verilerden oluşan yanıt örüntüleri incelendiğinde, doğru yanıt sayılarında özel okullarda öğrenim gören yanıtlayıcılar lehine, yanlış yanıt sayılarında ise devlet okullarında öğrenim gören yanıtlayıcılar lehine manidar fark bulunduğu anlaşılmaktadır. Özel okullarda öğrenim gören öğrencilerin doğru yanıt sayıları devlet okullarında öğrenim gören öğrencilerden daha fazla, yanlış yanıt sayıları ise daha azdır.

Yanıt örüntülerinde kayıp veri bulunan yanıtlayıcılar dikkate alındığında, gözlenen verilere benzer şekilde, doğru yanıt sayılarında özel okullar lehine, yanlış yanıt sayılarında ise devlet okulları lehine manidar fark bulunduğu anlaşılmaktadır. Özel okullarda öğrenim gören öğrencilerin doğru yanıt sayıları devlet okullarında öğrenim gören öğrencilerden daha fazla, yanlış yanıt sayıları ise daha azdır. Kayıp veri içeren madde sayıları ise devlet okulları lehine manidar düzeyde farklılaşmaktadır. Bu durum devlet okullarında öğrenim gören yanıtlayıcıların yanıt örüntülerinde, özel okullarda öğrenim gören yanıtlayıcılara göre daha fazla sayıda kayıp veri bulunduğunu göstermektedir.

Test genelinde de benzer bir durum söz konusudur. Test genelinde doğru yanıt sayılarında özel okullar lehine, yanlış yanıt sayıları ve kayıp veri içeren madde sayılarında ise devlet okulları lehine manidar fark bulunmaktadır. Özel okullarda öğrenim gören yanıtlayıcıların doğru yanıt sayıları, devlet okullarında öğrenim görenlere göre daha fazla, yanlış yanıt sayıları ve kayıp veri içeren madde sayıları ise daha azdır.

Tüm yanıtlayıcılar ile yanıt örüntüleri kayıp veri içermeyen yanıtlayıcıların profilleri, okul türü açısından benzerlik göstermektedir. Her iki grupta da özel okullarda öğrenim gören yanıtlayıcıların doğru yanıt sayısı devlet okullarında öğrenim gören öğrencilere göre daha fazladır. Yanlış yanıt sayıları açısından ise tam tersi bir durum söz konusudur.

Çizelge 3.8’de verilen çıktılar ve bu çıktılara bağlı değerlendirmelere bağlı olarak, testi alan yanıtlayıcılar arasında özel okullarda öğrenim görenlerin yanıt örüntülerinde kayıp veri oluşma olasılığının, devlet okullarında öğrenim görenlere göre daha düşük olduğu anlaşılmaktadır. Bu durum okul türü değişkeninin, kayıp veri oluşma olasılığı ile ilişkili olabileceğini göstermektedir.

Kayıp veri mekanizmasına yönelik yürütülen bir diğer analizde ‘bölge’ değişkeni dikkate alınmıştır. SBS 2011 Matematik Testi A Kitapçığını alan yanıtlayıcıların doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının, bölgelere göre manidar düzeyde fark içerip içermediği test edilmiştir. Hacimleri dolayısıyla yanıt örüntüleri gözlenen verilerden oluşan yanıtlayıcılara yönelik analiz çıktıları Çizelge 3.9’da, yanıt örüntüleri madde düzeyinde kayıplar içeren yanıtlayıcılara yönelik analiz çıktıları Çizelge 3.10’da, yanıt örüntüleri birim düzeyinde kayıplardan oluşan yanıtlayıcılara yönelik analiz çıktıları Çizelge 3.11’de ve test geneline yönelik çıktılar ise Çizelge 3.12’de verilmektedir. Gözlenen yanıt örüntülerinde kayıp veri, birim düzeyinde kayıp veri içeren yanıt örüntülerinde ise doğru ve yanlış yanıtlar bulunmadığı için bu örüntülere yönelik parametrelerden, ortalama, standart sapma ve Kruskal-Wallis H Testi değerleri elde edilememiştir.

Çizelge 3.9’da görüldüğü gibi yanıt örüntüleri gözlenen verilerden oluşan 144553 yanıtlayıcının büyük çoğunluğu (%19,4) İstanbul Bölgesinde yer almaktadır. İkinci sırada (%14.5) Akdeniz Bölgesi, üçüncü sırada (%13.0) Güneydoğu Anadolu Bölgesi gelmektedir. En az (%2.3) yanıtlayıcının bulunduğu bölge ise Doğu Karadeniz Bölgesidir.

Yanıt örüntüleri gözlenen verilerden oluşan yanıtlayıcılar dikkate alındığında doğru yanıt sayılarının ortalaması 8.89 ve standart sapması 6.24’tür. Yanlış yanıt ortalamaları 11.11 ve standart sapması 6.24’tür. Bu yanıtlayıcıların doğru ve yanlış yanıt sayıları, bölgelere göre manidar düzeyde farklılık göstermektedir.

Çizelge 3.9. Gözlenen Verilerde Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Yaşadıkları Bölgelere Göre Dağılımı

	Bölge	N	μ	SH (μ)	σ	Kruskal- Wallis H Testi*
Doğru	TR1 İstanbul	28106	8.58	.04	6.17	463.72
	TR2 Batı Marmara	5383	9.85	.09	6.15	
	TR3 Ege	15554	9.76	.05	6.50	
	TR4 Doğu Marmara	9286	11.02	.07	6.67	
	TR5 Batı Anadolu	15939	9.42	.05	6.35	
	TR6 Akdeniz	20939	8.88	.04	6.26	
	TR7 Orta Anadolu	7848	9.48	.07	6.35	
	TR8 Batı Karadeniz	6568	10.12	.08	6.52	
	TR9 Doğu Karadeniz	3364	11.19	.11	6.58	
	TRA Kuzeydoğu Anadolu	4602	7.53	.08	5.64	
	TRB Ortadoğu Anadolu	8132	7.44	.06	5.51	
	TRC Güneydoğu Anadolu	18733	6.78	.04	4.97	
	Bilgi Yok	99	7.53	.54	5.33	
	Toplam	144553	8.89	.02	6.24	
Yanlış	TR1 İstanbul	28106	11.42	.04	6.17	463.72
	TR2 Batı Marmara	5383	10.15	.09	6.55	
	TR3 Ege	15554	10.24	.05	6.50	
	TR4 Doğu Marmara	9286	8.98	.07	6.65	
	TR5 Batı Anadolu	15939	10.58	.05	6.35	
	TR6 Akdeniz	20939	11.12	.04	6.26	
	TR7 Orta Anadolu	7848	10.52	.07	6.35	
	TR8 Batı Karadeniz	6568	9.88	.08	6.52	
	TR9 Doğu Karadeniz	3364	8.81	.11	6.58	
	TRA Kuzeydoğu Anadolu	4602	12.47	.08	5.64	
	TRB Ortadoğu Anadolu	8132	12.56	.06	5.51	
	TRC Güneydoğu Anadolu	18733	13.22	.04	4.97	
	Bilgi Yok	99	12.47	.54	5.33	
	Toplam	144553	11.11	.02	6.24	
Kayıp	TR1 İstanbul	28106	.00	.00	.00	-
	TR2 Batı Marmara	5383	.00	.00	.00	
	TR3 Ege	15554	.00	.00	.00	
	TR4 Doğu Marmara	9286	.00	.00	.00	
	TR5 Batı Anadolu	15939	.00	.00	.00	
	TR6 Akdeniz	20939	.00	.00	.00	
	TR7 Orta Anadolu	7848	.00	.00	.00	
	TR8 Batı Karadeniz	6568	.00	.00	.00	
	TR9 Doğu Karadeniz	3364	.00	.00	.00	
	TRA Kuzeydoğu Anadolu	4602	.00	.00	.00	
	TRB Ortadoğu Anadolu	8132	.00	.00	.00	
	TRC Güneydoğu Anadolu	18733	.00	.00	.00	
	Bilgi Yok	99	.00	.00	.00	
	Toplam	144553	.00	.00	.00	

* p<0.05

Çizelge 3.10. Madde Düzeyinde Kayıp İçeren Yanıt Örüntülerinde Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Yaşadıkları Bölgelere Göre Dağılımı

	Bölge	N	μ	SH (μ)	σ	Kruskal- Wallis H Testi*
Doğru	TR1 İstanbul	65974	5.63	.02	3.88	282.35
	TR2 Batı Marmara	14713	6.06	.03	4.08	
	TR3 Ege	47767	5.80	.02	4.04	
	TR4 Doğu Marmara	40085	5.60	.02	4.03	
	TR5 Batı Anadolu	35396	6.28	.02	4.11	
	TR6 Akdeniz	49736	5.82	.02	3.99	
	TR7 Orta Anadolu	21248	5.83	.03	3.96	
	TR8 Batı Karadeniz	24395	5.67	.03	4.00	
	TR9 Doğu Karadeniz	14807	5.44	.03	3.96	
	TRA Kuzeydoğu Anadolu	11109	5.07	.03	3.50	
	TRB Ortadoğu Anadolu	20293	5.11	.03	3.55	
	TRC Güneydoğu Anadolu	37135	5.08	.02	3.45	
	Bilgi Yok	279	5.80	.23	3.88	
	Toplam	382937	5.66	.01	3.92	
Yanlış	TR1 İstanbul	65974	9.58	.02	3.79	742.19
	TR2 Batı Marmara	14713	9.19	.03	3.78	
	TR3 Ege	47767	9.28	.02	3.77	
	TR4 Doğu Marmara	40085	8.83	.02	3.68	
	TR5 Batı Anadolu	35396	9.57	.02	3.86	
	TR6 Akdeniz	49736	9.70	.02	3.82	
	TR7 Orta Anadolu	21248	9.73	.03	3.78	
	TR8 Batı Karadeniz	24395	9.22	.02	3.75	
	TR9 Doğu Karadeniz	14807	8.96	.03	3.70	
	TRA Kuzeydoğu Anadolu	11109	10.35	.04	3.62	
	TRB Ortadoğu Anadolu	20293	10.39	.03	3.61	
	TRC Güneydoğu Anadolu	37135	10.61	.02	3.61	
	Bilgi Yok	279	9.53	.22	3.64	
	Toplam	382937	9.59	.01	3.78	
Kayıp	TR1 İstanbul	65974	4.79	.01	3.31	677.01
	TR2 Batı Marmara	14713	4.75	.03	3.24	
	TR3 Ege	47767	4.92	.02	3.34	
	TR4 Doğu Marmara	40085	5.57	.02	3.56	
	TR5 Batı Anadolu	35396	4.15	.02	2.97	
	TR6 Akdeniz	49736	4.48	.01	3.15	
	TR7 Orta Anadolu	21248	4.44	.02	3.12	
	TR8 Batı Karadeniz	24395	5.11	.02	3.36	
	TR9 Doğu Karadeniz	14807	5.59	.03	3.60	
	TRA Kuzeydoğu Anadolu	11109	4.57	.03	3.11	
	TRB Ortadoğu Anadolu	20293	4.50	.02	3.11	
	TRC Güneydoğu Anadolu	37135	4.31	.02	3.09	
	Bilgi Yok	279	4.67	.20	3.25	
	Toplam	382937	4.75	.01	3.28	

* $p < 0.05$

Çizelge 3.11. Yanıt Örüntüleri Birim Düzeyinde Kayıp Verilerden Oluşan Yanıtlayıcıların Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Bölgelere Göre Dağılımı

	Bölge	N
Doğru	TR1 İstanbul	6
	TR3 Ege	4
	TR4 Doğu Marmara	5
	TR5 Batı Anadolu	1
	TR6 Akdeniz	4
	TR8 Batı Karadeniz	2
	TR9 Doğu Karadeniz	2
	TRB Ortadoğu Anadolu	1
	TRC Güneydoğu Anadolu	2
	Toplam	27
Yanlış	TR1 İstanbul	6
	TR3 Ege	4
	TR4 Doğu Marmara	5
	TR5 Batı Anadolu	1
	TR6 Akdeniz	4
	TR8 Batı Karadeniz	2
	TR9 Doğu Karadeniz	2
	TRB Ortadoğu Anadolu	1
	TRC Güneydoğu Anadolu	2
	Toplam	27
Kayıp	TR1 İstanbul	6
	TR3 Ege	4
	TR4 Doğu Marmara	5
	TR5 Batı Anadolu	1
	TR6 Akdeniz	4
	TR8 Batı Karadeniz	2
	TR9 Doğu Karadeniz	2
	TRB Ortadoğu Anadolu	1
	TRC Güneydoğu Anadolu	2
	Toplam	27

Çizelge 3.12. Doğru ve Yanlış Yanıt Sayıları ile Kayıp Veri İçeren Madde Sayılarının Yanıtlayıcıların Yaşadıkları Bölgelere Göre Dağılımı

	Bölge	N	μ	SH (μ)	σ	Kruskal- Wallis H Testi*
Doğru	TR1 İstanbul	94086	6.51	0.02	4.87	408.58
	TR2 Batı Marmara	20096	7.08	0.04	5.15	
	TR3 Ege	63325	6.77	0.02	5.06	
	TR4 Doğu Marmara	49376	6.62	0.02	5.10	
	TR5 Batı Anadolu	51336	7.26	0.02	5.13	
	TR6 Akdeniz	70679	6.73	0.02	4.97	
	TR7 Orta Anadolu	29096	6.81	0.03	4.99	
	TR8 Batı Karadeniz	30965	6.61	0.03	4.99	
	TR9 Doğu Karadeniz	18173	6.51	0.04	5.08	
	TRA Kuzeydoğu Anadolu	15711	5.79	0.04	4.38	
	TRB Ortadoğu Anadolu	28426	5.77	0.03	4.33	
	TRC Güneydoğu Anadolu	55870	5.65	0.02	4.11	
	Bilgi Yok	378	6.25	0.22	4.37	
	Toplam	527517	6.54	0.01	4.89	
Yanlış	TR1 İstanbul	94086	10.13	0.02	4.71	1479.58
	TR2 Batı Marmara	20096	9.45	0.03	4.70	
	TR3 Ege	63325	9.51	0.02	4.61	
	TR4 Doğu Marmara	49376	8.85	0.02	4.39	
	TR5 Batı Anadolu	51336	9.89	0.02	4.80	
	TR6 Akdeniz	70679	10.12	0.02	4.72	
	TR7 Orta Anadolu	29096	9.94	0.03	4.63	
	TR8 Batı Karadeniz	30965	9.36	0.03	4.49	
	TR9 Doğu Karadeniz	18173	8.94	0.03	4.38	
	TRA Kuzeydoğu Anadolu	15711	10.97	0.04	4.42	
	TRB Ortadoğu Anadolu	28426	11.01	0.03	4.36	
	TRC Güneydoğu Anadolu	55870	11.48	0.02	4.30	
	Bilgi Yok	378	10.30	0.22	4.34	
	Toplam	527517	10.01	0.01	4.64	
Kayıp	TR1 İstanbul	94086	3.36	0.01	3.54	1098.20
	TR2 Batı Marmara	20096	3.47	0.03	3.48	
	TR3 Ege	63325	3.72	0.01	3.60	
	TR4 Doğu Marmara	49376	4.52	0.02	3.88	
	TR5 Batı Anadolu	51336	2.86	0.01	3.12	
	TR6 Akdeniz	70679	3.15	0.01	3.34	
	TR7 Orta Anadolu	29096	3.25	0.02	3.31	
	TR8 Batı Karadeniz	30965	4.02	0.02	3.63	
	TR9 Doğu Karadeniz	18173	4.56	0.03	3.92	
	TRA Kuzeydoğu Anadolu	15711	3.23	0.03	3.35	
	TRB Ortadoğu Anadolu	28426	3.21	0.02	3.33	
	TRC Güneydoğu Anadolu	55870	2.87	0.01	3.24	
	Bilgi Yok	378	3.45	0.18	3.47	
	Toplam	527517	3.45	0.01	3.51	

* p<0.05

Çizelge 3.9'da görüldüğü gibi yanıt örüntüleri gözlenen verilerden oluşan yanıtlayıcılarda bölgelere göre belirlenen manidar farkın kaynağı çoklu karşılaştırma testlerinden Tamhane testi ile belirlenmeye çalışılmıştır. Tamhane testi çıktıları, Batı Marmara ve Ege, Doğu Marmara ve Doğu Karadeniz, Batı Anadolu ve Orta Anadolu, Kuzeydoğu Anadolu ve Ortadoğu Anadolu bölgeleri arasındaki farklar dışında diğer tüm farkların istatistiksel olarak manidar olduğunu göstermektedir. Buna göre doğru yanıt sayıları en fazla olan yanıtlayıcılar Doğu Karadeniz (11.19) ve Doğu Marmara (11.02) bölgelerinden, en az olan yanıtlayıcılar ise Güneydoğu Anadolu Bölgesi (6.78)'ndendir. Yanlış yanıt sayılarının en yüksek olduğu bölge Güneydoğu Anadolu Bölgesi (13.22), en az olduğu bölgeler ise Doğu Karadeniz (8.81) ve Doğu Marmara (8.98) bölgeleridir.

Çizelge 3.10'da görüldüğü gibi yanıt örüntüleri madde düzeyinde kayıp veri içeren 382937 yanıtlayıcının büyük çoğunluğu (%17.2) İstanbul Bölgesindendir. İkinci sırada (%13.0) Akdeniz Bölgesi, üçüncü sırada (%12.5) Ege Bölgesi gelmektedir. En az yanıtlayıcının (%3.8) bulunduğu bölge ise Batı Marmara Bölgesidir.

Yanıt örüntüleri madde düzeyinde kayıp veri içeren yanıtlayıcıların doğru yanıt sayılarının ortalaması 5.66 ve standart sapması 3.92'dir. Yanlış yanıt ortalaması 9.59 ve standart sapması 3.78'dir. Kayıp veri içeren madde sayısı ortalaması 4.75 ve standart sapması 3.28'dir. Yanıt örüntüleri madde düzeyinde kayıp veri içeren yanıtlayıcıların doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayıları, bölgelere göre manidar düzeyde farklılık göstermektedir. Çoklu karşılaştırma testi çıktıları, İstanbul ve Doğu Marmara, İstanbul ve Batı Karadeniz, Ege ve Akdeniz, Ege ve Orta Anadolu, Akdeniz ve Orta Anadolu, Kuzeydoğu Anadolu ve Ortadoğu Anadolu, Kuzeydoğu Anadolu ve Güneydoğu Anadolu bölgeleri arasındaki farklar dışında diğer tüm farkların istatistiksel olarak manidar olduğunu göstermektedir. Buna göre doğru yanıt sayıları en fazla olan yanıtlayıcılar Batı Anadolu Bölgesi (6.28)'nden, en az olan yanıtlayıcılar ise Güneydoğu Anadolu (5.08), Ortadoğu Anadolu (5.11) ve Kuzeydoğu Anadolu (5.07) bölgelerindendir. Yanlış yanıt sayılarının en yüksek olduğu bölgeler Güneydoğu Anadolu (10.61), Ortadoğu Anadolu (10.39) ve Kuzeydoğu Anadolu (10.35) bölgeleri, en az olduğu bölge ise Doğu Marmara Bölgesi (8.83)'dir. Kayıp veri içeren

madde sayılarının en yüksek olduğu bölge Doğu Karadeniz (5.59), en az olduğu bölge ise Batı Anadolu Bölgesi (4.15)'dir.

Çizelge 3.11'de görüldüğü gibi yanıt örüntüleri birim düzeyinde kayıp verilerden oluşan 27 yanıtlayıcının 6'sı İstanbul Bölgesinden, 5'i Doğu Marmara Bölgesindendir. Batı Marmara, Orta Anadolu ve Kuzeydoğu Anadolu bölgelerinden ise yanıt örüntüsü birim düzeyinde kayıp verilerden oluşan yanıtlayıcı bulunmamaktadır.

Çizelge 3.12'de görüldüğü gibi göre testin genelinde, 527517 yanıtlayıcının büyük çoğunluğu (%17.8) İstanbul Bölgesindendir. Bunu sırasıyla Akdeniz Bölgesi (%13.4) ve Ege Bölgesi (%12.0) izlemektedir. En az yanıtlayıcının bulunduğu bölgeler ise Kuzeydoğu Anadolu Bölgesi (%3.0) ve Doğu Karadeniz Bölgesi (%3.5)'dir.

Testin genelinde doğru yanıt ortalaması 6.54 ve standart sapması 4.89'dur. Yanlış yanıt sayılarının ortalaması 10.01 ve standart sapması 4.64'tür. Kayıp veri içeren madde sayısı ortalaması 3.45 ve standart sapması 3.51'dir. Test genelinde yanıtlayıcıların doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayıları, bölgelere göre manidar düzeyde farklılık göstermektedir. Çoklu karşılaştırma testlerinin çıktıları, İstanbul ve Doğu Karadeniz, Ege ve Akdeniz, Ege ve Orta Anadolu, Doğu Marmara ve Batı Karadeniz, Kuzeydoğu Anadolu ve Ortadoğu Anadolu bölgeleri arasındaki farklar dışında diğer tüm farkların istatistiksel olarak manidar olduğunu göstermektedir. Buna göre test genelinde doğru yanıt sayıları en fazla olan yanıtlayıcılar Batı Anadolu Bölgesi (7.26)'nden, en az olan yanıtlayıcılar ise Güneydoğu Anadolu Bölgesi (5.65)'ndendir. Yanlış yanıt sayılarının en yüksek olduğu bölge Güneydoğu Anadolu Bölgesi (11.48), en az olduğu bölge ise Doğu Marmara Bölgesi (8.85)'dir. Kayıp veri içeren madde sayılarının en yüksek olduğu bölge Doğu Karadeniz Bölgesi (4.56), en az olduğu bölge ise Batı Anadolu Bölgesi (2.86)'dir.

Tüm yanıtlayıcılar ile yanıt örüntüleri kayıp veri içermeyen yanıtlayıcıların profilleri, bölge değişkeni açısından farklılık göstermektedir. Her iki grupta da en az doğru yanıt ve en fazla yanlış yanıtlama yapılan bölge Güneydoğu Anadolu Bölgesidir. Diğer taraftan test genelinde en fazla doğru yanıtlama yapılan bölgeler Batı Marmara ve Batı Anadolu bölgeleri iken yanıt örüntüleri kayıp veri içermeyen yanıtlayıcılar dikkate alındığında, en fazla

doru yanıtlama yapılan bölgeler Doğu Karadeniz ve Doğu Marmara bölgeleridir.

Bölge değişkeni dikkate alınarak yürütülen analizler, doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarının, gerek test genelinde gerekse kayıp veri içermeye durumu dikkate alınarak gruplandırılmış yanıt örüntüleri düzeyinde, bölgelere göre manidar düzeyde farklılaştığını göstermektedir. Bu durum bölge değişkeninin, kayıp veri oluşma olasılığı ile ilişkili olabileceğini göstermektedir.

Bu çalışmada ele alınan veri setinde yer alan kayıp verilerin konumu üzerine yürütülen önceki çalışmalara bağlı olarak, kayıp veri örüntüsünün genel örüntü yapısına daha uygun olduğu belirlenmişti. Bu yapı, kayıp verilerin tam seçkisiz ya da en azından seçkisiz olduğu yönünde bir kanıt olarak değerlendirilebilir görülmektedir. Diğer taraftan kayıp veri mekanizması üzerine üç temel değişken kullanılarak yürütülen analizler, tüm yanıtlayıcılar ile yanıt örüntüleri kayıp veri içermeyen yanıtlayıcıların profilleri arasında belirgin farklar olmamakla birlikte, bu değişkenler ile kayıp veri oluşma olasılığı arasında manidar ilişkiler bulunabileceğini göstermektedir. Dikkate alınamamış olmakla birlikte, test ortamı, düzeltme kuralı, maddelerin teste yerleştirilme durumu gibi değişkenler ile kayıp veri oluşması olasılığı arasında da manidar ilişkilerin belirlenme olasılığı bulunmaktadır. Bu durum kayıp veri mekanizmasının tam seçkisiz olarak oluşmadığı yönünde şüpheyandırmakta ve TSK varsayımının ihlâl edilmekte olabileceğini göstermektedir.

Zayıf ve yeterli olmayan bir kanıt olmakla birlikte TSK varsayımının test edilmesinde yaygın şekilde kullanılan testlerden biri, 'Little's MCAR Test' olarak bilinmektedir. Bu çalışmada kullanılan veri seti üzerinde yürütülen 'Little's MCAR Test' kestirimleri ($\chi^2=33719.52$ ve $p<0.05$) de TSK varsayımının ihlâl edilmiş olabileceğini göstermektedir.

Kayıp veri mekanizmasını açıklayan diğer varsayımlar olan SK ve SOK varsayımlarının istatistiksel olarak test edilmesi mümkün değildir. Bununla birlikte her bir madde düzeyinde kayıp veri bulunması, her bir maddede hem 1 hem 0 puanlamaların bulunması, kayıp veri miktarının madde sırasına göre düzenli bir değişim göstermemesi, kayıp verilerin genel örüntü yapısında ve dağınık şekilde konumlanmaları, tüm yanıtlayıcılar ile

yanıt örüntüleri kayıp veri içermeyen yanıtlayıcıların profilleri arasında benzerlikler bulunması gibi bulgular dikkate alınarak, kayıp veri mekanizmasının tam seçkisiz olmasa da seçkisiz olduğu ileri sürülebilir. Bu durumda TSK varsayımı karşılanmamakla birlikte SK varsayımının karşılanmakta olduğu belirlenebilir görülmektedir.

TSK ve SK varsayımlarının karşılanmaması, kayıp veri mekanizmasının SOK varsayımı ile açıklanmakta olduğunun bir göstergesidir. SOK varsayımı karşılandığında ise ancak çok özel durumlara yönelik oldukça sınırlı kayıp veri yöntemlerinin kullanılması mümkündür. Ayrıca bu durumda yapılan ölçmenin ve kullanılan ölçme aracının yeniden gözden geçirilmesi gerekli görülmektedir.

SBS 2011 Matematik Testi A Kitapçığı verilerinde, kayıp veri mekanizmasına yönelik analiz ve incelemeler, kayıp verilerin oluşmasını açıklayan mekanizmanın TSK varsayımını karşılamadığı fakat SK varsayımını karşılama olasılığının yüksek olduğunu göstermektedir. Rubin (1987), Little ve Rubin (1987), Schaffer (1997), Allison (2002) gibi birçok araştırmacı, TSK varsayımının ihlâli durumunda, özellikle silme ve basit atamaya dayalı yöntemlerin kullanımının, kestirimlerde yanlılığa yol açtığını belirlemiştir. Diğer taraftan en çok olabilirlik ve çoklu atamaya dayalı kayıp veri yöntemlerinin tamamına yakınında TSK varsayımı aranmamakta, bu yöntemlerin kullanılabilmesi için SK varsayımının karşılanması yeterli görülmektedir.

Yapı Geçerliliği ve Tek Boyutluluk Özelliği

Bu çalışmada SBS 2011 Matematik Testi A Kitapçığının yapı geçerliliği ve tek boyutluluk özelliğine yönelik analizler, Açıklayıcı Faktör Analizi (AFA) ve Doğrulayıcı Faktör Analizi (DFA) kullanılarak gerçekleştirilmiştir. Farklı kayıp veri yöntemlerine göre elde edilen eksiksiz veri setleri üzerinde yürütülen AFA ve DFA çıktılarına bağlı olarak testin, tek boyutluluk özelliğine sahip olup olmadığının ve madde puanlarının toplanabilir olup olmadığının belirlenmesi amaçlanmıştır. AFA çıktıları Çizelge 3.13 ve Çizelge 3.14'te, DFA ölçme modeli Şekil 3.1'de ve DFA çıktıları ise Çizelge 3.15'te verilmektedir.

Çizelge 3.13. Kayıp Veri Yöntemlerine Göre Yapı Geçerliliği Parametreleri (Açımlayıcı Faktör Analizi - 1)

Kayıp Veri Yöntemi	KMO	Bartlett* (χ^2)	Oransal Değişim Değerleri		Faktör Sayısı	Faktörler Arası Korelasyon**	Faktör Özdeğerleri (>1.00)	Faktör Madde Sayıları	Toplam Açıklanan Varyans (%)		Faktör Yükleri		Model Uyumu İyiliği (χ^2 /sd)
			Aralık	k (>.30)					Aralık	Madde Sayısı			
Silme	.976	10773.3	.140-.521	15	1	-	8.286	20	41.43	.37-.76	0	20	1.97
0 Atama	.955	21402.7	.082-.364	5	2	.62	f_1 :5.813 f_2 :1.261	f_1 :20 f_2 :0	35.37	.30-.64	2	19	2.41
							f_1 :7.247 f_2 :1.042	f_1 :20 f_2 :0					
BO	.970	9073.0	.100-.466	13	2	.09	f_1 :5.478 f_2 :1.208	f_1 :19 f_2 :1	41.44	.32-.71	0	20	1.11
SO	.952	19160.2	.070-.360	3	2	.59	f_1 :5.338 f_2 :1.213	f_1 :19 f_2 :1	33.43	.28-.64	1	17	2.13
YNO	.950	18267.1	.070-.350	3	2	.59	f_1 :5.332 f_2 :1.482	f_1 :18 f_2 :2	32.75	.29-.63	1	16	2.00
YNM	.944	19229.5	.053-.362	5	2	.44	f_1 :5.185 f_2 :1.209	f_1 :19 f_2 :1	34.07	.19-.65	6	18	2.67
Dİ	.947	17298.6	.067-.335	2	2	.60	f_1 :5.478 f_2 :1.208	f_1 :19 f_2 :1	31.97	.28-.62	1	16	1.95
DE	.952	119158.4	.070-.360	3	2	.59	f_1 :5.627 f_2 :1.255	f_1 :19 f_2 :1	33.43	.27-.64	1	17	2.13
RA	.951	20459.9	.069-.387	5	2	.59	f_1 :6.280 f_2 :1.279	f_1 :19 f_2 :1	34.41	.27-.66	1	17	2.73
BM	.957	25092.2	.079-.433	10	2	.59	f_1 :5.412 f_2 :1.250	f_1 :20 f_2 :0	37.80	.28-.69	1	17	3.75
VÇ	.952	193209.3	.067-.346	3	2	.61	f_1 :5.597 f_2 :1.259	f_1 :19 f_2 :1	33.31	.26-.63	0	17	14.28
Çoklu Atama	.950	101528.4	.069-.380	5	2	.59			34.28	.29-.65	1	17	14.54

* sd=190, p=.000

** Oblique döndürme ve Kaiser normalleştirilmesinden sonra faktörler arası korelasyon değerleri.

Çizelge 3.14. Kayıp Veri Yöntemlerine Göre Yapı Geçerliliği Parametreleri (Açımlayıcı Faktör Analizi - 2)*

Kayıp Veri Yöntemi	KMO	Bartlett** (χ^2)	Oransal Değişim Değerleri		Özdeğer	Açıklanan Varyans (%)	Faktör Yükleri		Model Uyumu iyiliği (χ^2/sd)
			Aralık	k (>.30)			Aralık	k (>.30)	
Silme	DS	.976	10773.3	.140 - .521	15	8.286	.37 - .76	20	1.97
Basit Atama	0 Atama	.955	21402.7	.082 - .364	5	5.813	.30 - .64	19	6.92
	BO	.970	9073.0	.100 - .466	13	7.247	.32 - .71	20	1.84
	SO	.952	19160.2	.070 - .360	3	5.478	.27 - .64	17	2.81
	YNO	.950	18267.1	.070 - .350	3	5.338	.27 - .63	16	5.05
	YNM	.944	19229.5	.053 - .362	5	5.332	.19 - .65	18	9.23
	DI	.947	17298.6	.067 - .335	2	5.185	.29 - .62	16	4.82
	DE	.952	119158.4	.070 - .360	3	5.478	.27 - .64	17	5.30
	RA	.951	20459.9	.069 - .387	5	5.627	.27 - .66	17	7.08
En Çok Olabilirlik	BM	.957	25092.2	.079 - .433	10	6.280	.28 - .69	17	9.41
	VÇ	.952	193209.3	.067 - .346	3	5.412	.27 - .63	17	54.93
Çoklu Atama	MCMC	.950	101528.4	.068 - .380	5	5.597	.27 - .66	17	35.76

* Faktör sayısı 1 ile sınırlandırılmıştır.

** sd=190, p=.000

Çizelge 3.13 ve Çizelge 3.14'te yer alan KMO ve Bartlett testlerine yönelik değerlendirmeler önceki bölümlerde yapılmıştı. Oransal değişim (communalities) değerleri, her bir maddenin varyansını vermektedir. Bu değerlerin .30 ve üzerinde olması beklenir. Oransal değişim değerleri, faktör yükleri kestirimlerinden bağımsız olarak hesaplanmaktadır. Bu nedenle KMO ve Bartlett istatistiklerinde olduğu gibi bu istatistiklerde de faktör sınırlaması olduğunda ya da olmadığında aynı kestirimler üretilmektedir. Bu çalışmada, Çizelge 3.13 ve Çizelge 3.14'te yer alan kestirimler dikkate alındığında, oransal değişim değerleri .30 ve üzerinde olan madde sayısının, DS yönteminde 15, BO yönteminde 13, BM yönteminde 10, YNM ve RA yöntemlerinde 5, Dİ yönteminde 2, diğer yöntemlerde ise 3 olduğu görülmektedir. Oransal değişim değeri .30 ve üzerinde olan madde sayısı, en yüksek DS yönteminde en düşük ise Dİ yönteminde elde edilmektedir.

Faktör sayısı sınırlandırılmadığında (Çizelge 3.13), DS yönteminde 1 faktörlü, diğer yöntemlerde ise 2 faktörlü yapı elde edilmektedir. DS yöntemi dışındaki yöntemlerde, 2 faktörün özdeğerleri arasındaki fark 4 katın üzerindedir. Özdeğerler arasındaki farkın en yüksek olduğu yöntemler BO ve BM yöntemleridir. Faktörler arası korelasyonun en düşük olduğu yöntem BO yöntemi, en yüksek olduğu yöntemler ise 0 Atama ve VÇ yöntemleri olarak belirlenmektedir.

Faktör sayısı sınırlandırılmadığında (Çizelge 3.13), DS yöntemi dışındaki diğer yöntemlerden 0 Atama, BO ve VÇ yöntemlerinde tüm maddeler birinci faktörde toplanmaktadır. Bu yöntemlerde 13, 6, 8 ve 4. maddeler, her iki faktörde de .30'a yakın faktör yüklerine sahip çıkmaktadır. Bunların dışındaki diğer yöntemlerin tamamında ise birinci faktörde 19 madde, 2 faktörde sadece 1 madde yer almaktadır. İkinci faktörlerde yer alan bu madde, her bir yöntem için testin 13. maddesidir ve faktör yükleri açısından binişiklik oluşturmaktadır.

Faktör sayısı 1 ile sınırlandırıldığında (Çizelge 3.14), kayıp veri yöntemlerine göre faktör özdeğerleri 5.00'in üzerinde elde edilmektedir. En yüksek faktör özdeğeri DS yöntemi (8.29), BO yöntemi (7.25) ve BM yöntemi (6.28) kullanılarak kestirilmektedir. En düşük faktör özdeğeri ise Dİ yöntemi (5.185) ile elde edilmektedir.

Faktör sayısı sınırlandırılmadığında (Çizelge 3.13), kayıp veri yöntemlerine göre açıklanan toplam varyans yüzdeleri %31 ile %42 arasında değişmektedir. En yüksek açıklanan varyans DS yöntemi (%41.43) ve BO yöntemi (%41.44) kullanılması durumunda elde edilmektedir. Bunları BM yöntemi (%37.80) izlemektedir. En düşük açıklanan varyans ise Dİ yöntemi (%31.97) ile elde edilmektedir.

Faktör sayısı 1 ile sınırlandırıldığında (Çizelge 3.14), beklenildiği gibi, DS yöntemi dışında diğer yöntemlerin kullanılması durumunda, açıklanan toplam varyanslarının düştüğü görülmektedir. Açıklanan toplam varyans yüzdeleri, DS yöntemi dışında %25 ile %37 arasında değişmektedir. En yüksek açıklanan varyans, maddelerin tek faktörde toplandığı DS yönteminde (%41.43). İkinci sırada BO yöntemi (%36.23) ve üçüncü sırada ise BM yöntemi (%31.40) gelmektedir. En düşük açıklanan varyans kestirimi ise yine Dİ yöntemi (%25.92) ile elde edilmektedir.

AFA uygulamasında maddeler düzeyinde faktör yük değerlerinin .30 üzerinde olması ve ayrıca maddelerin binişiklik oluşturmaması aranan özelliklerdir.

Faktör sınırlandırması yapılmadığında (Çizelge 3.13), kullanılan kayıp veri yöntemlerine göre maddelerin faktör yük değerlerinin farklılaştığı görülmektedir. En yüksek faktör yük değerleri DS ve BO yöntemlerinde, en düşük faktör yük değerleri ise Dİ ve YNM yöntemlerinde elde edilmektedir. Faktör sayısı 1 ile sınırlandırıldığında (Çizelge 3.14) ise MCMC yöntemi dışında diğer yöntemlerde faktör yük değerlerinin yaklaşık benzer kaldığı görülmektedir. MCMC yönteminde faktör yükü alt sınırı, .29'dan .26'ya düşmektedir. Bu değişim veri setinin büyüklüğü ile ilişkili görülmektedir.

Faktör yük değerleri, faktör sayısı ve sınırlandırmasından bağımsız olarak hesaplanmaktadır. Bu nedenle gerek faktör sayısı sınırlandırılmadığında gerekse 1 ile sınırlandırıldığında, faktör yük değeri .30 ve üzerinde olan madde sayısı DS ve BO yöntemlerinde 20, 0 Atama yönteminde 19, YNM yönteminde 18, YNO ve Dİ yöntemlerinde 16 ve diğer yöntemlerde 17 olarak belirlenmektedir.

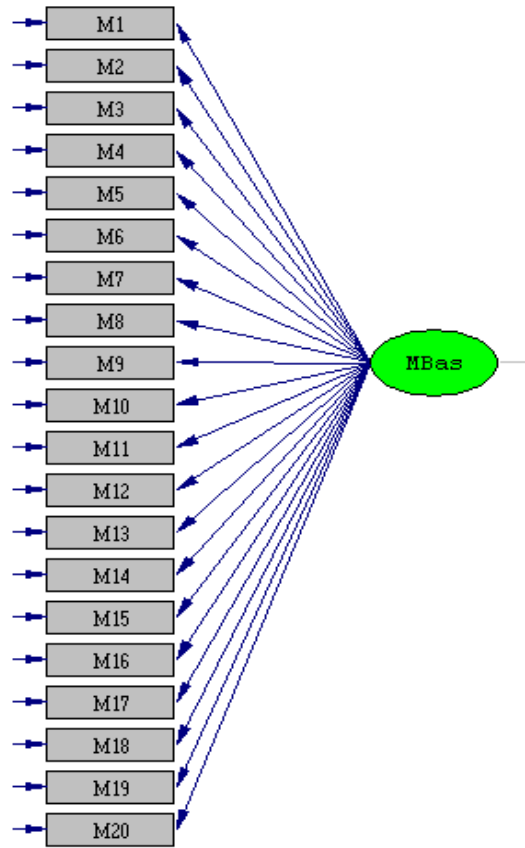
Faktör sınırlandırması yapılmadığında (Çizelge 3.13), DS, BO ve VÇ yöntemleri dışındaki yöntemlerin kullanılması durumunda binişik madde oluştuğu görülmektedir. YNM yönteminde 6, 0 Atama yönteminde 2, diğer

yöntemlerde ise 1 binişik madde bulunmaktadır. Faktör sayısı 1 ile sınırlandırıldığında ise maddeler tek faktör altında toplanmaya zorlandığı için doğal olarak birden fazla faktörde faktör yük değeri görülmesi mümkün olmamaktadır.

Faktör sayısı sınırlandırılmadığında (Çizelge 3.13), VÇ ve MCMC yöntemlerinin model uyumunu sağlamadığı görülmektedir. BM yönteminde orta düzeyde diğer yöntemlerde ise mükemmel ya da mükemmele yakın bir uyum belirlenebilmektedir. Faktör sayısı 1 ile sınırlandırıldığında (Çizelge 3.14), DS yöntemi dışında diğer yöntemlerde uyum iyiliği düzeyi düşmektedir. Bu durumda DS, BO, SO ve Dİ yöntemlerinde mükemmele yakın bir uyum görülmektedir. VÇ ve MCMC yöntemlerine göre model uyumu sağlanmamaktadır. Diğer yöntemlerde ise model uyumu orta ya da zayıf düzeydedir.

Şekil 3.1’de görüldüğü gibi bu çalışmada DFA için üretilen ölçme modeli, ‘matematik başarısı’ olarak tanımlanan bir gizil değişken ve test maddeleri ile tanımlanan 20 gözlenen değişken içermektedir.

Çizelge 3.15’te verilen DFA kestirimleri, standartlaştırılmış yol katsayılarının, kullanılan kayıp veri yöntemine göre değişmekle birlikte genel olarak .26 ile .76 arasında olduğunu göstermektedir. Uyum iyiliği yüksek düzeyde olan bir modelde yol katsayılarının .30 ve üzerinde olması beklenir. Bu değerlerin yüksek olması, gözlenen değişkenin gizil değişken üzerindeki açıklanan varyans değerinin ve buna bağlı olarak açıklama gücünün yüksek olduğunu göstermektedir. Bu çalışmada DS, O Atama, BO ve YNM yöntemlerinin kullanılması durumunda standartlaştırılmış yol katsayıları .30 ve üzerinde kestirilmektedir. Diğer yöntemlerin kullanılması durumunda ise modelde yol katsayısı .30’un altında olan gözlenen değişkenler bulunduğu anlaşılmaktadır. En yüksek yol katsayıları DS yönteminde, en düşük yol katsayıları ise MCMC ve VÇ yöntemlerinde elde edilmektedir. Yol katsayılarının manidarlığına yönelik t değerlerinin tamamı 2.56’nın üzerindedir. Bu durum gözlenen ve gizil değişkenler arasındaki tüm yolların istatistiksel olarak .01 düzeyinde manidar olduğunu göstermektedir.



Şekil 3.1. SBS 2011 Matematik Testi A Kitapçığı Verilerine Yönelik Ölçme Modeli (Conceptual Diagram)

Çizelge 3.15. Kayıp Veri Yöntemlerine Göre DFA Model Uyum İyiliği Parametreleri

Kayıp Veri Yöntemi	Yol Katsayıları Aralığı**	Hata Varyansları Aralığı**	χ^2/sd^{***}	NFI	GFI	AGFI	RMSEA	RMR	SRMR
DS	.37 - .76	.42 - .86	2.09	.99	.98	.97	.028	.0051	.021
0 Atama	.30 - .64	.59 - .91	8.36	.98	.97	.97	.038	.0056	.028
BO	.32 - .71	.50 - .90	2.02	.99	.98	.97	.026	.0050	.022
SO	.27 - .63	.60 - .93	6.15	.98	.98	.98	.032	.0042	.024
YNO	.29 - .64	.59 - .92	7.89	.98	.97	.97	.037	.0055	.027
YNM	.30 - .65	.58 - .96	11.76	.97	.96	.95	.046	.0073	.036
Dİ	.27 - .62	.62 - .93	5.58	.98	.98	.98	.030	.0046	.024
DE	.27 - .63	.60 - .93	6.13	.98	.98	.98	.032	.0042	.024
RA	.26 - .65	.57 - .93	8.41	.98	.97	.97	.038	.0058	.028
BM	.28 - .69	.53 - .92	11.34	.98	.96	.96	.045	.0056	.031
VÇ*	.26 - .63	.60 - .93	20.30	1.00	.99	.99	.019	.0031	.013
MCMC*	.26 - .66	.57 - .93	21.81	.99	.99	.98	.029	.0044	.020

* Model için önerilen modifikasyonlar dikkate alınarak hata varyansları ilişkilendirilmiştir.

** Standartlaştırılmış değerler dikkate alınmıştır ($t > 2.56$, $p < .01$).

*** $p = .000$ hesaplanmıştır.

Uyum iyiliği yüksek bir DFA modelinde standartlaştırılmış artık değerlerin yani hata varyanslarının .90'ın altında, tercihen .60 civarında olması beklenir. Çizelge 3.15'te görüldüğü gibi bu çalışmada kullanılan ölçme modelinde, kayıp veri yöntemlerine göre değişmekle birlikte, hata varyansları .42 ile .93 arasında kestirilmektedir. En düşük hata varyansları DS ve BO yöntemlerinde, en yüksek hata varyansları ise Dİ, DE, RA, VÇ ve MCMC yöntemlerinde elde edilmektedir.

Kurulan ölçme modeline yönelik olarak her bir kayıp veri yönteminde hesaplanan χ^2 değerleri istatistiksel olarak manidardır. Uyum iyiliğini değerlendirmek için kullanılan χ^2/sd oranlarının 2 civarında ve 3'ten küçük olması mükemmel derecede uyumu, 5 civarında ve 5'ten küçük olması orta düzeyde uyumu, 10'dan küçük olması ise zayıf düzeyde bir uyumu göstermektedir. Bu çalışmada, kestirilen χ^2/sd oranlarına göre DS ve BO yöntemlerinin kullanılması durumunda mükemmel düzeyde bir model-veri uyumu elde edilmektedir. YNM, BM, VÇ ve MCMC yöntemlerinde model-veri uyumu bulunmamaktadır. Diğer yöntemlerin ise zayıf düzeyde model-veri uyumu sağlayabildiği görülmektedir.

Uyum iyiliği parametrelerinden NFI, GFI ve AGFI değerleri, genel olarak .97 ve üzerindedir. Bu değerler her bir kayıp veri yönteminde mükemmel düzeyde model-veri uyumunu sağlandığına yönelik kanıt oluşturmaktadır. En yüksek kestirimler VÇ yönteminde, en düşük kestirimler ise YNM yönteminde elde edilmektedir.

Uyum iyiliğine yönelik hata parametrelerinden RMSEA, RMR ve SRMR değerleri genel olarak .05'in altında ve 0'a oldukça yakındır. Bu değerler de her bir kayıp veri yöntemine göre mükemmel düzeyde model-veri uyumunun sağlandığını göstermektedir. En düşük kestirimler VÇ yönteminde, en yüksek kestirimler ise YNM yönteminde elde edilmektedir.

SBS 2011 Matematik Testi A Kitapçığı verileri üzerinde farklı kayıp veri yöntemlerine göre yürütülen AFA ve DFA sonuçlarına göre silme ve basit atamaya dayalı yöntemlerin kullanılması durumunda açıklanan varyans yüzdeleri ve veri-model uyum düzeyi, EÇÖ ve VÇA yöntemlerine göre daha yüksek kestirilmektedir. Faktör yük değerleri açısından ele alındığında ise

basit atama yöntemleri, uç kestirimleri yani en düşük ve en yüksek kestirimleri üretmektedir.

Çokluk ve Kayrı (2011), basit atama yöntemlerinin, faktör yapılarını açıklamada yeterli performans gösterdiğini fakat açıklanan varyans yüzdelerinde, faktör yüklerine yönelik özdeğerlerde ve güvenirlik katsayılarında, gerçek değerlere göre düşüşe yol açtığını belirlemiştir. Bu çalışmada ise silme ve basit atama yöntemlerinin, açıklanan varyans yüzdelerini ve buna bağlı olarak veri-model uyum düzeyini artırma eğiliminde olduğu görülmektedir. Bu farklılığın, bu çalışmada kullanılan veri setinin iki kategorili puanlanan verilerden oluşması ve basit atama yöntemlerinin, genel olarak ortalama ve standart sapma değerlerini düşürerek verilerin benzeşiklik düzeyini artırmasından kaynaklandığı düşünülmektedir.

SBS 2011 Matematik Testi A Kitapçığına yönelik veri seti üzerinde yapı geçerliği çalışmaları olarak yürütülen AFA ve DFA sonuçları, her bir kayıp veri yönteminde de maddelerin tek bir faktör altında toplanabileceğini, tek faktörlü yapının mükemmel yakın düzeyde model-veri uyumunu sağladığını göstermektedir. Bu durum testin tek boyutluluk özelliğine sahip olduğuna ve ileri analizlerde toplam puanların kullanılabileceğine dair bir kanıt oluşturmaktadır. Bununla birlikte genel olarak silme ve basit atamaya dayalı yöntemlerin, en çok olabilirlik ve çoklu atamaya dayalı yöntemlere göre AFA ve DFA'da daha iyi performans gösterdiği ileri sürülebilir. Ayrıca AFA sonuçlarına bağlı olarak dört maddenin (sırasıyla 13, 6, 8 ve 4. maddeler) gözden geçirilmesinde gerektiği düşünülmektedir.

Madde Parametreleri

Bu çalışmada SBS 2011 Matematik Testi A Kitapçığına yönelik olarak, farklı kayıp veri yöntemleri ile elde edilen eksiksiz veri setleri üzerinde, klasik test kuramı temelinde madde parametre kestirimleri gerçekleştirilmiştir. Madde parametreleri olarak, madde güçlük indeksi, madde ayırıcılık indeksi, nokta çift serili ve çift serili korelasyon katsayıları ile madde güvenirlikleri dikkate alınmıştır.

İlk olarak madde güçlük indeksi kestirimleri gerçekleştirilmiştir. Bir testte yer alacak maddelerin güçlük düzeyinin ne olması gerektiği, testin amacına göre değerlendirilmektedir. Genellikle .20 ile .80 arasındaki güçlük değerleri, maddenin teste alınabilir olduğuna yönelik bir kanıt sağlamaktadır. Yanıtlayıcının maksimum performans göstermesi beklenen testlerde, güçlük indeksleri görece daha düşük yani ‘zor’ ya da ‘çok zor’ maddeler kullanılabilmektedir. Eğitim alanında özellikle öğrenci başarısını belirlemeye yönelik testlerde ise maddelerin çoğunluğunun güçlüklerinin orta düzeyde ve .50 civarında olması önerilmektedir (Lord ve Novick, 1968; Özçelik, 1981; Crocker ve Algina, 1986; Baykul, 2000; Tekin, 2007). Madde güçlük indeksi kestirimleri Çizelge 3.16’da verilmektedir.

Çizelge 3.16. Kayıp Veri Yöntemlerine Göre Madde Güçlük İndeksleri

Kayıp Veri Yöntemleri												
	0											
Madde	DS	Atama	BO	SO	YNO	YNM	Dİ	DE	RA	BM	VÇ	MCMC
1	.38	.25	.28	.25	.25	.25	.32	.25	.29	.27	.29	.29
2	.40	.25	.28	.25	.25	.25	.31	.25	.28	.26	.28	.28
3	.46	.28	.35	.28	.29	.28	.44	.28	.39	.34	.39	.39
4	.53	.30	.38	.30	.32	.30	.49	.30	.45	.38	.44	.45
5	.54	.45	.47	.67	.48	.67	.61	.67	.55	.53	.55	.55
6	.31	.23	.25	.23	.23	.23	.26	.23	.25	.23	.26	.25
7	.52	.41	.45	.63	.43	.63	.56	.63	.52	.50	.52	.52
8	.55	.45	.47	.59	.46	.59	.55	.59	.52	.51	.52	.52
9	.41	.29	.34	.29	.30	.29	.35	.29	.32	.31	.32	.32
10	.46	.36	.41	.36	.37	.36	.45	.36	.42	.40	.42	.42
11	.49	.35	.40	.35	.37	.35	.50	.35	.44	.42	.44	.45
12	.44	.31	.34	.31	.32	.31	.40	.31	.36	.34	.36	.36
13	.32	.17	.20	.17	.17	.17	.21	.17	.20	.17	.19	.19
14	.43	.30	.32	.30	.30	.30	.38	.30	.35	.31	.35	.35
15	.47	.39	.39	.39	.40	.39	.45	.39	.42	.40	.42	.42
16	.43	.34	.35	.34	.34	.34	.39	.34	.37	.35	.37	.37
17	.59	.61	.62	.70	.63	.70	.69	.70	.66	.67	.66	.66
18	.42	.27	.35	.27	.27	.27	.41	.27	.36	.31	.37	.36
19	.35	.16	.24	.16	.17	.16	.25	.16	.22	.18	.22	.21
20	.46	.39	.40	.39	.39	.39	.40	.39	.39	.39	.39	.39

Çizelge 3.16’da görüldüğü gibi bu çalışmada farklı kayıp veri yöntemleriyle elde edilen madde güçlük indeksi kestirimleri arasında sayısal farklılıklar bulunmaktadır. Genel olarak en yüksek kestirimler DS yönteminde, en düşük kestirimler ise 0 Atama yönteminde elde edilmiştir. Diğer bir deyişle

DS yönteminin kullanılması durumunda maddelerin 'kolay madde', 0 Atama yönteminin kullanılması durumunda ise maddelerin 'zor madde' olarak değerlendirilme olasılığı yüksektir.

Her bir kayıp veri yönteminde de madde güçlük indeksi en yüksek olan madde 17. maddedir. Bu maddeye yönelik güçlük indeksi değerleri, farklı kayıp veri yöntemlerine göre .61 ve .70 arasında değişmektedir. Madde güçlük indeksi en düşük olan madde ise DS yönteminde 6. madde, diğer yöntemlerde 13. ve 19. maddelerdir.

Madde güçlük indekslerinin .20 ile .80 aralığında olması ölçütü dikkate alındığında bu beklentiyi karşılayan yöntemlerin DS, BO, Dİ ve RA yöntemleri olduğu görülmektedir. Hiçbir kayıp veri yönteminde, madde güçlük indeksi .80'in üzerinde olan madde bulunmamaktadır. Diğer taraftan indeks değeri .20'nin altında olan madde sayısı VÇ ve MCMC yöntemlerinde 1, 0 Atama, SO, YNO, YNM, DE ve BM yöntemlerinde ise 2'dir.

Madde güçlük indekslerinin .50 civarında olması ölçütü dikkate alındığında, bu beklentiyi en iyi karşılayan yöntemlerin DS ve Dİ yöntemleri olduğu görülmektedir. Madde güçlük indeksleri .16 ile .70 arasında değerler almaktadır. İndeks değerleri aralığı en düşük olan yöntem DS yöntemi (.31 - .59 arası), en yüksek olan yöntemler ise SO, YNM ve DE yöntemleri (.19 - .70)'dir.

Parametrik özelliklere yönelik diğer inceleme ve analizlerde olduğu gibi madde güçlük indeksleri de dikkate alındığında, silme ve basit atamaya dayalı yöntemlerin genellikle uç değerler ürettiği, EÇÖ ve ÇVA yöntemlerinin ise görece daha ılımlı kestirimler sağladığı anlaşılmaktadır.

Madde Ayırıcılık İndeksi

Madde parametrelerine yönelik analizlerde ikinci olarak madde ayırıcılık indeksi kestirimleri gerçekleştirilmiştir. Testte yer alan her bir maddenin, madde ile ölçülmek istenen özelliğe sahip olan bireylerle sahip olmayan bireyleri ayırt edebilmesi beklenir. Madde ayırıcılık indeksi olarak tanımlanan bu parametre, ayırıcılık (discriminative) anlamında geçerlik, dolayısıyla yapı geçerliği ile ilgilidir. Madde ayırıcılık indeksinin .40 ve

üzerinde olması idealdir. Bununla birlikte .20 ve üzeri değerler de kabul edilebilir düzeyde bir ayırıcılığın varlığına kanıt olabilmektedir (Lord ve Novick, 1968; Özçelik, 1981; Baykul, 2000; Tekin, 2007). Madde ayırıcılık indeksi kestirimleri Çizelge 3.17’de verilmektedir.

Çizelge 3.17. Kayıp Veri Yöntemlerine Göre Madde Ayırıcılık İndeksleri

Madde	Kayıp Veri Yöntemleri											
	DS	Atama	BO	SO	YNO	YNM	Dİ	DE	RA	BM	VÇ	MCMC
1	.80	.55	.68	.50	.52	.50	.60	.50	.58	.57	.59	.60
2	.87	.57	.71	.52	.54	.52	.59	.52	.60	.59	.61	.61
3	.79	.53	.81	.48	.50	.48	.60	.48	.63	.68	.62	.64
4	.62	.36	.66	.30	.33	.30	.40	.30	.43	.52	.41	.42
5	.79	.72	.80	.45	.65	.45	.62	.45	.66	.72	.67	.68
6	.47	.34	.39	.31	.30	.31	.30	.31	.31	.29	.31	.31
7	.80	.70	.83	.53	.64	.53	.64	.53	.72	.78	.69	.73
8	.59	.43	.54	.33	.38	.33	.40	.33	.39	.46	.41	.40
9	.88	.65	.75	.60	.63	.60	.67	.60	.69	.67	.70	.69
10	.88	.74	.87	.70	.71	.70	.76	.70	.80	.83	.79	.81
11	.81	.68	.84	.61	.64	.61	.68	.61	.72	.78	.71	.73
12	.87	.63	.78	.57	.59	.57	.65	.57	.67	.68	.68	.68
13	.52	.24	.32	.21	.22	.21	.23	.21	.23	.21	.23	.23
14	.64	.48	.60	.43	.44	.43	.46	.43	.47	.47	.46	.47
15	.83	.70	.71	.66	.65	.66	.65	.66	.71	.69	.71	.70
16	.77	.61	.72	.58	.58	.58	.61	.58	.63	.63	.64	.63
17	.74	.63	.67	.51	.54	.51	.54	.51	.61	.55	.60	.61
18	.74	.41	.68	.37	.37	.37	.47	.37	.50	.50	.49	.50
19	.78	.31	.60	.27	.29	.27	.38	.27	.40	.35	.39	.39
20	.83	.60	.64	.59	.57	.59	.59	.59	.62	.59	.62	.62

Çizelge 3.17’de görüldüğü gibi kullanılan kayıp veri yöntemine göre madde ayırıcılık indeksleri kestirimlerinde sayısal farklılıklar oluşmaktadır. Genel olarak en yüksek madde ayırıcılık indeksleri DS yönteminde, en düşük ayırıcılık indeksleri ise SO ve YNM yöntemlerinde elde edilmektedir.

Madde ayırıcılık indeksleri, DS yönteminde .47 ile .88 arasında, SO ve YNM yöntemlerinde ise .21 ile .70 arasında değişmektedir. Ayırıcılık indeksleri aralığı en dar olan yöntem DS yöntemidir. Aralığın en geniş olduğu yöntem ise BM yöntemidir. BM yönteminde madde ayırıcılık indeksleri .21 ile .86 arasında değerler almaktadır.

Ayırıcılık indeksi en yüksek olan madde, DS yönteminde 9 ve 10. maddeler, diğer yöntemlerde 10. maddedir. İndeks değeri en düşük olan madde ise DS yönteminde 6. madde diğer yöntemlerde 13. maddedir. DS

yöntemi dışındaki diğer yöntemlerin kullanılması durumunda testte yer alan 20 madde arasında ayırıcılık gücü en yüksek olan 10. madde, en düşük olan ise 13. madde olarak belirlenmektedir.

Çizelge 3.17’de görüldüğü gibi hiç bir kayıp veri yönteminde, ayırıcılık indeksi .20’nin altında olan madde bulunmamaktadır. Madde ayırıcılık indeksi .40 ve üzerinde olan madde sayısı ise DS yönteminde 20, BO yönteminde 18, Dİ, RA, BM, VÇ ve MCMC yöntemlerinde 17, 0 Atama yönteminde 16, diğer yöntemlerde 14’tür.

Maddenin ayırıcılık gücüne yönelik olarak kullanılabilen diğer parametreler, madde puanları ile toplam puanlar arasındaki korelasyonlardır. İki kategorili puanlanan maddelerde, bu amaçla nokta çift serili korelasyon katsayısı (NÇKK) ve çift serili korelasyon katsayısı (ÇKK) kullanılması uygun görülmektedir. Madde güçlüğü’nün .50 civarında olması durumunda NÇKK, diğer durumlarda ise ÇKK’nin daha fazla bilgi sağlamaktadır (Lord ve Novick, 1968; Özçelik, 1981; Baykul, 2000; Tekin, 2007). Madde nokta çift serili korelasyon katsayıları kestirimleri Çizelge 3.18’de, çift serili korelasyon katsayıları kestirimleri ise Çizelge 3.19’da verilmektedir.

Çizelge 3.18. Kayıp Veri Yöntemlerine Göre Madde (Nokta Çift Serili) Korelasyon Katsayıları*

Madde	Kayıp Veri Yöntemleri											
	DS	Atama	BO	SO	YNO	YNM	Dİ	DE	RA	BM	VÇ	MCMC
1	.70	.61	.65	.61	.60	.61	.56	.61	.58	.63	.58	.60
2	.75	.63	.69	.63	.63	.63	.57	.63	.61	.66	.61	.63
3	.66	.57	.70	.57	.57	.57	.51	.57	.54	.65	.54	.56
4	.51	.41	.56	.39	.40	.39	.36	.39	.37	.51	.36	.39
5	.65	.58	.62	.38	.55	.38	.49	.38	.52	.58	.52	.53
6	.42	.36	.37	.36	.35	.36	.32	.36	.33	.34	.33	.33
7	.67	.57	.64	.43	.56	.43	.51	.43	.56	.63	.54	.57
8	.49	.38	.44	.29	.37	.29	.34	.29	.34	.40	.34	.34
9	.76	.65	.67	.60	.65	.65	.59	.65	.64	.67	.64	.64
10	.75	.64	.72	.70	.64	.64	.60	.64	.64	.70	.64	.66
11	.68	.60	.68	.61	.59	.59	.54	.59	.58	.67	.58	.60
12	.73	.61	.68	.57	.61	.61	.56	.61	.60	.65	.60	.61
13	.48	.37	.37	.21	.37	.37	.31	.37	.32	.33	.33	.32
14	.53	.48	.55	.43	.48	.47	.41	.47	.43	.49	.43	.44
15	.69	.58	.59	.58	.58	.58	.52	.58	.57	.61	.58	.58
16	.66	.55	.63	.55	.55	.55	.52	.55	.55	.59	.55	.55
17	.61	.47	.51	.41	.46	.41	.46	.41	.47	.47	.47	.48
18	.62	.47	.59	.47	.47	.47	.43	.47	.47	.54	.46	.47
19	.67	.50	.63	.50	.50	.50	.44	.50	.49	.52	.49	.50
20	.68	.54	.54	.55	.54	.55	.51	.55	.54	.55	.54	.54

*p=.000

Çizelge 3.19. Kayıp Veri Yöntemlerine Göre Madde (Çift Serili) Korelasyon Katsayıları*

Kayıp Veri Yöntemleri												
	0											
Madde	DS	Atama	BO	SO	YNO	YNM	Dİ	DE	RA	BM	VÇ	MCMC
1	.89	.83	.87	.82	.82	.82	.73	.82	.77	.85	.77	.80
2	.95	.87	.92	.86	.86	.86	.75	.86	.81	.90	.82	.83
3	.84	.77	.90	.75	.75	.75	.64	.75	.69	.84	.68	.72
4	.64	.54	.71	.52	.53	.52	.45	.52	.47	.65	.46	.48
5	.82	.73	.78	.50	.69	.50	.62	.50	.65	.73	.66	.67
6	.56	.49	.51	.49	.49	.49	.43	.49	.45	.47	.45	.45
7	.84	.72	.81	.55	.70	.55	.64	.55	.70	.78	.68	.71
8	.61	.48	.55	.37	.46	.37	.43	.37	.42	.51	.43	.43
9	.97	.86	.87	.86	.86	.86	.76	.86	.83	.89	.83	.84
10	.94	.82	.91	.82	.82	.82	.76	.82	.81	.89	.81	.83
11	.85	.77	.87	.76	.76	.76	.68	.76	.73	.84	.73	.75
12	.92	.80	.88	.80	.80	.80	.71	.80	.76	.84	.77	.78
13	.63	.55	.52	.54	.55	.54	.44	.54	.46	.49	.47	.47
14	.67	.64	.72	.63	.63	.63	.52	.63	.55	.65	.56	.57
15	.86	.74	.75	.74	.74	.74	.66	.74	.72	.77	.73	.73
16	.83	.72	.81	.71	.71	.71	.66	.71	.70	.76	.71	.71
17	.77	.60	.65	.54	.59	.54	.60	.54	.61	.61	.61	.62
18	.78	.64	.76	.64	.63	.64	.54	.64	.61	.71	.59	.60
19	.86	.75	.87	.75	.74	.75	.61	.75	.69	.76	.69	.70
20	.85	.69	.68	.69	.69	.69	.65	.69	.69	.70	.69	.69

*p=.000

Çizelge 3.18 ve Çizelge 3.19’da görüldüğü gibi kestirilen NÇKK ve ÇKK değerleri, istatistiksel olarak manidardır ve ayırıcılık indeksi kestirimlerine paralel değerler üretmektedir. Toplam puanlar ile korelasyonları en düşük ve en yüksek olan maddeler, kestirim aralığı en dar ve en geniş olan yöntemler birbirini teyit eder nitelikte örtüşmektedir. Kullanılan kayıp veri yöntemine göre farklılaşmakla birlikte genel olarak NÇKK değerleri .21 ile .76 aralığında, ÇKK değerleri ise .37 ile .97 aralığında değişmektedir.

Madde güçlüklerinin .50 civarından uzaklaştığı ve sağa çarpık bir dağılım gösterdiği bulgusuna ulaşılmıştı. Bu durumda ÇKK kestirimlerinin daha fazla bilgi sağladığı ileri sürülebilir. Buna göre Çizelge 3.19’da yer alan kestirimler dikkate alındığında, DS ve BO yöntemleri dışındaki yöntemlerde, maddeler ile toplam puanlar arasında orta düzeyde bir korelasyonun bulunduğu görülmektedir. En düşük korelasyon katsayısı kestirimleri RA yönteminde elde edilmiştir. DS ve BO yöntemlerinde ise korelasyonlar, diğer yöntemlere göre daha yüksek kestirilmektedir. Bu yöntemlerin kullanılması durumunda maddeler ile toplam puanlar arasında orta ve yüksek düzeyde

korelasyon belirlenmektedir. Testte yer alan 4, 6, 8 ve 13. maddeler ile toplam puanlar arasındaki korelasyonlar, diğer maddelere göre daha düşüktür.

Madde Güvenirlik Katsayıları

Madde güvenirlik katsayıları, bu çalışmada dikkate alınan bir diğer madde parametresidir. İki kategorili puanlanan maddelerde bu parametre, puan aralığının dar olması nedeniyle çok anlamlı görülmemektedir. Güvenirlik kanıtı olarak toplam puanlara yönelik güvenirlik kestirimlerinin dikkate alınmasının daha doğru olacağı ifade edilmektedir. İki kategorili puanlanan maddelerde madde güvenirlik katsayısı, basitçe madde standart sapması ile madde ayırıcılık indeksi değerleri çarpılarak elde edilebilmektedir (Baykul, 2000; Atılğan, 2009). Farklı kayıp veri yöntemlerine göre kestirilen madde güvenirlik katsayıları, Çizelge 3.20'de verilmektedir.

Çizelge 3.20. Kayıp Veri Yöntemlerine Göre Madde Güvenirlik Katsayıları

Madde	Kayıp Veri Yöntemleri											
	DS	Atama	BO	SO	YNO	YNM	Dİ	DE	RA	BM	VÇ	MCMC
1	.39	.24	.31	.21	.23	.22	.26	.21	.26	.24	.27	.27
2	.43	.25	.32	.22	.23	.22	.26	.22	.27	.25	.28	.27
3	.39	.24	.39	.20	.23	.22	.27	.20	.30	.29	.32	.31
4	.31	.17	.32	.12	.15	.14	.18	.12	.21	.21	.21	.21
5	.39	.36	.40	.20	.33	.21	.29	.20	.32	.32	.34	.34
6	.22	.14	.17	.13	.13	.13	.13	.13	.13	.12	.14	.13
7	.40	.34	.41	.23	.32	.26	.30	.23	.35	.35	.36	.37
8	.29	.21	.27	.15	.19	.16	.19	.15	.20	.21	.21	.20
9	.43	.30	.35	.27	.29	.27	.31	.27	.32	.30	.33	.32
10	.44	.36	.43	.32	.34	.34	.36	.32	.39	.38	.40	.40
11	.41	.33	.41	.27	.31	.29	.32	.27	.36	.35	.36	.36
12	.43	.29	.37	.25	.27	.26	.30	.25	.32	.31	.33	.33
13	.24	.09	.13	.08	.08	.08	.09	.08	.09	.08	.09	.09
14	.32	.22	.28	.19	.20	.20	.21	.19	.22	.21	.23	.22
15	.41	.34	.35	.31	.32	.32	.31	.31	.35	.33	.35	.35
16	.38	.29	.34	.27	.27	.27	.29	.27	.30	.29	.32	.30
17	.36	.31	.33	.23	.26	.23	.25	.23	.29	.25	.29	.29
18	.37	.18	.32	.15	.16	.16	.21	.15	.24	.21	.25	.24
19	.37	.12	.26	.10	.11	.10	.15	.10	.16	.13	.17	.16
20	.41	.29	.31	.29	.28	.29	.29	.29	.30	.29	.30	.30

Çizelge 3.20’de görüldüğü gibi her bir kayıp veri yönteminde de madde güvenilirlik katsayıları düşüktür. En yüksek kestirimler DS yönteminde (.22 - .44), en düşük kestirimler ise SO ve DE yöntemlerinde (.08 - .32) elde edilmiştir. Her bir kayıp veri yönteminde de güvenilirlik kestirimi en yüksek olan madde 10. maddedir. Güvenirlik kestirimi en düşük olan madde ise DS yönteminde 6. madde, diğer yöntemlerde 13. maddedir.

Test Parametreleri

SBS 2011 Matematik Testi A Kitapçığına yönelik test parametreleri kestirimlerinde öncelikle farklı kayıp veri yöntemlerine kullanılarak test toplam puanları elde edilmiş ve bu puanların dağılımları incelenmiştir. Farklı kayıp veri yöntemlerine göre elde edilen toplam puanlarının dağılımlarına ilişkin kestirimler Çizelge 3.21’de verilmektedir.

Çizelge 3.21. Kayıp Veri Yöntemlerine Göre Toplam Puanların Dağılımı

Kayıp Veri Yöntemleri	Min.	Max.	μ	SE (μ)	σ	Basıklık	Çarpıklık
DS	0.0	20.0	8.95	.17	6.26	.63	-1.19
0 Atama	0.0	20.0	6.57	.07	4.87	1.21	.63
BO	0.0	20.0	7.30	.14	5.61	.90	-0.56
SO	0.0	20.0	7.93	.06	4.44	1.08	.42
YNO	0.0	20.0	6.72	.07	4.82	1.20	.61
YNM	0.0	20.0	7.25	.06	4.57	1.23	.71
DI	0.0	20.0	7.93	.06	4.56	.99	.20
DE	0.0	20.0	7.93	.06	4.44	1.08	.42
RA	-2.1	20.0	7.75	.07	4.89	.91	-0.09
BM	0.0	20.0	7.75	.07	4.82	.95	-0.04
VÇ	-5.0	21.0	7.72	.07	4.94	.87	-0.12
MCMC	0.0	20.0	7.77	.07	4.93	.89	-0.18

Çizelge 3.21’de görüldüğü gibi toplam puanlar, RA ve VÇ yöntemleri dışında 0 ile 20 puan aralığında değişmektedir. RA ve VÇ yöntemlerinde kayıp veriler yerine olasılıklı ve ötelemeli atama ile [0, 1] aralığı dışında değer ataması yapılabilmektedir. DS yönteminde, kayıp veriler yerine değer

ataması yapılmamakta, bu veriler silinerek analiz dışı bırakılmaktadır. 0 Atama, YNM, BM ve MCMC yöntemlerinde kayıp veriler yerine 1 ve 0 değerleri atanmaktadır. Diğer yöntemlerde ise $[0, 1]$ aralığında değer ataması yapılmaktadır.

Farklı kayıp veri yöntemlerine göre elde edilen toplam puan ortalamaları, 6.57 ve üzerindedir. En yüksek ortalama DS yönteminde, en düşük ortalama ise 0 Atama yönteminde elde edilmektedir. 0 Atama yönteminin kullanılması durumunda toplam puan ortalamasının düşmesi, beklenen bir durumdur. DS yönteminde ortalamanın yükselmesi ise ihmal edilen verilerde 0 değerinin 1 değerine göre daha fazla olduğunu göstermektedir. Kayıp veri içeren yanıt örüntülerinde 0 puanlanan yani yanlış yanıtlanan madde sayısının 1 puanlanan yani doğru yanıtlanan madde sayısından yüksek olduğu anlaşılmaktadır.

Ortalama kestirime yönelik standart hata değerleri, .06 ve üzerindedir. En yüksek hatanın DS ve BO yöntemlerinin kullanılması durumunda ortaya çıktığı görülmektedir.

Verilerin yayılma ölçülerinden biri olarak standart sapma değerleri, 4.44 ve üzerindedir. En yüksek standart sapma DS ve BO yöntemlerinde elde edilmektedir. SO ve DE yöntemleri ise en düşük standart hata kestirimlerini üretmektedir.

Basıklık ve çarpıklık katsayılarına göre toplam puanlar sağa çarpık bir dağılıma sahiptir. Test puanlarının dağılımlarının, normal dağılımdan gelmediği iddia edilebilir görülmektedir. Bu bulgu SBS 2011 Matematik Testi A Kitapçığının, izleme testi yapısından çok maksimum performans testi yapısında olduğuna yönelik bir kanıt sağlamaktadır.

SBS 2011 Matematik Testi A Kitapçığına yönelik verilerde DS yöntemi, test ortalama, standart hata ve standart sapmalarının olduğundan daha yüksek kestirime yol açmakta, standart hata ve standart sapma değerlerini yükseltmektedir. Little ve Rubin (1987), Schaffer (1997), Allison (2002, 2009) gibi birçok araştırmacı, özellikle TSK varsayımının karşılanmadığı ve kayıp veri miktarının yüksek olduğu verilerde, DS gibi silmeye dayalı yöntemlerin yanlış kestirimler ürettiğini belirlemiştir. Bu yöntemler, evrene yönelik kestirimlerde kullanılabilir görülmemektedir.

SBS 2011 Matematik Testi A Kitapçığına yönelik verilerde 0 Atama yöntemi, toplam puanlara yönelik standart sapma ve standart hata kestirimlerinde belirgin bir farklılık oluşturmamakla birlikte, ortalama kestirimlerinde dikkate değer bir düşüşe yol açmaktadır. Lord (1974), Finch (2008), Hohensin ve Kubinger (2011), Culbertson (2011) gibi birçok araştırmacı, özellikle iki kategorili puanlanan testlerde, kayıp verilerin 'yanlış' kabul edilmesine dayalı 0 atama yönteminin, TSK ve SK varsayımları karşılanırsa bile yanlış kestirimlere yol açtığını ve optimal bir yöntem olarak kabul edilemeyeceğini belirlemiştir.

SBS 2011 Matematik Testi A Kitapçığına yönelik verilerde basit atama yöntemleri, toplam puanlara yönelik ortalama ve standart sapma kestirimlerini olduğundan daha düşük gösterme eğilimindedir. Bu yöntemler veri setinin homojenleşmesine yol açmaktadır. Lord (1974), Finch (2008), Hohensin ve Kubinger (2011) gibi birçok araştırmacı, özellikle iki kategorili puanlanan testlerde, kayıp verilerin ihmal edilebilir olmaması durumunda ve kayıp veri miktarının yüksek olması durumunda, silmeye dayalı yöntemlerin yanlış kestirimler ürettiğini belirlemiştir. Linacre (2004), bu tür durumlarda silmeye dayalı yöntemlerin ya da basit atama yöntemlerinin, ancak EÇO ve ÇVA yöntemleri gibi daha karmaşık yöntemler için başlangıç kestirimlerinin elde edilmesinde kullanılabileceğini belirtmektedir.

Little ve Rubin (1987), Allison (2002, 2009) gibi birçok araştırmacı, basit atama yöntemlerinde, değişkenler düzeyinde, standart sapma kestirimlerinin olduğundan daha düşük, ortalama kestirimlerinin ise olduğundan daha yüksek elde edildiğini belirlemiştir. Bu çalışmada ise basit atama yöntemlerinin kullanılması durumunda test ortalamaları, olduğundan daha düşük kestirilmektedir. Bu farklılığın, yanıtlayıcıların yanıt örüntülerinde 0 puanlanan madde sayısının ve maddeler düzeyinde kayıp veri miktarının fazlalığından kaynaklandığı düşünülmektedir. Basit atama yöntemleri, ortalama, medyan gibi merkezi eğilim ölçülerini ya da yakın nokta değerlerini dikkate almaktadır. Bu durum kayıp veriler yerine benzer ya da yakın değerler atanma olasılığını artırmakta ve veri setini daha homojen bir hale getirmektedir. Bu çalışmada, basit atama yöntemlerinin kullanılması, yanıt örüntülerinde 0 ya da 0'yakın değerlerin artmasına yol açmaktadır. Dolayısıyla maddeler düzeyinde ortalama kestirimleri, olduğundan daha

düşük elde edilmektedir. Maddeler düzeyinde kayıp veri miktarının fazlalığı da bu düşüşü daha belirgin hale getirmektedir.

SBS 2011 Matematik Testi A Kitapçığına yönelik kestirimlerde, EÇO ve ÇVA yöntemleri, silme ve basit atamaya dayalı yöntemlere göre genel olarak daha ılımlı kestirimler üretmektedir. Bununla birlikte VÇ yöntemi, kayıp veriler yerine değer atamada diğer yöntemlere göre daha geniş bir puan aralığını dikkate alabilmesine bağlı olarak, standart sapmanın daha yüksek kestirimine yol açmaktadır. MCMC yönteminde ise gerçek veri seti kullanılarak daha büyük bir veri seti üretilmekte ve kestirimler, bu büyük veri seti üzerinden elde edilmektedir. Dolayısıyla MCMC yönteminde standart hata, olduğundan daha düşük kestirilebilmektedir. Lord (1983), Schaffer (1997), DeMars (2002), Acock (2005) gibi birçok araştırmacı, iki kategorili puanlanan verilerde, özellikle ihmal edilebilirlik sağlanmadığı durumlarda, silme ve basit atamaya dayalı yöntemler yerine EÇO yöntemleri gibi olasılıklı yöntemlerin daha titiz kestirimler ürettiğini belirlemiştir. Rubin (1987), Little ve Rubin (1987) ve Schaffer (1997), bu tür verilerde ÇVA yöntemlerinin de titiz kestirimler ürettiğini belirtmektedir. Diğer taraftan Allison (2006), iki kategorili puanlanan verilerde ÇVA yöntemlerinin özellikle ortalama kestirimlerinde yanlışlık üretebileceğini belirlemiştir ki bu belirleme, bu çalışmada elde edilen bulgular ile örtüşmektedir.

Toplam puanların dağılımına yönelik olanlar dışında diğer test parametreleri Çizelge 3.22'de verilmektedir. Çizelge 3.22'de yer alan kestirimlerden minimum ve maksimum puanlar ile ayırıcılık indeksi değerleri, alt ve üst gruplar üzerinden elde edilmiştir. Alt ve üst gruplar, toplam gözlem birimlerinin yaklaşık %27'sini temsil etmektedir. Bu çalışmada alt ve üst gruplar, DS yönteminde 39000 civarı gözlem biriminden, MCMC yönteminde 700000 civarı gözlem biriminden, diğer yöntemlerde ise 12000 ile 15000 arasında gözlem biriminden oluşmaktadır. Toplam puanlar sağa çarpık bir dağılım gösterdiği için alt gruplarda yer alan gözlem birimi sayısı, üst gruplara göre daha fazladır.

Çizelge 3.22'de görüldüğü gibi alt gruplarda maksimum toplam puanlar 3 ve 5 arasında değişmektedir. En yüksek 'maksimum toplam puan' Dİ yönteminin kullanılması durumunda, en düşük 'maksimum toplam puan' 0 Atama ve BO yöntemlerinin kullanılması durumunda ortaya çıkmaktadır. Üst

gruplarda ise minimum puanlar 8 ile 15 arasında değişmektedir. En yüksek 'minimum toplam puan' DS yönteminde elde edilmektedir. Basit atamaya dayalı kayıp veri yöntemlerinin EÇO ve ÇVA yöntemlerine göre daha düşük kestirimler ürettiği görülmektedir.

Çizelge 3.22. Kayıp Veri Yöntemlerine Göre Test Parametreleri

Kayıp Veri Yöntemi	Test Parametreleri						
	Max. Puan (Alt Grup)	Min. Puan (Üst Grup)	Ortalama Güçlük	Ortalama Ayırıcılık	Ortalama NÇKK	Ortalama ÇKK	Güvenirlilik (α)
DS	4	15	.447	.636	.636	.804	.923
0 Atama	3	8	.329	.530	.529	.700	.866
BO	3	10	.365	.592	.592	.767	.903
SO	4	8	.363	.502	.492	.667	.844
YNO	4	8	.336	.523	.524	.690	.861
YNM	4	8	.363	.502	.503	.667	.844
DI	5	11	.421	.478	.478	.614	.823
DE	4	8	.363	.502	.503	.667	.844
RA	4	10	.388	.507	.508	.656	.849
BM	4	9	.364	.561	.560	.732	.886
VÇ	4	10	.389	.507	.507	.656	.849
MCMC	4	10	.388	.517	.517	.670	.857

Çizelge 3.22'de görüldüğü gibi testin ortalama güçlük değerleri, kayıp veri yöntemlerine göre .329 ile .447 arasında değişmektedir. En yüksek kestirimler DS ve DI yöntemlerinde, en düşük kestirim 0 Atama yönteminde elde edilmiştir. Basit atamaya dayalı kayıp veri yöntemlerinin EÇO ve ÇVA yöntemlerine göre genel olarak daha düşük test ortalama güçlüğü kestirimleri ürettiği görülmektedir.

Madde ve test parametreleri birlikte değerlendirildiğinde, kullanılan kayıp veri yöntemine göre madde güçlük indeksleri .16 ile .70 arasında, test ortalama güçlük indeksleri ise .33 ile .48 arasında değerler almaktadır. Bu farklılık, testin ağırlıklı olarak 'zor' maddelerden oluştuğunu göstermektedir. Bu kapsamda testte yer alan 6, 13 ve 19. maddelerin gözden geçirilmesi gerekli görülmektedir.

Testin ortalama ayırıcılık indeksi, kayıp veri yöntemlerine göre .502 ile .636 arasında değişmektedir. En yüksek ayırıcılık kestirimleri DS yönteminde, en düşük ayırıcılık kestirimi ise SO, YNM ve DE yöntemlerinde elde

edilmektedir. Bu durum DS yönteminin kullanılması durumunda alt ve üst gruplar arasındaki puan farklılıklarının olduğundan daha fazla belirlendiğini göstermektedir. Diğer taraftan SO, YNM ve DE yöntemlerinin kullanılması durumunda alt ve üst gruplar arasındaki puan farklılıkları azalmaktadır.

Ortalama korelasyon katsayılarından NÇKK ortalamaları, kayıp veri yöntemlerine göre .478 ile .636 arasında değişmektedir. Bu değerler, maddeler ile toplam puanlar arasındaki ilişkiler anlamında orta düzeyde bir ayırıcılığın bulunduğuna yönelik kanıt sağlamaktadır. ÇKK kestirimleri ise kayıp veri yöntemlerine göre .614 ile .804 arasında değişmektedir. Bu değerler, madde puanları ile test puanları arasında, kullanılan kayıp veri yöntemine göre, orta düzeyde ilişkilerin yanı sıra iyi düzeyde ilişkilerin de bulunduğunu göstermektedir. Orta ve iyi sayılabilecek düzeylerde ayırıcılık kanıtı sağlamaktadır. En yüksek ortalama NÇKK ve ÇKK kestirimleri DS yönteminde, en düşük kestirimler ise DI, SO, YNM ve DE yöntemlerinde elde edilmiştir.

Bu durum silmeye dayalı yöntemlerin, testin ayırıcılık düzeyini olduğundan daha yüksek belirlediğini göstermektedir. Diğer taraftan basit atamaya dayalı yöntemler ise genel olarak testin ortalama ayırıcılık düzeyini düşürmektedir.

Madde ve test parametrelerinin birlikte değerlendirilebilmesi amacıyla, kullanılan kayıp veri yöntemlerine göre madde parametrelerine yönelik aralıklar ve test parametreleri Çizelge 2.23'te verilmektedir.

Çizelge 3.23'te görüldüğü gibi kullanılan kayıp veri yöntemine göre madde ayırıcılık indeksleri .21 ile .88 arasında, test ortalama ayırıcılık indeksleri ise .50 ile .64 arasında değerler almaktadır. Bu farklılık, testte yer alan maddelerin çoğunluğunun ayırıcılık gücünün düşük olduğunu göstermektedir. Bu kapsamda testte yer alan 6 ve 13. maddelerin gözden geçirilmesi gerekli görülmektedir.

Madde parametrelerinden NÇKK ve ÇKK, kullanılan kayıp veri yöntemine göre .21 ile .97 arasında değişmektedir. Test ortalama NÇKK ve ÇKK değerleri ise kayıp veri yöntemine göre .48 ile .81 arasındadır. Bu farklılık, madde-test korelasyonu ortalamanın üzerinde olan madde sayısının daha fazla olduğunu göstermektedir. Bu kapsamda testte yer alan 6 ve 13. maddelerin gözden geçirilmesi gerekli görülmektedir.

Çizelge 3.23. Kayıp Veri Yöntemlerine Göre Madde Parametre Aralıkları ve Test Parametreleri

Kayıp Veri Yöntemi	Madde Güçlük	Test Ortalama	Güçlüğü	Madde Ayrıcılık	Test Ortalama	Ayrıcılığı	Madde NCKK	Test Ortalama	NCKK	Madde ÇKK Aralığı	Test Ortalama ÇKK	Madde Güvenirliliği	Aralığı	Test Güvenirliliği
DS	.31-.59	.447	.47-.88	.636	.42-.76	.636	.56-.97	.804	.24-.44	.923				
0 Atama	.16-.61	.329	.24-.74	.530	.36-.65	.529	.48-.87	.700	.09-.36	.866				
BO	.20-.62	.365	.32-.87	.592	.37-.72	.592	.51-.92	.767	.13-.43	.903				
SO	.16-.70	.363	.21-.70	.502	.21-.70	.492	.37-.86	.667	.08-.32	.844				
YNO	.17-.63	.336	.22-.71	.523	.35-.65	.524	.43-.76	.690	.08-.34	.861				
YNM	.16-.70	.363	.21-.70	.502	.36-.65	.503	.37-.86	.667	.08-.34	.844				
DI	.25-.69	.421	.23-.76	.478	.31-.60	.478	.43-.76	.614	.09-.36	.823				
DE	.16-.70	.363	.21-.70	.502	.36-.65	.503	.37-.86	.667	.08-.32	.844				
RA	.20-.66	.388	.23-.80	.507	.32-.64	.508	.42-.83	.656	.09-.39	.849				
BM	.17-.67	.364	.21-.83	.561	.33-.70	.560	.51-.90	.732	.08-.38	.886				
VÇ	.22-.66	.389	.23-.79	.507	.33-.64	.507	.43-.83	.656	.09-.40	.849				
MCMC	.21-.66	.388	.23-.81	.517	.32-.66	.517	.43-.84	.670	.09-.40	.857				

Çizelge 3.23'te görüldüğü gibi testin güvenilirliğine yönelik olarak Cronbach'ın α katsayısı kestirimleri, iç tutarlılık anlamında bilgi sağlamaktadır. Kestirilen α katsayıları, kayıp veri yöntemlerine göre .823 ve .923 arasında değişmektedir. Bu durum testin iç tutarlılık anlamında güvenilirliğinin yüksek sayılabilecek düzeyde olduğu yönünde kanıt sağlamaktadır. En yüksek α değeri DS yönteminde, en düşük α değeri ise Dİ, SO, YNM ve DE yöntemlerinde elde edilmiştir. Bu durum silmeye dayalı yöntemlerin, testin iç tutarlılık anlamında güvenilirlik düzeyini olduğundan daha yüksek belirlediğini göstermektedir. Diğer taraftan basit atamaya dayalı yöntemler ise genel olarak testin iç tutarlılık anlamında güvenilirlik düzeyini düşürmektedir.

Madde ve test parametreleri birlikte değerlendirildiğinde, madde güvenilirlik katsayıları, kullanılan kayıp veri yöntemine göre .08 ile .44 arasında değişmektedir ve oldukça düşük güvenilirlik düzeyini göstermektedir. Test güvenilirlik katsayıları olarak Cronbach α değerleri ise kullanılan kayıp veri yöntemine göre .82 ile .92 arasında değişmektedir. Test güvenilirlikleri, her bir kayıp veri yönteminde de yüksek düzeyde belirlenmektedir. Madde ve test güvenilirlikleri arasındaki bu farklılığın, puanlama ölçeği ve buna göre belirlenen hesaplama yönteminden kaynaklandığı düşünülebilir. SBS 2011 Matematik Testi A Kitapçığında maddeler, iki kategorili puanlanan sınıflama ölçeğindeki değişkenlerdir. Madde güvenilirlikleri, madde standart sapması ve ayıricılık indeksinin çarpımı ile elde edilmektedir. Test güvenilirliğinin hesaplanmasında dikkate alınan toplam puanlar ise eşit aralık ölçeğindedir. Test güvenilirliğinin hesaplanmasında Cronbach α katsayısı kullanılmaktadır. Bu nedenle Baykul (2000), özellikle iki kategorili puanlanan maddelerden oluşan testlerde, maddeler düzeyinde güvenilirlik kestirimleri yerine test güvenilirliğine yönelik kestirimlerinin dikkate alınmasının daha uygun olacağını belirtmektedir.

BÖLÜM 5

SONUÇ VE ÖNERİLER

Bu bölümde, çalışmada ulaşılan sonuçlar ve öneriler yer almaktadır. Sonuçlar, araştırma sorularına uygun olarak ayrı başlıklar altında sunulmaktadır.

Sonuç

Belirlenen araştırma soruları çerçevesinde ulaşılan sonuçlar, aşağıda başlıklar halinde sunulmaktadır.

Kayıp Verilerin İhmal Edilebilirliği

Bu çalışmada SBS 2011 Matematik Testi A Kitapçığına yönelik veri setinde yer alan kayıp verilerin ihmal edilebilir olup olmadığı, kayıp veri örüntüsü ve kayıp veri mekanizmasına yönelik analiz ve incelemelerle belirlenmeye çalışılmıştır.

Kayıp Veri Örüntüsü

Testi alan yanıtlayıcıların yanıt örüntülerinde yer alan kayıp veriler, genel örüntü yapısı göstermektedir. Buna göre kayıp verilerin veri setinde düzenli, planlı ya da monoton bir şekilde konumlanmadığı anlaşılmaktadır.

Testi alan yanıtlayıcıların %27.4'ü testte yer alan tüm maddelerde, doğru ya da yanlış, yanıtlama yapmıştır. Yanıtlayıcıların %49.3'ü, testte yer alan 20 maddeden en az 17'sinde yanıtlama yapmıştır. Diğer taraftan yanıtlayıcıların %72.6'sının yanıt örüntülerinde en az bir maddede kayıp veri bulunmaktadır. Test maddelerinin hiçbirinde yanıtlama yapmayan yanıtlayıcılar ise %0.01 gibi oldukça düşük bir yüzde oluşturmaktadır.

Maddeler düzeyinde incelendiğinde doğru yanıtlayıcı, yanlış yanıtlayıcı ve kayıp verilerinde düzenli bir artış ya da azalış bulunmadığı görülmüştür. Testte yer alan 20 maddenin her biri kayıp veri içermektedir. En az kayıp veri içeren madde 20. maddedir. Bu maddede yanıtlayıcıların %2'si yanıtlayıcı yapmamıştır. En fazla kayıp veri içeren madde ise 4. maddedir. Bu maddede yanıtlayıcıların %34.4'ü yanıtlayıcı yapmamıştır. Test maddelerinin 19'unda %5'in üzerinde, 14'ünde %10'un üzerinde, 7'sinde ise %20'nin üzerinde kayıp veri bulunmaktadır.

Kayıp Veri Mekanizması

SBS 2011 Matematik Testi A Kitapçığı verilerinde, kayıp veri mekanizmasına yönelik analiz ve incelemeler, yanıt örüntülerinde kayıp verilerin tam seçkisiz olmasa da seçkisiz oluşmuş olma olasılığının yüksek olduğunu göstermiştir. Kayıp verilerin oluşmasını açıklayan mekanizmanın TSK varsayımını karşılamadığı fakat SK varsayımını karşılama olasılığının yüksek olduğu belirlenmiştir.

Kayıp veri mekanizmasına yönelik inceleme ve analizlerde cinsiyet, okul türü ve bölge değişkenleri dikkate alınmıştır. Yanıtlayıcıların toplam doğru ve yanlış yanıt sayıları ile kayıp veri içeren madde sayılarında bu değişkenlere göre manidar farklar bulunduğu belirlenmiştir. Ayrıca tüm yanıtlayıcılar ile yanıt örüntülerinde kayıp veri bulunmayan yanıtlayıcıların profilleri, bu değişkenlere göre karşılaştırılarak incelenmiştir.

Testi alan yanıtlayıcıların %48.7'si kızlardan, %51.3'ü ise erkeklerden oluşmaktadır. Cinsiyet değişkeni dikkate alındığında kızlar lehine bir durum bulunduğu görülmüştür. Kızlar, erkeklere göre daha fazla sayıda doğru yanıtlayıcı ve daha az sayıda yanlış yanıtlayıcı yapmıştır. Kızların yanıt örüntülerinde, erkeklerin yanıt örüntülerinde olduğundan daha fazla sayıda kayıp veri bulunmaktadır. Buna bağlı olarak kızların yanıt örüntülerinde kayıp veri oluşma olasılığının, erkeklere göre daha düşük olduğu ileri sürülebilir görülmektedir.

Testi alan yanıtlayıcıların %96.7'si devlet okullarında, %3.3'ü ise özel okullarda öğrenim görmüştür. Okul türü değişkeni dikkate alındığında özel okullar lehine bir durum bulunduğu görülmüştür. Özel okullarda öğrenim

gören yanıtlayıcıların doğru yanıt sayıları, devlet okullarında öğrenim görenlere göre daha fazla, yanlış yanıt sayıları ve kayıp veri içeren madde sayıları ise daha azdır. Buna bağlı olarak özel okullarda öğrenim görenlerin yanıt örüntülerinde kayıp veri oluşma olasılığının, devlet okullarında devlet okullarında öğrenim görenlere göre daha düşük olduğu ileri sürülebilir görülmektedir.

Testi alan yanıtlayıcıların %17.8'i İstanbul, %13.4'ü Akdeniz, %12.0'ı Ege, %10.6'sı Güneydoğu Anadolu, %9.7'si Batı Anadolu, %9.4'ü Doğu Marmara, %5.9'u Batı Karadeniz, %5.5'i Orta Anadolu, %5.4'ü Ortadoğu Anadolu, %3.8'i Batı Marmara, %3.4'ü Doğu Karadeniz ve %3.0'ı Kuzeydoğu Anadolu bölgesindendir. Bölge değişkeni dikkate alındığında test genelinde doğru yanıt sayıları en fazla olan yanıtlayıcılar Batı Anadolu Bölgesinden, en az olan yanıtlayıcılar ise Güneydoğu Anadolu Bölgesindendir. Yanlış yanıt sayılarının en yüksek olduğu bölge Güneydoğu Anadolu Bölgesi, en az olduğu bölge ise Doğu Marmara Bölgesidir. Kayıp veri içeren madde sayılarının en yüksek olduğu bölge Doğu Karadeniz Bölgesi, en az olduğu bölge ise Batı Anadolu Bölgesidir. Buna bağlı olarak bölge değişkeninin kayıp veri oluşma olasılığı ile ilişkili olduğu ileri sürülebilir görülmektedir.

Bu çalışmada kayıp veri mekanizmasına yönelik inceleme ve analizlerde cinsiyet, okul türü ve bölge değişkenlerine göre belirlenen manidar farkların yanı sıra 'Little's MCAR Test' kestirimleri gerçekleştirilmiştir. 'Little's MCAR Test' kestirimleri ($\chi^2=33719.52$ ve $p<0.05$), kayıp verilerin tam seçkisiz oluşmadığı ve TSK varsayımının ihlâl edilmiş olabileceği yönünde kanıt sağlamıştır.

SBS 2011 Matematik Testi A Kitapçığı verilerinde, her bir madde düzeyinde kayıp veri bulunması, her bir maddede hem 1 hem 0 puanlamaların bulunması, kayıp veri miktarının madde sırasına göre düzenli bir değişim göstermemesi, kayıp verilerin genel örüntü yapısında ve dağınık şekilde konumlanmaları, tüm yanıtlayıcılar ile yanıt örüntüleri kayıp veri içermeyen yanıtlayıcıların profilleri arasında benzerlikler bulunması gibi bulgular, kayıp veri mekanizmasının tam seçkisiz olmasa da seçkisiz olduğu yönünde kanıt sağlamıştır. Buna bağlı olarak kayıp veri mekanizmasının SK varsayımını karşılama olasılığının yüksek olduğu belirlenmiştir.

Bu çalışmada, SBS 2011 Matematik Testi A Kitapçığına yönelik veri setinde, kayıp veri örüntüsü ve kayıp veri mekanizması üzerine yapılan inceleme ve analizler, kayıp verilerin ihmal edilebilir olmadığını göstermiştir. Sadece gözlenen verilerden hareketle gerçekleştirilecek analizlerin, manidar düzeyde yanlı kestirimler üretme olasılığı yüksektir. Buna bağlı olarak testin psikometrik özelliklerinin incelenmesinde, öncelikle, kayıp verilerle başa çıkmak için uygun yöntemlerin belirlenmesi ve bu yöntemlere göre istatistiksel düzeltmeler yapıldıktan sonra analizlere devam edilmesi gerekliliği ortaya çıkmıştır.

Bu çalışmada, SBS 2011 Matematik Testi A Kitapçığının psikometrik özelliklerine yönelik inceleme ve analizler, 12 farklı kayıp veri yöntemi dikkate alınarak gerçekleştirilmiştir. Belirlenen yöntemler arasında silmeye ve basit atamaya dayalı yöntemler bulunduğu gibi EÇO ve ÇVA yöntemleri de bulunmaktadır. Kullanılan kayıp veri yöntemlerinden DS yöntemi silmeye dayalı, 0 Atama, SO, BO, YNO, YNM, Dİ ve DE yöntemleri basit atamaya dayalı, RA, BM ve VÇ yöntemleri en çok olabilirlik yaklaşımına dayalı, MCMC ise çoklu atama yaklaşımına dayalı yöntemlerdir.

Yapı Geçerliliği ve Tek Boyutluluk Özelliği

Bir veri setinde kayıp verilerin varlığı, ideal ve istenen bir durum değildir. Bu nedenle kayıp veriler, genel olarak testin yapı geçerliliğine olumsuz katkı sağlamaktadır. Dahası ihmal edilebilir olmayan kayıp verilerin varlığında, özellikle değişkenler düzeyinde kayıp veri miktarının fazla olduğu durumlarda, yapı geçerliliğine yönelik analiz yöntemlerinin kullanılması mümkün olmamaktadır. Bu nedenle kayıp veriler içeren bir testin yapı geçerliliğine yönelik analizlerde, öncelikle uygun bir kayıp veri yöntemine göre eksiksiz veri setinin elde edilmesi gerekmektedir.

Bu çalışmada SBS 2011 Matematik Testi A Kitapçığının yapı geçerliliği ve tek boyutluluk özelliği, farklı kayıp veri yöntemleri kullanılarak elde edilen eksiksiz veri setleri üzerinde yürütülen Açımlayıcı Faktör Analizi (AFA) ve Doğrulayıcı Faktör Analizi (DFA) ile incelenmiştir.

SBS 2011 Matematik Testi A Kitapçığına yönelik veri seti üzerinde yapı geçerliği çalışmaları olarak yürütülen AFA ve DFA, testte yer alan bazı maddelerin (13, 6, 8 ve 4. maddeler) ölçülmek istenen yapının tanımlanmasını ve tek boyutluluk özelliğinin sağlanmasını güçleştirdiğini göstermiştir. Bununla birlikte kullanılan her bir kayıp veri yönteminde de maddeler tek bir faktör altında toplanabilmektedir. Elde edilen tek faktörlü yapı, mükemmele yakın düzeyde model-veri uyumu sağlamaktadır. Bu durum testin tek boyutluluk özelliğine sahip olduğuna ve ileri analizlerde toplam puanların kullanılabileceğine dair bir kanıt oluşturmaktadır.

Farklı kayıp veri yöntemleri karşılaştırıldığında AFA ve DFA'da silme ve basit atamaya dayalı yöntemlerin, EÇO ve ÇVA yöntemlerine göre görece daha iyi performans gösterdiği belirlenmiştir. Bu durum, silme ve basit atamaya dayalı yöntemlerin, özellikle iki kategorili puanlanan maddelerden oluşan testlerde, veri setini homojenleştirme eğiliminin doğal bir sonucu olarak değerlendirilmiştir.

Madde ve Test Parametreleri

Bu çalışmada SBS 2011 Matematik Testi A Kitapçığına yönelik olarak, farklı kayıp veri yöntemleri ile elde edilen eksiksiz veri setleri üzerinde, klasik test kuramı temelinde madde ve test parametre kestirimleri gerçekleştirilmiştir.

Madde güçlük indeksi kestirimleri, her bir kayıp veri yönteminde de .80'in altında elde edilmiştir. Madde güçlük indeksi en yüksek olan madde 17. madde, en düşük olan maddeler ise 6, 13 ve 19. maddelerdir. Güçlük indeksi .20'nin altında olan 6, 13 ve 19. maddelerin gözden geçirilmesi gerektiği belirlenmiştir.

Kullanılan kayıp veri yöntemleri karşılaştırıldığında, DS yönteminin madde güçlüklerini artırma, 0 Atama yönteminin ise madde güçlüklerini azaltma eğiliminde olduğu görülmüştür. DS yönteminin kullanılması durumunda veri setinde 1 puanlanan, 0 Atama yönteminin kullanılması durumunda ise 0 puanlanan madde sayısı artmaktadır. Buna bağlı olarak her iki yöntemin de madde güçlük indeksi kestirimlerinde yanlılığa yol açma

olasılığının yüksek olduğu değerlendirilmiştir. 0 Atama yöntemi dışındaki diğer basit atama yöntemleri, madde güçlüklerini .50 civarında yoğunlaştırma eğilimi göstermektedir. EÇO ve ÇVA yöntemleri ise güçlük indekslerinde düşüşe yol açmakla birlikte bu düşüş 0 Atama yöntemi kadar belirgin değildir.

Madde ayırıcılık indeksi kestirimleri dikkate alındığında her bir kayıp veri yönteminde de ayırıcılık indeksleri .20'nin üzerinde kestirilmiştir. Ayırıcılık indeksi en yüksek olan maddeler 9 ve 10. maddeler, en düşük olan maddeler ise 6 ve 13. maddeler olarak belirlenmiştir. Bu belirlemeye bağlı olarak ayırıcılık indeksi .40'ın altında elde edilen 6 ve 13. maddelerin gözden geçirilmesi gerektiği değerlendirilmiştir.

Kullanılan kayıp veri yöntemleri karşılaştırıldığında genel olarak silmeye dayalı yöntemlerin en yüksek, basit atama yöntemlerinin ise en düşük madde ayırıcılık indeksi kestirimleri ürettiği görülmüştür. EÇO ve ÇVA yöntemleri ise diğer yöntemlere göre görece daha ılımlı ve ortada kestirimler üretmektedir.

Madde ayırıcılığının bir diğer kanıtı olarak, madde ve toplam puanlar arasındaki korelasyonlar, NÇKK ve ÇKK kestirimleri ile incelenmiştir. Genel olarak silme ve basit atamaya dayalı yöntemlerin, EÇO ve ÇVA yöntemlerine göre daha yüksek korelasyon katsayıları ürettiği belirlenmiştir.

Madde güvenilirlik katsayısı kestirimleri, her bir kayıp veri yönteminde de madde güvenilirliklerinin düşük olduğunu göstermiştir. Güvenirlik katsayıları genel olarak .40'ın altındadır. Güvenirlik katsayısı en yüksek olan madde 10. madde, en düşük olan maddeler ise 6 ve 13. maddeler olarak belirlenmiştir. Bu belirlemeye bağlı olarak 6 ve 13. maddelerin gözden geçirilmesi gerektiği değerlendirilmiştir.

Farklı kayıp veri yöntemleri karşılaştırıldığında genel olarak silmeye dayalı yöntemlerin madde güvenilirlik katsayısını artırma, basit atama yöntemlerinin ise düşürme eğiliminde olduğu görülmüştür. EÇO ve ÇVA yöntemleri bu yöntemlere göre daha ılımlı ve ortada kestirimler üretmektedir.

SBS 2011 Matematik Testi A Kitapçığına yönelik olarak farklı kayıp veri yöntemlerine göre elde edilen test toplam puanları, her bir kayıp veri yönteminde de normal dağılımdan manidar şekilde sapmakta ve sağa çarpık bir dağılım göstermektedir. DS ve BO yöntemleri dışında diğer kayıp veri

yöntemleri ile elde edilen toplam puanlar arasındaki korelasyonlar manidardır ve pozitif yönlü güçlü bir ilişkinin varlığını göstermektedir.

En yüksek test ortalaması DS yönteminde elde edilmiştir. DS yönteminin kullanılması durumunda Testi alan yanıtlayıcıların test ortalaması 8.95 olarak belirlenmektedir. En düşük test ortalaması ise beklendiği gibi 0 Atama yönteminde elde edilmiştir. 0 Atama yönteminin kullanılması durumunda test ortalaması 6.57 olarak belirlenmektedir.

Farklı kayıp veri yöntemleri, test ortalaması, standart hata ve standart sapma kestirimleri açısından karşılaştırıldığında genel olarak silmeye dayalı yöntemlerin kestirilen değerleri artırma, basit atama yöntemlerinin ise düşürme eğiliminde olduğu görülmüştür. Silmeye dayalı yöntemlerin veri setini homojenleştirmekle birlikte veri kaybından dolayı kestirilen değerlerde artışa yol açtığı, basit atama yöntemlerinin ise veri setini homojenleştirme eğiliminden dolayı kestirilen değerlerde düşüşe yol açtığı değerlendirilmiştir. EÇO ve ÇVA yöntemleri, silme ve basit atama yöntemlerine göre daha ılımlı ve ortada kestirimler üretmektedir.

Test ortalama güçlük indeksi, ortalama ayırıcılık indeksi, NÇKK ve ÇKK kestirimleri dikkate alındığında genel olarak silmeye dayalı yöntemlerin en yüksek, basit atamaya dayalı yöntemlerin ise en düşük kestirimleri ürettiği görülmüştür. EÇO ve ÇVA yöntemleri, silme ve basit atama yöntemlerine göre daha ılımlı ve ortada kestirimler üretmektedir.

Farklı kayıp veri yöntemlerine göre test ortalama güçlük indeksleri .329 ile .447 arasında, test ortalama ayırıcılık indeksleri .502 ile .636 arasında, NÇKK değerleri 478 ile .636 arasında ve ÇKK değerleri .614 ile .804 arasında değişmektedir. Bu kestirimlere bağlı olarak SBS 2011 Matematik Testi A Kitapçığının, güçlük ve ayırıcılık düzeyi düşük maddelerden oluştuğu değerlendirilmiştir.

SBS 2011 Matematik Testi A Kitapçığının güvenirliği, farklı kayıp veri yöntemleri ile elde edilen toplam puanlara yönelik Cronbach α kestirimleri ile incelenmiştir. Her bir kayıp veri yönteminde de .80'in üzerinde güvenirlik katsayısı kestirimi elde edilmiştir. Buna bağlı olarak testin iç tutarlılık anlamında güvenirlik düzeyinin yüksek olduğu değerlendirilmiştir.

Farklı kayıp veri yöntemleri karşılaştırıldığında genel olarak silmeye dayalı yöntemlerin test güvenilirliğini artırma, basit atama yöntemlerinin ise düşürme eğiliminde olduğu görülmüştür. EÇO ve ÇVA yöntemleri bu yöntemlere göre daha ılımlı ve ortada kestirimler üretmektedir.

Bu çalışmanın sonuçları, ihmal edilebilir olmayan kayıp verilerin varlığında, iki kategorili puanlanan maddelerden oluşan testlerin psikometrik özelliklerine yönelik inceleme ve analizlerde, uygun bir kayıp veri yönteminin kullanılmasının gerekli olduğunu göstermektedir. Silmeye dayalı yöntemler ve 0 Atama yöntemi, bu tür verilerde kullanılabilir değildir. Basit atama yöntemlerinin ise yanlış kestirimler üretme olasılığı yüksektir. EÇO ve ÇVA yöntemleri, bu tür verilerde kullanılması en uygun kayıp veri yöntemleri olarak değerlendirilmektedir.

Öneriler

Bu çalışmada ulaşılan sonuçlar doğrultusunda, kayıp veriler içeren ve iki kategorili puanlanan maddelerden oluşan testlerin psikometrik özelliklerine yönelik inceleme, analiz ve kestirimlerde;

1. 0 Atama yönteminin kesinlikle kullanılmaması,
2. Ancak kayıp verilerin ihmal edilebilir olduğuna yönelik sağlam kanıtlar elde edilebilmişse silmeye dayalı yöntemlerin kullanılması,
3. Kayıp verilerin ihmal edilebilir olmadığı durumlarda;
 - a. Kayıp veri miktarı çok fazla (araştırmanın tür ve modeline göre değişebilmekle birlikte, genel olarak çok değişkenli istatistiksel tekniklerin kullanımında %5, diğer istatistiksel tekniklerin kullanımında ise %15'in üzerinde) ise veri toplama aracının gözden geçirilmesi ya da gerekli tedbirler alınarak veri toplama sürecinin yenilenmesi,
 - b. Kayıp veri miktarı makul düzeyde ya da az ise en çok olabilirlik ve çoklu veri atama yöntemlerinin tercih edilmesi

önerilmektedir.

Test geliştirme ve uygulama süreçlerinde, kayıp verilerin doğasına yönelik daha sağlıklı ön bilgiler edinilebilmesi için 'bilgisi olmama', 'erişilememe', 'boş bırakma', 'test uygulamasını erken bitirme ya da terk etme' gibi durumlara yönelik değişkenlerin tanımlanması ve bu değişkenlerin analizlerde dikkate alınması önerilmektedir.

İleri araştırmalara yönelik olarak;

1. Benzer incelemelerin, SBS'de yer alan diğer alt testler için de gerçekleştirilmesi,
2. Yanıt örüntülerinin yanı sıra,
 - a. Yanıtlayıcılara yönelik olarak okul başarı puanları, matematik ders notları, matematiğe karşı tutum düzeyleri,
 - b. Test ve test uygulamasına yönelik olarak, (üç ya da dört yanlış yanıtın bir doğru yanıtı götürmesi gibi şekillerde ifade edilen) düzeltme kuralı, test ortamı, uygulama süresi
3. Kayıp veri yöntemlerinin ve bu yöntemlere göre testin psikometrik özelliklerinin, MTK temelinde incelenmesi

önerilmektedir.

KAYNAKÇA

- Acock, A.A. (2005). Working with Missing Values. *Journal of Marriage and Family*, s.65, sy.1012-1028.
- Aiken, L.R. (1995). *Rating Scales and Checklist: Evaluating Behaviour, Personality and Attitudes*. New York: John Wiley&Sons, Inc.
- Allison, P.D. (2002). *Missing Data*. California: Sage Publication, Inc.
- Allison, P.D. (2006). Imputation of Categorical Variables with PROC MI. *Annual Meeting of the SAS Users Group International*, San Francisco, CA.
- Allison, P.D. (2009). Missing Data. Ed. Roger E. Millsao ve Alberto Maydeu-Olivares, *Quantitative Methods in Psychology*, sy.72-89. London: SAGE Publication.
- Arbuckle, J.L. (2007). *Amos 16 User's Guide*. Chicago:SPSS.
- Bal, C. (2003). Çok Gruplu Veri Setlerinde Eksik Gözlem Sorununun Çözülmesi ve Sağlık Alanında Bir Uygulama. *Yayımlanmamış doktora tezi*. Eskişehir: Osmangazi Üniversitesi, Sağlık Bilimleri Enstitüsü, Biyoistatistik ABD.
- Baykul, Y. (2000). *Eğitimde ve Psikolojide Ölçme: Klasik Test Teorisi ve Uygulaması*. Ankara: ÖSYM Yayınları.
- Bentler, P.M. ve Wu, E.J. (2002). *EQS 6 for Windows User's Guide*. Encino, CA: Multivariate Software.
- Brick, J.M. (2001). Analysis of Potential Nonresponse Bias. *Annual Meeting of the American Statisical Association*, August 5-9, 2001.
- Brown, C.H. (1983). Asymptotic Comparision of Missing Data Procedures for Estimating Factor Loadings. *Biometrika*, s.48, n.2, sy.269-291.
- Crocker, L. ve Algina, J. (1986). *Introduction to Classical and Modern Test Theory*. Belmont CA: Wadsworth Thomson Learning Company.
- Culbertson, M.J. (2011). Is It Wrong? Handling Missing Responses in IRT. *Annual Meeting of the National Council on Measurement in Education*, April 2011.
- Çokluk, Ö. ve Kayrı, M. (2011). The Effects of Methods of Imputation for Missing Values on the Validity and Reliability of Scales. *Educational Sciences: Theory&Practice*, s.11, sy.303-309.

- De Ayala, R.J., Plake, B.S. ve Impara, J.C. (2011). The Impact of Omitted Responses on the Accuracy of Ability Estimation in Item Response Theory. *Journal of Educational Measurement*, s.38, sy.213-234.
- DeMars, C. (2002). Incomplete Data and Item Parameter Estimates Under JMLE and MML Estimation. *Applied Measurement in Education*, s.15, sy.15-31.
- Demir, E. ve Parlak, B. (2012). Türkiye’de Eğitim Araştırmalarında Kayıp Veri Sorunu. *Eğitimde ve Psikolojide Ölçme ve Değerlendirme Dergisi*, s.3(1), sy.20-241.
- Dempster, A.P., Laird, N.M. ve Rubin, D.B. (1977). Maximum Likelihood Estimation from Incomplete Data via the EM Algorithm. *Journal of the Royal Statistical Society*, seri B, s.39, sy.1-38.
- Dural, S. (2010). Farklı Kayıp Veri Tekniklerinin Çok Göstergeli Örtük Gelişme Modelleri Üzerindeki Etkisi. *Yayımlanmamış doktora tezi*. İzmir: Ege Üniversitesi, Sağlık Bilimleri Enstitüsü, Psikometri ABD.
- Eberle, W. ve Toutenburg, H. (1999). Handling of Missing Values in Statistical Software Packages for Windows. *Paper: Institut für Statistik, Sonderforschungsbereich 386*, s.170.
- Elashoff, J.D. ve Elashoff, R.M. (1971). Missing Data Problems for Two Samples on a Dichotomous Variable. *Research and Development Memorandum*, no:73, sy.40-86. California, Stanford: Stanford Center for Research and Development in Teaching.
- Enders, C.K. (2006). A Primer on the Use of Modern Missing-Data Methods in Psychomatic Medicine Research. *Psychomatic Medicine*, s.68, sy.427-436.
- Enders, C.K. (2010). *Applied Missing Data Analysis*. New York: The Guilford Press.
- Erkuş, A. (2012). *Psikolojide Ölçme ve Ölçek Geliştirme: Temel Kavramlar ve İşlemler*. Birinci baskı. Ankara: Pegem Akademi Yayıncılık.
- Finch, H. (2008). Estimation of Item Response Theory Parameters in the Presence of Missing Data. *Journal of Educational Measurement*, s.45, no:3, sy.225-245.
- Graham, J.W. (2009). Missing Data Analysis: Making It Work in the Real World. *Annual Review of Psychology*, s.60, n.4, sy.549-576.
- Groves, R.M. (2006). Nonresponse Rates and Nonresponse Bias in Household Surveys. *Public Opinion Quartely*, s.70, n.5, sy.646-675.

- Hohensinn, C. ve Kubinger K.D. (2011). On the Impact of Missing Values on Item Fit and the Model Validness of the Rasch Model. *Psychological Test and Assesment Modeling*, s.53, sy.380-393.
- Holman, D.J. (1991-2003). *Mle: A Progammig Language for Building Likelihood Models, Version 2.1., User's Manual*. Seattle, WA: The University of Washington, Center for Statistics and Social Sciences.
- Jöreskog, K.G. ve Sörbom, D. (2006). *LISREL 8.8 for Windows*. Lincolnwood, IL: Scientific Software International.
- Kanuk, L. ve Berenson, C. (1975). Mail surveys and Response Rates. *Journal of Marketing Research*, s.12, sy.440-453.
- Linacre, J.M. (2004). Rasch Model Estimation: Further Topics. *Journal of Applied Measurement*, s.5, no:1, sy:95-110.
- Little, R.J.A ve Rubin, D.B. (1987). *Statistical Analysis with Missing Data, 2nd ed.* New York: John Wiley & Sons, Inc.
- Lord, F.M. (1973). Estimation of Latent Ability and Item Parameters When There are Omitted Responses. *Report Bulletin*, 73-37. Princeton: Educational Testing Service.
- Lord, F.M. (1974). Estimation of Latent Ability and Item Parameters When There are Omitted Responses. *Psychometrika*, s.39, sy.247-264.
- Lord, F.M. (1980). *Application of Item Response Theory to Practical Testing Problems*. Hillsdale, NJ: Erlbaum.
- Lord, F.M. (1983). Maximum Likelihood Estimation of Item Response Parameters When Some Responses are Omitted. *Psychometrika*, s.48, sy.477-482.
- Lord, F.M. ve Novick, M.R. (1968). *Statistical Theory of Mental Test Score*. California: Addison-Wesley Publishing.
- McKinley, R.L. ve Reckase, M.D. (1983). The Use of IRT Analysis on Dichotomous Data from Multidimesional Tests. *Annual Meeting of American Educational Research Association*, April 1983, Montreal.
- Mislevy, R.J. ve Wu, P.K. (1996). *Missing Responses and IRT Ability Estimation: Omits, Choice, Time Limits, Adaptive Testing*. ETS Research Report RR-96-30-ONR, Princeton, NJ: Educational Testing Service.
- Muthen, L.K. ve Muthen, B.O. (2009). *Mplus User's Guide (5. Baskı)*. Los Angeles: Muthen&Muthen.

- Özçelik, D. A. (1981). *Okullarda Ölçme ve Değerlendirme*. Ankara: ÜSYM-Eğitim Yayınları:3.
- Peng, C.Y.J., Harwell,M., Liou, S.M. ve Ehman L.H. (2007). *Advances in Missing Data Methods and Implication for Educational Research*. Ed. Sholomo S. Sawilowski, *Real Data Analysis*, sy.31-77. USA: IAP-Information Age Publishing.
- Puma, M.J., Olsen, R.B., Bell, S.H. ve Price, C. (2009). *What Do When Data are Missing in Group Randomized Controlled Trial* (NCE 2009-0049). Washington, DC: National Center for Education Evaluation and Regional Assistance, Institute of Educational Sciences, U.S. Department of Education.
- Raghunathan, T.E., Lepkowski, J.M., Hoewyk, J.V. ve Solenberger, P. (2001). A Multivariate Technique for Multiply Imputing Missing Values Using a Sequence of Regression Models. *Survey Methodology*, s.27, n.1, sy.85-95.
- Raghunathan, T.E., Solenberger, P.W. ve Van Hoewyk, J., (2002). *IVEware: Imputation and Variance Estimation Software User Guide*. University of Michigan, Survey Research Center.
- Rubin, D.B. (1976). Inference and Missing Data. *Biometrika*, s.63, n.3, sy.581-592.
- Rubin, D.B. (1987). *Multiple Imputation for Nonresponse in Surveys*. New York: John Wiley & Sons, Inc.
- Schaffer, J.L. (1997). *Analysis of Incomplete Multivariate Data*. London: Chapman&Hall.
- Schaffer, J.L. (1999). *NORM: Multiple Imputation of Incomplete Multivariate Data under a Normal Model*. University Park: Department of Statistics, Pennsylvania State University.
- Tavşancıl, E. (2010). *Tutumların Ölçülmesi ve SPSS ile Veri Analizi*. Dördüncü baskı. Ankara: Nobel Yayın Dağıtım.
- Tekin, H. (2007). *Eğitimde Ölçme ve Değerlendirme* (18. baskı). Ankara: Yargı Yayınları.
- Whipple, T.W. ve Muffo, J.A. (1982). Adjusting for Nonresponse Bias: The Case of an Alumni Survey. *22nd Annual Forum of the Association for Institutional Research*, May, 16-19. Denver: Cleveland State University.

EKLER

EK A
SBS 2011 MATEMATİK TESTİ A KİTAPÇIĞI

8. SINIF

MATEMATİK

A

1. $(-3)^{-2}$ sayısı aşağıdaki sayılardan hangisi ile çarpılırsa sonuç 3 olur?

A) 3^3 B) 3^{-1}
C) 3^2 D) $(-3)^{-3}$

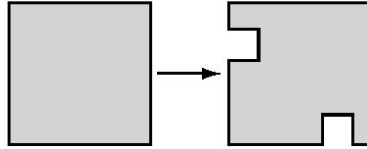
2.



Bir okulda yapılan atık pil toplama kampanyasında 2^{11} adet kalem pil toplanmıştır. Toplanan bu piller ile en fazla kaç metrekaarelik arazinin kirlenmesinin önüne geçilmiş olur?

A) 2^{13} B) 4^{22} C) 6^{11} D) 8^{11}

3. Efe, proje ödevi için alanı 484 cm^2 olan kare şeklindeki kartondan, alanları otuz altışar santimetrekaare olan iki kareyi şekildeki gibi kesip çıkarmıştır.



Kalan kartonun çevre uzunluğu kaç santimetredir?

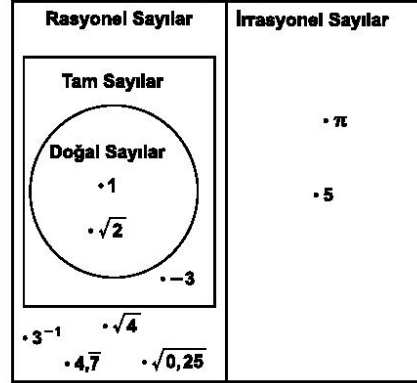
A) 88 B) 112 C) 124 D) 136

4. Bir basketbol sahasında orta yuvarlak denilen ve alanı $9,72 \text{ m}^2$ olan dairesel bölgenin çapı kaç metredir? (π yerine 3 alınız.)

A) 2,8 B) 3,2 C) 3,6 D) 4,1

5. Sayı kümeleri arasındaki ilişkiye örnek vermek amacıyla aşağıdaki şema çizilmiştir.

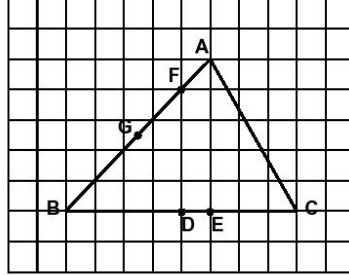
Gerçek Sayılar



Bu şemanın doğru olabilmesi için hangi iki sayının yer değiştirmesi gerekir?

A) π ile $\sqrt{0,25}$ B) π ile $4,\overline{7}$
C) 5 ile $\sqrt{4}$ D) 5 ile $\sqrt{2}$

6. Verilen ABC üçgeninde hangi iki noktadan geçen doğru, üçgenin bir kenarının orta dikmesidir?

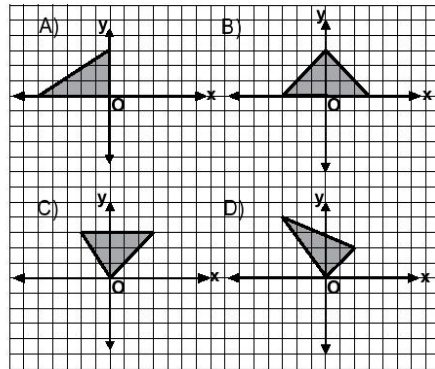


- A) A ile E B) F ile E
C) A ile D D) F ile D

7. Nisan ayında 1 ton inşaat demirinin fiyatı 1230 TL ile 1270 TL arasında değişmiştir. Nisan ayında 11 ton demir alan bir müteahhit kaç TL ödemiş olabilir?

- A) 13 420 B) 13 470
C) 13 680 D) 14 080

8. x eksenine göre yansıması ile orijin etrafında 180° dönme altındaki görüntüsü aynı olan şekil aşağıdakilerden hangisidir?

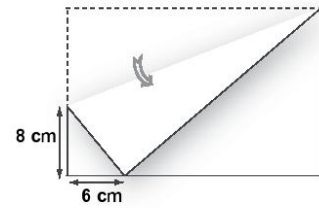


9. Aşağıdakilerden hangisi üçgen şeklinde yapraklardan oluşan yandaki takvimin bir yaprağıdır?



- A) B)
C) D)

10. Dikdörtgen biçimindeki bir kâğıt, şekildeki gibi bir köşesi uzun kenarının üzerine gelecek biçimde katlanıyor.



Şekilde verilen ölçülere göre, bu kâğıdın kısa kenarının uzunluğu kaç santimetredir?

- A) 14 B) 16 C) 18 D) 22

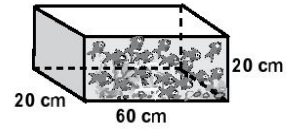
11. Atatürk Orman Çiftliği'ndeki Hayvanat Bahçesi'ne giriş bilet ücretleri aşağıda verilmiştir:



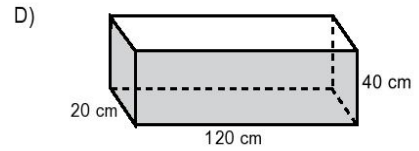
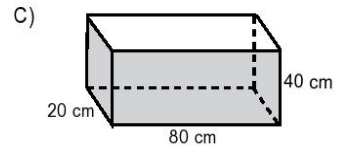
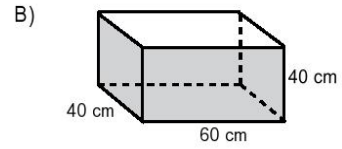
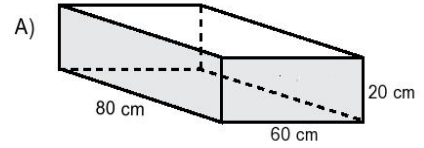
Bir günde 91 kişi bilet alarak Hayvanat Bahçesi'ni gezmiş ve 230 TL elde edilmiştir. Buna göre kaç öğrenci Hayvanat Bahçesi'ni gezmiştir?

- A) 78 B) 67 C) 57 D) 38

- 12.



Dikdörtgenler prizması şeklinde akvaryum imal edilen bir atölyede, yukarıdaki akvaryumun 4 katı hacminde imal edilecek yeni akvaryum hangisi olamaz?

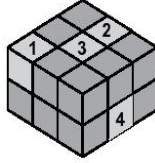


8. SINIF

MATEMATİK

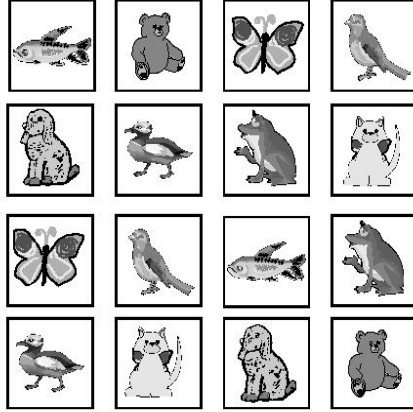
A

13. Birim küplerden oluşan yandaki yapıda, numaralandırılmış küplerden hangisi çıkarıldığında yapının yüzey alanı değişmez?



- A) 1 B) 2 C) 3 D) 4

14. Aşağıdaki kartlar ters çevrilip karıştırılıyor ve resimler görülmeyecek şekilde yeniden diziliyor.

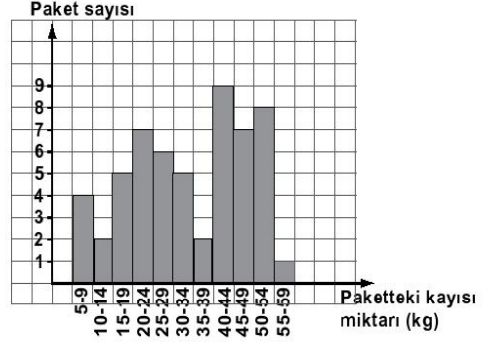


Rastgele açılan iki kartta da kelebek resminin bulunma olasılığı nedir?

- A) $\frac{1}{192}$ B) $\frac{1}{120}$ C) $\frac{1}{64}$ D) $\frac{1}{56}$

15. Aşağıda bir fabrikada hazırlanan kayısı paketlerinin kütlelerine göre dağılımı verilmiştir:

Grafik: Kayısı Paketlerinin Kütlelerine Göre Dağılımı



Grafığe göre, aşağıdakilerden hangisi doğrudur?

- A) Kütlesi en fazla olan paket 44 kilogramdır.
 B) Toplam 59 paket hazırlanmıştır.
 C) Her gruptan en az 2 paket hazırlanmıştır.
 D) 40 kg ve üzerinde toplam 25 paket hazırlanmıştır.
16. En çok hangi renk otomobil satışı yapıldığını araştıran biri, en çok beyaz otomobillerin satıldığı sonucuna ulaşmıştır. Bu sonucu elde etmek için hangi ölçü kullanılmıştır?

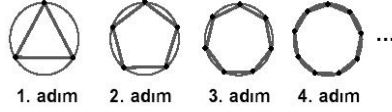
- A) Ortanca B) Tepe değer
 C) Aritmetik ortalama D) Açıklık

8. SINIF

MATEMATİK

A

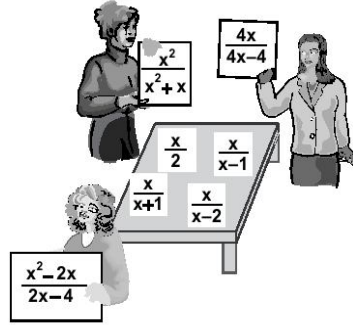
17.



Yukarıda verilen örüntü, aynı kurala göre devam ettirildiğinde 19. adımdaki çemberin içine çizilen çokgenin kenar sayısı kaçtır?

- A) 24 B) 33 C) 39 D) 42

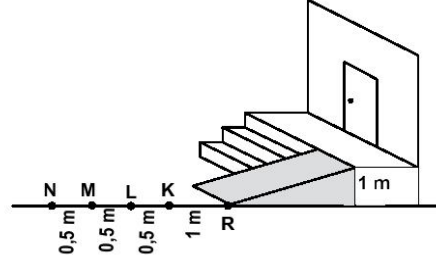
18. Aşağıda üç oyuncunun ellerindeki kartlara yazılmış rasyonel cebirsel ifadeler görülmektedir.



Her oyuncu elindeki karta yazılmış ifadenin en sade biçimi yazılı olan kartı aldığı anda masada hangi kart kalır?

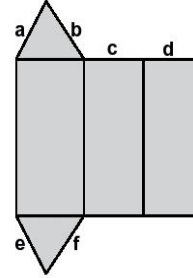
- A) $\frac{x}{2}$ B) $\frac{x}{x-1}$ C) $\frac{x}{x+1}$ D) $\frac{x}{x-2}$

19. Şekilde R noktasından başlayan engelli rampasının eğimi %10'dur. Yerden yüksekliği 1 m olan bu rampanın eğimi %8 olsaydı, rampa hangi noktadan başlardı?



- A) N B) M C) L D) K

20. Tabanı çeşitkenar üçgen şeklinde olan bir prizmanın açılımı aşağıda verilmiştir.



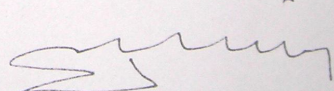
Aşağıdakilerin hangisindeki ayrıt uzunlukları birbirine eşittir?

- A) a ile c B) a ile f
C) f ile d D) e ile d

MATEMATİK TESTİ BİTTİ.
FEN VE TEKNOLOJİ TESTİNE GEÇİNİZ.

EK B

VERİ TALEP DİLEKÇESİ

T.C. MİLLÎ EĞİTİM BAKANLIĞINA (Yenilik ve Eğitim Teknolojileri Genel Müdürlüğü)	<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="text-align: center; padding: 2px;">Yenilik ve Eğitim Teknolojileri Genel Müdürlüğü</td> </tr> <tr> <td style="text-align: center; padding: 2px;">18 Temmuz 2012</td> </tr> <tr> <td style="text-align: center; padding: 2px;">NUMARA: 12070</td> </tr> </table>	Yenilik ve Eğitim Teknolojileri Genel Müdürlüğü	18 Temmuz 2012	NUMARA: 12070
Yenilik ve Eğitim Teknolojileri Genel Müdürlüğü				
18 Temmuz 2012				
NUMARA: 12070				
Ankara Üniversitesi Eğitim Bilimleri Fakültesi Ölçme ve Değerlendirme Anabilim Dalında doktora öğrenimi görmektedirim. Kayıp veri yöntemlerinin testlerin psikometrik özelliklerinin kestiriminde kullanımı üzerine hazırladığım tez önerim, enstitü tarafından belirlenen tez izleme jürisi tarafından kabul edilmiş olup danışmanım tarafından onaylanmış bir örneği ektedir. Tez çalışmamda söz konusu yöntemleri, gerçek ve yeterli büyüklükte bir veri seti üzerinde uygulamayı düşünmekteyim. Bu nedenle Türkiye’de ilköğretim 8 inci sınıf öğrencilerine yönelik Seviye Belirleme Sınavı (SBS)’nin 2009 ve/veya 2010 ve /veya 2011 ve/veya 2012 uygulamalarına yönelik sadece A-B-C-D yanıt örüntülerine ve cinsiyet, il gibi temel değişkenlere ihtiyaç duymaktayım. Başkaca bir kişisel bilgiye ihtiyacım bulunmamaktadır. Bilgilerinizi ve gereğini arz ederim.				
18../07/2012  Ergül DEMİR				
EK: Doktora Tez Önerisi (1 Adet-9 Sayfa)				
 Prof. Dr. Ali Zaimettin Koc				
Adres : MEB HBÖGM Açık Lise Müdürlüğü Teknik Okullar/Ankara Kurum Sicil No. : 304327 Tel : 0505 430 12 48 E-Posta : erguldemir@gmail.com				

EK C
ÖZGEÇMİŞ

Adı ve Soyadı : Ergül DEMİR

Doğum Tarihi : 15.11.1977

İletişim Bilgileri : (Adres) Ankara Üniversitesi Eğitim Bilimleri Fakültesi
Ölçme ve Değerlendirme Bölümü. Cebeci / ANKARA
(Tel.) 0312 363 33 50 - 3122

E-Posta Adresi : erguldemir@gmail.com

Öğrenim Durumu :

Derece	Bölüm/Program	Üniversite	Yıl
Yüksek Lisans	Eğitimde Ölçme ve Değerlendirme	Hacettepe Üniversitesi	2010
Yüksek Lisans	Kamu Yönetimi	TODAİE	2010
Lisans	Matematik Öğretmenliği	Dokuz Eylül Üniversitesi	1998
Lisans	İşletme	Anadolu Üniversitesi	2010

İş Deneyimi :

Unvan	Görev Yeri	Yıl
Öğretim Görevlisi	Ankara Üniversitesi Eğitim Bilimleri Fakültesi	2012 -
Matematik Öğretmeni	MEB Merkez ve Merkeze Bağlı Taşra Teşkilatları	2007 - 2012
Matematik Öğretmeni	MEB Taşra Teşkilatları	2002 - 2007
Matematik Öğretmeni	Özel Öğretim Kurumları	1999 - 2002