



**DOĞRUSAL KARIŞIK MODELLERDE PARAMETRELER İÇİN GÜVEN
ARALIKLARINA YENİ YAKLAŞIMLAR**

Hatice Tül Kübra AKDUR

**DOKTORA TEZİ
İSTATİSTİK ANABİLİM DALI**

**GAZİ ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ**

HAZİRAN 2017

Hatice Tül Kübra AKDUR tarafından hazırlanan “DOĞRUSAL KARIŞIK MODELLERDE PARAMETRELER İÇİN GÜVEN ARALIKLARINA YENİ YAKLAŞIMLAR” adlı tez çalışması aşağıdaki jüri tarafından OY BİRLİĞİ ile Gazi Üniversitesi İstatistik Anabilim Dalında DOKTORA TEZİ olarak kabul edilmiştir.

Danışman: Prof. Dr. Hülya BAYRAK

İstatistik, Gazi Üniversitesi

Bu tezin, kapsam ve kalite olarak Doktora Tezi olduğunu onaylıyorum.



İkinci Danışman: Doç. Dr. Fikri GÖKPINAR

İstatistik, Gazi Üniversitesi

Bu tezin, kapsam ve kalite olarak Doktora Tezi olduğunu onaylıyorum.



Başkan: Prof. Dr. Tahir HANALIOĞLU

Endüstri Mühendisliği, TOBB Ekonomi ve Teknoloji Üniversitesi

Bu tezin, kapsam ve kalite olarak Doktora Tezi olduğunu onaylıyorum.



Üye: Prof. Dr. Ensar BAŞPINAR

Zootekni, Ankara Üniversitesi

Bu tezin, kapsam ve kalite olarak Doktora Tezi olduğunu onaylıyorum.



Üye: Prof. Dr. Semra ERBAŞ

İstatistik, Gazi Üniversitesi

Bu tezin, kapsam ve kalite olarak Doktora Tezi olduğunu onaylıyorum.



Üye: Doç. Dr. Yaprak Arzu ÖZDEMİR

İstatistik, Gazi Üniversitesi

Bu tezin, kapsam ve kalite olarak Doktora Tezi olduğunu onaylıyorum.



Tez Savunma Tarihi: 09/06/2017

Jüri tarafından kabul edilen bu tezin Doktora Tezi olması için gerekli şartları yerine getirdiğini onaylıyorum.



Prof. Dr. Hadi GÖKÇEN
Fen Bilimleri Enstitüsü Müdürü

ETİK BEYAN

Gazi Üniversitesi Fen Bilimleri Enstitüsü Tez Yazım Kurallarına uygun olarak hazırladığım bu tez çalışmada;

- Tez içinde sunduğum verileri, bilgileri ve dokümanları akademik ve etik kurallar çerçevesinde elde ettiğimi,
- Tüm bilgi, belge, değerlendirme ve sonuçları bilimsel etik ve ahlak kurallarına uygun olarak sunduğumu,
- Tez çalışmada yararlandığım eserlerin tümüne uygun atıfta bulunarak kaynak gösterdiğimi,
- Kullanılan verilerde herhangi bir değişiklik yapmadığımı,
- Bu tezde sunduğum çalışmanın özgün olduğunu,

bildirir, aksi bir durumda aleyhime doğabilecek tüm hak kayıplarını kabullendiğimi beyan ederim.

Hatice Tül Kübra AKDUR
30/ 06/ 2017

DOĞRUSAL KARIŞIK MODELLERDE PARAMETRELER İÇİN GÜVEN
ARALIKLARINA YENİ YAKLAŞIMLAR

(Doktora Tezi)

Hatice Tül Kübra AKDUR

GAZİ ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ

Haziran 2017

ÖZET

Bu tezde, doğrusal karışık-etki modellerindeki parametrelerin tahmin ve çıkarım yöntemleri incelenmiştir. Model parametreleri ile ilgili çıkarım yapmak için literatürde hipotez testleri ve güven aralıkları yöntemlerine başvurulur. Tekrarlanan ölçümlü ve uzun vadeli çalışmalarda sıklıkla kullanılan bazı doğrusal karışık-etki modelleri incelenmiştir. İlgili modellerin sabit etki parametreleri için literatürde kullanılan güven aralığı metotları ve yeni uyarlanan güven aralığı metodu doğru sonuç verme açısından karşılaştırılmıştır. Güven aralıklarının kurulması bazı dağılımsal varsayımların belirlenmesini gerektirir. Bu dağılımsal varsayımlar sağlanmadığında veya küçük örnek çaplarında literatürde yer alan standart analitik metotlar doğru sonuç vermemektedir. Bu nedenle alternatif olarak parametrik bootstrap güven aralığı, profil olabilirlik ve yapay-skor (pseudo-score) güven aralığı yöntemine başvurulabilir. İç-içe hata regresyon modelinin sabit etki parametresi için gerçekleştirilen simülasyonlarda küçük örnek çaplarında parametrik bootstrap metodunun iyi bir alternatif olduğu görülmüştür. Ayrıca iç-içe hata regresyon modeli için deney birimi içi korelasyon orta dereceli olduğunda ve orta örnek çapı büyüklükleri için profil olabilirlik metodunun kapsama olasılığı bakımından iyi sonuç verdiği görülmüştür. Yapay-skor güven aralığı, ilk olarak bu tezde doğrusal karışık-etki modellerinin sabit etki parametresi için uyarlanmıştır. Yeni uyarlanan yapay-skor metodu ve literatürde var olan diğer güven aralığı metotlarının performanslarının bağımlı değişkenin farklı kovaryans yapılarında nasıl değiştiğini görebilmek amacıyla geniş bir simülasyon çalışması gerçekleştirilmiştir. Küçük örnek çaplarında ve bağımlı değişkenin otheregresif kovaryans yapısı durumunda yapay-skor metodunun diğer metotlara göre kapsama olasılığı ve ortalama uzunluk bakımından daha iyi sonuç verdiği görülmüştür. Yapay-skor metodu doğrusal karışık etki modellerinin sabit etki parametresi için yeni bir güven aralığı yöntemi olarak istatistik literatürüne kazandırılmıştır.

Bilim Kodu : 20503
Anahtar Kelimeler : Karışık-etkili model, güven aralığı, parametrik bootstrap, yapay-skor, profil olabilirlik, pivot
Sayfa Adedi : 119
Danışman : Prof.Dr. Hülya BAYRAK
İkinci Danışman : Doç.Dr. Fikri GÖKPINAR

NEW APPROACHES TO CONFIDENCE INTERVALS FOR PARAMETERS IN LINEAR MIXED MODELS

(Ph. D. Thesis)

Hatice Tül Kübra AKDUR

GAZİ UNIVERSITY
GRADUATE SCHOOL OF NATURAL AND APPLIED SCIENCES

June 2017

ABSTRACT

In this thesis, estimation and inference methods of linear mixed model parameters are examined. Hypothesis testing and confidence interval procedures are performed in the statistics literature in order to make inference about model parameters. Some set of linear mixed- effect models, which are frequently used in repeated measures and longitudinal studies, are investigated. The existing confidence interval methods and new adaptation method for fixed effect parameters of the models are compared in terms of providing the accurate results. Constructing a confidence interval requires to be determined some distributional assumptions. Standard analytical confidence interval methods do not produce consistent results when the distributional assumptions are not met or small sample sizes are used. For this reason, parametric bootstrap confidence interval, profile likelihood and pseudo-score confidence interval method can be applied as an alternative method. It is seen that parametric bootstrap is a good alternative for small sample sizes in the simulation study for fixed effect parameter of nested error regression model. In nested error regression model it is reported that profile likelihood confidence interval method provides better coverage probability results under medium sample sizes and moderate intra-class correlation. The pseudo-score confidence interval method is firstly adapted for fixed effect parameters of linear mixed models. An extensive simulation study is conducted in order to see how the performances of the new adaptive pseudo-score and the other existing methods change under the different covariance structure of response variable. In the cases of small samples and autoregressive covariance structure of response variable pseudo-score method provides better coverage rates and average expected lengths than the existing methods. Therefore, pseudo-score confidence interval method is introduced into the statistics literature as a new confidence interval method for fixed effect parameters of linear mixed-effect models.

Science Code	: 20503
Key Words	: Mixed-effect models, confidence interval, pseudo-score, parametric bootstrap, profile likelihood, pivot
Page Number	: 119
Supervisor	: Prof. Dr. Hülya BAYRAK
Co-Supervisor	: Assoc. Prof. Dr. Fikri GÖKPINAR

TEŞEKKÜR

Doktora ders dönemi ve yeterlilik sınavım dahil olmak üzere tez sürecim boyunca bana her türlü desteği veren, çıkmaza düştüğüm zamanlarda beni yeniden teşvik eden güler yüzünü hiçbir zaman esirgemeyen rehberliğine her zaman güvendiğim kıymetli danışmanım Prof. Dr. Hülya Bayrak'a en içten şükranlarımı sunarım. Tez izleme komitelerinde değerli zamanlarını bana ayıran, yapıcı eleştirileri ve önerileri için eş danışmanım Doç. Dr. Fikri Gökpınar, Prof. Dr. Ensar Başpınar ve Doç. Dr. Yaprak Arzu Özdemir'e teşekkürü bir borç bilirim. Washington Eyalet Üniversitesi'ndeki yüksek lisans danışmanım Prof. Dr. Nairanajana Dasgupta'ya doktora eğitimime devam etmem gerektiği konusundaki teşvikleri için çok teşekkür ederim. Calgary Üniversitesi'nde doktora yurt dışı araştırmam için 2214-A bursunu sağlayan Tübitak'a şükranlarımı sunarım. Tübitak burs sürecindeki danışmanım Doç. Dr. Alex de Leon'a doktoradan sonra yapacağım çalışmalar konusunda bana rehberlik ettiği için teşekkürlerimi sunarım. Birlikte çok güzel anıları paylaştığımız, her zorluğa birlikte göğüs gerdiğimiz canım dostum Deniz Özonur'a doktora sürecinde her zaman yanımda olduğu için çok teşekkür ederim. Burada istesem de adlarını sayamayacağım 2003 senesinde çıktığım bu akademik yolculukta bana destek olup rehberlik etmiş tüm hocalarıma ve yol arkadaşlarıma en içten şükranlarımı sunarım.

Hiç bir zaman kardeşliğini esirgemeyen kuzenim Gülran Aytaç'a teşekkürü bir borç bilirim. Sevgili eşim Deniz Akdur'a doktora sürecimde sonsuz desteği, sabrı ve tahammülüyle yanımda olduğu için en içten teşekkürlerimi sunarım. Annem Sultan Şenol ve babam Danyal Şenol'a her zaman bana güvendikleri ve en zor durumlarda yeniden motive olmamı sağladıkları için sonsuz minnetlerimi sunarım. Son olarak tüm aile bireylerime koşulsuz sevgi ve destekleri için en içten şükranlarımı sunarım. Tezimi geçen yıl aramızdan ayrılan sevgili babaannem Ayser Şenol'a gururla ithaf etmek isterim.

İÇİNDEKİLER

	Sayfa
ÖZET	iv
ABSTRACT.....	v
TEŞEKKÜR	vi
İÇİNDEKİLER	vii
ÇİZELGELERİN LİSTESİ.....	x
ŞEKİLLERİN LİSTESİ.....	xi
SİMGELER VE KISALTMALAR.....	xv
1. GİRİŞ	1
2. KLASİK YÖNTEMLERİN TANITILMASI	7
3. DOĞRUSAL KARIŞIK-ETKİ MODELİ	9
3.1. Tek Yönlü Rasgele Etkili ANOVA Modeli (RANOVA)	12
3.2. Fay-Herriot Model	13
3.3. İç içe Hata Regresyon Modeli	13
3.4. Doğrusal Karışık Etki Modellerinde Kovaryans Yapıları	16
3.4.1. Rasgele kesim modeli	18
3.4.2. Rasgele eğim modeli	21
3.4.3. Rasgele etkiler için kovaryans yapıları	26
3.4.4. Hata terimleri için kovaryans yapıları	27
4. DKE MODELLERİNDE TAHMİN VE KESTİRİM METODLARI	33
4.1. ANOVA Tahmini	33
4.2. En Çok Olabilirlik Tahmin Metodu	33
4.3. Kısıtlı En Çok Olabilirlik Tahmin Metodu	36
4.4. Karma Etkiler için Nokta Tahmini	39

	Sayfa
4.5. Sabit Etkiler için Nokta Tahmini	41
4.6. Hata Kareler Ortalaması için Yaklaşımlar	41
4.6.1. Naive hko yaklaşımı	41
4.6.2. Sabit etki tahminleri için Kenward-Roger varyans- kovaryans yaklaşımı	42
4.7. DKE Modelinde Sabit Etkiler için Analitik Yaklaşımlar	45
5. DKE MODELLERİNDE SAYISAL ALGORİTMALAR VE İSTATİSTİKSEL YAZILIMLAR	47
5.1. BM Algoritması için Bazı Tanımlamalar ve Gösterimler	47
5.2. Beklenti ve Maksimizasyon Algoritması	49
5.3. Karışık Etki Modelleri için BM Algoritması	58
5.3.1. İki varyans bileşenli karışık-etki modeli	60
5.3.2. İç içe hata regresyon modeli için BM algoritması	64
5.4. Doğrusal Karışık Etki Modellerinde Kullanılan İstatistiksel Yazılımlar	65
6. GÜVEN ARALIĞI	67
6.1. Pivot Ölçüm Metodu	70
7. DKE MODELİNDE GÜVEN ARALIĞI METODLARI	73
7.1. Profil Olabilirlik Güven Aralığı Metodu	73
7.2. Wald Güven Aralığı Metodu	75
7.3. t-Naive Güven Aralığı Metodu	77
7.4. Yapay-Skor (Pseudo-Score) Güven Aralığı Metodu	78
7.5. Parametrik Bootstrap Güven Aralığı Metodu	79
7.5.1. Bootstrap-t güven aralığı	81
8. SİMULASYON ÇALIŞMALARI VE VERİ ANALİZLERİ	83
8.1. Uygulama: 6 Şehir Çalışması	83

	Sayfa
8.2. İç-içe Hata Regresyon Modeli için Simülasyon Çalışması	87
8.3. Rasgele Eğim Modeli için Simülasyon Çalışması	95
8.4. Uygulama	103
9. SONUÇ VE ÖNERİLER	107
KAYNAKLAR	111
ÖZGEÇMİŞ	117



ÇİZELGELERİN LİSTESİ

Çizelge	Sayfa
Çizelge 5.1. Yapay dölleme veri seti	51
Çizelge 5.2. Tek faktörlü rastgele etki modeli için varyans analizi tablosu	55
Çizelge 5.3. Domuz çiftliğinde yavruların ağırlık artış veri seti	62
Çizelge 5.4. DKE modeli matris gösterimi	63
Çizelge 5.5. Beklenti ve maksimizasyon adımları	63
Çizelge 5.6. DKE modelinin sabit etkileri ve varyans bileşenleri tahminleri	64
Çizelge 5.7. İç içe hata regresyon modeli için üretilen veri seti	64
Çizelge 5.8. İç-içe hata regresyon modelinde parametre tahmini için başlangıç değerleri	65
Çizelge 5.9. İç-içe hata regresyon modeli parametre tahminleri BM adımları	65
Çizelge 5.10. İç-içe hata regresyon modelinde parametre tahminleri	65
Çizelge 8.1. Sabit etkilerin KEÇO tahminleri ve standart hataları	85
Çizelge 8.2. $\log(FEV1)$ tekrarlanan ölçümlerinin korelasyon matrisi	86
Çizelge 8.3. Rasgele etkilerin varyans, kovaryans ve standart hata tahminleri	87
Çizelge 8.4. Simülasyon şeması	88
Çizelge 8.5. Simülasyon şeması	96
Çizelge 8.6. Durum 1’de β_1 ’e ilişkin KO ve OA sonuçları	99
Çizelge 8.7. Durum 2’de $\rho = 0.25$ için β_1 ’e ilişkin KO ve OU sonuçları	99
Çizelge 8.8. Durum 2’de $\rho = 0.5$ için β_1 ’e ilişkin KO ve OU sonuçları	99
Çizelge 8.9. Durum 2’de $\rho = 0.75$ için β_1 ’e ilişkin KO ve OU sonuçları	100
Çizelge 8.10. Durum 3’de $\rho = 0.25$ için β_1 ’e ilişkin KO ve OU sonuçları	101
Çizelge 8.11. Durum 4’de $\rho = 0.5$ için β_1 ’e ilişkin KO ve OU sonuçları	101
Çizelge 8.12. Durum 5 için β_1 ’e ilişkin KO ve OU sonuçları	102

Çizelge	Sayfa
Çizelge 8.13. Durum 6 için β_1 'e ilişkin KO ve OU sonuçları	102
Çizelge 8.14. Durum 7 için β_1 'e ilişkin KO ve OU sonuçları	103
Çizelge 8.15. Uzun vadeli depresyon çalışması verisinin bir alt seti	104
Çizelge 8.16. Parametre ve güven aralığı tahminleri	105



ŞEKİLLERİN LİSTESİ

Şekil	Sayfa
Şekil 1.1. Fiyat-talep verisi için klasik regresyon ve DKE model yaklaşımı	3
Şekil 3.1. Zaman içinde yanıt değişkenin marjinal ve koşullu değişimi	19
Şekil 3.2. Hata terimi ile marjinal ve koşullu yanıt değişkenin değişimi	20
Şekil 3.3. Hata terimi ile marjinal ve koşullu yanıt değişkenin değişimi	23
Şekil 7.1. Profil olabirlik güven aralığı gösterimi	74
Şekil 8.1. $m=3$ ve $\rho=0.25$ kombinasyonu için kapsama olasılığı	89
Şekil 8.2. $m=9$ ve $\rho=0.25$ kombinasyonu için kapsama olasılığı	90
Şekil 8.3. $m=15$ ve $\rho=0.25$ kombinasyonu için kapsama olasılığı	91
Şekil 8.4. $m=3$ ve $\rho=0.5$ kombinasyonu için kapsama olasılığı	92
Şekil 8.5. $m=9$ ve $\rho=0.5$ kombinasyonu için kapsama olasılığı	93
Şekil 8.6. $m=15$ ve $\rho=0.5$ kombinasyonu için kapsama olasılığı	94
Şekil 8.7. $n=5$, $m=20$ için bantlı kovaryans yapısında yapay-skor değerleri histogramı	96

SİMGELER VE KISALTMALAR

Bu çalışmada kullanılmış simgeler ve kısaltmalar, açıklamaları ile birlikte aşağıda sunulmuştur.

Kısaltmalar

Açıklamalar

BM	Beklenti ve Maksimizasyon
BS	Bootstrap
BS-t-KR	Bootstrap-t-Kenward Roger
DEİDYTE	Deneyssel en iyi doğrusal yansız tahmin edici
DEİDYK	Deneyssel en iyi doğrusal yansız kestirim
DGEKK	Deneyssel geliştirilmiş en küçük kareler
DKE	Doğrusal karışık-etkili
EÇO	En çok olabilirlik
EİDYTE	En iyi doğrusal yansız tahmin edici
EİDYK	En iyi doğrusal yansız kestirim
EKK	En küçük kareler
GEKK	Genelleştirilmiş en küçük kareler
HKO	Hata Kareler Ortalaması
KEÇO	Kısıtlı en çok olabilirlik
KO	Kapsama Olasılığı
KTG	Kareler toplamı grup
KTH	Kareler toplamı hata
OU	Ortalama Uzunluk
PBS	Parametrik bootstrap
SKK	Sınıf içi korelasyon katsayısı

1. GİRİŞ

Doğrusal karışık-etki modelleri genetik, eğitim, mühendislik, biyoloji, sosyal bilimler gibi bir çok uygulama alanında oldukça güçlü ve popüler bir araçtır. Bir doğrusal model yanıt değişken ve açıklayıcı değişken arasındaki ilişkiyi belirler. Doğrusal model hem sabit hem de rasgele etki içerdiğinde, doğrusal karışık-etki modeli olarak adlandırılır.

Doğrusal karışık-etki (DKE) modelleri için birçok teorik ve uygulamalı bilgi ile istatistik literatürüne katkıda bulunulmuştur. DKE modelleri metodolojisi istatistik bilimini yeni bir seviyeye taşımıştır. DKE modellerinin teorik yapısı, parametre tahmin yöntemleri ile ilgili en geniş kaynaklar: Searle, Casella, ve McCulloch (1992); Molenberghs ve Verbeke (2000); Diggle (2002); Demidenko (2004); Fitzmaurice, Laird ve Ware (2004); McCulloch, Searle, ve Neuhaus (2001); Jiang (2007) olarak verilebilir. İstatistiksel yazılımlarla DKE modellerinin uygulanmasına yönelik en geniş kapsamlı kaynaklar: Verbeke ve Molenberghs (1997); Raudenbush ve Bryk (2002); Faraway (2016); West, Welch, ve Galecki (2014); Brown ve Prescott (2014); Littell, Milliken ve Wolfinger, (2006); Hox, Moerbeek ve van de Schoot (2010)'dur.

Doğrusal karışık-etki modelleri, yüzey (spatial) verisi, zaman içinde uzun-vadeli ölçüm verisi (longitudinal), tekrarlanan ölçümlü (repeated measures) veriler gibi verilerin bazı sınıflama faktörlerine göre gruplanmış fiziksel, biyolojik, ve sosyal bilimler araştırmalarında oldukça kullanışlıdır. Klasik istatistik modellerinde en önemli varsayımlar, gözlemlerin aynı kitleden gelmesi, bağımsız ve aynı dağılımlı olmasıdır. DKE modellerinin verisi daha kompleks, çok düzeyli ve hiyerarşik bir yapıda olmaktadır. Aynı sınıftan/ düzeyden/ deney biriminden olan birimler bağımlı, sınıflar/ düzeyler/ deney birimleri arasında gözlemler birbirinden bağımsız yapıda olmaktadır. Bu durumda iki değişkenlik ortaya çıkmaktadır: birisi sınıflar/ düzeyler/ deney birimleri içindeki değişkenlik, diğeri de sınıflar/ düzeyler/ deney birimleri arasındaki değişkenliktir.

DKE modellerinin ismindeki doğrusallık modeldeki parametrelerin doğrusallığı, karışık denmesi de modeldeki bağımsız değişkenlerin düzeylerinin sabit veya rastgele olmasından gelmektedir. DKE modelinin kullanım alanları aşağıdaki gibi özetlenebilir (Demidenko, 2004: 2);

- Kompleks olarak sınıflanmış ve zaman içinde ölçülmüş verileri modellemek için,
- Farklı değişken kaynaklarına sahip verileri modellemek için,
- Biyolojik çeşitliliği ve heterojenliği modellemek için,
- Bayesçi ve frekansçı yaklaşımlar arasında bir köprü kurmak için.

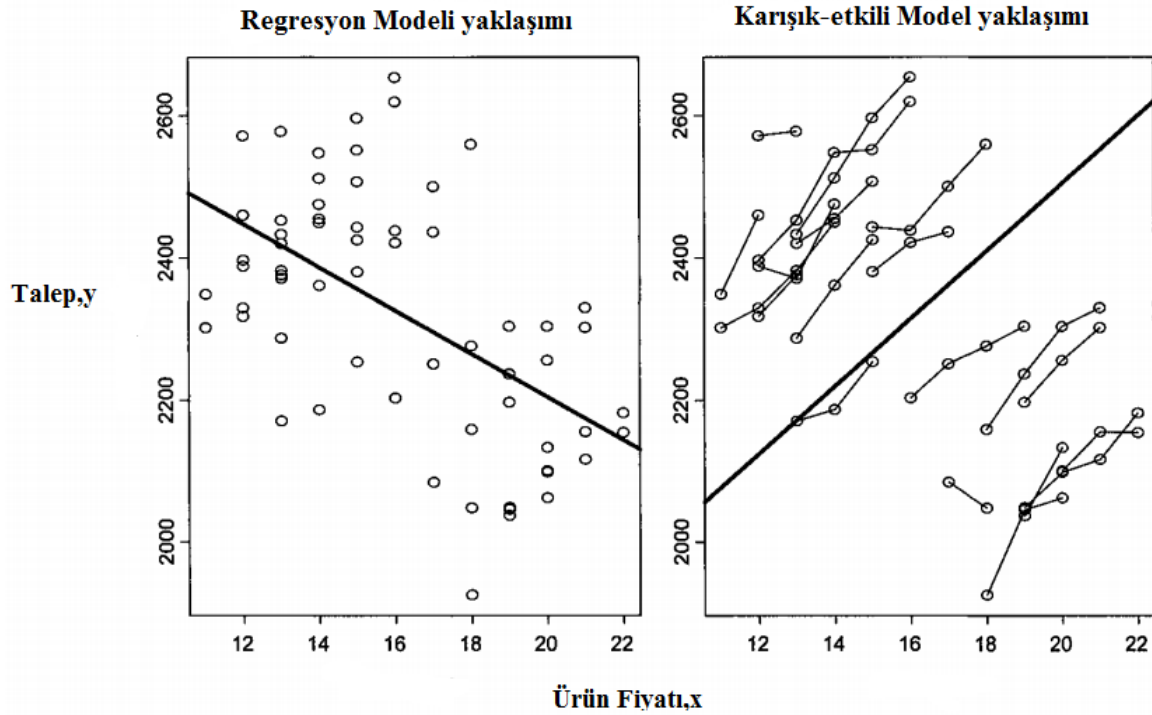
Zootekni alanında birçok çalışmada sığırcılık alanında toplanan veri setleri modellenerek sığır ırklarının ıslahı DKE modelleri ile belirlenmiştir. Ortaya çıkan birçok aşılama ve yabancı ırkların karışması gibi birçok zorluğa rağmen bu modellerin esnekliği sayesinde genetik ıslah çalışmaları başarıya ulaştırılmıştır (Onyango, 2009). Psiko-dil bilimciler ve diğer kavramsal psikologlar deneyleri için genellikle çalıştıkları üniversitedeki gönüllülerden gözlem toplarlar. Dil sadece o üniversitede konuşulmadığından; diğer insanlar ve dili yeni öğrenenler tarafından kullanıldığı için burada gönüllüler rastgele etki olarak ele alınır. Gönüllüler arasındaki genetik, çevresel, politik veya diğer farklılıklar rastgele etki olarak modellenmiş olur. Munofu, Moore, Payne, Walton ve Flint (2003) genler ve psikolojik kişisel özellikler arasındaki ilişkiyi modellemek için DKE modelinin özel hallerini kullanmışlardır. Genetik yapının kişisel özellikler arasında etkisinin olup olmadığı araştırılmıştır. Bir genin çekinik veya baskın olma özellikleri bu modeller sayesinde belirlenebilmiştir. Ayrıca, farklı işlerde parmakların güçlerinin koordinasyon yapılarını anlamak için biyomekanik veriler elde edilip DKE modelleri ile analiz edilmişlerdir (Bazzoli, Letué ve Martinez, 2014). Çalışmada bu modellerin kullanılma sebebi klasik varyans analizi metotlarının yeterli olmayacağının düşünülmesidir. Çünkü eş-zamanlı aynı bireyden alınan verilerin tekrarlı olarak alınmasından kaynaklanan bağımlı veri yapısı ancak DKE modelleri ile çözülebilecektir.

Bu modeller ayrıca genetik çalışmalarda oldukça yaygın kullanılır. Örneğin fenotip ve genotip arasındaki ilişkiyi modellemek için genom kapsamlı ilişki çalışmaları (Genome-wide association studies (GWASs)) olarak adlandırılan deneylerin modellenmesinde çok önemli bir yere sahiptir. DKE modellerinin en genel gösterimi Eş.1.1’de ifade edilir:

$$Y = X\beta + Zb + e \quad (1.1)$$

Klasik istatistiksel modellerle yapılan bir çıkarım iki değişken arasında negatif bir ilişkiyi işaret ederken diğer taraftan DKE modelleri ile analiz edildiğinde tam tersi bir durum ortaya çıkabilir. Örneğin farklı ürünlerin fiyatları ile ürüne olan talep arasındaki ilişkiyi

düşünelim. Veriye klasik regresyon modeli uygulandığında ve aşağıdaki grafik elde edildiğinde (Demidenko, 2004: 3) regresyon doğrusu negatif bir ilişkiyi göstermektedir. Bu veriler farklı ürünlerden toplandığı için aslında her ürünün gözlem değerleri kendi içinde değerlendirilmelidir. Bu ise DKE modeli metodolojisi kullanılarak yapılabilir. Şekil 1.1'deki grafikten görüldüğü gibi ürünler kendi içinde değerlendirildiğinde ürünün fiyatı ve talep arasında pozitif bir ilişki görülmektedir. İfade edildiği gibi model seçimi doğru yapılmadığında analiz araştırmacıyı yanlış sonuçlara ve yorumlara götürebilir.



Şekil 1.1. Fiyat-talep verisi için klasik regresyon ve DKE model yaklaşımı

Yukarıda bahsedilen çalışmalarda model seçildikten sonra o modelin parametreleri ve fonksiyonlarıyla ilgili çıkarımlar yapmak deneyi sonuçlandırıp yorumlayabilmek açısından çok önemli bir aşamadır. Birçok uygulama doğrusal karışık-etki modellerinin parametreleri ile ilgili çıkarım yapmayı gerektirir. Bu çıkarımlar doğru tahmin edici yöntemleri ve ilgili tahminlerin hata kareler ortalaması bulunarak gerçekleştirilebilir. Kackar ve Harville (1984) doğrusal karışık modellerde sabit ve rasgele etkinin tahminleri için hata kareler ortalamalarını araştırmışlardır. Harville ve Carriquiry (1992), klasik ve bayesgil tahmin yöntemlerini dengeli olmayan karışık doğrusal modellere uygulamışlardır. Lahiri ve Rao (1995), küçük alan tahminlerinde (small area estimation) hata kareler ortalamasının (HKO) tahmini için sağlam tahmin ediciler önermiştir. Das ve Krishen (1999) doğrusal olmayan

karışık-etki modellerinde parametrelerin standart hatalarını tahmin etmek için bootstrap metodlarını kullanmışlardır. Chatterjee, Lahiri ve Li (2008) doğrusal karışık modellerde kullanılan deneysel en iyi doğrusal yansız tahmin edicinin (DEİDYTE) (EBLUE-empirical best linear unbiased estimator) dağılımını ve onunla ilgili tahmin aralıklarını bulmak için parametrik bootstrap yaklaşımını kullanmışlardır. Kenward ve Roger (1997) küçük örneklemelerde kısıtlı en çok olabilirlik (KEÇO) fonksiyonundan sabit etkiler için çıkarım yöntemi araştırmışlardır. İstatistiksel çıkarımlar parametreler için hipotez testleri veya güven aralığı/ bölgesi kurularak gerçekleştirilmektedir. Harville (1976) sabit ve rasgele etkilerin doğrusal kombinasyonları için güven bölgeleri (confidence regions) oluşturmuştur. Bates, Maechler, Bolker ve Walker (2014) lme4 isimli R kütüphanesinde DKE modeli parametrelerinin EÇO ve KEÇO tahmin edicilerini elde etmiştir. Ayrıca Wald ve profil olabilirlik güven aralıklarını lme4 kütüphanesine dahil etmiştir.

Diciccio ve Efron (1996), BC, bootstrap-t, ABC ve calibration bootstrap güven aralığı oluşturma metodlarını, bu metodların arkasındaki teoriyi ve ayrıca bu metodların en çok olabilirlik (EÇO) tabanlı güven aralıkları ile arasındaki bağlantıyı göstermişlerdir. Özellikle fonksiyonlar doğrusal olmadığında güven aralıklarının doğrudan hesaplanması mümkün değildir veya hesaplamaları oldukça güçtür.

Doğrusal regresyon modelinde, regresyon parametrelerinin en küçük kareler tahmin edicileri (EKK) yansızdır, yani EKK tahmin edicisi tekrarlı örneklemede ilgili parametrenin değerini sistematik olarak fazla ya da daha az tahmin etmez. Doğrusal yansız tahmin ediciler arasında, EKK tahmin edicileri minimum varyansa sahiptir. Yani hiçbir doğrusal yansız tahmin edici EKK tahmin edicilerinin örnekleme dağılımı kadar parametre değeri etrafında yoğunlaşamaz. Minimum varyanslı doğrusal tahmin edici en iyi doğrusal yansız tahmin edici (EİDYTE) olarak bilinir. EKK tahmin edicileri bilinen standart hata formüllerine sahiptirler, test istatistiğinin gözlenen değerinin hesaplanması ve varyans tahmininin hesaplanması kolaydır. Ayrıca, herhangi bir regresyon katsayısının test istatistiği bilinen serbestlik dereceli Student t dağılımına sahiptir. Karışık etkili modellerde, durum daha zor olmaktadır. Sabit etkilerin en iyi yansız doğrusal tahmin edicileri modeldeki rasgele terimlerin varyanslarına bağlıdır. Rasgele terimlerin varyansları genellikle bilinmezdir ve tahminleri araştırılır. Bu varyanslar genellikle bilinmezdir ve yaklaşımlarını araştırmak gerekir. Modelin varyans tahminleri kullanılarak oluşturulmuş sabit etki tahmin edicilerinin kapalı formda yazılamayan örnekleme varyansları ya da

standart hataları vardır. Sabit etkilerden yararlanarak elde edilen test istatistikleri ve örnekleme varyanslarının yaklaşımları sonlu örnek çaplarında bilinmeyen bir dağılıma sahiptirler. Sonuç olarak, bu standart metodları kullanarak yapılacak olan çıkarımın iki potansiyel hata kaynağı olduğu söylenebilir.

Yukarıda bahsedilen bu güçlükler sabit etki parametresi için parametrik bootstrap güven aralıkları kurularak üstesinden gelinebileceği son yıllarda yapılan çalışmalarda gösterilmiştir. Bu yaklaşım altında, doğru çıkarımlar varyans tahminin yaklaşımına ve test istatistiğinin bilinen dağılımına ihtiyaç duymadan parametrik bootstrap yöntemleri ile yapılabilir. Yöntemin avantajlarına ve diğer analitik yöntemlere göre kolaylığına rağmen, bu metodun yayınlanmış uygulaması oldukça sınırlıdır. Bu yöntemi en son Staggs (2009) çalışmasında bazı doğrusal karışık etki modellerinin sabit etki parametreleri için kullanmıştır. Ayrıca Liu (2013) çalışmasında doğrusal karışık modeldeki parametrelerin fonksiyonları için parametrik bootstrap yöntemini kullanmıştır. Akdur, Özönur ve Bayrak (2016) DKE modellerinin özel bir hali olan iç içe hata regresyon modelinin sabit etki parametresi için parametrik bootstrap güven aralığı yöntemini analitik güven aralığı yöntemleriyle karşılaştırmıştır. Parametrik bootstrap güven aralığı yönteminin handikapı oldukça zaman ve maliyet gerektiren bir yöntem olmasıdır.

Wald-tipi olarak bilinen güven aralıkları en çok olabilirlik tahmin edicilerinin asimptotik özelliklerine bağlı olarak oluşturulmuş analitik güven aralığı metotlarıdır. Ayrıca analitik metod olarak gösterilen diğer bir metod da profil olabilirlik yöntemine dayanan güven aralığı metodudur. Bir diğeri ise Pearson Ki-kare testinin güven aralığı için uyarlanan metodudur. Bu metodun yayınlanmış uygulaması olumsuzluk tablolarının parametreleri ve lojistik regresyon parametreleri içindir. Lang (2008) skor ve profil olabilirlik metotlarını olumsuzluk tablolarının parametreleri için güven aralığı oluştururken kullanmıştır. Agresti ve Ryu (2010) yapay-skor (pseudo-score) ve profil olabilirlik güven aralıklarını kategorik veri analizinde kullanmak için oluşturmuştur. Bu modellerde yapay-skor metodunun güven aralığı yöntemi olarak iyi sonuç verdiği gösterilmesine rağmen DKE modellerinde uygulaması literatürde görülmemiştir.

Bu tez çalışmasında bu konuda çalışan araştırmacılara ve uygulamacılara katkı vereceği düşüncesiyle DKE modelleri üzerinde yapay-skor güven aralığı metodu uyarlanacaktır. Bu tez çalışmasının temel olarak iki önemli amacı vardır. Birincisi küçük örnek çaplarında ve

farklı kovaryans yapılarında literatürde var olan güven aralığı metodlarının sabit etki parametresini kapsama olasılığı bakımından performanslarını incelemektir. Bir diğer önemli amaç ise yapay-skor metodunun yeni bir analitik güven aralığı metodu olarak DKE modelleri için uyarlamaktır. Tezde değinilen güven aralığı metodlarının performanslarının modelin bağımlı değişkeninin kovaryans yapısının değişiminden nasıl etkilendiğini inceleyerek araştırmacılara en doğru yöntemi sunmak bu tezin en önemli katkılarından biri olacaktır. Ayrıca bu tezde tekrarlanan ölçümlü ve uzun vadeli çalışmalarda ortaya çıkan veri setlerinin bağımlılık yapısından oluşan kovaryans yapıları ve DKE modelleri ile analizi ele alınmıştır. Bu tür çalışmalarda ortaya çıkan veri setleri doğru kovaryans yapısının seçimiyle DKE modelleri ile analiz edilebilmektedir.



2. KLASİK YÖNTEMLERİN TANITILMASI

Doğrusal karışık-etki modeli tanıtılmadan önce aşağıda yaygın olarak kullanılan iki metod tanıtılıp bu metodların DKE modelleri ile ilişkisi açıklanacaktır. DKE modeli, ilk olarak Laird ve Ware (1982) tarafından geliştirilmiştir. Demidenko (2004, 4-5) DKE modelini; varyans analizi (ANOVA), varyans bileşenleri (VARCOMP) ve regresyon modellerinin bir kombinasyonu olarak tanımlamıştır. Örneğin, en basit haliyle, tek yönlü ANOVA modeli aşağıda regresyon formunda ifade edilmiştir.

$$y_{ij} = \beta_i + e_{ij}, \quad i = 1, \dots, n, \quad j = 1, \dots, m_i \quad (2.1)$$

n terimi modeldeki deney birimi, küme, grup veya sınıf sayısını, m_i terimi her bir deney birimden alınmış gözlem/ ölçüm sayısını, y_{ij} i . grubun j . ölçüm değerini göstermektedir. $\{e_{ij}\}$ i . grubun j . bireye ait hata terimi olup, $N(0, \sigma^2)$ parametrelili normal dağılıma sahiptir. Önemli ama genellikle çokça vurgulanmayan diğer bir ANOVA varsayımı da $\{\beta_1, \dots, \beta_n\}$ parametrelerinin sabit etki parametreleri olmasıdır. Her deney birimi için, gözlemler $\{y_{i1}, y_{i2}, \dots, y_{im_i}\}$ aynı dağılımlı bağımsız ve β_i ortalamalı oldukları görülebilir. Bilinen ANOVA hipotezi birimlerin/ deneklerin (units/subjects) aynı olması, veya $H_0 : \beta_1 = \dots = \beta_n$ olarak ifade edilebilir. ANOVA modeli ayrıca doğrusal regresyon modelinin özel bir halidir,

$$y = X\beta + e, \quad (2.2)$$

y vektörü $N \times 1$ 'lik gözlem vektörü, X matrisi $N \times p$ boyutunda tasarım matrisi β ise $p \times 1$ boyutunda parametre vektörüdür. Örneğin tek-yönlü ANOVA modeli regresyon formunda ifade edilebilir. gözlemleri y vektöründe $N = \sum_{i=1}^n m_i$ olmak üzere X tasarım matrisinin elemanları 0 ve 1 ve grup sayısı veya deney birimi sayısı modeldeki parametre sayısına eşit olur yani $p = n$ 'dir. Çeşitli ANOVA modelleri doğrusal model olarak Searle (1971)'de gösterilmiştir. Tüm ANOVA modellerinin iki önemli özelliği vardır:

parametreler $\{\beta_1, \dots, \beta_p\}$ en küçük kareler yöntemi ile tahmin edilir ve F-testi doğrusal hipotez testi için kullanılır. Eş. 2.2 ve 2.1 sabit etkili modellerdir.

$\{\beta_i\}$ parametrelerinin rasgele olduğu varsayılırsa farklı bir istatistik modeline gidilecektir. Bu parametrelerin ortak ortalaması β ve varyansları da σ_β^2 olarak normal dağılımlı olduğu varsayılırsa, $\beta_i \sim N(\beta, \sigma_\beta^2)$, $\beta_i = \beta + b_i$ olmak üzere varyans bileşenleri modeli aşağıdaki gibi yazılır;

$$y_{ij} = \beta + b_i + e_{ij} \quad (2.3)$$

b_i terimi rasgele etki olarak adlandırılır. ANOVA modeli sabit etki modeli varyans bileşenleri ise rasgele etki modelidir. Eş. 2.1 ve Eş. 2.3 benzer görünseler de, farklı istatistiksel özelliklere sahiptirler. ANOVA modelinde birimler içinde gözlemler bağımlı değildir, varyans bileşenleri modelinde ise birimler içinde gözlemler birbirlerine bağımlıdır ve korelasyon katsayısı $\sigma_\beta^2 / (\sigma^2 + \sigma_\beta^2)$ olarak formüle edilir. Gauss Markov teoremine göre, Eş. 2.1 EKK tahmin edicileri en çok olabilirlik (EÇO) tahmin edicileriyle aynıdır tutarlıdır ve etkindir, ama bu durum Eş. 2.3 modeli için sağlanmaz. Ayrıca, m_i 'ler farklı olduğunda EÇO tahmin edicilerinin kapalı formları için çözüm yoktur.

Karışık etkili modeller ANOVA ve varyans bileşenleri modellerinin bir kombinasyonudur. Örneğin, n ($i=1, \dots, n$) kişi için $t_{i1}, t_{i2}, \dots, t_{im_i}$ gibi farklı zamanlarda alınan kan basıncı değerlerinde; y_{ij} , i . kişinin t_{ij} . zamandaki kan basıncı değerini gösterebilir. t_{ij} zaman vektörü x_{ij} tam tasarım matrisinde içerilsin. Eş. 2.3'deki model kan basıncı değerinin kişiden kişiye değişiklik gösterip göstermediğini anlamak için yeterli olabilir. Aynı zamanda kan basıncı farklı zamanlarda kişinin farklı yaş aralıklarına denk geldiğini düşünülürse zaman değişkeninin etkisi de öğrenilmek istenebilir. Bu durumda, varyans bileşenleri modeli karışık-etki modeline dönüşür,

$$y_{ij} = x'_{ij}\beta + b_i + e_{ij} \quad (2.4)$$

3. DOĞRUSAL KARIŞIK-ETKİ MODELİ

Basit doğrusal regresyon modelinde, $Y = \beta_0 + \beta_1 x + e$, regresyon katsayıları sabit etkilidir. Bu parametreler tüm kitleyi temsil eder; sabit terim $x=0$ noktasında Y vektörünün ortalamasını, eğim parametresi x 'deki bir birimlik artışın Y 'deki ortalama değişimini gösterir.

Sabit etkinin tersine, rasgele etki kitleyi bütün olarak karakterize etmez. Etki rasgeledir çünkü, elde edilen etkinin değeri kitleden alınan rasgele bazı özel birimlere bağlıdır. Örneğin, tekrarlanan ölçümlü tasarımlarda, çalışmadaki yanıt değişkeninin değerlerinin bir kısmı deney birimleri arasındaki farklılıktan kaynaklanmaktadır. Bu kendine özgü deney birim etkisi rasgele etki yoluyla modellenenir. Benzer şekilde, öğrenci başarı skorunun belirlenmesi çalışmasında her bir sınıfa ait olan etki rasgele etki yoluyla modellenir. Doğrusal regresyon modelindeki hata terimi de aslında bir rasgele etkidir.

Aşağıdaki model ve varsayımlarla birlikte matris formunda genel DKE modeli gösterimidir;

$$\begin{cases} y_i = X_i \beta + Z_i b_i + e_i, & i = 1, \dots, n \\ b_i \sim N(0, G) \\ e_i \sim N(0, R_i) \\ b_1, \dots, b_n, e_1, \dots, e_n & \text{kov}(b, e) = 0 \end{cases} \quad (3.1)$$

Eş. 3.1 modelinde y_i i . deney birimi için m veya m_i boyutlu yanıt değişken vektörü, $1 \leq i \leq n$, n deney birimi sayısı, X_i ve Z_i sırasıyla $m \times p$ ve $m \times q$ boyutlu açıklayıcı değişken (tasarım) matrisleri, β p boyutlu sabit etki vektörü, b_i q boyutlu rasgele etki vektörü, e_i m boyutlu hata terimlerinin vektörüdür.

Eş. 3.1 modelinde $\text{kov}(b_i) = G$ olacak şekilde, b_i rasgele etki vektörü ve Z_i ilgili tasarım matrisidir. Z_i 'nin sütunları X_i 'nin sütunlarının bir alt seti olabilir. Rasgele farklılaşması istenen β regresyon parametrelerinin belli bir bölümü Z_i 'nin sütunlarını oluşturan X_i sütunları ile belirlenebilir.

X_i 'nin gerekli sütunları Z_i 'ye eklenerek β 'daki istenen parametrelerin rasgele farklılaşması sağlanabilir. Rasgele etkiler, b_i , G kovaryans matris ve sıfır ortalama vektörlü çok değişkenli normal dağılıma sahiptir. Prensipde, b_i için herhangi bir çok değişkenli dağılım varsayılabilir ama pratikte b_i 'nin çok değişkenli normal dağılıma sahip olduğu varsayılır. $\{y_i\}$ vektörlerini birleştirerek $\sum m_i \times 1$ boyutlu y vektörü ve $\{X_i\}$ matrislerini birleştirerek $\sum m_i \times p$ boyutlu X matrisine ve $\{Z_i\}$ matrisleri de $Z = \text{diag}\{Z_1, \dots, Z_n\}$ olarak tam bir matris eşitliğine dönüştürülebilir;

$$y = X\beta + Zb + e \quad (3.2)$$

Eş. 3.2'de verilen model bir doğrusal model gibi görünse de, varyans parametrelerinin (genellikle) bilinmez olması nedeniyle tahmin metodolojisinde bu model doğrusal olmayan istatistik kategorisinde incelenmektedir. Z matrisinin elemanları grup üyeliğini belirten 0 veya 1 değerlerinden ve X matrisinin sütunu kesim (intercept) terimini gösteren 1 değerlerinden oluşursa bu modele varyans bileşenleri modeli denir.

Eğer $R_i = \sigma^2 I_{m_i}$ olarak alınırsa Eş. 3.1 koşullu bağımsızlık modeli olarak bilinir. i . deney biriminden alınmış m_i gözlemin b_i ve β koşulu üzerinden bağımsız olduğunu ifade eder. R_i esnetilerek başka kovaryans modellerine sahip olması istenirse i . deney biriminden alınmış m_i gözlemin b_i ve β koşulu üzerinden bağımsız olduğunu ifade edilemez. y_i 'nin b_i üzerinden koşullu dağılımı $X_i\beta + Z_i b_i$ ortalama ve R_i kovaryans matrisi ile normal dağılımlıdır. $f(y_i | b_i)$ ve $f(b_i)$ ilgili olasılık yoğunluk fonksiyonlarını göstermek üzere; y_i 'nin marjinal yoğunluk fonksiyonu aşağıdaki gibidir;

$$f(y_i) = \int f(y_i | b_i) f(b_i) db_i$$

Böylece y_i , $X_i\beta$ ortalama vektörü ve $V_i = Z_i G Z_i' + R_i$ kovaryans matrisi ile m_i (veya dengeli durumda m) boyutlu normal dağılıma (marjinal dağılımı) sahip olur. Bu yaklaşım SAS PROC MIXED prosedüründe kullanılmaktadır. Çıkarımlar y_i 'nin marjinal dağılımına

göre yapılır. Tezde dengeli durum ele alındığı için deney birimi içinde gözlem sayısı genellikle m olarak ifade edilmiştir.

Rasgele etkiler ve hata terimleri için normallik varsayımları model kurulumunda gerekli olmasa da parametre tahminleri, hipotez testleri ve rasgele etkilerin kestirimleri için gerekli bir varsayımdır. α vektörü $V_i = \text{Var}(Y_i)$ ’deki varyans ve kovaryans parametrelerini içeren vektör olsun. α vektörü G matrisinin elemanlarını ve R_i ’deki tüm elemanları içerir. Y_i ’nin marjinal modelinde bulunan tüm parametrelerin vektörü $\theta = (\beta', \alpha')$ ile ifade edilsin. Marjinal olabilirlik fonksiyonundan elde edilen tahminlere göre yapılan yaklaşımlar klasik yaklaşımlardır. Eş. 3.1’de verilen modelin varsayımları altında DKE modelleri için marjinal olabilirlik ve log-olabilirlik fonksiyonunun genel yapısı aşağıdaki gibidir;

$$L_{E\zeta O}(\beta, \alpha) = \left\{ \prod_{i=1}^n (\pi)^{-m_i/2} |V_i(\alpha)|^{-\frac{1}{2}} \exp \left(-\frac{1}{2} (Y_i - X_i \beta)' V_i^{-1}(\alpha) (Y_i - X_i \beta) \right) \right\},$$

$$l(\beta, \alpha) = \ln L(\beta, \alpha) = -\frac{1}{2} N \ln(2\pi) - \frac{1}{2} \sum_i \ln(|V_i(\alpha)|) - \frac{1}{2} \sum_i (Y_i - X_i \beta)' V_i^{-1}(\alpha) (Y_i - X_i \beta) \quad (3.3)$$

İlk olarak modelin varyans parametrelerinin bilinir olduğu varsayılırsa β ’nın en çok olabilirlik tahmini Eş. 3.3 maksimize edilerek elde edilir (Laird ve Ware, 1982),

$$W_i = V_i^{-1}(\alpha) \text{ olmak üzere } \hat{\beta} = \left(\sum_{i=1}^n X_i' W_i X_i \right)^{-1} \sum_{i=1}^n X_i' W_i Y_i, \quad (3.4)$$

$$\text{Var}(\hat{\beta}) = \left(\sum_{i=1}^n X_i' W_i X_i \right)^{-1} \left(\sum_{i=1}^n X_i' W_i V(Y_i) W_i X_i \right) \left(\sum_{i=1}^n X_i' W_i X_i \right)^{-1} \quad (3.5)$$

$$\text{Var}(\hat{\beta}) = \left(\sum_{i=1}^n X_i' W_i X_i \right)^{-1} \quad (3.6)$$

$E(Y_i)$ ortalaması, $X_i \beta$ olarak doğru şekilde belirlendiğinde Eş. 3.4 için yeterlilik koşulu yansız olmasıdır. Eş. 3.5 ve 3.6 kovaryans matrisinin doğru olarak belirlendiğini varsayar. Böylece Eş. 3.6’ya göre yapılacak olan analizler kovaryans matrisinden sapmalardan etkilenecektir ve sağlam (robust) çıkarımlar olmayacaktır. Bu sapmalardan etkilenmeyecek

$Var(\hat{\beta})$ için bir tahmin edici sandviç tahmin edici adıyla Liang ve Zeger (1986) tarafından önerilmiştir. Eş. 3.5’de verilen $V(Y_i)$ yerine $r_i = y_i - X_i\hat{\beta}$ olmak üzere $r_i r_i'$ yazılır. Ortalama doğru belirlendiği takdirde bu tahmin edici tutarlı olacaktır. Modelin varyans parametreleri α bilinmediğinde eğer tahmini elde edilebiliyorsa $\hat{\alpha}$, $\hat{V}_i = V_i(\hat{\alpha}) = \hat{W}_i^{-1}$ olarak alınır. β sabit etki parametre vektörünün tahmini Eş. 3.4’de W_i yerine $\hat{W}_i$ kullanılarak elde edilir. $\hat{\beta}$ ’nın standard hata tahminleri Eş. 3.5 ve 3.6’de α yerine $\hat{\alpha}$ kullanılarak elde edilir. SAS PROC MIXED ve R’in lme4 paketinden ilgili tahminler elde edilebilir. Bu yazılımlar sayısal algoritmalar yoluyla EÇO tahminlerini bulur. Aşağıda sıklıkla kullanılan bazı DKE modelleri tanıtılacaktır.

3.1. Tek Yönlü Rasgele Etkili ANOVA Modeli (RANOVA)

Bu model DKE model tanıma tam olarak uymamaktadır, çünkü modelde sabit etki yoktur. Scheffe, (1956)’ya göre bu model ilk olarak 1886 yılında bir astronomun farklı gecelerde tekrarlı teleskopik gözlemler alması sonucu ortaya çıkmıştır. Tek yönlü rasgele sınıflama modeli olarak da bilinir. Tek faktörlü rastgele etki modeli zootekni, biyoloji, klinik deney çalışmaları, güvenilirlik analizi, çevresel deneyler gibi istatistiğin birçok uygulama alanında sıklıkla kullanılmaktadır. Özellikle son yıllarda laboratuvarlar arası yapılan çalışmalarda laboratuvar etkisini modellemek için kullanılmıştır. Model aşağıdaki gibi yazılır;

$$Y_{ij} = \mu + \alpha_i + e_{ij}, \quad (3.7)$$

Y_{ij} terimi i . sınıftaki j . gözlemi gösterir, α_i i . sınıftaki rasgele etkiyi, e_{ij} ise Y_{ij} ile ilgili rasgele hatayı göstermektedir. μ modelin genel ortalama ve sabit etki terimidir. e_{ij} terimleri aynı dağılımlı ve bağımsızdır, $N(0, \sigma_e^2)$. α_i terimleri aynı dağılımlı ve bağımsızdır, $\alpha_i \sim N(0, \sigma_\alpha^2)$. e_{ij} ve α_i terimleri bağımsızdır. Rasgele etkinin varlığını test etmek için $H_0: \sigma_\alpha^2 = 0$; $H_1: \sigma_\alpha^2 \neq 0$ hipotezi kurulur. Modeldeki toplam grup/ sınıf/ deney birimi sayısı n ve i . sınıftan alınan ölçüm/ gözlem sayısı m_i ile gösterilir. Dengeli veri setleri için her sınıftan aynı sayıda gözlem alınır ve m ile gösterilir.

Sınıf içi korelasyon katsayısı, hem birim hem de varyansları bakımından aynı gruba ait ölçümler arasındaki ilişkiyi belirlemek amacıyla geliştirilmiştir. Örneğin, bir batında doğan yavruların ağırlıkları arasındaki ilişki, ikizlerin IQ değerleri arasındaki ilişkinin ölçümünde sınıf içi korelasyon katsayısı (SKK) kullanılır. SKK, aynı denekten elde edilen ölçümlerin değişkenliğini, denekler ve gözlemler üzerinden elde edilen toplam değişkenlik ile karşılaştırır ve $\rho = \sigma_\alpha^2 / (\sigma_\alpha^2 + \sigma_e^2)$ olarak tanımlanmıştır. Ayrıca bu katsayı yerine bazen varyans oranı kullanılmaktadır, $\gamma = \sigma_\alpha^2 / \sigma_e^2$. Bu iki terimin aralarındaki ilişki $\rho = \gamma / \gamma + 1$ olarak tanımlanır.

3.2. Fay-Herriot Model

Küçük alan tahmini (small area estimation) için anket örneklemelerinde çok az örnek birimi içeren küçük coğrafi alanlarda çıkarım ve kestirim yapmak için kullanılan bir modeldir (Hulting ve Hartville, 1991; Fay ve Herriot, 1979) Model aşağıdaki gibi ifade edilir;

$$Y_i = X_i\beta + b_i + e_i, \quad (3.8)$$

b_i ve e_i 'ler birbirlerinden bağımsız dağılımları $b_i \sim N(0, \sigma_b^2)$, $e_i \sim N(0, \sigma_e^2)$ olan rasgele değişkenlerdir. σ_i^2 'lerin genellikle bilinir olduğu varsayılır veya daha önceki anket çalışmalarından tahminlerine genellikle ulaşılabilir.

3.3. İç içe Hata Regresyon Modeli

Fay-Herriot model ve tek-yönlü RANOVA modeli iç içe hata regresyon modelinin özel durumlarıdır, aşağıdaki gibi gösterilir;

$$Y_{ij} = X_{ij}\beta + b_i + e_{ij}, \quad i = 1, \dots, n, \quad j = 1, \dots, m \quad (3.9)$$

b_i ve e_{ij} 'ler birbirlerinden bağımsız dağılımları $b_i \sim N(0, \sigma_b^2)$, $e_{ij} \sim N(0, \sigma_e^2)$ olan rasgele değişkenlerdir. Simülasyon çalışmasında ele alınan $Y_{ij} = \beta_0 + \beta_1 x_i + b_i + e_{ij}$ iç içe hata regresyon modelinin özel bir durumu için modelin parametreleri ile ilgili tanımlamaları

yapılacaktır. Bu bölümde hata terimlerinin bağımsız ve sabit varyanslı olduğu varsayımı altında modelle ilgili tahmin ediciler verilmiştir.

$Y_i = (Y_{i1}, \dots, Y_{im})'$ olmak üzere bir genel kanonik model gösterimi aşağıdaki gibi verilir,

$$Y_i = X_i\beta + Z_ib_i + e_i, \quad i = 1, \dots, n, \quad b_i \sim N(0, G) \quad e_i \sim N(0, R_i = \sigma_e^2 I_m) \quad (3.10)$$

İç içe hata regresyon modeli için bu terimler; $b_i = (b_i)$ rastgele etki sabiti, $\beta = (\beta_0, \beta_1)'$ sabit etki vektörü, $X_i = (1, x_{ij})_{m \times 2}$ sabit etkiler için tasarım matrisi, $Z_i = (1)_{m \times 1}$ rastgele etki terimi için tasarım matrisi $e_i = (e_{i1}, \dots, e_{im})'$ gösterilebilir. İlgilenilen parametre vektörü $\theta = (\beta_0, \beta_1, \sigma_b^2, \sigma_e^2, \rho)$, ve genel parametre fonksiyonu $g(\theta)$ ile ifade edilecektir.

Varyanslar bilinir olduğunda veya elde edilen tahminleri aşağıdaki eşitliklerde yerine koyularak sabit etki tahminleri aşağıdaki gibi elde edilir (Jiang, 2007: 71). $d_i = m_i / (\sigma_e^2 + m_i \sigma_b^2)$, olmak üzere dengeli ve dengesiz veri-seti durumunda sabit etkilerin EİDYTE'leri aşağıdaki gibi ifade edilebilir;

$$\hat{\beta}_{BLUE,0} = \frac{(\sum_{i=1}^n d_i x_i^2)(\sum_{i=1}^n d_i \bar{y}_i) - (\sum_{i=1}^n d_i x_i)(\sum_{i=1}^n d_i x_i \bar{y}_i)}{(\sum_{i=1}^n d_i)(\sum_{i=1}^n d_i x_i^2) - (\sum_{i=1}^n d_i x_i)^2}$$

$$\hat{\beta}_{BLUE,1} = \frac{(\sum_{i=1}^n d_i)(\sum_{i=1}^n d_i x_i \bar{y}_i) - (\sum_{i=1}^n d_i x_i)(\sum_{i=1}^n d_i \bar{y}_i)}{(\sum_{i=1}^n d_i)(\sum_{i=1}^n d_i x_i^2) - (\sum_{i=1}^n d_i x_i)^2}$$

$m_i = m$ olarak yazılırsa dengeli tasarım olarak ifade edilir ve sabit etkilerin EİDYTE'leri aşağıdaki gibi olur:

$$\hat{\beta}_{E/DYTE,0} = \frac{(\sum_{i=1}^n x_i^2)(\sum_{i=1}^n \bar{y}_i) - (\sum_{i=1}^n x_i)(\sum_{i=1}^n x_i \bar{y}_i)}{n(\sum_{i=1}^n x_i^2) - (\sum_{i=1}^n x_i)^2}$$

$$\hat{\beta}_{E/DYTE,1} = \frac{\sum_{i=1}^n (x_i - \bar{x})(\bar{y}_i - \bar{y}_..)}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

Yukarıdaki eşitlikte görüldüğü gibi dengesiz tasarım durumunda sabit etki tahmin edicileri varyans bileşenlerine bağlıdır, ama dengeli durumda varyans bileşenlerinden bağımsızdır.

Sabit etki tahminlerinin $\hat{\beta}_{E/DYTE} = (\hat{\beta}_{E/DYTE,0}, \hat{\beta}_{E/DYTE,1})'$ varyansı da aşağıdaki gibidir;

$$\text{Var}(\hat{\beta}_{E/DYTE}) = (X'V^{-1}X)^{-1} \quad (3.11)$$

Bu model için tahmin edicilerin varyanslarının açık hali

$D = (\sum_{i=1}^n d_i)(\sum_{i=1}^n d_i x_i^2) - (\sum_{i=1}^n d_i x_i)^2$ olmak üzere aşağıdaki gibi yazılır;

$$\text{Var}(\hat{\beta}_{E/DYTE}) = \frac{1}{\sigma_\varepsilon^2} \begin{pmatrix} \sum_{i=1}^n d_i x_i^2 & -\sum_{i=1}^n d_i x_i \\ -\sum_{i=1}^n d_i x_i & \sum_{i=1}^n d_i \end{pmatrix}$$

EİDYTE içinde varyans bileşenleri mevcut olduğunda bunların tahminleri yerine yazabilir. Sonuçta elde edilen tahmin edicileri deneysel EİDYTE (DEİDYTE) olarak adlandırılır. Normal dağılım varsayımı altında, varyans bileşenlerinin tahminleri en çok olabilirlik tahminleri olduğunda, DEİDYTE en çok olabilirlik tahmin edicileri ile aynı olacaktır. DEİDYTE oldukça karmaşıktır. Özellikle y içinde doğrusal değildir. Ayrıca, Eş. 3.11'de V varyans bileşenleri tahminleriyle yer değiştirirse, DEİDYTE'daki gerçek varyasyonun az tahmin edileceği düşünülmektedir. Kackar ve Harville (1984), DEİDYTE'nin $\hat{\beta}_{EBLUE}$, β için yansız tahmin edici olduğunu varyans bileşenlerinin tahmin edicilerinin düzgün ve değişmez olduğunu normallik varsayımı altında göstermiştir, yani $E(\hat{\beta}_{DE/DYTE}) = \beta$ 'dir. Ek olarak, normallik varsayımı altında herhangi verilen bir a vektörü için aşağıdaki eşitlik gösterilmiştir;

$$\text{var}(a' \hat{\beta}) = \text{var}(a' \hat{\beta}_{E/DYTE}) + E \left\{ a' (\hat{\beta} - \hat{\beta}_{E/DYTE}) \right\}^2 \quad (3.12)$$

$\text{var}(a' \hat{\beta}_{E/DYTE}) = a' \text{Var}(\hat{\beta}_{E/DYTE}) a$ olduğundan, Eş. 3.12'nin sağındaki ilk terim Eş. 3.11'de varyans bileşenlerinin tahminleri yerleştirilerek tahmin edilebilir. Eş. 3.12'nin son terimi bu şekilde tahmin edilemez. Bu durum büyük örneklerde çok önemli bir sorun teşkil etmez. Küçük örneklerde ise Kackar ve Harville (1984) ve Kenward ve Roger (1997) bu tahmin edicilerin varyansları için düzeltmeler önermişlerdir. En azından büyük örneklem durumlarında sabit etkiler için güven aralığı kurulumunda bu zorluk bir probleme yol açmaz. ANOVA modelleri için Jiang (1998b) varyans bileşenleri kısıtlı en çok olabilirlik tahminleri olduğunda Eş.3.11 $\hat{\beta}$ 'nin asimptotik kovaryans matrisidir. Ayrıca varyans bileşenleri en çok olabilirlik tahminleri olduğunda Eş.3.11 $\hat{\beta}$ 'nin

asimptotik kovaryans matrisidir (Miller, 1977). $\hat{V}$ varyans bileşenlerinin en çok olabilirlik veya kısıtlı en çok olabilirlik tahminleri olmak üzere $a'\beta$ için büyük örneklerde asimptotik güven aralığı (Wald güven aralığı);

$$\left[a'\beta - z_{\frac{1-\alpha}{2}} \left\{ a'(X'\hat{V}^{-1}X)^{-1}a \right\}^{1/2}, a'\beta + z_{\frac{1-\alpha}{2}} \left\{ a'(X'\hat{V}^{-1}X)^{-1}a \right\}^{1/2} \right] \text{ olarak ifade edilir.}$$

β_1 için büyük örneklerde asimptotik güven aralığı $\hat{\sigma}_\varepsilon^2$ ve $\hat{\sigma}_b^2$ varyans bileşenlerinin kısıtlı en çok olabilirlik veya en çok olabilirlik tahminleri $a'=(0,1)'$, $\hat{d}_i = m_i / (\hat{\sigma}_\varepsilon^2 + m_i \hat{\sigma}_b^2)$, $\hat{D} = (\sum_{i=1}^n \hat{d}_i)(\sum_{i=1}^n \hat{d}_i x_i^2) - (\sum_{i=1}^n \hat{d}_i x_i)^2$ olmak üzere,

$$\left[\hat{\beta}_1 - z_{\frac{1-\alpha}{2}} \left(\frac{\sum_{i=1}^n \hat{d}_i}{\sigma_b^2 \hat{D}} \right)^{1/2}, \hat{\beta}_1 + z_{\frac{1-\alpha}{2}} \left(\frac{\sum_{i=1}^n \hat{d}_i}{\sigma_b^2 \hat{D}} \right)^{1/2} \right] \text{ ile ifade edilir.}$$

3.4. Doğrusal Karışık Etki Modellerinde Kovaryans Yapıları

Doğrusal karışık etki modelleri gruplanmış veri setlerinin yanı sıra uzun vadeli (longitudinal) çalışmalarda ve tekrarlı ölçüm verilerinin analizinde sıklıkla kullanılmaktadır. DKE modellerinin altında yatan temel dayanak bazı regresyon parametre setinin bir deney biriminden diğerine değişmesidir, bu sebeple kitledeki doğal değişkenliğin kaynağı hesaba katılır. Her deney birimi için kendi ortalama yanıt eğrisi oluşturulur ve regresyon parametre setinin rasgele olduğu düşünülür. DKE modellerinin ayırt edici özelliği; yanıt değişkeninin ortalama değeri tüm deney birimleri tarafından paylaşıldığı varsayılan kitle karakteristiği β ve deney birimine özgü rasgele etkiden oluşur. Littell, Pendergast ve Natarajan (2000) tekrarlı ölçüm verilerini analiz etmek için DKE modellerinin nasıl kullanılacağını farklı kovaryans yapıları ile birlikte SAS Proc Mixed fonksiyonu yardımıyla açıklamıştır. Dawson, Gennings ve Carter (1997) tekrarlı ölçümlerde kovaryans yapılarını belirlemek için grafiksel teknikleri kullanmışlardır.

Bu bölümde DKE modellerinde ortaya çıkabilecek olası kovaryans yapıları ve uzun vadeli DKE modelleri incelenecektir. Bir çok çalışmada, $R_i = \text{Var}(e_i)$ diyagonal matris olarak alınır ; I_m $m \times m$ boyutlu birim matris olmak üzere $R_i = \sigma^2 I_m$ gösterilir. Bu durumda hata terimlerinin deney birimi içinde bağımsız olduğu varsayılır; $\text{kov}(e_{ij}, e_{ik}) = 0$. Bu çok

sınırlayıcı bir durumdur, tekrarlanan ölçüm ve uzun vadeli çalışmalara uygun olmayan bir yapıdır. R_i için uygun kovaryans yapıları AIC, BIC, AICC gibi bilgi kriterlerine göre belirlenerek seçilmelidir (Gurka, 2006). Ayrıca teoride $G_i = Var(b_i)$ olarak yazılabilir, yani gruplar arası kovaryans matrisi farklı olabilir ama genellikle aynı olduğu varsayılır. İki farklı G_i 'nin uygun olacağı bir çalışmayı Lee ve Bryk (1989) literatüre kazandırmıştır.

Doğrusal modelde bir rasgele etki terimi olması yanıtlar arasında kovaryansa neden olur ve $kov(Y_i) = V_i$ rasgele etki için ayrı bir yapıyı barındırır. Doğrusal bir modele rasgele etkilerin dahil edilmesiyle zamanın bir fonksiyonu olarak tekrarlı ölçümlerin kovaryansı açıklanır. Deney birimleri içindeki değişkenlik ve deney birimleri arasındaki değişkenlik DKE modelleri ile açıkça ayırt edilir. Ölçüm veya tekrar sayısından bağımsız olarak modeldeki rasgele etkilerin kovaryans yapıları genellikle az sayıda parametreyle açıklanır. DKE modelleri, uzun vadede toplanmış yanıt değişken değerlerinin deney birimleri içi ve deney birimleri arası analizini gerçekleştirebilir. Kitledeki yanıt değişkeninin ortalama değişimini açıklayan parametreleri tahmin etmenin yanı sıra her bir deney biriminin zaman/ ölçüm içindeki kendi değişimini de açıklar. Bağımlı değişkenin kovaryans matrisi V_i ; R_i deney birimleri içi değişkenliği ölçen ve $Z_i G Z_i'$ deney birimleri arasında değişkenliği ölçen iki farklı yapıdan oluşur. Bazı kaynaklarda bağımlı değişkenin kovaryansı doğrudan R_i yoluyla modellenmiştir. Little ve diğerleri (2000) tekrarlı ölçümlerde bağımlı değişkenin kovaryans matrisi V_i 'yi R_i matrisini toeplitz, otoregresif, birleşik simetrik gibi değişik yapılar seçerek ve $G = 0$ alarak modellemiştir. Yine aynı çalışmada G ve R_i matrislerinin her ikisini de otoregresif yapıda seçerek aynı deney biriminden alınan gözlemler arasındaki kovaryansın iki farklı kaynaktan gelmesini sağlamıştır. Tezde bu çalışmadan esinlenerek gözlemler arası kovaryansın iki farklı kaynaktan gelmesi sağlanması amacıyla SAS PROC Mixed prosedürüne benzer olarak V_i , G ve R_i bileşenleri yoluyla modellenmiştir. Bunun nedeni V_i 'nin daha esnek bir yapıya sahip olabilmesidir. Ayrıca DKE modelleri dengeli olmayan örnek çaplarıyla kolaylıkla baş edebilir. Her deney biriminden aynı sayıda ölçüm alınması gerekmez. Sonuç olarak, DKE modelleri dengesiz uzun vadede toplanmış (longitudinal) veya tekrarlanan ölçümlü verileri (repeated measures) analiz etmede oldukça kullanışlıdır.

3.4.1. Rasgele kesim modeli

Rasgele kesim modeli (random intercept model) deney birimine göre değişimi içeren bir rasgele etkiye sahip doğrusal modeldir;

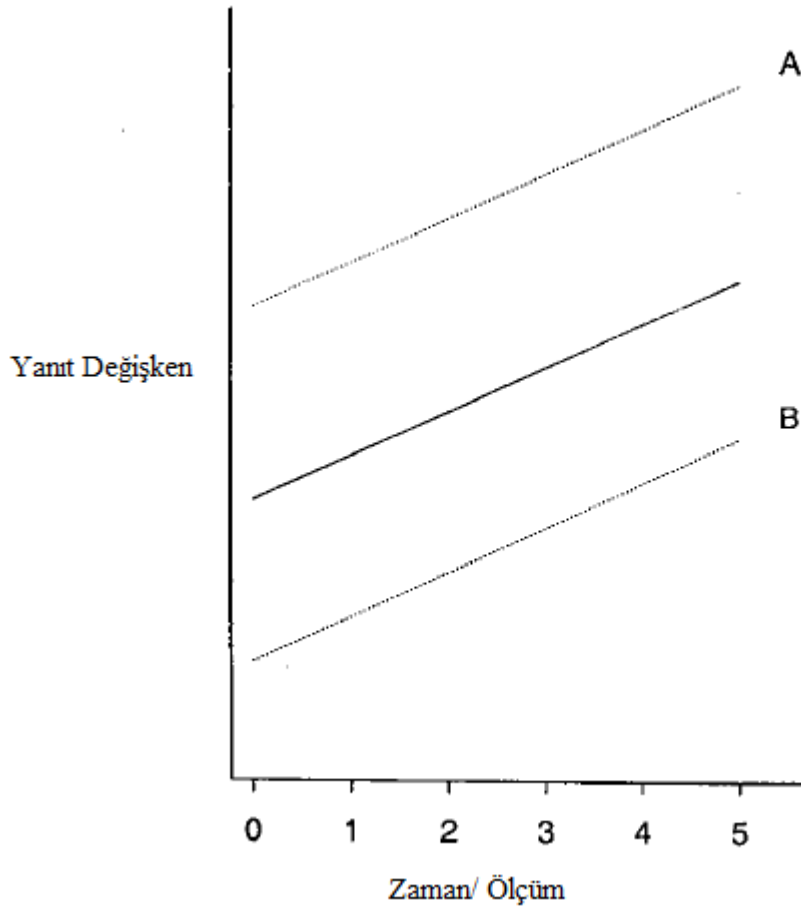
$$Y_{ij} = X'_{ij}\beta + b_i + e_{ij}, \quad (3.13)$$

burada b_i rasgele deney birimi etkisi ve e_{ij} rasgele hata terimidir. i . deney biriminin j . Durumdaki yanıt değişken değerinin kitle ortalaması $X'_{ij}\beta$ 'den $b_i + e_{ij}$ kadar farklı olduğu varsayılır. Rasgele etki terimi ve hata teriminin rasgele, ortalamalarının 0 ve varyanslarının sırasıyla $Var(b_i) = \sigma_b^2$ ve $Var(e_{ij}) = \sigma_e^2$ olduğu varsayılır. Bu model kitledeki ortalama yanıt profilinin $E(Y_{ij}) = X'_{ij}\beta$ yanı sıra her deney biriminin ortalama yanıt eğrisini $E(Y_{ij} | b_i) = X'_{ij}\beta + b_i$ ifade eder ve deney birimine özgü b_i etkisi verildiğinde Y_{ij} 'nin koşullu ortalamasını gösterir. Deney birimine özgü rasgele etkilerin dağılımı üzerinde $E(Y_{ij}) = X'_{ij}\beta$ terimi Y_{ij} 'nin marjinal ortalamasını gösterir. Her iki durumda da ortalama yanıt değişken değeri koşullu olarak açıklayıcı değişken X_{ij} değerlerine bağlıdır. Eş. 3.13 modelindeki parametrelerin izahı düşünülürse; regresyon parametresi β ortalama yanıtın ölçüm/ zaman içinde değişim yapısını ifade eder. b_i , i . deney biriminin ölçüm/ zaman içinde kitle ortalamasından farklılaşma trendini ifade eder. Açıklayıcı değişkenlerin etkisi hesaba katıldığında b_i deney biriminin kitle ortalama kesiminden (population mean intercept) ne kadar uzaklaştığını gösterir. Eş. 3.13 modeli aşağıdaki gibi açık olarak ifade edilir;

$$\begin{aligned} Y_{ij} &= X'_{ij}\beta + b_i + e_{ij} \\ &= \beta_1 X_{ij1} + \beta_2 X_{ij2} + \dots + \beta_p X_{ijp} + b_i + e_{ij} \\ &= \beta_1 + \beta_2 X_{ij2} + \dots + \beta_p X_{ijp} + b_i + e_{ij} \\ &= (\beta_1 + b_i) + \beta_2 X_{ij2} + \dots + \beta_p X_{ijp} + e_{ij}, \end{aligned} \quad (3.14)$$

burada tüm i ve j 'ler için $X_{ij1} = 1$ olduğunda β_1 sabit etki için kesim terimi olur. Model bu şekilde ifade edildiğinde i . deney birimi için kesim (intercept) $\beta_1 + b_i$ bir deney biriminden diğerine rasgele değişimini açıkça ifade edebilir. Çünkü b_i rasgele etkisinin ortalamasının 0 olduğu varsayılır, b_i i . deney biriminin kesiminin $\beta_1 + b_i$ kitle kesimi

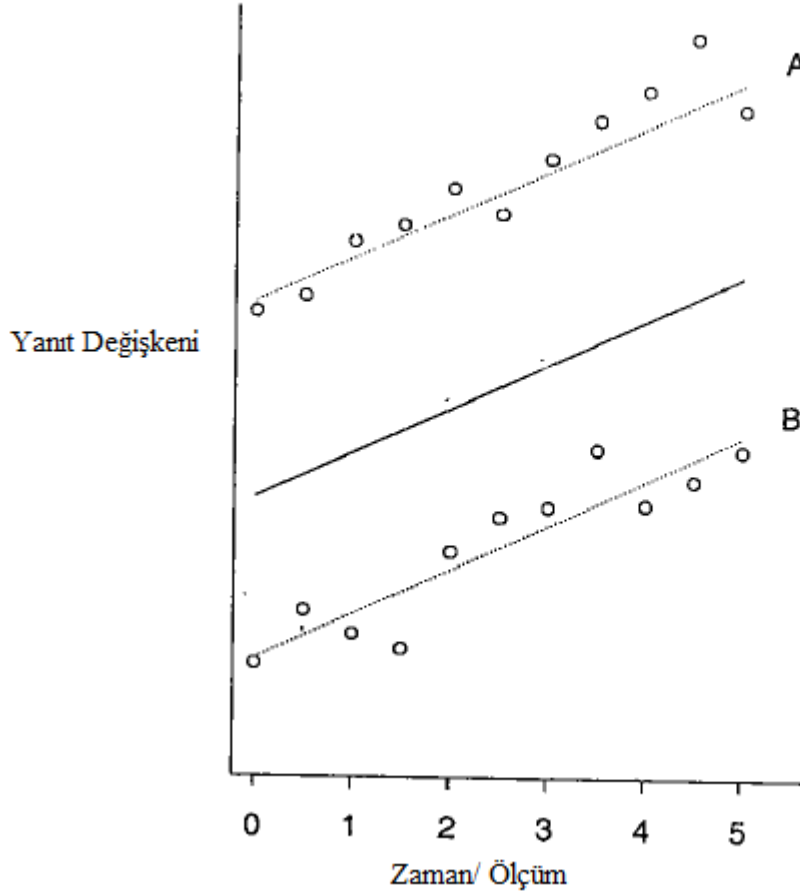
β_1 'den ne kadar farklı olduğunu ifade eder. Model denkleminin grafiksel gösterimi düşünülerek DKE modellerinin bu basit örneği daha iyi anlaşılabilir. Şekil 3.1 marjinal ortalama yanıt değişken değerlerinin zaman/ ölçüm içinde kitledeki doğrusal değişimini (kesiksiz çizgi ile ifade edilmiştir) ayrıca koşullu ortalama yanıt değişken değerlerinin iki deney birimi A ve B için kitle trendinden farklılığı (kesikli çizgilerle) ifade edilmiştir.



Şekil 3.1. Zaman içinde yanıt değişkeninin marjinal ve koşullu değişimi

Şekil 3.1'deki gösterimde deney birimi A kitle ortalamasından daha fazla yanıt vermiştir (responds) böylece pozitif bir b_i 'ye sahip olmalıdır. Bir diğer taraftan deney birimi B kitle ortalamasından düşük yanıt vermiştir ve negatif bir b_i 'ye sahiptir. Rasgele kesime sahip karışık etki modeli A ve B deney birimlerindeki tekrarlı ölçümler için mükemmel bir şekilde yayılan deney birimine ait eğrileri varsaymaz. e_{ij} hata teriminin eklenmesi deney birimine özel eğrinin altına ve üstüne düşen rasgele ölçümlere izin vermiş olur. Bu durum Şekil 3.2'de gösterilmiştir. Aynı deney birimindeki farklı tekrarlı ölçümler için marjinal

kovaryansı ele alınsın. Deney birimine özgü etki üzerinde ortalama alındığında Y_{ij} için marjinal ortalama $E(Y_{ij}) = \mu_{ij} = X'_{ij}\beta$ 'dır. Y_{ij} 'ler arası marjinal kovaryans, Y_{ij} 'nin μ_{ij} marjinal ortalamasından farklılığı bakımından tanımlanır.



Şekil 3.2. Hata terimi ile marjinal ve koşullu yanıt değişkeninin değişimi

Örneğin, Şekil 3.2’de bu farklılıklar A deney birimi için bütün ölçümlerde pozitif, B deney birimi için ise negatif olmuştur. Böylece yanıtlar arasında marjinal olarak zaman/ ölçüm içinde pozitif bir korelasyonu göstermektedir. Rasgele kesim modeli için her yanıt değişken değerinin marjinal varyansı aşağıdaki gibi gösterilir;

$$\begin{aligned}
 Var(Y_{ij}) &= Var(X'_{ij}\beta + b_i + e_{ij}) \\
 &= Var(b_i + e_{ij}) \\
 &= Var(b_i) + Var(e_{ij}) \\
 &= \sigma_b^2 + \sigma_e^2
 \end{aligned}$$

Herhangi aynı deney biriminden alınmış iki yanıt değişken ölçümü Y_{ij} ve Y_{ik} arasındaki kovaryans aşağıdaki gibidir;

$$\begin{aligned} \text{kov}(Y_{ij}, Y_{ik}) &= \text{kov}(X'_{ij}\beta + b_i + e_{ij}, X'_{ik}\beta + b_i + e_{ik}) \\ &= \text{kov}(b_i + e_{ij}, b_i + e_{ik}) \\ &= \text{kov}(b_i, b_i) \\ &= \text{Var}(b_i) = \sigma_b^2 \end{aligned}$$

Aynı deney birimi için tekrarlı ölçümlerin marjinal kovaryans matrisi aşağıdaki birleşik simetrik (compound symmetry) yapısına sahip olur;

$$\text{kov}(Y_i) = \begin{pmatrix} \sigma_b^2 + \sigma_\varepsilon^2 & \sigma_b^2 & \sigma_b^2 & \dots & \sigma_b^2 \\ \sigma_b^2 & \sigma_b^2 + \sigma_\varepsilon^2 & \sigma_b^2 & & \sigma_b^2 \\ \sigma_b^2 & \sigma_b^2 & \sigma_b^2 + \sigma_\varepsilon^2 & & \sigma_b^2 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \sigma_b^2 & \sigma_b^2 & \sigma_b^2 & \dots & \sigma_b^2 + \sigma_\varepsilon^2 \end{pmatrix}.$$

Aynı deney biriminden alınmış iki ölçüm arasındaki kovaryans σ_b^2 ve korelasyon

$$\text{korr}(Y_{ij}, Y_{ik}) = \frac{\sigma_b^2}{\sigma_b^2 + \sigma_\varepsilon^2} \text{ olacaktır. Korelasyonun bu basit tanımı karışık etki modellerinin}$$

önemli bir özelliğini vurgular: b_i rasgele etkisinin doğrusal modele dahil edilmesi tekrarlı ölçümler arasında korelasyona neden olmuştur. DKE modelinin en basit hali olan bu model genellikle uzun vadeli verilerin analizi için uygun olmaz. Bu nedenle diğer bölümde daha genelleştirilmiş DKE modellerine yer verilecektir. DKE modelleri için ortaya çıkabilecek farklı kovaryans yapıları ileriki bölümlerde tanıtılacaktır.

3.4.2. Rasgele eğim modeli

Bu bölümde rasgelelik gösteren ek regresyon parametrelerinin dahil edilmesiyle daha genel bir DKE modeli olan rasgele eğim (katsayı) modeli (random slope (coefficient) model) tanıtılacaktır. Ayrıca DKE modellerinin bazı önemli özellikleri üzerinde durulacaktır. DKE modelinin en önemli prensibi bazı regresyon parametreler setinin deney biriminden deney birimine değişmesidir. Yukarıdaki en basit örnekte yalnızca kesimin rasgele olmasına izin verilmiştir. Bu tek rasgele etkinin modele dahil edilmesi tekrarlı ölçümler arası sınırlı bir formda olmasına rağmen kovaryansa neden olmuştur. Bazı regresyon parametrelerinin

rasgele olmasına izin vererek uzun vadeli veriler için daha uygun bir kovaryans yapısı elde edilebilir. Örneğin zaman içinde alınan tekrarlanan ölçümlü bir çalışma ele alınsın. Yetişkinlikte değişen kişisel özelliklerin oranı tahmin edilmek için kişilik ölçümleri deneklere belli yıllarda belirli bir zaman ölçeğinde uygulanarak gerekli değerler toplanmıştır (Terracciano ve ark., 2005). Araştırmacı denekler içi skor sınıflandırmalarını görmezden gelip kişilik özellikleri ve deney birimi yaşına göre regresyon modeli kurulabilir. Bu sonuç, değişim oranının yansız bir tahmin edicisini verebilir, ama deney birimi içinde hataların ilişkili olması doğrusal regresyon modelinin EKK model varsayımını bozacaktır. Böylece eğim parametresinin tahmininin standart hatasını kullanarak yapılan çıkarımlar güvenilir olmayacaktır. Bu nedenle DKE modeli veri setinin analizinde uygun olur. Aşağıdaki büyüme eğrisi (growth curve) modeli veri setine uydurulmuştur;

$$Y_{ij} = \beta_1 + YAS_{ij}\beta_2 + b_{1i} + YAS_{ij}b_{2i} + e_{ij} \quad (3.15)$$

Y_{ij} : i. denek için j. zaman noktasında kişisel özellik skoru,

β_1 : sabit terim,

β_2 : eğim terimi (denekler arası ortalama eğim),

YAS_{ij} : i. denneğin j. zaman noktasındaki yaşı,

b_{1i} : i. denneğin sabit teriminin ortalama sabit etkiden, β_1 rasgele sapması,

b_{2i} : i. denneğin eğim teriminin ortalama eğim teriminden, β_2 rasgele sapması,

e_{ij} : i. denneğin j. zaman noktasında ilgili rasgele hata terimi.

e_{ij} hata terimlerinin bağımsız ve $N(0, \sigma_e^2)$ normal dağılımlı olduğu veya hata terimlerinin de tekrarlı ölçümlerden dolayı deney birimi içinde bağımlı olduğu bir kovaryans yapısı R_i ele alınabilir. b_{1i} ve b_{2i} terimlerinin e_{ij} 'lerden bağımsız ve ortak (joint) dağılımlarının $N(0, G)$ olduğu varsayılır. G ise 2×2 boyutlu bir varyans-kovaryans matrisidir.

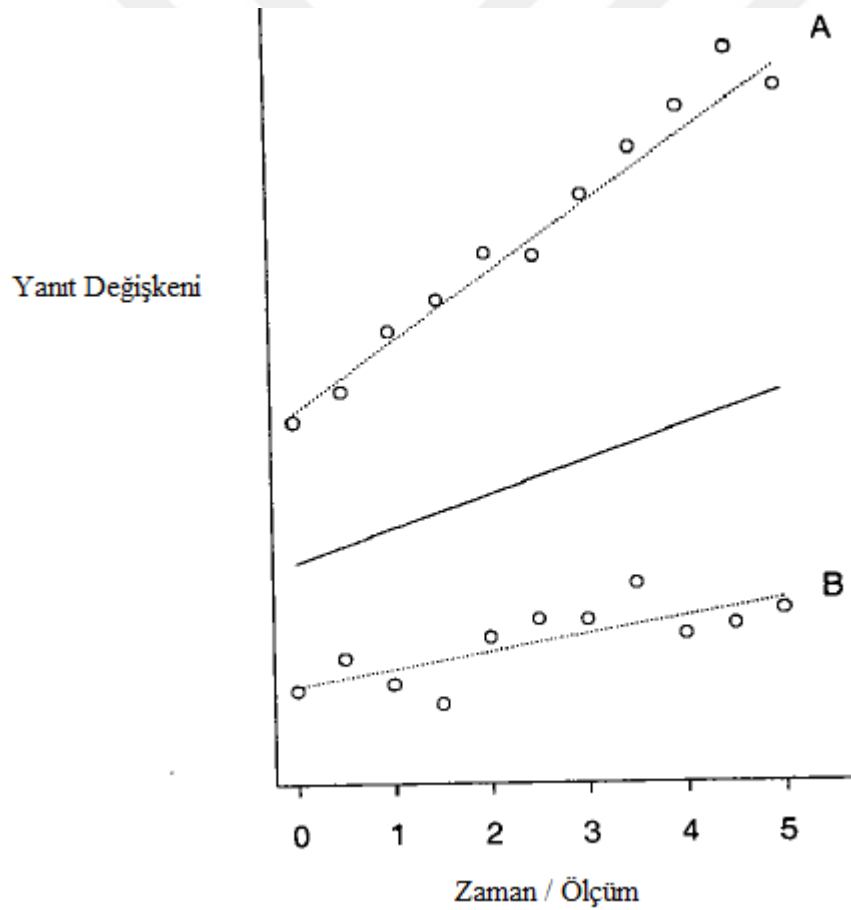
Bu modelde, sabit eğim terimi (fixed slope) için EİDYTE rasgele terimlerin bilinmeyen varyanslarına bağlıdır. Bu varyansların formüldeki parametrelerinin EİDYTE'leri için yerine kullanılacak tahminleri varsa, bulunan tahmin edicinin örnekleme varyansları kapalı formda yazılamaz ve yaklaşık olarak bulunmalıdır. Eğim parametresi için güven aralığı ve yokluk hipotezinin testi $\beta_2 = 0$ (kişisel özellikler yaşa bağlı olarak değişmez) uygun test

istatistiğinin yeterli şekilde yaklaşık olmasına ve varyans tahmininin doğruluğuna bağlı olacaktır.

Bu fikri daha iyi açıklamak için deney birimleri arası kesimi ve eğimi değişen i . deney biriminin j . ölçüm durumu için aşağıdaki model ele alınsın;

$$Y_{ij} = \beta_1 + \beta_2 t_{ij} + b_{1i} + b_{2i} t_{ij} + e_{ij}, \quad j=1, \dots, m_i \quad (3.16)$$

Bu modelde, her deney biriminin yanıt değişkeni kendi alt düzeyine göre farklılaşmasının yanı sıra zaman içinde yanıt değişkeninin değişimine göre de farklılaşmıştır.



Şekil 3.3. Hata terimi ile marjinal ve koşullu yanıt değişkeninin değişimi

Bu durumu içeren model denkleminin grafiksel gösterimi Şekil 3.3 ile daha iyi anlaşılır. Şekil 3.3 marjinal ortalama yanıtın zaman / ölçüm içinde doğrusal olarak nasıl değiştiğini

ayrıca iki farklı deney biriminin (A ve B) koşullu yanıt ortalamalarının kitle trendinden nasıl farklılaştığını göstermektedir.

DKE modelindeki en önemli farklılık Y_{ij} 'nin marjinal ortalaması ve koşullu ortalaması arasındaki farktır. b_i verildiğinde Y_i 'nin koşullu veya deney birimine özgü ortalaması $E(Y_i | b_i) = X_i\beta + Z_ib_i$ şeklindedir.

Y_i 'nin marjinal veya kitlenin hepsinden (b_i 'nin dağılımı üzerinden ortalananmış) elde edilmiş ortalaması,

$$\begin{aligned} E(Y_i) &= \mu_i \\ &= E\{E(Y_i | b_i)\} = E(X_i\beta + Z_ib_i) \text{ olarak ifade edilir.} \\ &= X_i\beta + Z_iE(b_i) = X_i\beta \end{aligned}$$

DKE modellerinde regresyon parametreleri β 'lar her deney birimi için aynı kabul edilir ve kitle-ortalananmış (averaged) yorumları vardır. β 'nın tersine b_i vektörü (ilgili sabit etki ile birleştirildiğinde) deney birimine-özümlü regresyon terimlerini oluşturur. Rasgele etkiler sabit etkilerle birleştğinde her deney birimi için ortalama yanıt profillerini oluştururlar. i . deney birimi için ortalama yanıt profili aşağıdaki gibidir;

$$E(Y_i | b_i) = X_i\beta + Z_ib_i.$$

Bir çok çalışmada, R_i diyagonal matris olarak alınır ; I_{m_i} $m_i \times m_i$ boyutlu birim matris olmak üzere $R_i = \sigma^2 I_{m_i}$ gösterilir. Bu durumda hata terimlerinin deney birimi içinde bağımsız olduğu varsayılır; $\text{kov}(e_{ij}, e_{ik}) = 0$. Bu çok sınırlı bir durumdur ve R_i için farklı kovaryans yapılarını düşünmek en doğru olana karar vermek daha doğrudur. Eş. 3.16'daki spesifik model denkleminin matris ve vektör gösterimleri aşağıdaki gibi seçilebilir;

$$X_i = Z_i = \begin{pmatrix} 1 & t_{i1} \\ 1 & t_{i2} \\ \vdots & \vdots \\ 1 & t_{in_i} \end{pmatrix}.$$

Burada, $q = p = 2$ ve Z_i X_i 'nin iki sütunundan oluşmuştur. Bu model her deney biriminin yanıt değişkenindeki alt düzeylerinin farklılığının yanı sıra zaman/ ölçüm içinde ortalama yanıt değişkenindeki değişimin de farklılaşabileceğini modellemiştir. Açıklayıcı değişkenlerin etkileri (faktörler, maruz kalmalar, deney birimlerinin arka plan özellikleri) bu açıklayıcı değişkenlere bağlı kesim ve eğim terimlerinin farklı grup ve faktör düzeylerinde farklılaşmasına izin verilerek kapsama alınabilir.

Örneğin, kontrol ve işlem grubu olan iki gruplu bir çalışmayı ele alınsın. Yanıt değişkeni zaman/ ölçüm içinde doğrusal bir eğilimde değişiyorsa kesimlerin ortalaması ve eğimler gruplara bağlıysa aşağıdaki DKE modeli uydurulabilir;

$$Y_{ij} = \beta_1 + \beta_2 t_{ij} + \beta_3 Grup_i + \beta_4 t_{ij} \times Grup_i + b_{1i} + b_{12} t_{ij} + e_{ij},$$

Eğer deney birimi işlem grubuna atanmışsa $Grup_i = 1$ değerini alır, ve eğer kontrol grubuna atanmışsa $Grup_i = 0$ değerini alır. Bu modelde kontrol grubu için X_i aşağıdaki gibi olur:

$$X_i = \begin{pmatrix} 1 & t_{i1} & 0 & 0 \\ 1 & t_{i2} & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 1 & t_{in_i} & 0 & 0 \end{pmatrix}$$

İşlem grubu için tasarım matrisi aşağıdaki gibi olur:

$$X_i = \begin{pmatrix} 1 & t_{i1} & 1 & 1 \\ 1 & t_{i2} & 1 & 1 \\ \vdots & \vdots & \vdots & \vdots \\ 1 & t_{in_i} & 1 & 1 \end{pmatrix}$$

Z_i tasarım matrisi kontrol ve işlem grubu için aynı kalacaktır:

$$Z_i = \begin{pmatrix} 1 & t_{i1} \\ 1 & t_{i2} \\ \vdots & \vdots \\ 1 & t_{in_i} \end{pmatrix}$$

3.4.3. Rasgele etkiler için kovaryans yapıları

Y_i yanıt değişken vektörünün kovaryans matrisi $V_i = Z_i G Z_i' + R_i$ olarak gösterilir. Tezde G rasgele etkiler için kovaryans matrisinin her deney birimi için aynı olduğu varsayılacaktır. Fei et al. (2016) çalışmasında farklı rasgele etkilerin her deney biriminde farklı kovaryansa sahip olabileceğini ifade etmiştir. Genellikle literatürde görülen kovaryans matrisleri varyans bileşenleri yapısında veya yapısız bir matris şeklinde görülmüştür. Fei ve diğerleri (2016) çalışmasında rasgele etkilerin kovaryans yapılarını varyans bileşenleri, yapısız ve AR(1) olarak seçmiştir. 3 rasgele etkiye sahip DKE modelinde $Y_{ij} = \beta_0 + \beta_1 x_{ij} + \beta_2 t_{ij} + b_{1i} + b_{2i} x_{ij} + b_{3i} t_{ij} + e_{ij} \quad i = 1, \dots, n; \quad j = 1, \dots, m$ aşağıda en çok görülen rasgele etkiler için kovaryans yapıları özetlenmiştir.

a. Genel simetrik pozitif tanımlı kovaryans yapısı (yapısız):

$$G = \text{Var}(b_i) = \begin{bmatrix} \sigma_{b_1}^2 & \sigma_{b_1, b_2} & \sigma_{b_1, b_3} \\ \sigma_{b_2, b_1} & \sigma_{b_2}^2 & \sigma_{b_2, b_3} \\ \sigma_{b_3, b_1} & \sigma_{b_3, b_2} & \sigma_{b_3}^2 \end{bmatrix}$$

b. Birim matrisle çarpılan aynı varyanslı (homojen) kovaryans yapısı:

$$G = \text{Var}(b_i) = \begin{bmatrix} \sigma^2 & 0 & 0 \\ 0 & \sigma^2 & 0 \\ 0 & 0 & \sigma^2 \end{bmatrix}$$

c. Birleşik simetrik (compound symmetry) kovaryans yapısı:

$$G = \text{Var}(b_i) = \begin{bmatrix} \sigma_1^2 + \sigma_2^2 & \sigma_1^2 & \sigma_1^2 \\ \sigma_1^2 & \sigma_1^2 + \sigma_2^2 & \sigma_1^2 \\ \sigma_1^2 & \sigma_1^2 & \sigma_1^2 + \sigma_2^2 \end{bmatrix}, G = \text{Var}(b_i) = \begin{bmatrix} \sigma^2 & \sigma_1^2 & \sigma_1^2 \\ \sigma_1^2 & \sigma^2 & \sigma_1^2 \\ \sigma_1^2 & \sigma_1^2 & \sigma^2 \end{bmatrix},$$

$$G = \text{Var}(b_i) = \sigma^2 \begin{bmatrix} 1 & \rho & \rho \\ \rho & 1 & \rho \\ \rho & \rho & 1 \end{bmatrix}.$$

- d. Birinci mertebeden geçici otokorelasyonu modellemek için kullanılan otoregresif kovaryans yapısı:

$$G = \text{Var}(b_i) = \sigma^2 \begin{bmatrix} 1 & \rho & \rho^2 \\ \rho & 1 & \rho \\ \rho^2 & \rho & 1 \end{bmatrix}.$$

- e. Heterojen birleşik simetrik kovaryans yapısı:

$$G = \text{Var}(b_i) = \begin{bmatrix} \sigma_1^2 & \sigma_1 \sigma_2 \rho & \sigma_1 \sigma_3 \rho \\ \sigma_2 \sigma_1 \rho & \sigma_2^2 & \sigma_2 \sigma_3 \rho \\ \sigma_3 \sigma_1 \rho & \sigma_3 \sigma_2 \rho & \sigma_3^2 \end{bmatrix}.$$

- f. Heterojen geçici otokorelasyonu modellemek için kullanılan otoregresif kovaryans yapısı:

$$G = \text{Var}(b_i) = \begin{bmatrix} \sigma_1^2 & \sigma_1 \sigma_2 \rho & \sigma_1 \sigma_3 \rho^2 \\ \sigma_2 \sigma_1 \rho & \sigma_2^2 & \sigma_2 \sigma_3 \rho \\ \sigma_3 \sigma_1 \rho^2 & \sigma_3 \sigma_2 \rho & \sigma_3^2 \end{bmatrix}.$$

Yukarıda gösterimi yapılan kovaryans yapıları rasgele etkiler için sıklıkla kullanılan yapılardır.

3.4.4. Hata terimleri için kovaryans yapıları

Bu bölümde, literatürde sıklıkla görülen hata terimlerinin sahip olabileceği kovaryans yapılarından bahsedilecektir. Bu yapılar özellikle uzun vadeli zaman içinde yapılan tekrarlı ölçüm çalışmalarında ortaya çıkmaktadır. Veri seti belli bir kovaryans yapısına uymadığında veya uydurulamadığında az veri ile çok sayıda parametre tahmin etmek zorunda kalınır. Bir yapıya uydurulmamış kovaryansın en önemli eksikliği az veri ile çok sayıda parametre tahmini yapılmasıdır. Bu durumda ilgilenilen parametre vektörü β 'nın tahmininin doğruluğu azalır. Böylece β için yapılan çıkarımlar yanlış olur. Aynı şekilde uygun olmayan bir kovaryans modeli seçilirse yine β çıkarımları doğru olmayacaktır. Hata terimlerinin kovaryans modeli seçilirken bu iki durum iyi dengelenmelidir.

Zaman serileri için geliştirilmiş bir çok yapı seri korelasyon içeren uzun vadeli çalışmalar için kullanılabilir. Uzun vadeli veri setleri daha az ölçüm ve deney birimi içermesi bakımından zaman serileri verisinden biraz farklı olsa da ortak bir karakteristiğe sahiptirler: tekrarlı ölçümler pozitif veya negatif korelasyonludur ve yakın zamanlarda alınan ölçümler arası korelasyon daha güçlüdür.

Birleşik simetrik kovaryans yapısı

Tekrarlı ölçüm verilerinde ilk kullanılan kovaryans yapısıdır. Bu yapı ölçümler arası varyansı ve korelasyonu sabit kabul eder, σ^2 , ve tüm i ve k için $\text{kor}(e_{ij}, e_{ik}) = \rho$ 'dır. Literatürde, uzun vadeli veri setlerinde genellikle korelasyonun zaman içinde gözlemler arasında azaldığı görülmüştür. Zaman içinde korelasyonun sabit olması bu çalışmalarda genellikle gerçekçi değildir çok az çalışmada bu duruma rastlanmıştır. $-1 \leq \rho \leq 1$ sınırlaması olmak üzere bu yapının matris formu aşağıdadır;

$$\text{kov}(e_i) = \sigma^2 \begin{pmatrix} 1 & \rho & \rho & \dots & \rho \\ \rho & 1 & \rho & \dots & \rho \\ \rho & \rho & 1 & \dots & \rho \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho & \rho & \rho & \dots & 1 \end{pmatrix}.$$

Bu yapı genellikle rasgele kesim modelinde görülmektedir. Bu en basit DKE modelinin kovaryans matrisi bir önceki bölümde açıklandığı gibi marjinal olarak birleşik simetrik yapıya sahip olmaktadır.

Toeplitz kovaryans yapısı

Zaman içinde eşit olarak ayrılmış ölçümler arasında aynı korelasyonu varsayar. Kovaryans Toeplitz yapısında olduğunda ölçümler arası varyans aynı olduğu varsayılır, σ^2 , ve tüm i ve k için $\text{kor}(e_{ij}, e_{ij+k}) = \rho_k$ 'dır. Toeplitz kovaryans yapısı birbirine en yakın olan ölçümler arası gözlemlerin korelasyonunun sabit olduğunu varsaydığından ölçümler eşit zaman aralıklarıyla yapılmalıdır.

$$kov(e_i) = \sigma^2 \begin{pmatrix} 1 & \rho_1 & \rho_2 & \cdots & \rho_{m-1} \\ \rho_1 & 1 & \rho_1 & \cdots & \rho_{m-2} \\ \rho_2 & \rho_1 & 1 & \cdots & \rho_{m-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho_{m-1} & \rho_{m-2} & \rho_{m-3} & \cdots & 1 \end{pmatrix}$$

Otoregresif kovaryans yapısı

Otoregresif yapı durumlar arası varyansın aynı olduğunu varsayar, σ^2 , ve tüm i ve k için $kov(e_{ij}, e_{ij+k}) = \rho^k$ ’dır. Oldukça yaygın kullanılan bir yapıdır ve ölçüm sayısından bağımsız sadece iki parametresi vardır. Otoregresif kovaryans yapısı Toeplitz’in özel bir formu olduğundan ölçümler eşit zaman aralıklarında olduğunda uygun olur. Tekrarlı ölçümler zaman uzaklığı arttıkça korelasyonun zaman içinde azalacağını varsayar.

Aynı deney biriminden alınan gözlemler zaman içinde korelasyonları genellikle çok hızlı bir şekilde düşmez. e_{ij} hata terimi $e_{i0} \sim N(0, \sigma^2)$ ile başlayan ve $e_{ij} \sim \rho e_{ij-1} + w_{ij}$, $w_{ij} \sim N(0, \sigma^2[1 - \rho^2])$ yapısına sahipse teorik olarak otoregresif bir proses olduğu söylenir. j . durumdaki hata terimi bir önceki durumdaki hata terimi ρe_{ij-1} ve ek bağımsız hata teriminin w_{ij} deterministik bir fonksiyonudur. Bu proses için aşağıdakiler yazılabilir;

$$Var(e_{ij}) = \sigma^2 \quad kov(e_{ij}, e_{ik}) = \sigma^2 \rho^{|j-k|}, \quad kov(e_i) = \sigma^2 \begin{pmatrix} 1 & \rho & \rho^2 & \cdots & \rho^{m-1} \\ \rho & 1 & \rho & \cdots & \rho^{m-2} \\ \rho^2 & \rho & 1 & \cdots & \rho^{m-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho^{m-1} & \rho^{m-2} & \rho^{m-3} & \cdots & 1 \end{pmatrix}.$$

Birleşik simetrik, Toeplitz ve otoregresif kovaryans zaman içinde varyansların sabit olduğunu varsayar. Bu varsayım esnetilebilir. Heterojen varyansın olduğu modelde, $Var(e_{ij}) = \sigma_j^2$ düşünülebilir. Bu yapı m varyans parametresi ve bir korelasyon parametresine sahiptir.

Böylece heterojen otoregresif kovaryans modeli aşağıdaki gibi verilebilir;

$$kov(e_i) = \begin{pmatrix} \sigma_1^2 & \rho\sigma_1\sigma_2 & \rho^2\sigma_1\sigma_3 & \dots & \rho^{m-1}\sigma_1\sigma_m \\ \rho\sigma_1\sigma_2 & \sigma_2^2 & \rho\sigma_2\sigma_3 & \dots & \rho^{m-2}\sigma_2\sigma_m \\ \rho^2\sigma_1\sigma_3 & \rho\sigma_2\sigma_3 & \sigma_3^2 & \dots & \rho^{m-3}\sigma_3\sigma_m \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho^{m-1}\sigma_1\sigma_m & \rho^{m-2}\sigma_2\sigma_m & \rho^{m-3}\sigma_3\sigma_m & \dots & \sigma_m^2 \end{pmatrix}$$

Bantlı kovaryans yapısı

Belirlenmiş bir aralıktan sonra korelasyonun sıfır olacağını varsayan bir kovaryans modelidir. Örneğin bant derecesi üç olan kovaryans modeli $k \geq 3$ için $kor(e_{ij}, e_{ij+k}) = 0$ olduğunu varsayar. Bu bölümde bahsedilen kovaryans yapılarına bantlı yapı uygulanabilir. Örneğin iki bantlı Toeplitz yapısı aşağıdaki gibi gösterilebilir;

$$kov(e_i) = \sigma^2 \begin{pmatrix} 1 & \rho_1 & 0 & \dots & 0 \\ \rho_1 & 1 & \rho_1 & \dots & 0 \\ 0 & \rho_1 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & 1 \end{pmatrix}$$

Bantlı kovaryans yapısı tekrarlı ölçümler arası zaman aralığı arttıkça korelasyonun hızlı bir şekilde sıfıra ulaştığını varsayar. Genellikle tıbbi araştırmalarda çalışma çok uzun bir zaman aralığı ile yapılsa da korelasyonun sıfır olması genellikle gözlenmez.

Üstel kovaryans yapısı

Ölçümler zaman içinde eşit aralıklarla alınmadıysa otoregresif kovaryans modeli genelleştirilebilir;

$\{t_{i1}, \dots, t_{im}\}$ i. deney biriminden alınan gözlemlerin zamanlarını temsil etmek üzere tüm ölçüm durumlarında sabit σ^2 varyans varsayılırsa $\rho \geq 0$ için $kor(e_{ij}, e_{ik}) = \rho^{|t_{ij} - t_{ik}|}$ olur. İki tekrarlı ölçüm arasında zaman aralığı arttıkça korelasyon üstel olarak azalır. Bu yapı üstel kovaryans olarak aşağıdaki gibi ifade edilir;

$$Kov(e_{ij}, e_{ik}) = \sigma^2 \rho^{|t_{ij} - t_{ik}|} = \sigma^2 \exp(-\theta |t_{ij} - t_{ik}|), \quad \theta = -\log(\rho) \text{ veya } \rho = \exp(-\theta), \quad \theta \geq 0$$

Zaman ölçeğinin doğrusal dönüşümleri altında üstel kovaryans modeli değişmezdir. Eğer t_{ij} , $(a + bt_{ij})$ ile yer değiştirilirse (örneğin zaman ölçeği günden haftaya değişirse) kovaryans matrisi yine aynı forma sahip olur. Üstel kovaryans yapısında, aynı durumdan tekrarlı olarak yapılan ölçümlerin (veya aynı durumda deney biriminden alınmış tekerrürlerin) korelasyonun bir olduğu ve ölçümler arası zaman aralığı arttıkça korelasyonun hızlıca sıfır olduğu varsayılır. İlk özellik ölçümlerin hatasız olduğu varsayımını ortaya çıkarır; bu genellikle tıbbi çalışmalarda gerçekçi olmayan bir varsayımdır. İkinci özellik de korelasyonların tekrarlı ölçümler arasında sıfıra yaklaşması uzun vadeli tıbbi çalışmalarda pek rastlanmayan bir durumdur.

Hibrid kovaryans yapıları

Otoregresif ve birleşik simetrik kovaryans yapıları bir araya getirilerek uzun vadeli çalışmalarda ortaya çıkan bir çok istenmeyen durum hibrid yapılarla üstesinden gelinebilir. Örneğin ölçümler arası korelasyon zaman aralıkları arttıkça otoregresif durumda hızlıca sıfıra yaklaşır ama bu korelasyonun belli bir değerde sabitlenmesi istenirse otoregresif yapı birleşik simetrik yapı ile birleştirilebilir. Böylece bu istenmeyen durum hibrid bir kovaryans modeli ile çözülebilir. Aşağıdaki kovaryans modelleri ele alınsın;

$$\Sigma_1 = \sigma_1^2 \begin{pmatrix} 1 & \rho_1 & \rho_1 & \dots & \rho_1 \\ \rho_1 & 1 & \rho_1 & \dots & \rho_1 \\ \rho_1 & \rho_1 & 1 & \dots & \rho_1 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho_1 & \rho_1 & \rho_1 & \dots & 1 \end{pmatrix}, \quad \Sigma_2 = \sigma_2^2 \begin{pmatrix} 1 & \rho_2^{|t_{i1}-t_{i2}|} & \rho_2^{|t_{i1}-t_{i3}|} & \dots & \rho_2^{|t_{i1}-t_{im}|} \\ \rho_2^{|t_{i2}-t_{i1}|} & 1 & \rho_2^{|t_{i2}-t_{i3}|} & \dots & \rho_2^{|t_{i2}-t_{im}|} \\ \rho_2^{|t_{i3}-t_{i1}|} & \rho_2^{|t_{i3}-t_{i2}|} & 1 & \dots & \rho_2^{|t_{i3}-t_{im}|} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho_2^{|t_{im}-t_{i1}|} & \rho_2^{|t_{im}-t_{i2}|} & \rho_2^{|t_{im}-t_{i3}|} & \dots & 1 \end{pmatrix}$$

olmak üzere $kov(e_i) = \Sigma_1 + \Sigma_2$ yazılır. Bu kovaryans modeli için aşağıdaki ifadeler elde edilir:

$$Var(e_{ij}) = \sigma_1^2 + \sigma_2^2$$

$$kov(e_{ij}, e_{ik}) = \rho_1 \sigma_1^2 + \rho_2^{|t_{ij}-t_{ik}|} \sigma_2^2$$

$$kor(e_{ij}, e_{ik}) = \frac{\rho_1 \sigma_1^2 + \rho_2^{|t_{ij}-t_{ik}|} \sigma_2^2}{\sigma_1^2 + \sigma_2^2}.$$

Aynı durumdan deney biriminden alınmış tekrarlı ölçümler arası korelasyonu

$\frac{\rho_1 \sigma_1^2 + \rho_2 \sigma_2^2}{\sigma_1^2 + \sigma_2^2}$ 'dur. $\rho_1 < 1$ olduğunda bu korelasyon birden az olur. Ayrıca, zaman aralığı

arttıkça, korelasyonun birden sıfırlanması yerine en az $\frac{\rho_1 \sigma_1^2}{\sigma_1^2 + \sigma_2^2}$ olur (Fitzmaurice ve diğerleri, 2004: 173).

4. DKE MODELLERİNDE TAHMİN VE KESTİRİM METODLARI

Bu bölümde DKE modelinin istatistikleri üzerinde durulacaktır. DKE modelinin rasgele ve sabit etkilerinin tahmin edilebilmesi için öncelikle modeldeki varyans parametrelerinin tahmin edilmesi gerekir. İstatistik literatüründe çeşitli varyans tahmin metodları vardır, bunlar ANOVA, en çok olabilirlik, kısıtlı veya artık en çok olabilirlik metodlarıdır. DKE modelindeki sabit etki parametrelerinin tahminleri EÇO ve KEÇO yöntemlerine göre sayısal algoritmalar yoluyla bulunur.

4.1. ANOVA Tahmini

Fisher, varyans ve varyans analizi terimlerini 1918’de istatistik literatürüne kazandırmıştır (Scheffe, 1956). Dengeli verilerde varyans parametreleri için ANOVA tahmini, gözlenen değerlerden ortalamaları çıkarılıp kareleri alınarak oluşturulan denklemlerden elde edilir. Örneğin, n sınıf ve her sınıfta m gözlemi olan tek yönlü ANOVA modelinde grup içi hata kareler ortalaması $E(KTH) = n(m-1)\sigma_e^2$. ve gruplar arası hata kareler ortalaması $E(KTG) = (n-1)\sigma_e^2 + m(n-1)\sigma_\alpha^2$ olarak ifade edilir. σ_α^2 terimi dengeli olmayan veri setlerinde negatif olabilir. ANOVA varyans tahminlerinin dengeli olmayan durumlarda genellikle elde edilmesi zordur. Henderson (1953) bu durumda kendi adını taşıyan üç metod önermiştir (Searle ve diğerleri, 1992: 34).

4.2. En Çok Olabilirlik Tahmin Metodu

En çok olabilirlik metodunun Fisher’e dayanan bir tarihi vardır. Hartley ve Rao (1967)’ye kadar karışık etkili modellerde kullanılmamışlardır. Bölüm 3’de bahsedildiği üzere Eş. 3.1 marijnal model olarak bilinir ve $Y_i \sim N(X_i\beta, Z_iGZ_i' + R_i)$ olduğunu ifade eder.

β yerine olabilirlik fonksiyonunda Eş. 3.4 yazıldıktan sonra, α ’nın en çok olabilirlik tahminleri Eş. 3.10’daki olabilirlik fonksiyonu α ’ya göre maksimize edilerek elde edilir. α ve β ’nin aynı anda (simultaneously) tahmini Eş. 3.3’deki birleşik (joint) log-olabilirlik fonksiyonunun maksimize edilmesi ile mümkün olur.

Sonuçta elde edilen tahminler $\hat{\alpha}_{E\check{O}}$ ve $\hat{\beta}_{E\check{O}}$ olarak ifade edilir. Standard hatalar Fisher bilgi matrisinin tersi yoluyla elde edilir. Bu yaklaşım α 'nın tahmini sonucu ortaya çıkan ekstra değişkenliği dikkate alan $\hat{\beta}_{E\check{O}}$ 'nin kovaryans matrisini ortaya çıkarır. Bu yöntem SAS'ın PROC MIXED prosedüründe kullanılmaz. Eş. 3.3'de verilen birleşik (joint) olabilirlik fonksiyonu kompakt model Eş.3.2 için aşağıdaki gibi yazılır;

$$f(y) = \frac{1}{(2\pi)^{N/2} |V|^{1/2}} \exp \left\{ -\frac{1}{2} (y - X\beta)' V^{-1} (y - X\beta) \right\}$$

α vektörü V 'deki tüm varyans bileşenlerini içermek üzere log-olabilirlik fonksiyonu aşağıdaki gibi yazılır;

$$l(\beta, \alpha) = c - \frac{1}{2} \log |V| - \frac{1}{2} (y - X\beta)' V^{-1} (y - X\beta), \quad (4.1)$$

α_r r . varyans bileşeni ve qv α vektörünün boyutu olmak üzere parametrelere göre log-olabilirlik fonksiyonu türevlenerek aşağıdakiler elde edilir;

$$\frac{\partial l}{\partial \beta} = X' V^{-1} y - X' V^{-1} X \beta, \quad (4.2)$$

$$\frac{\partial l}{\partial \alpha_r} = \frac{1}{2} \left\{ (y - X\beta)' V^{-1} \frac{\partial V}{\partial \alpha_r} V^{-1} (y - X\beta) - iz \left(V^{-1} \frac{\partial V}{\partial \alpha_r} \right) \right\} \quad r = 1, \dots, qv. \quad (4.3)$$

EÇÖ tahmini bulmak için kullanılan standart prosedür $\partial l / \partial \beta = 0$, $\partial l / \partial \alpha_r = 0$ EÇÖ denklemlerini sıfıra eşitlemektir. Eş. 4.2 ve 4.3'de verilen çözümler her zaman EÇÖ tahminleri olmayabilir. p , β vektörünün boyutudur ve X tam (sütun) ranklı matris olduğundan $rank(X) = p$ 'dir. Varyans bileşenlerinin EÇÖ tahminleri $\hat{V} = V(\hat{\alpha})$ içermek üzere Eş. 4.2'den aşağıdaki elde edilir;

$$\hat{\beta} = (X' \hat{V}^{-1} X)^{-1} X' \hat{V}^{-1} y \quad (4.4)$$

Böylece, α yani varyans parametreleri için EÇO tahmini bulunursa, β 'nın kapalı-form ifadesi Eş. 4.4'deki gibi elde edilir. Eş. 4.3 ve 4.2 α 'nın EÇO tahmini için aşağıdaki ifadeyi göstermeye olanak sağlar;

$$P = V^{-1} - V^{-1}X(X'V^{-1}X)^{-1}X'V^{-1},$$

$$y'P \frac{\partial V}{\partial \alpha_r} Py = iz \left(V^{-1} \frac{\partial V}{\partial \alpha_r} \right). \quad (4.5)$$

$\hat{\beta}$ 'nin hesaplanabilmesi için öncelikle Eş.4.5 çözülmelidir. EÇO tahminlerinin asimptotik özellikleri tutarlılık asimptotik normallik gibi ilk olarak Hartley ve Rao(1967) tarafından incelenmiştir.

Uygun koşullar altında EÇO tahminleri tutarlı ve Fisher bilgi matrisinin tersine eşit olan asimptotik kovaryans matrisi ile asimptotik olarak normal dağılır. $\theta = (\beta', \alpha')$ olmak üzere düzenleyici koşullar altında Fisher bilgi matrisi aşağıdaki gibi olur;

$$Var \left(\frac{\partial l}{\partial \theta} \right) = -E \left(\frac{\partial^2 l}{\partial \theta \partial \theta'} \right) \quad (4.6)$$

Eş. 4.2 ve 4.3 yoluyla Eş. 4.6'nın elemanları için sonraki ifadeler elde edilebilir. Örneğin V 'nin iki kez sürekli olarak parametrelere göre türevlenebilir olduğu varsayılırsa;

$$E \left(\frac{\partial^2 l}{\partial \beta \partial \beta'} \right) = -X'V^{-1}X,$$

$$E \left(\frac{\partial^2 l}{\partial \beta \partial \alpha_r} \right) = 0, \quad 1 \leq r \leq qv,$$

$$E \left(\frac{\partial^2 l}{\partial \alpha_r \partial \alpha_s} \right) = -\frac{1}{2} iz \left(V^{-1} \frac{\partial V}{\partial \alpha_r} V^{-1} \frac{\partial V}{\partial \alpha_s} \right), \quad 1 \leq r, s \leq qv \quad (4.7)$$

Eş. 4.6'nın β 'ya bağlı olmadığı görülmüştür. Bu nedenle Eş. 4.6 $I(\alpha)$ olarak ifade edilebilir.

4.3. Kısıtlı En Çok Olabilirlik Tahmin Metodu

DKE modelleri genellikle bir çok sabit etki içerir. Sabit etkiler tahmin edilirken ortaya çıkan serbestlik derecesinin kaybolması problemi açık bir şekilde ele alınarak varyans bileşenlerini tahmin etmek önemlidir. Kısıtlı en çok olabilirlik tahmini bu kaybı önlemeye olanak sağlar. KEÇO yönteminde varyans parametrelerinin EÇO tahmin edicileri bulunmadan önce sabit etkiler EÇO fonksiyonundan çıkarılır (Patterson ve Thompson,1971). X üzerinden Y 'nin doğrusal regresyonundan artıklar hesaplanarak EKK regresyonundan elde edilen artıklar vektörü Y vektörü olarak karışık etkili modele uydurulurken kullanılır, bu yeni vektör üzerinden EÇO tahmin edicileri hesaplanır (EKK artıkları bağımlı gözlemler olarak kullanılır). Orijinal model için son bulunan EÇO varyans parametre tahminleri KEÇO olur. EÇO ve KEÇO tahmin edicileri tutarlı, asimptotik olarak yansızdır. Sabit etki tahmin edilirken serbestlik derecesi kaybı KEÇO tekniğinde ek bir avantaj olarak düzeltilir ama EÇO tekniğinde bu düzeltme yoktur. Sabit etkileri yok etmek için veride dönüştürme yapılır daha sonra dönüştürülmüş veri varyans bileşenlerinin tahmini için kullanılır. Dengeli veri için örnek çapından bağımsız olarak KEÇO tahminleri yansızdır (Searle ve diğerleri, 1992: 91). EÇO tekniğinde bu düzeltme gözardı edildiğinden tahminler yanlıdır. Eş. 3.1'deki modeldeki tüm n adet deney birimine özgü regresyon modelleri bir araya getirilerek Eş. 3.2'de kompakt model oluşturulur. Z matrisi ana diyagonalinde Z_i matrislerini içeren ve diğer elemanlarının sıfır olduğu blok-diyagonal bir matristir. Y 'nin boyutu $\sum_{i=1}^n m_i$ 'dir ve N ile ifade edilir.

Y gözlem değerlerinin vektörü $X\beta$ ortalama vektörü ve ana diyagonal bloklarında V_i matrislerinin olduğu diğer elemanları sıfır olan $V(\alpha)$ blok-diyagonal kovaryans matrisli normal dağılıma sahiptir. α varyans bileşenlerinin kısıtlı en çok olabilirlik tahminleri (KEÇO) X matrisinin sütunlarına ortogonal sütunları olan $N \times (N - p)$ boyutlu tam-rank K matrisinin olduğu hata kontrastlarını içeren $U = K'Y$ setinin bir fonksiyonu olan olabilirlik fonksiyonunun maksimize edilmesi ile elde edilir. $K'Y = K(X\beta + e) = Ke \sim N(0, K'VK)$ yani U vektörü sıfır ortalama vektörlü ve $K'V(\alpha)K$ kovaryans matrisli normal dağılımlıdır. U böylece β 'ya bağımlılığını yitirir. V 'deki parametrelerin tahmini için $K'Y$ kullanılır. Harville (1974) hata kontrastlarının olabilirlik fonksiyonunu $\hat{\beta}$ Eş. 3.4'de verilmek üzere aşağıdaki gibi göstermiştir;

$$L(\alpha) = (2\pi)^{-(N-p)/2} \left| \sum_{i=1}^n X_i' X_i \right|^{1/2} \left| \sum_{i=1}^n X_i' V_i^{-1} X_i \right|^{1/2} \prod_{i=1}^n |V_i|^{-1/2} \times \exp \left\{ -\frac{1}{2} \sum_{i=1}^n (Y_i - X_i \hat{\beta})' V_i^{-1} (Y_i - X_i \hat{\beta}) \right\} \quad (4.8)$$

$\hat{\alpha}$ varyans bileşenlerinin KEÇO tahmini hata kontrastlarına bağlı değildir. DKE modellerinde sabit etkiler vektörünün EÇO tahminleri KEÇO tahminleri ile aynı olmaz. Eş. 4.8'deki olabilirlik fonksiyonu $L_{EÇO}(\beta, \alpha) = L_{EÇO}(\theta)$, C α 'ya bağlı olmayan bir sabit ve $W_i(\alpha) = V_i^{-1}(\alpha)$ olmak üzere Eş. 3.3'de verilen en çok olabilirlik fonksiyonu aşağıdaki eşitliğe denktir;

$$L(\alpha) = C \left| \sum_{i=1}^n X_i' W_i(\alpha) X_i \right|^{-1/2} L_{EÇO}(\hat{\beta}(\alpha), \alpha) \quad (4.9)$$

Eş. 4.9'da $\left| \sum_{i=1}^n X_i' W_i(\alpha) X_i \right|$ β 'ya bağlı değildir. Aşağıda verilen KEÇO olabilirlik fonksiyonunun α ve β parametrelerine göre maksimize edilmesiyle α ve β 'nın KEÇO tahminleri bulunur;

$$L_{KEÇO}(\theta) = \left| \sum_{i=1}^n X_i' W_i(\alpha) X_i \right|^{-1/2} L_{EÇO}(\theta) \quad (4.10)$$

Ayrıca $\hat{\beta}_{KEÇO} = (X' \hat{V}_{KEÇO}^{-1} X)^{-1} X' \hat{V}_{KEÇO}^{-1} y$ ve $Var(\hat{\beta}_{KEÇO}) = (X' \hat{V}_{KEÇO}^{-1} X)^{-1}$ olur.

Eş. 3.2'deki DKE modellerinin kompakt matris formu için KEÇO yöntemi aşağıdaki gibi ifade edilir :

$U \sim N(0, K'VK)$ böylece U 'nun birleşik olasılık yoğunluk fonksiyonu,

$$f_k(U) = \frac{1}{(2\pi)^{(N-p)/2} |K'VK|^{1/2}} \exp \left\{ -\frac{1}{2} U' (K'VK)^{-1} U \right\} \quad (4.11)$$

$$l_k(\theta) = c - \frac{1}{2} \log(|K'VK|) - \frac{1}{2} U' (K'VK)^{-1} U \quad (4.12)$$

KEÇO türevlenerek y bakımından aşağıdaki ifade elde edilir,

$$\frac{\partial l_k}{\partial \alpha_i} = \frac{1}{2} \left\{ y' P \frac{\partial V}{\partial \alpha_i} P y - i z \left(P \frac{\partial V}{\partial \alpha_i} \right) \right\}, \quad i = 1, \dots, qv \quad (4.13)$$

$$P = K(K'VK)^{-1} K' = V^{-1} - V^{-1} X (X' V^{-1} X)^{-1} X' V^{-1} \quad (4.14)$$

Eş. 4.12 maksimize edilerek α için KEÇO tahminleri elde edilir. EÇO durumunda olduğu gibi KEÇO eşitliği $\partial l_K / \partial \alpha = 0$ sağlayan değerler bulunur. KEÇO tahminleri K matrisinin dönüşümüne dayalı olsa da tahminler K'ya bağlı değildir. Eş. 4.13 ve 4.14'den görülebilir. Bu sonuç önemlidir, çünkü K'nın seçimi tek (unique) değildir ve tahminin K'nın seçimine bağlı olması istenmeyen bir durumdur.

Uygun koşullar altında KEÇO tahminleri tutarlı ve Fisher bilgi matrisinin tersine eşit olan asimptotik kovaryans matrisi ile asimptotik olarak normal dağılır. $\theta = (\beta', \alpha')$ olmak üzere düzenleyici koşullar altında kısıtlayıcı Fisher bilgi matrisi aşağıdaki gibidir;

$$\text{Var}\left(\frac{\partial l_K}{\partial \alpha}\right) = -E\left(\frac{\partial^2 l_K}{\partial \alpha \partial \alpha'}\right) \quad (4.15)$$

V'nin iki kez sürekli olarak (continuously) parametrelere göre türevlenebilir olduğu varsayılırsa;

$$E\left(\frac{\partial^2 l_K}{\partial \alpha_r \partial \alpha_s}\right) = -\frac{1}{2} \text{iz} \left(V^{-1} \frac{\partial V}{\partial \alpha_r} V^{-1} \frac{\partial V}{\partial \alpha_s} \right), \quad 1 \leq i, j \leq qv \quad (4.16)$$

yazılır.

En çok olabilirlik tahminlerinin bazı kötü özellikleri;

- EÇO tahminleri genellikle yanlıdır.
- EÇO tahminleri tek olmak zorunda değildir.
- EÇO tahminleri her zaman elde edilemeyebilirler.
- EÇO tahminlerini genellikle analitik olarak elde etmek zordur. Sayısal yöntemlerle (Beklenti ve Maksimizasyon ve Newton-Raphson algoritması gibi) elde edilebilirler. Bu yöntemlerin tek handikabı tahmin edicilerin başlangıç değerlerinden iyi veya kötü bir şekilde kolayca etkilenmeleridir. Algoritmalara iyi başlangıç değerleriyle başlanırsa iyi parametre tahminleri elde edilebilir.

En çok olabilirlik tahminlerinin bazı iyi özellikleri;

- $\hat{\theta}_n$ “büyük” örnek çaplarında θ için tutarlı tahmin edicidir,

$$\lim_{n \rightarrow \infty} |\hat{\theta}_n - \theta| = 0 \quad .$$

Yani EÇO tahmini örnek çapı arttıkça gerçek parametre değerine yaklaşmaktadır.

- $\hat{\theta}_n$, asimptotik olarak yansızdır, $\lim_{n \rightarrow \infty} E(\hat{\theta}_n) = \theta$. Yeterince “büyük” örnekler için, EÇO tahmininin örnekleme dağılımı gerçek popülasyon değeri etrafında yoğunlaşır. EÇO tahminleri yanlı olsa da, örnek çapı arttıkça bu yan önemsiz olacaktır. Örnek olarak, $x_1, \dots, x_n$ rasgele örneği parametreleri μ ve σ^2 olan normal dağılımdan çekilmiş olsun, σ^2 için EÇO tahmini $\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$ olacaktır. Bu tahmin edici yanlıdır. Bu nedenle genellikle yansız hale gelmiş $s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$ tahmin edicisi kullanılmaktadır.
- $\hat{\theta}_n$, asimptotik olarak etkindir. Yani, bütün yansız tahmin ediciler içinde minimum varyansa sahiptir. Diğer bir deyişle en kesin tahmin olmaya eğilimlidir.
- $\hat{\theta}_n$ 'in varyansı asimptotik olarak büyük örneklerde bilinirdir. $I_n(\theta)^{-1}$, Fisher bilgi matrisinin tersi olmak üzere (n örnek çapı için) $\text{var}(\hat{\theta}_n) = I_n(\theta)^{-1}$ olacaktır. Bu çok önemli bir sonuçtur çünkü EÇO tahmininin standart hatasını hesaplamamıza olanak sağlamaktadır. Bu durum genellikle büyük örneklerde sağlanmaktadır. Küçük örneklerde bu tahmin edici güvenilir olmadığı için literatürde birçok düzeltme önerilmiştir.
- $\hat{\theta}_n$, asimptotik olarak normal dağılır. EÇO tahmininin örnekleme dağılımı en azından büyük örnek çaplarında bilinirse bir önceki özellikle birlikte hipotez testlerinin ve güven aralıklarının oluşturulması için bir temel sağlanmış olur.

4.4. Karma Etkiler için Nokta Tahmini

Modeldeki varyans parametre tahminleri biliniyorsa, karma modeldeki sabit ve/veya rasgele etkilerin tahminleri yapılabilir. Karma etki, sabit ve rasgele etkinin tahmin

edilebilir bir fonksiyonudur. k ve l sabit vektörler olsun. Karma etki $T = k'\beta + l'b$ olarak gösterilsin. Örneğin, Fay-Herriot modelinde i . küçük alan ortalaması $T = x_i\beta + b_i$ olarak yazılır. T 'nin gözlenen değerinin her küçük alan için gözlenen veriler yardımıyla tahmin etmek küçük alan çalışmalarının amacıdır. Karma etki kestirimi $T = k'\beta + l'b$ hem tahmin ($k'\beta$, sabit bileşen olduğu için) hem de kestirim (prediction) ($l'b$, rasgele bileşen) içerir. Bu tezde tahmin sabit etkiler için, kestirim de rasgele ve karma etkiler için kullanılacaktır.

Karma etki T için en iyi doğrusal yansız kestirim (EİDYK) (BLUP), $t(\theta) = k'\beta_{GEKK} + l'GZ'V^{-1}(Y - X\beta_{GEKK})$ olarak Goldberger (1962) ve Henderson (1953) tarafından gösterilmiştir. DKE modellerinin her deney birimi ya da sınıfın kendine özgü parametreye sahip olma varsayımından ortaya çıkmış önemli bir özellik de ilgili deney biriminin gözlem sayısı tahmin edilecek (sabit etki) parametresinden az olsa da deney birimine özgü kestirimler yapılabilir. β_{GEKK} genelleştirilmiş en küçük kareler denklemlerinden elde edilir (Aitken, 1936).

$X'V^{-1}Y = X'V^{-1}\beta$ ve buradan da $\beta_{GEKK} = (X'V^{-1}X)^{-1}X'V^{-1}Y$ olarak elde edilir. β_{GEKK} , β 'nın en iyi yansız doğrusal tahmin edicisi olmanın yanı sıra, β_{GEKK} , herhangi verilen bir V için β 'nın EÇO tahmin edicisidir (Birkes ve Wulf, 2003).

Uygulamada, genellikle θ vektöründeki varyans parametreleri bilinmezdir, β 'nın tahmini ve T 'nin kestirimi iki adımlı bir yaklaşım gerektirir. Öncelikle varyans parametre tahminleri hesaplanır. G ve V varyans-kovaryans matrislerinin tahminleri $\hat{G}$ ve $\hat{V}$ 'yi oluşturmak için kullanılır. Oluşturulan bu tahminler $t(\theta)$ ve β_{GEKK} formüllerinde yerine konulur. β_{GEKK} formülünde V yerine $\hat{V}$ 'yi kullanmak, β 'nın deneysel (tahmin edilmiş) genelleştirilmiş en küçük kareler (DGEKK, EGLS) tahminini verir, $\hat{\beta}_{DGEKK} = (X'\hat{V}^{-1}X)^{-1}X'\hat{V}^{-1}Y$ olarak gösterilir. $t(\theta)$ formülünde $\hat{G}$, $\hat{V}$ ve $\hat{\beta}_{DGEKK}$ yerine konularak elde edildiğinde, sonuç T için deneysel en iyi yansız doğrusal kestirim (DEİDYK) (EBLUP) olur, $t(\hat{\theta}) = k'\beta_{DGEKK} + m'\hat{G}Z'\hat{V}^{-1}(Y - X\beta_{DGEKK})$.

Formülde $\hat{G}$, $\hat{V}$ ve $\hat{\beta}_{DGEKK}$ ANOVA, EÇO veya KEÇO tahmini yerine koymak $t(\theta)$ 'yı hesaplarırken beklenen değeri değiştirmez (Das, Jiang ve Rao, 2004; Kackar ve Harville, 1984). Yani, DEİDYK, EİDYK gibi yansızdır.

4.5. Sabit Etkiler için Nokta Tahmini

Birçok karışık etkili model uygulamasında, araştırmacı öncelikle sabit etkilerle ilgili çıkarımlarla ilgilenir. Sabit etkilerin doğrusal kombinasyonları yukarıda anlatılan karma etkilerin kestirimi probleminin özel bir durumudur, $T = k'\beta + l'b$ formülünde l sıfır olarak alınırsa $T = k'\beta$ formülüne yani sabit etkiyi tahmin etmeye yönelik olur. Karma etki için önerilen teorik sonuçlar sabit etki durumu için yeniden düzeltilebilir. Rasgele etkilerin olmadığı durum $T = k'\beta$ kestirim yerine tahmin gerektirir. Sabit etki, $T = k'\beta$ için, EİDYTE $t(\theta) = k'\beta_{GEKK}$ 'dir. EİDYK tahmin edicisinin özel durumu olan EİDYTE, yansız ve minimum varyanslı tahmin edicidir. Rasgele etkilerin varyans-kovaryans matrisi bilinir olduğunda genelleştirilmiş EKK ile sabit etkilerin tahminleri bulunur. EİDYK durumunda olduğu gibi varyans parametreleri bilinir olmadığında β_{GEKK} hesaplanamaz ve β_{GEKK} formülünde V yerine $\hat{V}$ yazılarak β 'nın tahmin edilmiş genelleştirilmiş en küçük kareler tahmin edicisi β_{DGEKK} tahmin edilir. $t(\theta)$ için formülde β_{GEKK} yerine β_{DGEKK} konularak T için DEİDYK, $t(\hat{\theta}) = k'\beta_{DGEKK}$ bulunur. Birçok DKE modelinde DEİDYK ve DEİDYTE ile ilgili problem hata kareler ortalamalarının (HKO) kapalı formunun olmamasıdır (Hulting ve Harville, 1991).

4.6. Hata Kareler Ortalaması için Yaklaşımlar

4.6.1. Naive hko yaklaşımı

Kackar ve Harville (1984), bir çok popüler varyans parametre tahmin edicisinin (KEÇO ve EÇO dahil olmak üzere), DEİDYK veya DEİDYTE'nin varyansı, $HKO(t(\hat{\theta}))$, aşağıdaki gibi göstermiştir;

$$HKO(t(\hat{\theta})) = HKO(t(\theta)) + E(t(\hat{\theta}) - t(\theta))^2.$$

Bilinen varyans parametreleri için, EİDYK'in varyansı, $HKO(t(\theta))$ bilinirdir. Geleneksel naıve yaklaşımda $HKO(t(\theta))$ için olan formülde V yerine $\hat{V}$ yazılarak $HKO(t(\hat{\theta}))$ elde edilir. Bu naıve yaklaşım $HKO[t(\theta)|\hat{\theta}]$ olarak gösterilsin.

Naıve yaklaşım iki potansiyel hata kaynağı içerir. İlk hata kaynağı, $\hat{V}$ 'nin, $HKO(t(\theta))$ 'yı tahmin etmek için V yerine kullanılmasıdır. İkinci hata kaynağı, $\hat{V}$ 'nin $t(\hat{\theta})$ 'yı tahmin etmek için V yerine kullanılmasıdır, bu durum $t(\theta)$ ve $t(\hat{\theta})$ arasında uyumsuzlığa neden olur, bu uyumsuzluk $E(t(\hat{\theta}) - t(\theta))^2$ naıve yaklaşımda tamamen görmezden gelinmektedir. Bir başka deyişle, DEİDYK 'deki varyans, varyans parametreleri tahminlerini hesaplamak için kullanılan örneklemdaki değişkenlik tarafından şişirilir. Naıve yaklaşım 'varyans yayılma' olarak adlandırılan bu durumu hesaba almakta başarısızdır (Littell, 2002).

EÇO ve KEÇO metodları tutarlı tahmin ediciler sağladığı için, örnek çapı arttıkça V 'nin tahmini olan $\hat{V}$ gelişmektedir, naıve yaklaşımın beklenen hatası sıfıra yaklaşmaktadır (Littell, 2002). Küçük veya ortalama sayıda sınıf/ deney birimi için, $HKO(t(\hat{\theta}))$ değeri ciddi biçimde az tahmin edilebilir, bu durum küçük örneklem için şaşırtıcı değildir. Ayrıca, naıve yaklaşımlarla gerçekleştirilen hipotez testlerinin küçük örneklemde 1. Tip hata oranı artabilir (Catellier ve Muller, 2000; Schaalje, McBride, ve Fellingham, 2002; Kenward ve Roger, 1997).

4.6.2. Sabit etki tahminleri için Kenward-Roger varyans- kovaryans yaklaşımı

Naıve yaklaşımların eksikliklerini bilen bir çok araştırmacı Taylor serisi açılımlarını çeşitli DKE modellerinin varyans parametreleri tahmininde, DEİDYK'nin HKO'sının analitik yaklaşımlarını geliştirmek için çalışmışlardır (Das ve diğerleri, 2004; Datta ve Lahiri, 2000; Prasad ve Rao, 1990). Birinci mertebeden Taylor serisi açılımına göre, Kackar ve Harville (1984) yılında HKO yaklaşımları önermiştir. Harville ve Jeske (1992), Kackar ve Harville yaklaşımını 2. mertebeden Taylor açılımını kullanarak, $HKO(t(\theta))$ 'nın tahmin edicisi $HKO[t(\theta)|\hat{\theta}]$ 'nın yanını düzeltmek için modifiye etmişlerdir. DKE modellerinde sabit etki tahminlerinin kesinlikleri/ hassasiyetleri asimptotik dağılımlarına göre tahmin edilmektedir. Aslında, bir önceki verilen bölümde bahsedilen bu tahmin genellikle yanlı

olup ve sabit etki tahmin edicisinin varyansını normalden az veya fazla tahmin etmektedir. Kenward ve Roger (1997) sabit etki parametrelerinin KEÇO tahmin edicilerinin varyansları için bir düzeltme önermiştir. Bu yaklaşımın orjinal makaledeki hali daha genel bir açıklama olması açısından aşağıda verilmiştir. DKE modelleri için özelleştirilerek simülasyon çalışmasında bu yaklaşım sabit etki tahminlerinin kovaryans matrisini elde etmek için kullanılmıştır.

Y $N \times 1$ boyutlu normal dağılıma sahip doğrusal modelin gözlem vektörü olsun. Y için $N \times p$ boyutlu ve p tam ranklı tasarım matrisi X , $p \times 1$ bilinmeyen sabit etki parametre vektörü β ve $N \times N$ boyutlu ve elemanları $\sigma_{rx1} = (\sigma_1, \dots, \sigma_r)'$ olan Σ kovaryans matrisi olmak üzere $Y \sim N(X\beta, \Sigma)$ olduğu varsayalım. Kenward ve Roger (1997) verinin bağımsız grupları içeren gözlemlerden oluştuğunu dolayısıyla Σ kovaryans matrisinin DKE modellerinde olduğu gibi blok-diyagonal bir matris olduğunu varsaymıştır. Tekrarlanan ölçümlü, gruplanmış veya yüzey (spatial) modellerinde kullanılan veri setlerine uygun olarak düşünülmüştür.

β 'nin GEKK tahmin edicisi $\tilde{\beta} = (X' \Sigma^{-1} X)^{-1} X' \Sigma^{-1} y$, ve bu tahmin edicinin varyans-kovaryans matrisi $\Phi(\sigma) = (X' \Sigma(\sigma)^{-1} X)^{-1}$ olarak gösterilir, yani buradaki $\tilde{\beta}$ tahmininde varyans parametrelerinin bilinir olduğu varsayılır. σ 'nin KEÇO tahmin edicisi $\hat{\sigma}$, ve β 'nin KEÇO tahminine dayalı genelleştirilmiş tahmini EKK tahmini $\hat{\beta} = (X' \Sigma^{-1}(\hat{\sigma}) X)^{-1} X' \Sigma^{-1}(\hat{\sigma}) y$ veya $\hat{\beta} = \Phi(\hat{\sigma}) X' \Sigma^{-1}(\hat{\sigma}) y$ olarak gösterilmiştir. Kackar ve Harville (1984) $\hat{\beta}$ 'nin β için yansız bir tahmin olduğunu göstermişlerdir. β 'nin tahmininin kesinliği üzerinde farklı kovaryans yapılarının, modelin ve örnek çapının etkisi olur. Ancak geleneksel asimptotik özelliklere dayalı yaklaşım Σ 'nin tahmininden kaynaklanan değişkenliği hesaba katmaz. Böyle durumlarda asimptotik yaklaşıma dayalı Wald-tipi güven aralıklarının güvenilirliği azalır.

$\hat{\Phi} = \Phi(\hat{\sigma})$, $Var(\hat{\beta})$ için küçük örnek çaplarında tahmin edici olarak kullanıldığında iki potansiyel hata kaynağı içerir: İlki, $\hat{\Phi}$, Φ 'yı olduğundan az tahmin eder, nedeni ise $\Phi = Var(\tilde{\beta})$, $\hat{\sigma}$ 'nin $\hat{\beta} = \tilde{\beta}(\hat{\sigma})$ içindeki değişkenliğini hesaba katmaz. İkincisi de $\hat{\Phi} = \Phi(\hat{\sigma})$ tahmin edicisi $\Phi(\sigma)$ için yanlıdır. Kackar ve Harville (1984) $\hat{\beta}$ 'nin içindeki

değişkenliği $Var(\hat{\beta}) = \Phi + \Lambda$ şeklinde iki parçaya bölmüştür Λ 'yı yani ikinci hata kaynağını, ele almışlardır. $\Lambda = Var(\hat{\beta} - \tilde{\beta})$ olarak gösterilir. Λ geleneksel asimptotik kovaryans matrisinin $Var(\hat{\beta})$ içindeki değişkenliği ne kadar az tahmin ettiğini gösteren miktarı içeren bileşendir. İlk hata kaynağının düzeltilmesi için bir yaklaşım Harville ve Jeske (1992) tarafından önerilmiştir. Kenward ve Roger (1997) ilk hata kaynağını düzeltmek için bir yaklaşım önermiştir ve $\hat{\Phi}_A$ ile gösterdikleri $Var(\hat{\beta})$ tahminini iki düzeltmeyi de birleştirerek hesaplamışlardır. Kackar ve Harville, σ etrafında Taylor serisi açılımını kullanmışlardır. $P_i = X' \frac{\partial \Sigma^{-1}}{\partial \sigma_i} X$; $Q_{ij} = X' \frac{\partial \Sigma^{-1}}{\partial \sigma_i} \Sigma \frac{\partial \Sigma^{-1}}{\partial \sigma_j} X$ ve W_{ij} ,

$W = V[\hat{\sigma}]$ 'nin (i, j) . elemanı olmak üzere, Λ 'yı aşağıdaki gibi ifade etmişlerdir;

$$\Lambda \sim \Phi \left\{ \sum_{i=1}^r \sum_{j=1}^r W_{ij} (Q_{ij} - P_i \Phi P_j) \right\} \Phi \quad (4.17)$$

Kenward ve Roger, $\hat{\Phi} = \Phi(\hat{\sigma})$ tahmin edicisinin $\Phi(\sigma)$ için yanlılığını dikkate alıp, σ etrafında Taylor serisi açılımını kullanarak, $\hat{\Phi} = \Phi(\hat{\sigma})$ 'yı

$$\hat{\Phi} \sim \Phi + \sum_{i=1}^r (\hat{\sigma}_i - \sigma_i) \frac{\partial \Phi}{\partial \sigma_i} + \frac{1}{2} \sum_{i=1}^r \sum_{j=1}^r (\hat{\sigma}_i - \sigma_i) (\hat{\sigma}_j - \sigma_j) \frac{\partial^2 \Phi}{\partial \sigma_i \partial \sigma_j} \quad \text{olarak ifade etmişlerdir.}$$

Buradan, $\Sigma = \Sigma(\hat{\sigma})$ olmak üzere,

$$R_{ij} = X' \Sigma^{-1} \frac{\partial^2 \Sigma}{\partial \sigma_i \partial \sigma_j} \quad \text{ve} \quad \frac{\partial^2 \Phi}{\partial \sigma_i \partial \sigma_j} = \Phi (P_i \Phi P_j + P_j \Phi P_i - Q_{ij} - Q_{ji} + R_{ij}) \Phi$$

$$E[\hat{\Phi}] \sim \Phi + \frac{1}{2} \sum_{i=1}^r \sum_{j=1}^r W_{ij} \frac{\partial^2 \Phi}{\partial \sigma_i \partial \sigma_j} \quad (4.18)$$

Böylece, Kenward ve Roger küçük örneklem için düzeltilmiş $\hat{\beta}$ 'nin varyans-kovaryans matrisini Eş. 4.17 ve 4.18'den aşağıdaki gibi elde etmişlerdir;

$$\hat{\Phi}_A = \hat{\Phi} + 2\hat{\Phi} \left\{ \sum_{i=1}^r \sum_{j=1}^r W_{ij} \left(Q_{ij} - P_i \hat{\Phi} P_j - \frac{1}{4} R_{ij} \right) \right\} \hat{\Phi} \quad (4.19)$$

Kovaryans yapısının doğrusal bir formda tanımlanmış önemli bir sınıfı aşağıdaki gibi gösterilir;

$$\Sigma = \sum_{j=1}^r \sigma_j G_j$$

Bu sınıf karışık etkili modellerden ortaya çıkan yapıları da içerir. Böyle kovaryans yapılarında tüm (i, j) çiftleri için $\frac{\partial^2 \Sigma}{\partial \sigma_i \partial \sigma_j} = 0$ ve $R_{ij} = 0$ olur. Böylece Eş. 4.19 aşağıdaki ifadeye dönüşür;

$$\hat{\Phi}_A = \hat{\Phi} - \sum_{i=1}^r \sum_{j=1}^r W_{ij} \frac{\partial^2 \Phi}{\partial \sigma_i \partial \sigma_j} \quad (4.20)$$

$$= \hat{\Phi} + 2\hat{\Lambda} \quad (4.21)$$

W , $\hat{\sigma}$ için yaklaşık varyans kovaryans matrisi olmak üzere gözlenen veya beklenen bilgi matrisinin tersinden elde edilebilir. Kenward ve Roger yaklaşımının teorik genişletilmiş elde edilişi ve farklı bir versiyonu ile ilgili geniş bir açıklama Alnosaier (2007) tezinde bulunabilir. Bu yaklaşım Bootstrap-t yöntemiyle aralık oluştururken kullanılmıştır ve Bootstrap-t-Kenward Roger olarak ifade edilmiştir (BS-t-KR). Ayrıca yapay-skor güven aralığı oluşturulurken de bu yaklaşımdan yararlanılmıştır.

4.7. DKE Modelinde Sabit Etkiler için Analitik Yaklaşımlar

Sabit etki durumunda, β_{DGEKK} içindeki tüm tahmin ediciler Y 'nin normal dağılıma sahip doğrusal kombinasyonlarıdır. Bu nedenle, DEİDYTE, $t(\hat{\theta}) = k' \beta_{DGEKK}$ normal dağılıma sahip olacaktır. DEİDYTE, yansız olduğunda varyansı, σ_t^2 HKO'suna eşit olur. σ_t^2 'nin bilindiği durumda, test istatistiği $[t(\hat{\theta}) - T] / \sigma_t$ standart normal dağılıma sahip olur. $T = k' \beta$ için $1 - \alpha$ güven aralığı $P(-z_{\alpha/2} < [t(\hat{\theta}) - T] / \sigma_t < z_{1-\alpha/2}) = 1 - \alpha$ eşitliğinden elde edilir. Yani T için aralık $[t(\hat{\theta}) - z_{1-\alpha/2} \sigma_t, t(\hat{\theta}) + z_{1-\alpha/2} \sigma_t]$ olarak ifade edilir.

σ_t^2 'nin bilindiği durumda, test istatistiği $[t(\hat{\theta}) - T] / \sigma_t$ pivotudur. Pivot, parametre(ler)in ve verinin bir fonksiyonu olmakla birlikte dağılımı parametreden bağımsızdır. σ_t^2 bilinmediğinde pivot $[t(\hat{\theta}) - T] / s_t$, s_t^2 'nin σ_t^2 'ye yaklaşımı kullanılır. $[t(\hat{\theta}) - T] / s_t$

yapay-pivotunun dağılımı genellikle Student-t dağılımına yaklaşıır. Verilen bir t istatistiğinin serbestlik derecesi için, $1-\alpha$ güven aralığı aşağıdaki gibi yazılır;

$$[t(\hat{\theta}) - t_{1-\alpha/2} s_t, t(\hat{\theta}) + t_{1-\alpha/2} s_t] . \quad (4.22)$$

Bu yaklaşım altında güven alt ve üst sınırlarını hesaplamak için bir seçenek, s_t^2 'nin Naive yaklaşımını alıp N toplam gözlem sayısı ve p toplam sabit etki parametre sayısı olmak üzere $N - p$ 'yi t dağılımının serbestlik derecesi olarak kullanmaktır (Demidenko, 2004: 131). Naive yaklaşımın eksi yanlılığı veya serbestlik derecesinin yanlış olması sonucu yapay-pivot dağılımı t dağılımından daha ağır kuyruklara sahip olabilir (Harville ve Carriquiry, 1992). Aralık kapsama oranı belirlenen nominal seviyeden daha az olabilir, bu durum McLean ve Sanders (1988) simulasyon çalışmasında gösterilmiştir. Giesbrecht ve Burns (1985), $N - p$ serbestlik derecesini kullanmak yerine Satterwaite'ın (1941, 1946) serbestlik derecesi yaklaşım metodunu adapte etmiştir. Kenward ve Roger (1997) sabit etkilerin doğrusal kombinasyonlarını test etmek için, Prasad-Rao (1991) HKO yaklaşımına dayalı F test istatistiğini önerip bu F istatistiğinin paydasının serbestlik derecesi için de ayrı bir yaklaşım önermişlerdir. Bu yeni önerilen serbestlik derecesi aslında Giesbrecht ve Burns (1985) tarafından kullanılan Satterthwaite (1941,1946)'ın adapte edilen serbestlik derecesi yaklaşımı ile aynıdır. $T = k' \beta$ doğrusal kombinasyonu için Kenward-Roger F istatistiğinin kare kökü yapay-pivot (s_t^2 Prasad-Rao yaklaşımından elde edilmek üzere) $[t(\hat{\theta}) - T] / s_t$ ile aynıdır. Bu yapay-pivotun dağılımı t dağılımına yaklaşıır ve Eş. 4.22'deki formülde güven limitleri oluşturmak için kullanılabilir. Kenward-Roger metodu, serbestlik dereceleri ve Prasad-Rao yaklaşımı DDFM=KR seçeneği ile SAS PROC MIXED'de hesaplanabilir. Bu bölümde σ_t^2 bilinmediği durum için anlatılan metodların doğruluğu seçilen HKO yaklaşımının doğruluğuna ve yapay-pivot dağılımının student-t dağılımına ne kadar benzediğine bağlıdır.

5. DKE MODELLERİNDE SAYISAL ALGORİTMALAR VE İSTATİSTİKSEL YAZILIMLAR

EÇO ve KEÇO tahminlerinin hesaplanması DKE modellerinin analizindeki en zorlu aşamadır. Normal dağılımlı DKE modellerinde olabilirlik ya da kısıtlı olabilirlik fonksiyonları kapalı-formda ifade edilebilir. Bu fonksiyonlar Newton-Raphson, Fisher–Scoring, beklenti ve maksimizasyon (BM) (expectation and maximization-EM) gibi sayısal yöntemlerle maksimize edilir. En doğru algoritma DKE modelinin özel yapısına uygun olarak seçilmelidir.

5.1. BM Algoritması için Bazı Tanımlamalar ve Gösterimler

Klasik Üstel Aile (Regular Exponential Family)

$s(y)$: $k \times 1$ boyutlu yeterli istatistik vektörü,

$c(\theta)$: $k \times 1$ boyutlu parametre vektörü,

$\theta_{d \times 1} \in \Omega$, d -boyutlu bir konveks parametre uzayında ve $f(y; \theta)$ olasılık yoğunluk fonksiyonu; $b(y)$ ve $a(\theta)$ scalar değerler olmak üzere bir üstel ailenin bileşik olasılık yoğunluk fonksiyonu aşağıdaki formda yazılır;

$$f(y; \theta) = b(y) \exp \{ c(\theta)^T s(y) \} / a(\theta).$$

$c(\theta)$ doğal veya kanonik parametre vektörü olarak adlandırılır. $k = d$ ve $c(\theta)$ 'nin Jakobyeni, $\frac{\partial c}{\partial \theta}$, $k \times k$ boyutlu tam ranklı bir matrisse, $f(y; \theta)$ 'nin üstel ailede olduğu bilinir. Bu durumda,

$$f(y; \theta) = b(y) \exp \{ \theta^T s(y) \} / a(\theta).$$

Tam-veri seti skor vektörü

$$S(\theta; y) = \frac{\partial l(\theta; y)}{\partial \theta}$$

Gözlemlenmiş-veri seti vektörü

$$S_{göz}(\theta; y_{göz}) = \frac{\partial l(\theta; y_{göz})}{\partial \theta}$$

Ayrıca beklenen değer operatörünün ve türev operatörünün değişiminin olabileceği şartlar için aşağıdaki ifade verilebilir;

$$S_{göz}(\theta; y_{göz}) = E_{\theta} [S(\theta; y) | y_{göz}]$$

Tam-veri seti Bilgi Matrisi (Fisher Information Matrix)

$$I(\theta; y) = \frac{-\partial^2 l(\theta; y)}{\partial \theta \partial \theta^T}$$

Tam-veri seti Beklenen Bilgi Matrisi

$$I(\theta; y) = E_{\theta} [I(\theta; y)]$$

Gözlemlenmiş-veri seti Bilgi Matrisi

$$I_{göz}(\theta; y_{göz}) = \frac{-\partial^2 l(\theta; y_{göz})}{\partial \theta \partial \theta^T}$$

Gözlemlenmiş-veri seti beklenen Bilgi Matrisi

$$I_{göz}(\theta; y_{göz}) = E_{\theta} [I(\theta; y_{göz})]$$

Gözlemlenmiş-veri seti beklenen koşullu Bilgi Matrisi

$$I_{göz}(\theta; y_{göz}) = E_{\theta} [I(\theta; y) | y_{göz}]$$

Kayıp Bilgi Prensibi

$l_{göz}(\theta; y_{göz}) = l(\theta; y) - \log f_2(y_{kay} | y_{göz}; \theta)$ olduğu hatırlanırsa θ 'ya göre türev alınırsa aşağıdaki ifade bulunur,

$$I_{göz}(\theta; y_{göz}) = I(\theta; y) + \frac{\partial^2 \log f_2(y_{kay} | y_{göz}; \theta)}{\partial \theta \partial \theta^T}$$

$y | y_{göz}$ üzerinde koşullu beklenen değer alınır:

$$I_{göz}(\theta; y_{göz}) = I_{Koşullu}(\theta; y_{göz}) - I_{kay}(\theta; y_{göz}) \text{ elde edilir.}$$

Bu ifadeden kayıp bilgi matrisi aşağıdaki gibi elde edilir:

$$I_{göz}(\theta; y_{göz}) = -E_{\theta} \left\{ \frac{\partial^2 f_2(y_{kay} | y_{göz}; \theta)}{\partial \theta \partial \theta^T} | y_{göz} \right\}$$

Kayıp bilgi prensibi aşağıdaki eşitliğin olduğunu gösterir:

$$Gözlemlenmiş Bilgi = Tam Bilgi - Kayıp Bilgi$$

5.2. Beklenti ve Maksimizasyon Algoritması

EÇO ve KEÇO tahminleri profil log-olabilirlik fonksiyonunun maksimize edilmesi ile elde edilir. Bu doğrusal olmayan bir optimizasyon problemidir. Newton-Raphson gibi standart sayısal prosedürler olsa da bunun gibi prosedürlerin sıkıntısı genellikle başlangıç değerlerine oldukça duyarlı algoritmalar olmasıdır. Başlangıç değerleri kötü olduğunda yakınsama problemleri ortaya çıkar. Bu nedenle bu tezde en çok olabilirlik fonksiyonunun maksimize edilmesi probleminde daha iyi bir sayısal algoritma olan beklenti ve maksimizasyon sayısal yöntemi ele alınmıştır. Beklenti ve Maksimizasyon algoritması diğer yöntemlere göre yavaş olsa da başlangıç değerlerine karşı daha sağlamdır ve başlangıç değerlerinden fazla etkilenmez. Kötü başlangıç değerleri algoritmaya verildiğinde yakınsaması yavaşlar. İki yöntemin avantajları bir algortmada birleştirilebilir. Örneğin kötü başlangıç değerleri seçilirse ilk olarak BM algoritması ile başlanır belli adımdan sonra Newton-Raphson algoritmasına geçilebilir. Bu tezde BM algoritması bazı DKE modelleri için tanıtılacaktır.

Beklenti ve maksimizasyon algoritması bazı rastgele değişkenlerin gözlemlenemediği eksik olduğu veya tamamlanamadığı durumlarda en çok olabilirlik yöntemiyle parametre tahmini için kullanılan oldukça genel, yinelemeli bir algortmadır. BM algoritması veri setinde bazı eksiklikler olduğunda parametre tahmini için sezgisel bir fikri formülleştirmiştir;

1. Eksik değerler tahmin değerleriyle yer değiştirilir,
2. Parametreler tahmin edilir,

Yukarıdaki iki adım aşağıdakiler uygulanarak yakınsama sağlanana kadar tekrarlanır;

1. adımda tahmini parametre değerlerini gerçek değerler gibi kullanılarak,
2. adımda tahmini değerleri gözlenmiş değerler olarak alınarak, yakınsama gerçekleşene kadar döngü devam ettirilir.

Genel teorik çerçevesi Orchard ve Woodbury (1972)'nin kayıp bilgi prensibinde verilmeden önce de BM algoritması kullanılmaktaydı. Beklenti ve maksimizasyon terimi ilk olarak algoritmanın farklı durumlarda davranışı ve oldukça fazla sayıda uygulamasıyla birlikte Dempster, Laird ve Rubin (1977) tarafından istatistik literatürüne kazandırılmıştır.

p -boyutlu parametre vektörüne $\theta \in \Theta \subseteq \mathcal{R}^p$ sahip y rastgele vektörü için bileşik olasılık yoğunluk fonksiyonu $f(y; \theta)$ olarak gösterilsin. y rastgele vektörü için tam bir veri seti y gözlemlendiyse, y 'nin dağılımına bağlı olarak θ için en çok olabilirlik tahmin edicilerini elde etmek en önemli amaçtır. y için log-olabilirlik fonksiyonu aşağıdaki gibidir ve maksimize edilmesi gerekir;

$$\log L(\theta; y) = l(\theta; y) = \log f(y; \theta)$$

Kayıp veri olması durumunda tam veri setinin sadece bir fonksiyonu gözlemlenmiş olur. $y_{göz}$ gözlemlenmiş ama “tamamlanmamış” veriyi ve y_{kay} gözlemlenmemiş kayıp veriyi temsil etmek üzere bu fonksiyon $(y_{göz}, y_{kay})$ olarak tanımlanır. Tanımın sadeliği bakımından, kayıp verinin rastgele olduğu varsayılсын f_1 , $y_{göz}$ 'ün bileşik olasılık yoğunluk fonksiyonu olmak üzere,

$$\begin{aligned} f(y; \theta) &= f(y_{kay}, y_{göz}; \theta) \\ &= f_1(y_{göz}; \theta) f_2(y_{kay} | y_{göz}; \theta) \end{aligned}$$

yazılabilir. Böylece $l_{göz}(\theta; y_{göz})$ gözlemlenmiş-verinin log-olabilirliği olmak üzere aşağıdaki eşitlik yazılır;

$$l_{göz}(\theta; y_{göz}) = l(\theta; y) - \log f_2(y_{kay} | y_{göz}; \theta).$$

$l_{göz}$ 'ün maksimize edilmesinin zor ama tam-veri setini maksimize etmenin kolay olduğu durumlarda BM algoritması faydalıdır. y tam-veri seti olarak gözlemlenemediği için,

$l(\theta; y)$ tam olarak değerlendirilemez dolayısıyla maksimize edilemez. BM algoritması, $y_{göz}$ verildiğinde koşullu bekleneniyle yer değiştirerek $l(\theta; y)$ 'yı yinelemeli olarak maksimize etmeyi amaçlar. Bu beklenen θ 'nın o yineleme için geçerli değerindeki tam-veri setinin dağılımına göre hesaplanır. Yani, $\theta^{(0)}$, θ için verilen ilk değerse, ilk yinelemede $Q(\theta, \theta^{(0)}) = E_{\theta^{(0)}} [l(\theta; y) | y_{göz}]$ 'nın hesaplanmasını gerektirir. $Q(\theta, \theta^{(0)})$, θ 'ya göre maksimize edilmiş olur ve aşağıdaki eşitsizliği sağlayacak şekilde $\theta^{(1)}$ 'i bulmak için kullanılır; $Q(\theta^{(1)}, \theta^{(0)}) \geq Q(\theta, \theta^{(0)})$; $\theta \in \Theta$. Sonuç olarak, BM algoritması, Tahmin adımı (T-adım) ve onu takip eden Maksimizasyon-adımını (M-adım) içerir ve aşağıdaki gibi tanımlanır.

Tahmin-Adımı (T-Adım): $Q(\theta, \theta^{(t)})$ hesaplanır;

$$Q(\theta, \theta^{(t)}) = E_{\theta^{(t)}} [l(\theta; y) | y_{göz}]$$

Maksimizasyon-Adımı (M-adım): Θ içinde $Q(\theta^{(t+1)}, \theta^{(t)}) \geq Q(\theta, \theta^{(t)})$ olacak şekilde $\theta^{(t+1)}$ bulunur. $L(\theta^{(t+1)}) - L(\theta^{(t)})$ arasındaki fark önceden belirlenen çok küçük bir δ 'dan küçük olana kadar yinelemeye devam edilir. Log-olabilirlik fonksiyonun θ için bulunan bir yeterli istatistikte doğrusal olduğu gösterilebildiğinde bu iki adımın hesaplanması oldukça basitleşir. Tam-veri setinin dağılımının üstel aile formatında olması bu durumu olası hale getirir. Bu durumda T-adımı, gözlemlenmiş veri verildiğinde tam verinin yeterli istatistiğinin beklenen değerini hesaplamaya indirgenir. Tam-veri setinin dağılımı üstel aileden olduğunda M-adımı da kolaylaşır. M-adımı, T-adımında hesaplanan log-olabilirlik fonksiyonunun beklenen değerini maksimize etmeyi içermektedir. Tam-veri setinin dağılımı üstel olduğunda, diğer yinelemeye devam etmek için tüm log-olabilirlik fonksiyonunun beklenen değerinin maksimize edilmesi yerine T-adımında hesaplanan yeterli istatistiğin beklenen değeri maksimize edilmeye çalışılır. Böylece M-adımı ve T-adımı oldukça kolaylaşmış olur. Üstel dağılıma ait bazı örnekler bu adımlarla birlikte aşağıda incelenecektir.

Karmaşık en çok olabilirlik fonksiyonu hesaplamaları için sayısal algoritmalar içinde BM algoritması Newton-Raphson algoritması gibi algoritmalara göre daha üstün özelliklere sahiptir. En önemli özelliklerinden ilki, tam-veri setinin hesaplamalarına dayalı olduğu için uygulaması genellikle daha kolaydır. T-adımının her yinelemesi sadece tam-veri setinin

koşullu dağılımının beklenen değerini bulmayı içerir üstel aile formatı algoritmayı daha kolaylaştırır. M-adımının her yinelemesi tam-veri seti için uygun olan en çok olabilirlik tahmin edicilerinin basit kapalı formlarını içerir. Ayrıca sayısal olarak dengeli bir algoritmadır: her yineleme en çok olabilirlik fonksiyonunun $l_{göz}(\theta; y_{göz})$ değerini artırmayı gerektirir. $l_{göz}(\theta; y_{göz})$ eğer sınırlıysa, $l_{göz}(\theta^{(t)}; y_{göz})$ dizisi bir değere yakınsayacaktır. $\theta^{(t)}$ dizisi yakınsarsa, bu $l_{göz}(\theta; y_{göz})$ 'nun ya bir yerel maksimum noktasına ya da bir eyer noktasına (saddle point) sahip olacaktır ve $l_{göz}(\theta; y_{göz})$ tek modlu bir fonksiyonsa en çok olabilirlik tahmin edicisi tek olacaktır.

Örnek. $y = (y_1, \dots, y_n)^T$ tam-veri seti vektörünün normal dağılımlı $N(\mu, \sigma^2)$ bir rastgele örneklem olduğu varsayalım.

$$\begin{aligned} f(y; \mu, \sigma^2) &= \left(\frac{1}{2\pi\sigma^2} \right)^{\frac{n}{2}} \exp \left\{ -\frac{1}{2} \sum_{i=1}^n \frac{(y_i - \mu)^2}{\sigma^2} \right\} \\ &= \left(\frac{1}{2\pi\sigma^2} \right)^{\frac{n}{2}} \exp \left\{ -1/2\sigma^2 \left(\sum y_i^2 - 2\mu \sum y_i + n\mu^2 \right) \right\} \end{aligned}$$

$\left(\sum y_i, \sum y_i^2 \right)$ 'nin normal dağılımın parametreleri $\theta = (\mu, \sigma^2)$ için yeterli istatistikler

olduğu bilinmektedir. Tam-veri setinin en çok olabilirlik fonksiyonu;

$$\begin{aligned} l(\mu, \sigma^2; y) &= -\frac{n}{2} \log(\sigma^2) - \frac{1}{2} \sum_{i=1}^n \frac{(y_i - \mu)^2}{\sigma^2} + \text{sabit} \\ &= -\frac{n}{2} \log(\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^n y_i^2 + \frac{\mu^2}{\sigma^2} \sum_{i=1}^n y_i - \frac{n\mu^2}{\sigma^2} + \text{sabit} \end{aligned}$$

Böylece, tam-veri setine dayalı en çok olabilirlik fonksiyonu tam-veri seti yeterli istatistiklerinde doğrusaldır. y_i vektörünün $N(\mu, \sigma^2)$ sahip aynı dağılımlı normal olduğu, $y_i, i=1, \dots, m$ gözlemlenmiş ve $y_i, i=m+1, \dots, n$ kayıp veri (rastgele) vektörü olduğu varsayalım. Gözlemlenmiş veri vektörü $y_{göz} = (y_1, \dots, y_m)^T$ olarak gösterilsin. Tam-veri seti y 'nin vektörü üstel aileden olduğu için, T-adımı aşağıdaki ifadelerin hesaplanmasını gerektirir;

$$E_{\theta} \left(\sum_{i=1}^n y_i \mid y_{göz} \right) \text{ ve } E_{\theta} \left(\sum_{i=1}^n y_i^2 \mid y_{göz} \right)$$

Yani, tüm-veri setinin en çok olabilirlik fonksiyonunun hesaplanması yerine yeterli istatistiklerin koşullu beklenen değerlerinin hesaplanması yeterli olacaktır (Little ve Rubin, 1987: 131). Böylece, T-adımın t . yinelenmesinde $\mu^{(t)}$ ve $\sigma^{2(t)}$, (μ, σ^2) için iterasyonun yeni tahminleri $E_{\mu^{(t)}, \sigma^{2(t)}}(y_i) = \mu^{(t)}$ ve $E_{\mu^{(t)}, \sigma^{2(t)}}(y_i^2) = \mu^{2(t)} + \sigma^{2(t)}$ olmak üzere aşağıdakiler hesaplanır;

$$\begin{aligned} s_1^{(t)} &= E_{\mu^{(t)}, \sigma^{2(t)}} \left(\sum_{i=1}^n y_i \mid y_{göz} \right) \\ &= \sum_{i=1}^m y_i + (n-m)\mu^{(t)} \end{aligned} \quad (5.1)$$

$$\begin{aligned} s_2^{(t)} &= E_{\mu^{(t)}, \sigma^{2(t)}} \left(\sum_{i=1}^n y_i^2 \mid y_{göz} \right) \\ &= \sum_{i=1}^m y_i^2 + (n-m) \left[\sigma^{2(t)} + \mu^{(t)} \right] \end{aligned} \quad (5.2)$$

M-adımında, (μ, σ^2) için tam-veri setinin en çok olabilirlik tahmin edicileri: $\hat{\mu} = \frac{\sum_{i=1}^n y_i}{n}$ ve

$$\hat{\sigma}^2 = \frac{\sum_{i=1}^n y_i^2}{n} - \left(\frac{\sum_{i=1}^n y_i}{n} \right)^2, \text{ dir.}$$

μ ve σ^2 'nin yeni yinelenmelerini elde edebilmek için yukarıdaki ifadelerin sağ tarafındaki tam-veri setinin yeterli istatistikler için T-adımında hesaplanan beklenen değerler yerine konularak M-adımı tanımlanacaktır. $y = (y_{m+1}, \dots, y_n)^T$ gözlemlenemeyen kısmı olduğu için tam-veri setinin yeterli istatistikleri doğrudan hesaplanamayacaktır. Bu durumda aşağıdaki ifadelerden yararlanılacaktır;

$$\mu^{(t+1)} = \frac{s_1^{(t)}}{n} \quad (5.3)$$

$$\sigma^{2(t+1)} = \frac{s_2^{(t)}}{n} - \mu^{(t+1)2} \quad (5.4)$$

Böylece, $\mu^{(0)}$ ve $\sigma^{2(0)}$ başlangıç değerleriyle başlayarak Eş. 5.1 ve 5.2 T-adımında hesaplanır. Bir sonraki adımdaki $\mu^{(1)}$ ve $\sigma^{2(1)}$ değerleri Eş. 5.3 ve 5.4 bir önceki

başlangıç değerleri yerine konularak M-adımında hesaplanır ve yakınsama gerçekleşene kadar yinelemelere devam edilir. Böylece, BM algoritması Eş. 5.1, 5.2, 5.3 ve 5.4 arasında yenilenir. Aslında (μ, σ^2) için en çok olabilirlik tahmin edicileri açıkça $\hat{\mu} = \sum_{i=1}^m y_i / m$ ve $\hat{\sigma}^2 = \sum_{i=1}^m y_i^2 / m - \hat{\mu}^2$ ile ifade edilebildiğinden, BM algoritmasının kullanımı gerekli değildir. Bu örnek sadece BM algoritmasını basitçe tanıtmak için verilmiştir.

Örnek. Varyans Bileşenleri Tahmini (Little ve Rubin, 1987: 150)

İneklerin yapay döllemesiyle ilgili bir çalışmada, rastgele seçilen boğalardan alınan sperm örnekleri, gebelik üretme kabiliyetlerine göre test edilmiştir. Alınan sperm örnekleri sayısı boğadan boğaya değişmiştir ve yanıt değişkeni her örnekten elde edilen gebe bırakma yüzdesidir. Bu çalışmada boğa etkisi rasgele etki olarak varsayılmıştır. Veri seti Çizelge 5.1’de verilmiştir.

Çizelge 5.1. Yapay dölleme veri seti

Boğa(<i>i</i>)	Gebe bırakma yüzdesi	m_i
1	46, 31, 37, 62, 30	5
2	70, 59	2
3	52, 44, 57, 40, 67, 64, 70	7
4	47, 21, 70, 46, 14	5
5	42, 64, 50, 69, 77, 81, 87	7
6	35, 68, 59, 38, 57, 76, 57, 29, 60	9
Toplam		35

Böyle bir veri setini analiz etmek için tek faktörlü rastgele etki modeli uygun olacaktır;

$$y_{ij} = a_i + e_{ij}, \quad j = 1, \dots, m_i, \quad i = 1, \dots, n,$$

Boğa etkisinin normal dağılımlı, bağımsız; $a_i \sim N(\mu, \sigma_a^2)$, boğa etkisi-içi hata terimleri $e_{ij} \sim N(0, \sigma^2)$ normal dağılımlı, bağımsız olduğu varsayılır. Ayrıca a_i ve e_{ij} ’nin birbirinden bağımsız rastgele değişkenler olduğu varsayılın. Bu model varyans bileşenleri modelidir, yeniden parameterize edilerek DKE modeline uygun hale getirilir. Dempster, Laird ve Rubin (1977) varyans bileşenlerinin EÇO tahminlerini BM algoritmasına dayalı olarak elde etmişlerdir.

Çizelge 5.2. Tek faktörlü rastgele etki modeli için varyans analizi tablosu

Değişimin Kaynağı	Serb.Der.	Kareler Toplamı	Kareler Ortalaması	F değeri	Beklenen Kareler Ortalaması
Boğalar Arası	5	3222,059	664,41	2,68	$\sigma^2 + 5,67\sigma_a^2$ $(5\sigma^2 + 28,34 \sigma_a^2)/5$
Boğalar İçi (Hata)	29	7200,341	248,29		σ^2
Toplam	34	10522,400			

Hata teriminin varyansının ANOVA tahmini Çizelge 5.2’de $\hat{\sigma}^2 = 248,29$ olarak elde edilmiştir. Boğa rastgele etkisinin varyansının ANOVA tahmini $(664,41 - 248,29)/5,67 = \hat{\sigma}_a^2 = 73,39$ olarak elde edilir.

En çok olabilirlik tahmin edicileri $\theta = (\mu, \sigma_a^2, \sigma^2)$ elde etmede BM algoritması kullanmak için öncelikle $y^* = (y, a)^T$ ’nin bileşik olasılık yoğunluk fonksiyonu y^* -tam-veri seti varsayılarak bulunmalıdır. Bu bileşik olasılık yoğunluk fonksiyonu iki faktörün çarpımı olarak yazılabilir: ilk parça a_i verildiğinde y_{ij} ’nin bileşik olasılık yoğunluk fonksiyonu ve ikinci parça da a_i ’nin bileşik olasılık yoğunluk fonksiyonu olarak seçilmelidir.

$$f(y^*; \theta) = f_1(y | a; \theta) f_2(a; \theta)$$

$$= \prod_i \prod_j \left\{ \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2\sigma^2}(y_{ij} - a_i)^2} \right\} \prod_i \left\{ \frac{1}{\sqrt{2\pi}\sigma_a} e^{-\frac{1}{2\sigma_a^2}(a_i - \mu)^2} \right\}$$

Böylece, aşağıda verilen tam-veri seti için yeterli istatistiklerde en çok olabilirlik fonksiyonu doğrusal olacaktır.

$$T_1 = \sum a_i$$

$$T_2 = \sum a_i^2$$

$$T_3 = \sum_i^n \sum_j^{m_i} (y_{ij} - a_i)^2 = \sum_i^n \sum_j (y_{ij} - \bar{y}_i)^2 + \sum_i^n m_i (\bar{y}_i - a_i)^2$$

Tam-veri olması y ve a ’nın tamamen ulaşılabilir olduğunu varsayar. Sadece y gözlendiği için, $y_{göz}^* = y$ olarak yazılır.

BM algoritmasının T-adımı $y_{göz}^*$ verildiğinde T_1, T_2, T_3 'ün beklenen değerlerinin hesaplanmasını gerektirir. Bu beklenen değerleri hesaplamak için y verildiğinde a 'nın koşullu dağılımına ihtiyaç vardır. $y^* = (y, a)^T$ 'nin bileşik olasılık yoğunluk fonksiyonu, $N = \sum_i^n m_i$ boyutlu birlerin vektörü j_N , $\mu^* = (\mu, \mu_a)^T$, $\mu = \mu j_N$, $\mu_a = \mu j_n$ ve Σ^* , $(N+n) \times (N+n)$ boyutlu matris olmak üzere $(N+n)$ boyutlu çok değişkenli normal dağılıma $N(\mu^*, \Sigma^*)$ sahiptir. J_{m_i} $m_i \times m_i$ boyutlu birlerin matrisi olmak üzere $\Sigma_i = \sigma^2 I_{m_i} + \sigma_a^2 J_{m_i}$, $m_i \times m_i$ boyutlu bir matristir.

$$\Sigma = \begin{bmatrix} \Sigma_1 & & 0 \\ & \Sigma_2 & \\ 0 & & \ddots \\ & & & \Sigma_n \end{bmatrix}, \quad \Sigma_{12} = \sigma_a^2 \begin{bmatrix} j_{m_1} & & 0 \\ & j_{m_2} & \\ 0 & & \ddots \\ & & & j_{m_n} \end{bmatrix}$$

$$\Sigma^* = \begin{pmatrix} \Sigma & \Sigma_{12} \\ \Sigma_{12}^T & \sigma_a^2 I \end{pmatrix}$$

y 'nin bileşik olasılık yoğunluk fonksiyonunun kovaryans matrisi Σ ; aynı boğa için olan gözlemler arasındaki kovaryans σ_a^2 ve farklı boğalar arasındaki kovaryans 0, y_{ij} 'ler ortak normal dağılımlı ve genel ortalaması μ ve genel varyansı $\sigma^2 + \sigma_a^2$ parametrelerine sahiptir. Kovaryanslar aşağıda özetlenmiştir;

$$\begin{aligned} kov(y_{ij}, y_{i'j'}) &= kov(a_i + e_{ij}, a_{i'} + e_{i'j'}) \\ &= \sigma^2 + \sigma_a^2 \quad \text{eğer } i = i', j = j', \\ &= \sigma_a^2 \quad \text{eğer } i = i', j \neq j', \\ &= 0 \quad \text{eğer } i \neq i'. \end{aligned}$$

$kov(y_{ij}, a_i) = \sigma_a^2$ eğer $i = i'$; $kov(y_{ij}, a_i) = 0$ eğer $i \neq i'$ olduğundan Σ_{12} ; y ve a 'nın kovaryansıdır. y verildiğinde a 'nın koşullu dağılımını bulmak için Σ 'nın tersine ihtiyaç

vardır ve $\Sigma_i^{-1} = \frac{1}{\sigma^2} \left[I_{m_i} - \frac{\sigma_a^2}{\sigma^2 + m_i \sigma_a^2} J_{m_i} \right]$ olmak üzere aşağıdaki gibidir;

$$\Sigma^{-1} = \begin{bmatrix} \Sigma_1^{-1} & & 0 \\ & \Sigma_2^{-1} & \\ 0 & & \ddots \\ & & & \Sigma_n^{-1} \end{bmatrix}$$

Çok değişkenli normal teoriden bilindiği üzere, y verildiğinde a 'nın koşullu dağılımı $\alpha = \mu_a + \Sigma'_{12}\Sigma^{-1}(y - \mu)$ ve $A = \sigma_a^2 I - \Sigma'_{12}\Sigma^{-1}\Sigma_{12}$ olmak üzere $a \sim N(\alpha, A)$ 'dır.

$$w_i = \frac{\sigma_a^2}{\sigma^2 + m_i \sigma_a^2}, \quad \bar{y}_{i.} = \left(\sum_{j=1}^{m_i} y_{ij} \right) / m_i \quad \text{ve} \quad v_i = w_i \sigma_a^2 \quad \text{olmak üzere}$$

$a_i | y \stackrel{a.d.b.}{\sim} N(w_i \mu + (1 - w_i) \bar{y}_{i.}, v_i)$ olur. y (veya $y_{göz}^*$) verildiğinde T_1, T_2, T_3 'ün beklenen değerlerinin hesaplanması için bu koşullu dağılıma ihtiyaç olacaktır. T -adımının t . yinelemesi aşağıdaki gibi gerçekleşir;

$$T_1^{(t)} = \sum \left[w_i^{(t)} \mu^{(t)} + (1 - w_i^{(t)}) \bar{y}_{i.} \right]$$

$$T_2^{(t)} = \sum \left[w_i^{(t)} \mu^{(t)} + (1 - w_i^{(t)}) \bar{y}_{i.} \right]^2 + \sum v_i^{(t)}$$

$$T_3^{(t)} = \sum_i \sum_j \left(y_{ij} - \bar{y}_{i.} \right)^2 + \sum_i m_i \left[w_i^{(t)2} (\mu^{(t)} - \bar{y}_{i.})^2 + v_i^{(t)} \right]$$

Böylece tam-veri seti en çok olabilirlik tahmin edicileri aşağıdaki gibi olacaktır;

$$\hat{\mu} = \frac{T_1}{n}$$

$$\hat{\sigma}_a^2 = \frac{T_2}{n} - \hat{\mu}^2,$$

$$\hat{\sigma}^2 = \frac{T_3}{N}$$

T -adımındaki ifadeler kullanılarak M -adımında yeterli istatistiklerin beklenen değerleri yerine konularak en çok olabilirlik tahmin edicileri elde edilir;

$$\hat{\mu}^{(t+1)} = \frac{T_1^{(t)}}{n}$$

$$\hat{\sigma}_a^{2(t+1)} = \frac{T_2^{(t)}}{n} - \hat{\mu}^{(t+1)2}$$

$$\hat{\sigma}^{2(t+1)} = \frac{T_3^{(t+1)}}{N}$$

Bu iki denklem seti arasında yinelemeler yapılarak BM algoritması işletilir. Başlangıç değerleri olarak parametrelerin ANOVA tahminlerine yakın değerler kullanılır: $\mu^{(0)} = 54$, $\sigma_a^{2(0)} = 54,827$, $\sigma^2(0) = 70$. 30 yinelemenin sonunda en çok olabilirlik tahminleri $\hat{\mu} = 53,3184$, $\hat{\sigma}_a^2 = 54,827$ ve $\hat{\sigma}^2 = 249,22$ olarak elde edilir.

5.3. Karışık Etki Modelleri için BM Algoritması

BM algoritması karışık etkili modeller için kullanılırken rasgele etkilere kayıp gözlem olarak bakılır. $y_{N \times 1}$ gözlemlenmiş rastgele vektör $X_{N \times p}$ ve $Z_{N \times q_i}$ bilinen değerli tasarım matrisleri, bilinmeyen parametre vektörü $\beta_{p \times 1}$ ve b_i gözlemlenemeyen rastgele etkilerin $q_i \times 1$ boyutlu vektörleri olmak üzere BM algoritmasının daha açık ifade edilebilmesi için varyans bileşenleri modelinin aşağıdaki formu verilmiştir bu form DKE modelleri için de kullanılabilir (Searle ve diğerleri, 1992: 298);

$$y = X\beta + \sum_{i=1}^r Z_i b_i + e$$

$$e_{N \times 1} \sim N(0, \sigma_0^2 I_N), \quad b_i \sim N_{q_i}(0, \sigma_i^2 \Sigma_i) \quad i = 1, \dots, r$$

Hata terimlerinin N boyutlu çok değişkenli normal dağılımlı olduğu, rastgele etkilerin her birinin de q_i boyutlu çok değişkenli normal dağılıma sahip olduğu ve hata terimlerinden ve birbirlerinden bağımsız olduğu varsayılın. Tam-veri seti vektörü $(y, b_1, \dots, b_r)$ ve

y gözlemlenmiş (tam olmayan) veri vektörü olarak gösterilsin. $q = \sum_{i=0}^r q_i$ ve $q_0 = N$ olmak

üzere y için kovaryans matrisi aşağıdaki gibi ifade edilir;

$$\text{kov}(y, b_j') = \text{kov}(X\beta + \sum Z_i b_i, b_j') = Z_j \text{kov}(b_j, b_j') = \sigma_j^2 Z_j$$

$$V = \text{var}(y) = \sum_{i=0}^r Z_i Z_i^T \sigma_i^2 = \sum_{i=1}^r Z_i Z_i^T \sigma_i^2 + \sigma_0^2 I_N$$

$$\mu_{q \times 1} = \begin{bmatrix} X\beta \\ 0 \\ \vdots \\ 0 \end{bmatrix} \text{ ve } \sum_{q \times q} = \begin{bmatrix} V & {}_r\{\sigma_i^2 Z_i\}_{i=1}^r \\ {}_c\{\sigma_i^2 Z_i^T\}_{i=1}^r & {}_d\{\sigma_i^2 I_{q_i}\}_{i=1}^r \end{bmatrix} \text{ olmak üzere } y \text{ ve } b_1, \dots, b_r \text{ için bileşik}$$

dağılım da q -boyutlu çok değişkenli normal dağılıma sahip olacaktır, $N(\mu, \Sigma)$. Böylece

$$w = \left[(y - X\beta)^T, b_1^T, \dots, b_r^T \right] \text{ olmak üzere } f(y, b_1, \dots, b_r) = (2\pi)^{-\frac{1}{2}q} |\Sigma|^{-\frac{1}{2}} \exp\left(-\frac{1}{2} w^T \Sigma^{-1} w\right),$$

$(y, b_1, \dots, b_r)$ 'nin olasılık yoğunluk fonksiyonu olarak yazılır.

$$\sum_{q \times q} = \begin{bmatrix} V & {}_r\{\sigma_i^2 Z_i\}_{i=1}^r \\ {}_c\{\sigma_i^2 Z_i^T\}_{i=1}^r & {}_d\{\sigma_i^2 I_{q_i}\}_{i=1}^r \end{bmatrix} \text{ parçalı bir matristir ve 4 parçanın da tersi mevcuttur. } r$$

alt indisi satır vektörü, c alt indisi sütun vektörü, d alt indisi diyagonal matrisler olarak yazılacağını gösterir. Bu matrisin determinant ve tersi V 'nin Schur bileşeni yoluyla bulunur (Searle ve diğerleri, 1992:445, 453). Yukarıda verilen bileşik oyfu $f(y, b_1, \dots, b_r)$ daha basitleştirmek için:

$$\begin{aligned} |\Sigma| &= \left| {}_d\{\sigma_i^2 I_{q_i}\}_{i=1}^r \right| \left| V - {}_r\{\sigma_i^2 Z_i\}_d \left\{ (\sigma_i^2)^{-1} I_{q_i} \right\}_c \left\{ \sigma_i^2 Z_i^T \right\} \right| \\ &= \prod_{i=1}^r (\sigma_i^2)^{q_i} \left| V - \sum_i \sigma_i^2 Z_i Z_i^T \right| \\ &= \prod_{i=1}^r (\sigma_i^2)^{q_i} \left| \sigma_0^2 I_N \right| \\ &= \prod_{i=0}^r (\sigma_i^2)^{q_i} \end{aligned}$$

$$\Sigma^{-1} = \begin{bmatrix} 0 & 0 \\ 0 & {}_d\{\sigma_i^{-2} I_{q_i}\}_{i=1}^r \end{bmatrix} + \begin{bmatrix} I & \\ {}_c\{Z_i^T\} \end{bmatrix} \sigma_0^{-2} I_N \begin{bmatrix} I & {}_r\{Z_i\} \end{bmatrix} \text{ (Searle ve diğerleri, 1992: 450).}$$

Tam-veri seti log-olabilirlik fonksiyonu $b_0 = y - X\beta - \sum_{i=1}^r Z_i b_i = e$ olmak üzere aşağıda verildiği gibidir;

$$l = -\frac{1}{2} q \log(2\pi) - \frac{1}{2} \sum_{i=0}^r q_i \log \sigma_i^2 - \frac{1}{2} \sum_{i=0}^r \frac{b_i^T b_i}{\sigma_i^2}.$$

Buradan yeterli istatistikler: $b_i^T b_i$, $i=1, \dots, r$ ve $y - \sum_{i=1}^r Z_i b_i$ en çok olabilirlik tahmin edicileri;

$$\hat{\sigma}_i^2 = \frac{b_i^T b_i}{q_i} \quad i = 0, 1, \dots, r$$

$$\hat{\beta} = (X^T X)^{-1} X^T \left(y - \sum_{i=1}^r Z_i b_i \right).$$

5.3.1. İki varyans bileşenli karışık-etki modeli

$y = X\beta + Z_1 b_1 + e$ ve $e_{N \times 1} \sim N(0, \sigma_0^2 I_N)$; $b_1 \sim N(0, \sigma_1^2 Z_1)$ ve y 'nin kovaryans matrisi $V = Z_1 Z_1^T \sigma_1^2 + \sigma_0^2 I_n$ olarak yazılır.

$b_0 = y - X\beta - Z_1 b_1 = e$ olmak üzere tam-veri seti log-olabilirlik fonksiyonu aşağıdaki gibidir;

$$l = -\frac{1}{2} q \log(2\pi) - \frac{1}{2} \sum_{i=0}^1 q_i \log \sigma_i^2 - \frac{1}{2} \sum_{i=0}^1 \frac{b_i^T b_i}{\sigma_i^2}$$

En çok olabilirlik tahmin edicileri aşağıdaki gibi olacaktır;

$$\hat{\sigma}_i^2 = \frac{b_i^T b_i}{q_i} \quad i = 0, 1$$

$$\hat{\beta} = (X^T X)^{-1} X^T (y - Z_1 b_1)$$

Yeterli istatistiklerin, $b_i^T b_i, i=0, 1$ ve gözlemlenmiş veri-seti vektörü y üzerine koşullanmış $y - Z_1 b_1$ beklenen değerlerinin bulunması gerekir: $b_i | y$, q_i -boyutlu çok değişkenli normal dağılımlı olduğu için, $b_i | y \sim N_{q_i}(\sigma_i^2 Z_i^T V^{-1}(y - X\beta), \sigma_i^2 I_{q_i} - \sigma_i^4 Z_i^T V^{-1} Z_i)$ olarak yazılır.

$$E(b_0 | y) = \sigma_0^2 V^{-1}(y - X\beta),$$

$$E(b_0^T b_0 | y) = \sigma_0^4 (y - X\beta)^T V^{-1} Z_0 Z_0^T V^{-1} (y - X\beta) + \text{tr}(\sigma_0^2 I_{q_0} - \sigma_0^4 Z_0^T V^{-1} Z_0), \quad Z_0 = I_N.$$

Yeterli istatistiklerin beklenen değerleri aşağıdaki gibi olacaktır;

$$E(b_i^T b_i | y) = \sigma_i^4 (y - X\beta)^T V^{-1} Z_i Z_i^T V^{-1} (y - X\beta) + iz(\sigma_i^2 I_{q_i} - \sigma_i^4 Z_i^T V^{-1} Z_i)$$

$$E(y - Z_1 b_1 | y) = X\beta + \sigma_0^2 V^{-1} (y - X\beta).$$

BM Algoritması adımları

Adım 1 (Tahmin-Adımı);

$V^{(t)} = Z_1 Z_1^T \sigma_1^{2(t)} \sigma_0^{2(t)} I_n$, $i = 0, 1$ aşağıdakiler hesaplanır:

$$\begin{aligned} \hat{s}_i^{(t)} &= E(b_i^T b_i | y) \big|_{\beta=\beta^{(t)}, \sigma_i^2=\sigma_i^{2(t)}} \\ &= \sigma_i^{4(t)} (y - X\beta^{(t)})^T V^{(t)-1} Z_i Z_i^T V^{(t)-1} (y - X\beta^{(t)}) + iz(\sigma_i^{2(t)} I_{q_i} - \sigma_i^{4(t)} Z_i^T V^{(t)-1} Z_i) \end{aligned}$$

$$\hat{w}^{(t)} = E(y - Z_1 b_1 | y) \big|_{\beta=\beta^{(t)}, \sigma_i^2=\sigma_i^{2(t)}} = X\beta^{(t)} - \sigma_0^{2(t)} V^{(t)-1} (y - X\beta^{(t)})$$

Adım 2 (Maksimizasyon Adımı);

$$\sigma_i^{2(t+1)} = \hat{s}_i^{(t)} / q_i$$

$$\beta^{(t+1)} = (X^T X)^{-1} X^T \hat{w}^{(t)} \text{ elde edilir.}$$

Örnek. Domuz üretilen bir çiftlikte 5 erkek domuz için 2 farklı dişi domuz rastgele seçilerek her biri ile eşleştirilmiştir. Bir defada doğan her iki domuz yavrusunun günlük aldıkları kilo miktarı pound cinsinden ölçülmüştür. İlgili veri seti Çizelge 5.3’de verilmiştir (Snedecor ve Cochran, 1967). α_i i . erkek domuzla ilişkili sabit etki, β_{ij} i . erkek domuz ve j . dişi domuzla ilişkili rastgele etki, e_{ijk} hata terimidir. Tahmin edilmesi gereken model iki varyans bileşeni ve bir sabit etki terimine sahiptir, burada bir rasgele etki terimi vardır yani $b_1 = \beta_{ij}$, $r = 1$ ve β_{ij} rasgele etkisinin $5 \times 2 = 10$ düzeyi vardır dolayısıyla $q_1 = 10$ olduğu için $\beta_{ij} \sim N_{10}(0, \sigma_\beta^2 I_{10})$ yazılabilir; Model: $y_{ijk} = \mu + \alpha_i + \beta_{ij} + e_{ijk}$

Çizelge 5.3. Domuz çiftliğinde yavruların ağırlık artış veri seti

Erkek Domuz	Dişi Domuz	Ağırlık Artımı (pound)		Erkek Domuz	Dişi Domuz	Ağırlık Artımı
1	1	2,77		3	2	2,72
1	1	2,38		3	2	2,74
1	2	2,58		4	1	2,87
1	2	2,94		4	1	2,46
2	1	2,28		4	2	2,31
2	1	2,22		4	2	2,24
2	2	3,01		5	1	2,74
2	2	2,61		5	1	2,56
3	1	2,36		5	2	2,50
3	1	2,71		5	2	2,48

Veri setinin analizi yazılan program koduna uygun olması amacıyla varyans bileşenleri modelinde α_5 ile ifade edilen son domuzun kategorisi referans kategorisi olarak alınarak aşağıdaki DKE modeline uydurulmuştur. Bu örnekle birlikte varyans bileşenleri modeli DKE modeline dönüştürebilmek için yeniden parametrelendirilmiştir. Yukarıdaki tablolarda verilen veri setinin sabit etkilerle ilişkili tasarım matrisi X ve rasgele etkiler ile ilişkili tasarım matrisi Z , sabit etkilerin tahminleri ve varyans bileşenlerinin tahminleri ve DKE model denklemi aşağıda verilmiştir.

Çizelge 5.4. DKE modeli matris gösterimi

X Tasarım Matrisi	Z Tasarım Matrisi	Model Denklemleri
$\begin{bmatrix} 1 & 1 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 \end{bmatrix}$	$\begin{bmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}$	$Y = X \begin{bmatrix} \alpha_0 \\ \alpha_1 \\ \alpha_2 \\ \alpha_3 \\ \alpha_4 \end{bmatrix} + Z \begin{bmatrix} \beta_{11} \\ \beta_{12} \\ \beta_{21} \\ \beta_{22} \\ \beta_{31} \\ \beta_{32} \\ \beta_{41} \\ \beta_{42} \\ \beta_{51} \\ \beta_{52} \end{bmatrix} + \begin{bmatrix} e_{111} \\ e_{112} \\ e_{121} \\ e_{122} \\ e_{211} \\ e_{212} \\ e_{221} \\ e_{222} \\ e_{311} \\ e_{312} \\ e_{321} \\ e_{322} \\ e_{411} \\ e_{412} \\ e_{421} \\ e_{422} \\ e_{511} \\ e_{512} \\ e_{521} \\ e_{522} \end{bmatrix}$

Çizelge 5.5. Beklenti ve maksimizasyon adımları

İterasyon No	Log-Olabilirlik Değeri	Hata Terimleri Varyans Tahmini	Rasgele Etkilerin Varyans Tahmini
256	20,639	0,03869933	0,0088285
257	20,639	0,03869935	0,008828467
260	20,64	0,03869941	0,008828376
261	20,64	0,03869943	0,008828347
262	20,64	0,03869945	0,00882832

Çizelge 5.6. DKE modelinin sabit etkileri ve varyans bileşenleri tahminleri

$\hat{\alpha}_0 = 2,57$	$\hat{\alpha}_1 = 0,0975$	$\hat{\alpha}_2 = -0,04$	$\hat{\alpha}_3 = 0,0625$	$\hat{\alpha}_4 = -0,10$
$\hat{\sigma}_e^2 = 0,0387$		$\hat{\sigma}_\beta^2 = 0,0088$		

5.3.2. İç içe hata regresyon modeli için BM algoritması

İç-içe hata regresyon modeli (nested error regression model) iki varyans bileşenli doğrusal karışık etki modelinin özel bir halidir. Yukarıdaki örnekte sabit etkiler için bağımsız değişkenler kategoriktir bu modelde ise sürekli. Tezin simülasyon aşamasında bu modelle ilgilenilmiştir ve modelin sabit etki parametresi için farklı güven aralıkları incelenmiştir. Bu örnekte BM algoritması kullanılarak belli bir kombinasyon için ve normallik varsayımı altında modelin parametreleri tahmin edilecektir. BM algoritması için kodlar R programında yazılarak aşağıdaki model için parametrelerin BM yöntemi ile en çok olasılık tahmin edicileri elde edilmiştir ;

$$Y_{ij} = \beta_0 + \beta_1 x_{ij} + b_i + e_{ij}; \quad i = 1, 2; \quad j = 1, 2, 3; \quad n = 2; m = 3$$

$$i = 1, \dots, n; \quad j = 1, \dots, m \quad b_i \sim N(0, \sigma_b^2 I_2) \quad e_i \sim N(0, \sigma_e^2 I_3)$$

Çizelge 5.7. İç içe hata regresyon modeli için üretilen veri seti

Y	X	Z	Model Denklemini
$\begin{bmatrix} -2,660 \\ 0,034 \\ 1,047 \\ 2,372 \\ 1,911 \\ 1,443 \end{bmatrix}$	$\begin{bmatrix} 1 & 0,506 \\ 1 & -4,561 \\ 1 & -7,447 \\ 1 & -19,468 \\ 1 & -11,493 \\ 1 & 6,126 \end{bmatrix}$	$\begin{bmatrix} 1 & 0 \\ 1 & 0 \\ 1 & 0 \\ 0 & 1 \\ 0 & 1 \\ 0 & 1 \end{bmatrix}$	$Y = X \begin{bmatrix} \beta_0 \\ \beta_1 \end{bmatrix} + Z \begin{bmatrix} b_1 \\ b_2 \end{bmatrix} + \begin{bmatrix} e_{11} \\ e_{12} \\ e_{13} \\ e_{21} \\ e_{22} \\ e_{23} \end{bmatrix}$

R programında üretilen Çizelge 5.7'deki veri seti kullanılarak model için BM algoritması yoluyla iç-içe hata regresyon modelinin sabit etkileri ve varyans bileşenlerinin parametre tahminleri elde edilecektir.

Modeldeki sabit etki parametresinin ilk değeri olarak modelde rastgele etki yokmuş gibi düşünülerek veri seti doğrusal regresyon modeline uydurularak doğrusal regresyon tahmin edicisi ilk değer olarak alınmıştır.

Çizelge 5.8. İç-içe hata regresyon modelinde parametre tahmini için başlangıç değerleri

$\beta_0^{(0)} = 0,03746$	$\beta_1^{(0)} = -0,10711$
$\sigma_b^{2(0)} = 2,317$	$\sigma_e^{2(0)} = 1,94$

Rastgele etki ve hata teriminin varyanslarının ilk değerleri yine benzer bir mantıkla modelde sabit etki yokmuş gibi düşünülerek veri seti ANOVA modeline uydurulmuştur. Buradan ANOVA tahminleri varyans bileşenlerinin ilk değerleri olarak seçilerek algoritmada kullanılmıştır.

Çizelge 5.9. İç-içe hata regresyon modeli parametre tahminleri BM adımları

İterasyon No	Log-Olabilirlik Değeri	Hata Terimleri Varyans Tahmini	Rasgele Etkilerin Varyans Tahmini
106	-4,914936	1,470034	0,5568394
107	-4,914936	1,470036	0,5568322
108	-4,914936	1,470039	0,5568256
109	-4,914936	1,470041	0,5568195
110	-4,914936	1,470043	0,5568138
111	-4,914936	1,470045	0,5568085
112	-4,914936	1,470047	0,5568037
113	-4,914936	1,470048	0,5567992
115	-4,914936	1,470051	0,5567912
117	-4,914936	1,470054	0,5567843
118	-4,914936	1,470055	0,5567813

Çizelge 5.10. İç-içe hata regresyon modelinde parametre tahminleri

$\hat{\beta}_0 = 0,14459$	$\hat{\beta}_1 = -0,08956$
$\hat{\sigma}_b^2 = 0,5567$	$\hat{\sigma}_e^2 = 1,4700$

5.4. Doğrusal Karışık Etki Modellerinde Kullanılan İstatistiksel Yazılımlar

Normal dağılımlı DKE modellerinin analizi için prosedürler ve paketler SPSS, S-Plus, SAS, STATA ve R gibi istatistiksel yazılımlarda mevcuttur.

Bir çok yeni yazılım prosedürü gruplanmış veya uzun-vadede toplanmış tekrarlanan ölçümlü veri setlerinin her grup veya deney birimi için gözlem sayısı kadar satıra sahip “uzun” veya “dikey” formatta olmasını gerektirir. Genellikle tekrarlanan ölçümlü çalışmalarda yaygın olan geniş formatta girilmiş veri setleri SPSS (VARSTOCASES komutuyla), Stata gibi hazır menü yazılımlarıyla yeniden şekillendirilerek dikey formata dönüştürülür. Model denklemini bazı yazılımlar menü yoluyla (örneğin SPSS MIXED) bazı yazılımlar (SAS PROC MIXED, R lme paketi gibi) komutlar yardımıyla belirler. Yeni geliştirilen yazılımlar kullanıcıya tahmin yöntemini seçme özelliği vermektedir. Model uydurulma aşamasında EÇO veya KEÇO tahmin yöntemi olarak seçilebilir. SAS PROC MIXED fonksiyonu kovaryans parametrelerinin tahmininde kullanıcıya algoritmanın başlangıç değerlerini belirleyebilme özelliği sunmuştur. Ayrıca Newton-Raphson ve BM algoritması gibi sayısal yöntemler birleştirilerek ve olabilirlik fonksiyonu optimize edilerek tahmin değerlerine ulaşan istatistiksel yazılımlar son yıllarda geliştirilmiştir. Bir çok yazılım çapraz (crossed) rasgele etkilerle model uydurulmasına olanak sağlamasına rağmen, DKE modellerinin en önemli özelliklerinden biri olan veri setinin farklı gruplarında/ deney birimlerinde farklı varyans-kovaryans yapılarına izin vermekte ciddi boşluklar vardır. R lme4 paketi kullanıcının rasgele etkilerin kovaryans yapısını belirlemesine izin verir. SAS PROC MIXED prosedürü hem rasgele etkilerin hem de hata terimlerinin kullanıcı tarafından belirlenmesine izin vermiştir. Bununla birlikte gruba/ deney birimine özel rasgele etki ve hata terimlerinin kovaryans matrislerinin belirlenmesine de SAS PROC MIXED, SAS PROC GLIMMIX, WinBugs, SPSS: MIXED / GENLINMIXED yazılımları olanak verir. Kovaryans matrislerinin G , R , V varyans elemanlarının tahminleri negatif olmamalıdır. Bu nedenle bir çok yeni yazılım prosedürünün içinde varyans bileşenlerinin tahminlerinin negatif olmaması için tahmin yöntemleri algoritmalarına parametre uzayları için kısıtlar yerleştirilmiştir.

6. GÜVEN ARALIĞI

Parametrenin değerinin nokta tahmini ile birlikte, tahminin gerçek değere ne kadar yakın olduğunu bilmek gereklidir. Tahmin edicinin varyansı veya hata kareler ortalaması bulunarak bu soruya az da olsa bilgi sağlanabilir. Başka bir yaklaşım da aralık tahminlerini düşündürmektedir; gerçek parametre değerini bu aralığın hangi olasılıkla içereceği önemli bir problemidir.

Örnek 6.1. $n=40$ rasgele örneği $X_i \sim \exp(\theta)$ üstel dağılımdan çekilmiş olsun ve θ parametresi ortalama yaşam süresini temsil etsin. Örnek ortalaması, $\bar{X}$, θ parametresi için UMVUE (düzgün minimum varyanslı yansız tahmin edici) olsun. Verilmiş bir veri seti için θ tahmini $\bar{x} = 93.1$ ay olarak elde edilsin. Bu tahmin edicinin en iyi ve optimal özelliklere sahip bir tahmin edici olduğunu bilmemize rağmen, nokta tahmini kendi başına doğruluğu hakkında bir bilgi sağlamamaktadır. Bu probleme çözüm, uç noktaları rastgele değişken olan ve bu noktalar arasında 1'e yakın bir olasılıkla, 0.95 gibi, gerçek parametre değerini içeren bir aralık oluşturmak olacaktır.

$2n\bar{X} / \theta \sim \chi^2(2n)$ olmak üzere ki-kare dağılımının yüzdeliklerini kullanarak yukarıdaki örnek için güven aralığı oluşturulsun;

$$n = 40, v = 80, \chi_{0.025}^2(80) = 57.15, \chi_{0.975}^2(80) = 106.63,$$

$$P[57.15 < 80\bar{X} / \theta < 106.63] = 0.975 - 0.025 \text{ ve sonuç olarak } \theta \text{ için güven aralığı}$$

$$P[80\bar{X} / 106.63 < \theta < 80\bar{X} / 57.15] = 0.95 \text{ olacaktır.}$$

Genel olarak, rasgele uç noktalara sahip olan aralık, rasgele aralık olarak adlandırılır. $(80\bar{X} / 106.63, 80\bar{X} / 57.15)$ aralığı gerçek parametreyi 0.95 olasılıkla içeren rasgele bir aralıktır. $\bar{X}$ yerine $\bar{x} = 93.1$ tahmini konulursa aralık $(69.9, 130.3)$ olarak bulunacaktır. Bu aralığa θ için % 0.95 güven aralığı denir. Tahmin edilmiş aralık bilinen uç noktaları içerdiğinden, 0.95 olasılıkla gerçek parametre değeri θ 'yı içerdiğini söylemek uygun olmayacaktır. θ parametresi bilinmez olmasına rağmen bir sabittir ve bu spesifik aralık θ 'yı içerebilir de içermeyebilir de. Tahminden önce rasgele aralığın % 0.95 olasılıkla

gerçek parametreyi içereceğini söylemek, $69.9 < \theta < 130.3$ aralığının gerçek parametre değerinin içerilmesinde % 0.95 güvende olduğunu iddia etme sonucuna götürür.

Güven aralıkları bu bölümde formüsel olarak tanımlanacaktır (Bain ve Engelhardt, 1992: 360);

$X_1, \dots, X_n$ rasgele değişkenler Ω bir aralık ve $\theta \in \Omega$ olmak üzere bileşik olasılık yoğunluk fonksiyonu $f(x_1, \dots, x_n; \theta)$ şeklindedir. L ve U istatistikleri $L = l(X_1, \dots, X_n)$ ve $U = u(X_1, \dots, X_n)$ şeklinde tanımlansın. $x_1, \dots, x_n$ veri seti bir denemenin sonucunda gözlenmiş değerler olsun, l ve u için gözlenen değerler $l(x_1, \dots, x_n)$ ve $u(x_1, \dots, x_n)$ şeklinde yazılabilir.

6.1. Tanım $P[(l(x_1, \dots, x_n) < \theta < u(x_1, \dots, x_n))] = \gamma$, $0 < \gamma < 1$ olarak gösterilen aralığa θ için $(l(x_1, \dots, x_n), u(x_1, \dots, x_n))$ % 100γ güven aralığı denir. Gözlenen değerler $l(x_1, \dots, x_n)$ ve $u(x_1, \dots, x_n)$ alt ve üst güven sınırları olarak adlandırılır.

İstatistik literatüründe θ_l ve θ_u ile sırasıyla alt ve üst güven limitleri gösterilebilir. Bazen kısaltılmış ifadeler $l(x) = l(x_1, \dots, x_n)$ ve $u(x) = u(x_1, \dots, x_n)$ gözlenen limitleri göstermek için kullanılabilir.

Rastgele aralık ve gözlemlenmiş aralık arasındaki ayrım iyi yapılmalıdır. Bu ayrım nokta tahmini konusundaki tahmin edici ve tahmin arasındaki ayrıma karşılık gelir. (L, U) aralık tahmin edicisini, $(l(x), u(x))$ de aralık tahminini gösterir. γ olasılık düzeyine güven düzeyi denir. Güven aralığı kavramının en yaygın tanımı olasılığın frekans özelliğine dayalıdır. Aralık tahminleri birçok farklı örnekten hesaplanırsa, uzun vadede yaklaşık % 100γ olasılıkla aralığın gerçek parametre değeri θ 'yı kapsayacağı beklenir.

Bazı araştırmalarda alt ve üst sınırlardan herhangi birisi ilgi konusu olabilir. Aşağıda tek taraflı güven sınırı tanımı verilmiştir.

6.2. Tanım. Tek Taraflı Güven Sınırları;

1. $P[l(x_1, \dots, x_n) < \theta] = \gamma$ ve $l(x) = l(x_1, \dots, x_n)$ θ için tek taraflı % 100 γ güven düzeyinde alt güven sınırı denir.
2. $P[(\theta < u(x_1, \dots, x_n))] = \gamma$ ve $u(x) = u(x_1, \dots, x_n)$ θ için tek taraflı % 100 γ güven düzeyinde üst güven sınırı denir.

Yukarıdaki tanımları sağlayan güven sınırlarını elde etmek her zaman çok açık değildir. Yeterlilik kavramı genellikle bu problem için iyi bir çözümdür. Yeterli istatistik bulunduğunda, bu istatistiğin fonksiyonu olan güven sınırları elde edilebilir.

Örnek 6.2. $X_i \sim \exp(\theta)$ dağılımından n örnek hacimli rasgele örnek üstel dağılımdan seçilmiş olsun ve θ için tek taraflı alt güven sınırı elde edilsin. $\bar{X}$ istatistiğinin θ için yeterli bir istatistik ve $2n\bar{X} / \theta \sim \chi^2(2n)$ olduğu da bilinmektedir. Böylece,

$$\begin{aligned} \gamma &= P[2n\bar{X} / \theta < \chi^2(2n)] \\ &= P[2n\bar{X} / \chi^2(2n) < \theta] \end{aligned}$$

$\bar{x}$ gözlenirse, tek taraflı % 100 γ güven düzeyinde alt güven limiti $l(x) = 2n\bar{x} / \chi^2_\gamma(2n)$, ve üst güven limiti de benzer olarak $u(x) = 2n\bar{x} / \chi^2_{1-\gamma}(2n)$ şeklinde gösterilir.

θ için % 100 γ güven aralığı oluşturulmak istensin. $\alpha_1 > 0$ ve $\alpha_2 > 0$ değerleri $\alpha_1 + \alpha_2 = \alpha = 1 - \gamma$ olacak şekilde seçilir ve böylece güven aralıkları aşağıdaki gibi oluşturulur;

$$\begin{aligned} P[\chi^2_{\alpha_1}(2n) < 2n\bar{X} / \theta < \chi^2_{1-\alpha_2}(2n)] &= 1 - \alpha_1 - \alpha_2, \\ P[2n\bar{X} / \chi^2_{1-\alpha_2}(2n) < \theta < 2n\bar{X} / \chi^2_{\alpha_1}(2n)] &= \gamma. \end{aligned}$$

Genellikle uygulamada $\alpha_1 = \alpha_2$ olarak belirlenir ve eşit kuyruklu seçim olarak bilinir, yani $\alpha_1 = \alpha_2 = \alpha / 2$. Bu forma uyan güven aralığı $(2n\bar{x} / \chi^2_{1-\alpha/2}(2n), 2n\bar{x} / \chi^2_{\alpha/2}(2n))$ şeklindedir.

Belirlenmiş bir güven düzeyi için, en küçük uzunluk gibi bazı optimal özellikleri içeren metotların kullanılması gerekir. $U-L$ ilgili güven aralığının uzunluğu da bir rastgele

değişken olacaktır, bu uzunluğun minimum beklenen uzunluk olması istenir. Bazı problemler için, α_1 ve α_2 'nin eşit kuyruklu seçimi minimum beklenen uzunluğu sağlar, ama diğer seçimlerde sağlanmaz.

Örnek 6.3. Normal dağılımdan bir rasgele örnek seçilmiş olsun, $X_i \sim N(\mu, \sigma^2)$ ve σ^2 bilindiği varsayalım. Bu durumda $\bar{X}$, μ için yeterli bir istatistik olur, ve $Z = \sqrt{n}(\bar{X} - \mu) / \sigma \sim N(0,1)$ olduğu da bilinir. Dağılımın simetrik olması özelliğinden $z_{\alpha/2} = -z_{1-\alpha/2}$ μ için güven aralığı,

$$1 - \alpha = P\left[-z_{1-\alpha/2} < \sqrt{n}(\bar{X} - \mu) / \sigma < z_{1-\alpha/2}\right]$$

$$= P\left[\bar{X} - z_{1-\alpha/2}\sigma / \sqrt{n} < \mu < \bar{X} + z_{1-\alpha/2}\sigma / \sqrt{n}\right]$$

μ için % 100 γ güven aralığı da $(\bar{x} - z_{1-\alpha/2}\sigma / \sqrt{n}, \bar{x} + z_{1-\alpha/2}\sigma / \sqrt{n})$ olarak ifade edilir.

σ^2 'nin bilinmediği durumda bu çözüm farklılaşır, çünkü güven sınırları bilinmeyen parametreye bağlı olacağından hesaplanamaz. Güven aralığını oluşturma yönteminde küçük bir değişiklik yaparak μ için σ^2 'nin bilinmeyen bir parametre olması durumunda da güven aralığı oluşturulur. Aslında en büyük zorluk birden çok bilinmeyen parametrenin olduğu durumlardır. Genellikle bu tür durumlarda birden çok parametreye aynı anda uygulanan bileşik güven alanı oluşturulur. Pivot kavramı bu bilinmeyen parametre olduğunda nasıl güven aralığı kurulacağı sorusuna cevap olarak ortaya çıkmıştır. Bundan sonraki bölümde pivot kavramı üzerinde durulacaktır.

6.1. Pivot Ölçüm Metodu

$X_1, \dots, X_n$ 'in bileşik olasılık yoğunluk fonksiyonu $f(x_1, \dots, x_n; \theta)$ olsun. Bir çok bilinmeyen parametrenin de içinde olduğu θ için güven limitlerini elde etmek isteyebiliriz.

6.1.1. Tanım: Sadece $X_1, \dots, X_n$ ve θ 'nın bir fonksiyonu olan $Q = q(X_1, \dots, X_n; \theta)$ rastgele değişkenine pivot ölçüm denir. $Q = q(X_1, \dots, X_n; \theta)$ ölçüsünün dağılımı θ veya bilinmeyen diğer parametrelere bağlı değildir (Bain ve Engelhardt, 1992: 363).

Örnek 6.1.1. *Örnek 6.1'*de ki-kare dağılımına sahip rastgele değişken incelendi. Buradaki $Q = 2n\bar{X} / \theta$ olarak tanımlanan Q pivot ölçüm tanımına tamamen uymaktadır. θ için güven sınırları elde ederken Q hakkında bir olasılık ifadesinden yararlanılmıştır.

Genel olarak , θ için bir pivot ölçüm varsa ve bu pivot ölçümün yüzdelikleri q_1 ve q_2 elde edilebiliyorsa, güven aralığı elde etmek olası hale gelmektedir.

$$P[q_1 < q(X_1, \dots, X_n; \theta) < q_2] = \gamma$$

Gözlenen örneklem $x_1, \dots, x_n$ olmak üzere θ için $\% 100\gamma$ güven alanı $q_1 < q(x_1, \dots, x_n; \theta) < q_2$ eşitsizliğini sağlayan $\theta \in \Omega$ setidir. Böyle bir güven alanı bir aralık olmak zorunda değildir ve genellikle daha karmaşıktır. Ama bazı durumlarda güven aralıkları elde edilebilir. $x_1, \dots, x_n$ sabit değerlerinin her bir seti için, $q(x_1, \dots, x_n; \theta)$ θ 'nın monoton artan (veya azalan) bir fonksiyonu ise her zaman bir güven aralığını oluşturmak mümkündür. Pivot ölçümleri ortaya çıkaran dağılımların belli özelliklerini saptamak ayrıca mümkündür. Olasılık yoğunluk fonksiyonu (oyf) $f(x; \theta) = f_0(x - \theta)$ formundaysa θ konum parametresi olur. Oyf $f(x; \theta) = (1/\theta) f_0(x/\theta)$ formundaysa θ ölçek parametresidir. θ_1 ve θ_2 parametreleri için oyf $f(x; \theta_1, \theta_2) = (1/\theta_2) f_0((x - \theta_1)/\theta_2)$ şeklinde ise konum-ölçek parametreleri sırasıyla θ_1 ve θ_2 olacaktır. Bu durumların herhangi birinde EÇO tahmini elde edilebilirse, pivot ölçüm oluşturmak için kullanılır.

Yardımcı Teorem 6.1.1. (Bain ve Engelhardt, 1992: 363) $X_1, \dots, X_n$ 'in oyf $f(x; \theta)$ $\theta \in \Omega$ olmak üzere ve θ için en çok olabilirlik tahmin edicisi de $\hat{\theta}$ olsun. θ konum parametresi olmak üzere, $Q = \hat{\theta} - \theta$ pivot ölçümüdür. θ ölçek parametresi olmak üzere, $Q = \hat{\theta} / \theta$ pivot ölçümüdür.

Örneğin, $X_i \sim N(\mu, \sigma^2)$ ve σ^2 bilinsin, μ konum parametresinin EÇO tahmin edicisi $\bar{X}$ ve böylece $\bar{X} - \mu$ pivot ölçüm olacaktır. $X_i \sim \exp(\theta)$ olduğunda θ için en çok olabilirlik tahmin edicisi $\bar{X}$ ve $Q = \bar{X} / \theta$ da pivot ölçüm olacaktır. Bazen küçük değişiklikler yaparak (bilinen bir ölçek faktörü ile çarpmak gibi) pivot ölçümün dağılımı

bilinir hale getirilebilir. Örneğin, üstel dağılım için, $Q = \bar{X} / \theta$ pivot ölçümü yerine yüzdeliklerinin bilinmesi nedeniyle $2n\bar{X} / \theta \sim \chi^2(2n)$ olarak seçmek daha iyi olacaktır.

Yardımcı Teorem 6.1.2. (Bain ve Engelhardt, 1992: 364) $X_1, \dots, X_n$ 'in oyf konum-ölçek parametrelili bir dağılıma $f(x; \theta_1, \theta_2) = \left(\frac{1}{\theta_2}\right) f_0\left(\frac{x - \theta_1}{\theta_2}\right)$ sahip rasgele örnek olsun. θ_1 ve θ_2 parametreleri için EÇO tahmin edicileri $\hat{\theta}_1$ ve $\hat{\theta}_2$ mevcutsa $(\hat{\theta}_1 - \theta_1) / \hat{\theta}_2$ ve $\hat{\theta}_2 / \theta_2$ sırasıyla θ_1 ve θ_2 parametreleri için pivot ölçümlerdir.

Örnek 6.1.2. $X_i \sim N(\mu, \sigma^2)$ μ ve σ^2 bilinmeyen parametrelili normal dağılımdan bir rasgele örnek olsun. $\hat{\mu}$ ve $\hat{\sigma}$, μ ve σ için en çok olabilirlik tahmin edicileri olsun. Birinden biri bilinmeyen parametre olarak düşünülüp her parametre için güven aralığı kurmak için gerekli pivot ölçümler sırasıyla $(\hat{\mu} - \mu) / \hat{\sigma}$ ve $\hat{\sigma} / \sigma$ olacaktır. Bazı bilinen dağılımsal özelliklerden yararlanmak için $S^2 = n\hat{\sigma}^2 / (n-1)$ yansız tahmin edicisi bakımından sonuçlar aşağıdaki gibi ifade edilir,

$$\frac{\bar{X} - \mu}{S / \sqrt{n}} \sim t(n-1) \text{ ve } \frac{(n-1)S^2}{\sigma^2} \sim \chi^2(n-1).$$

$(n-1)$ serbestlik dereceli t dağılımının $1-\alpha/2$ yüzdeliği $t_{1-\alpha/2} = t_{1-\alpha/2}(n-1)$ olmak üzere

$$1-\alpha = P\left[-t_{1-\alpha/2} < \frac{\bar{X} - \mu}{S / \sqrt{n}} < t_{1-\alpha/2}\right] = P\left[\bar{X} - t_{1-\alpha/2}S / \sqrt{n} < \mu < \bar{X} + t_{1-\alpha/2}S / \sqrt{n}\right]$$

$\bar{x}$ ve s gözlenen değerler olmak üzere μ için $\% 100(1-\alpha)$ güven aralığı

$\left(\bar{x} - t_{1-\alpha/2}s / \sqrt{n}, \bar{x} + t_{1-\alpha/2}s / \sqrt{n}\right)$ 'dir.

7. DKE MODELİNDE GÜVEN ARALIĞI METODLARI

Bu bölümde DKE modellerinin sabit etki parametreleri için kullanılan standart güven aralığı metotlarıyla birlikte son yıllarda sabit etki ve karma etki için kullanılan parametrik bootstrap güven aralığı yöntemi açıklanacaktır. Ayrıca daha önce kategorik modellerin parametrelerinde kullanılan yapay-skor (pseudo-score) metodu DKE modellerinin sabit etki parametresi için ilk olarak bu tezde uyarlanacaktır.

7.1. Profil Olabilirlik Güven Aralığı Metodu

Profil olabilirlik güven aralığı (POGA) olabilirlik oran istatistiğinin asimptotik olarak ki-kare dağılımı özelliğine dayalıdır. Örnek çapının küçük olduğu durumlarda Wald güven aralığı metodundan daha iyi sonuç verdiği görülmüştür. Bu çalışmada DKE modellerinin sabit etki parametresi için POGA metodu kullanılacağı için notasyonlar DKE modelleri için uyarlanmıştır.

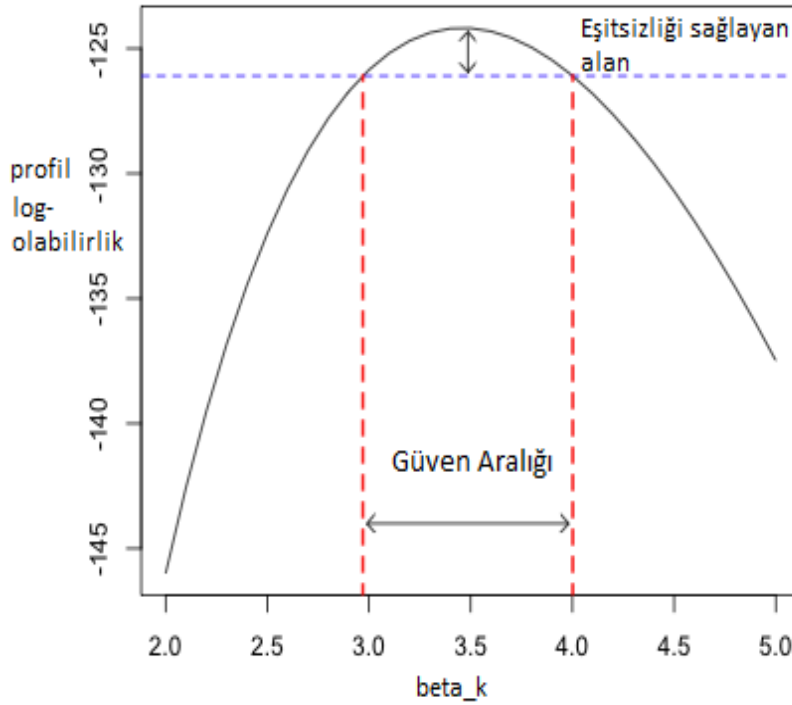
Herhangi bir DKE modelinin bilinmeyen parametre vektörü $\theta = (\beta_k, \delta)$, β_k ilgilenilen sabit etki parametresi, δ modeldeki ilgilenilmeyen parametre vektörü ve $\beta_{k(0)}$ ilgili parametrenin aldığı sabit bir değer olmak üzere $\beta_k = \beta_{k(0)}$ hipotezi test edilmek istensin. $L(\beta_k, \delta)$ modelin olabilirlik fonksiyonu ve β_k için profil olabilirlik fonksiyonu $L_1(\beta_k) = \max_{\delta} L(\beta_k, \delta)$ olarak gösterilsin. β_k için alınan her sabit $\beta_{k(0)}$ değerinde bu indirgenmiş modeldeki diğer parametrelerin EÇO tahminleri üzerinden $L_1(\beta_{k(0)})$ fonksiyonu maksimize edilir. Yani profil olabilirlik fonksiyonu (POF) bir olabilirlik fonksiyonu değildir; POF'daki her nokta olabilirlik fonksiyonunun maksimum değeridir.

$\hat{\beta}_k$, β_k parametresi için $\hat{\delta}$ da δ parametre vektörü için tam modelde EÇO tahminleri, χ^2_1 1 serbestlik dereceli ki-kare dağılımı ve $\beta_k = \beta_{k(0)}$ indirgenmiş modelde δ parametre vektörü için EÇO tahminleri $\hat{\delta}_0$ olmak üzere bu test için olabilirlik oran (likelihood ratio, LR) istatistiği aşağıdaki gibidir;

$$LR = 2 \log \left[\frac{L(\beta_k, \hat{\delta})}{L(\beta_{k(0)}, \delta_0)} \right] \sim \chi_1^2. \quad (7.1)$$

$\chi_{1,\alpha}^2$ 1 serbestlik dereceli ki-kare dağılımının $1-\alpha$ çeyrekliği olmak üzere β_k için $\%100(1-\alpha)$ profil olabilirlik güven aralığı aşağıdaki eşitsizliği sağlayan ($\beta_k = \beta_{k(0)}$) hipotezinin reddedilemediği $\beta_{k(0)}$ değerlerini içerir:

$$2 \left[\log L(\hat{\beta}_k, \hat{\delta}) - \log L(\beta_{k(0)}, \hat{\delta}_0) \right] < \chi_{1,\alpha}^2 \quad \text{veya} \quad \log L(\beta_{k(0)}, \hat{\delta}_0) > \log L(\hat{\beta}_k, \hat{\delta}) - \frac{\chi_{1,\alpha}^2}{2}.$$



Şekil 7.1. Profil olabilirlik güven aralığı gösterimi

$\log L(\hat{\beta}_k, \hat{\delta})$ ve $\frac{\chi_{1,\alpha}^2}{2}$ değerleri sabit değerler olduğu için $\log L(\beta_{k(0)}, \hat{\delta}_0)$ değerinin $\log L(\hat{\beta}_k, \hat{\delta}) - \frac{\chi_{1,\alpha}^2}{2}$ değerinden büyük olmasını sağlayan $\beta_{k(0)}$ değerleri Şekil 7.1’de gösterildiği gibi β_k için güven aralığı oluşturacaktır.

Profil olabilirlik güven aralığı algoritması

Tam modelden log-olabilirlik fonksiyonunun değeri ve parametrelerin tahmin değerleri elde edilir. İlgilenilen sabit etki β_k için hesaplanan tahmin değeri ve standart hatası alt sınır belirlemek için kullanılır.

Profil log-olabilirlik fonksiyonu maksimum değerinin solunda kalan alanın Şekil 7.1’de verildiği gibi artan bir fonksiyon olduğu varsayılınsın. Alt güven limitini hesaplamak için aşağıdaki algoritma verilmiştir;

1. $c = 1.96$ gibi sabit bir değer olmak üzere alt güven limitini hesaplamak için mantıklı bir alt sınır belirlenir, $\beta'_k = \hat{\beta}_k - c * SH(\hat{\beta}_k)$.
2. $\hat{\beta}_k$ ’den β'_k ’e kadar olan aralık eşit olarak, örneğin 100 noktaya bölünür.
3. Her $\beta_{k(i)}$ değeri için, $\hat{\delta}_{(i)}$ değerlerinde $\ln L(\beta_{k(i)}, \hat{\delta}_{(i)})$ maksimize edilerek profil olabilirlik $\ln L_1(\beta_{k(i)})$ değeri hesaplanır.
4. $\ln L_1(\beta_{k(i)}) \geq \ln L(\hat{\beta}_k, \hat{\delta}) - 1.92$ eşitsizliğini sağlayan en küçük $\beta_{k(i)}$ değeri β_k ’nin alt güven limiti olarak belirlenmiş olur. Üst güven limiti bu algoritmayla benzer şekilde elde edilir.

7.2. Wald Güven Aralığı Metodu

Wald istatistiğine dayalı güven aralığı hesaplama basitliği nedeniyle sıklıkla kullanılmaktadır. Wald güven aralığı EÇO tahmininin asimptotik olarak normal dağılıma sahip olması özelliğine göre oluşturulmuştur. Genel olarak, Wald metodu en çok olabilirlik metodunun kuadratik bir yaklaşımına dayalıdır. Y_i ’nin marjinal modelinde bulunan tüm parametrelerin vektörü $\theta = (\beta', \alpha')$ ile ifade edilsin. θ vektörü, α vektörü G matrisinin ve $R_i = \text{Var}(e_i)$ ’deki tüm elemanları içerir. $\hat{\theta}$, herhangi bir DKE parametresi θ için EÇO tahmini ve $SH(\hat{\theta})$, $\hat{\theta}$ için standart hata olmak üzere Wald istatistiği aşağıdaki gibi yazılır,

$$W = \frac{\hat{\theta} - \theta}{SH(\hat{\theta})} \quad (7.2)$$

$\hat{\theta}$, EÇO tahmin edicisi olduğu için büyük örneklerde dağılımı yaklaşık olarak normal dağılıma uymaktadır. Hessian matrisinin negatif bekleneni Fisher bilgi matrisi $\hat{\theta}$ 'nin standart hatasını elde etmek için kullanılabilir. Eş. 3.3'de marjinal $l(\beta, \alpha)$ logaritmik en çok olabilirlik fonksiyonu, β sabit etki parametrelerinin vektörü ve α varyans parametrelerinin vektörü olmak üzere Hessian matrisi aşağıdaki gibi ifade edilir;

$$H = \begin{pmatrix} H_{\beta\beta} & H_{\beta\alpha} \\ H_{\alpha\beta} & H_{\alpha\alpha} \end{pmatrix} = \begin{pmatrix} \frac{\partial^2 l}{\partial \beta \partial \beta} & \frac{\partial^2 l}{\partial \beta \partial \alpha} \\ \frac{\partial^2 l}{\partial \alpha \partial \beta} & \frac{\partial^2 l}{\partial \alpha \partial \alpha} \end{pmatrix}$$

Jenrich ve Schluchter (1986) doğrusal karışık modellerde Hessian matrisinin elemanları için aşağıdaki ifadeleri vermiştir:

$$H_{\beta\beta} = -\sum_{i=1}^n X_i' V_i^{-1} X_i$$

$$[H_{\alpha\beta}]_{rj} = [H_{\beta\alpha}]_{jr} = -\sum_{i=1}^n x_{ij}' V_i^{-1} \dot{V}_{ir} V_i^{-1} e_i, \quad j=1, \dots, p; \quad r=1, \dots, qv$$

$$[H_{\alpha\alpha}]_{rs} = -\frac{1}{2} \sum_{i=1}^n iz \left(V_i^{-1} \dot{V}_{ir} V_i^{-1} (2e_i e_i' - V_i) V_i^{-1} \dot{V}_{is} \right) \\ + \frac{1}{2} \sum_{i=1}^n iz \left(V_i^{-1} (e_i e_i' - V_i) V_i^{-1} \ddot{V}_{i,rs} \right), \quad r, s = 1, \dots, qv$$

Burada $e_i = Y_i - X_i \beta - Z_i b_i$, x_{ij} X_i 'nin j. sütunu $\dot{V}_{ir} = \partial V_i / \partial \alpha_r$, ve $\ddot{V}_{i,rs} = \partial^2 V_i / \partial \alpha_r \partial \alpha_s$. $\dot{V}_{ir}$ ve $\ddot{V}_{i,rs}$ matrislerinin elemanları sırasıyla $\alpha_1, \dots, \alpha_{qv}$ vektörüne göre V_i 'nin birinci ve ikinci türevlerinin elemanlarıdır. V_i kovaryans matrisinde qv toplam varyans parametre sayısını gösterir.

Fisher bilgi matrisi negatif Hessian matrisinin beklenen değeridir. T_i , $qv \times qv$ boyutlu matris ve elemanları $T_{jr}^i = \frac{1}{2} iz(V_i^{-1} \dot{V}_{ij} V_i^{-1} \dot{V}_{ir})$, $j=1, \dots, p; \quad r=1, \dots, qv$ olmak üzere, Fisher bilgi matrisi aşağıdaki gibidir;

$$I(\theta) = E(-H) = \begin{pmatrix} \sum_{i=1}^n X_i' V_i^{-1} X_i & 0 \\ 0 & \sum_{i=1}^n T_i \end{pmatrix}$$

Gözlenen bilgi matrisi, $\hat{I}(\theta)$, $\hat{I}(\theta) = -\hat{H}$ olarak gösterilir. $\hat{H}$, EÇO tahmin edicisi $\hat{\theta}$ 'dan elde edilmiştir. EÇO tahmin edicisi $\hat{\theta}$ 'nın varyansı bilgi matrisinin tersi hesaplanarak elde edilir:

$$\text{var}(\hat{\theta}) = [\hat{I}(\theta)^{-1}] = \left(-E[\hat{H}(\theta)] \right)^{-1} \big|_{\hat{\theta}=(\hat{\beta}, \hat{\alpha})} \text{ veya } \text{var}(\hat{\theta}) = [\hat{I}(\theta)^{-1}] = \left(-[\hat{H}(\theta)] \right)^{-1} \big|_{\hat{\theta}=(\hat{\beta}, \hat{\alpha})}$$

Yani, $\hat{\theta}$ 'nın standart hatası yukarıdaki beklenen veya gözlenen bilgi matrislerinden elde edilmiş varyans-kovaryans matrisinin köşegen elemanlarının karekökü olacaktır.

DKE modelinde β_k ($k=1, \dots, p$) bir sabit etki parametresi ve β_k onun EÇO tahminini ve $SH(\beta_k)$ de β_k tahminin standart hatasını gösterebilir. EÇO tahminlerinin standart hataları Hessian matrisinin negatif beklenen veya gözlenen değer matrisi (Fisher Bilgi Matrisi) kullanılarak elde edilir. Sabit etki parametresine ilişkin Wald istatistiği aşağıdaki gibidir;

$$\frac{\beta_k - \hat{\beta}_k}{SH(\hat{\beta}_k)}, \quad (7.3)$$

Eş.7.3'deki güven aralığı kurmak için gerekli olan pivottur ve dağılımının asimptotik standart normal olduğu varsayılır. Yeterli örnek çaplarında bu varsayımın doğruluğu

altında $z_{1-\frac{\alpha}{2}}$ standart normal dağılımın $1-\frac{\alpha}{2}$ çeyrekliği olmak üzere $\%100(1-\alpha)$

asimptotik Wald güven aralığı $\beta_k \pm z_{1-\frac{\alpha}{2}} SH(\hat{\beta}_k)$ olarak yazılır.

7.3. t-Naive Güven Aralığı Metodu

Wald güven aralığı oluşturulurken Eş.7.3'deki pivotun dağılımının asimptotik standart normal dağılımlı olduğu varsayılmıştır. Bu durum yeterli örnek çapları için doğru olsa da küçük örnek çaplarında güvenilir bir varsayım değildir. Demidenko (2004: 131) z-

dağılımının çeyreklikleri yerine student-t dağılımının çeyrekliklerini kullanmayı önermiştir. $t_{n^*m-p;1-\frac{\alpha}{2}}$, n^*m-p (dengeli veri seti durumunda) serbestlik dereceli student-t dağılımının $1-\frac{\alpha}{2}$ çeyrekliği olmak üzere β_k için $\%100(1-\alpha)$ t-naive güven aralığı $\hat{\beta}_k \pm t_{n^*m-p;1-\frac{\alpha}{2}} SH(\hat{\beta}_k)$ olarak elde edilir.

7.4. Yapay-Skor (Pseudo-Score) Güven Aralığı Metodu

Diğer bölümlerde bahsedilen analitik güven aralığı metodlarına ek olarak ilk defa bu tezde yapay-skor güven aralığı metodu doğrusal karışık-etki modellerindeki sabit etki parametreleri için tanıtılacaktır. Yapay-skor güven aralığı metodu ilk olarak kesikli bağımsız değişkenlere sahip istatistiksel modellerin parametreleri için kullanılmıştır (Agresti ve Ryu, 2010). Modeldeki ilgilenilen sabit etki parametresi farklı sabitlenmiş değerler aldığıında bu modeller için oluşturulan bağımlı değişkenin tahmini değerleri (fitted values) ile tam (full) modelin tahmini değerlerini karşılaştıran Pearson Ki-Kare test istatistiğine dayalı bir metottur. Profil-olabilirlik güven aralığı metodunun çalışmadığı durumlarda bu güven aralığı metodu kullanılabilir. İstatistiksel çıkarımlar profil olabilirlik metodunun çıkarımlarına asimptotik olarak alternatiftir.

DKE modelindeki bağımlı (yanıt) değişkeninin gözlenen değer vektörü Y_i $i = 1, \dots, n$ ve $\hat{Y}_i$ 'ler EÇO (veya KEÇO) tahmini değer vektörünü ifade etsin. İlgili sabit etki parametresinin β_k belli bir aralıkta sabit değeri $\beta_{k(0)}$ alınarak elde edilmiş indirgenmiş modelin EÇO tahmini değer vektörü $\hat{Y}_{i0}$ olarak ifade edilsin. Bu tez çalışmasında Pearson istatistiğinin genelleştirilmiş formu (Lovison, 2005) kullanılarak bu tezde R programı içinde geliştirilen bir sayısal algoritma yoluyla DKE modellerinin sabit etki parametreleri için yeni bir güven aralığı metodu oluşturulmuştur. Genelleştirilmiş Pearson istatistiğinin yapısı aşağıdaki gibidir;

$$X^2 = \sum_i \frac{(\hat{Y}_i - \hat{Y}_{i0})^2}{\text{var}(\hat{Y}_{i0})} = (\hat{Y} - \hat{Y}_0)' \hat{V}_0^{-1} (\hat{Y} - \hat{Y}_0), \quad (7.4)$$

burada $var(\hat{Y}_{i0})$ indirgenmiş model altında Y_i vektörünün tahmini varyans değerlerini ve $\hat{V}_0$ da ilgili değerleri içeren varyans-kovaryans matrisini göstermektedir. Tez çalışmasında bu metod için tahmini varyans-kovaryans matrisleri Kenward-Roger (1997) yaklaşımı kullanılarak elde edilmiştir. $\hat{Y}_{i0}$ değerleri $\beta_k = \beta_{k(0)}$ kısıtı altında belli bir aralıkta seçilen $\beta_{k(0)}$ değerleri için elde edilen indirgenmiş modellerden elde edilir. $\chi^2_{1,\alpha}$ 1 serbestlik dereceli Ki-Kare dağılımının $(1-\alpha)$ yüzdeliğini gösterebilir. β_k sabit etki parametresi için $\% 100(1-\alpha)$ güven aralığı Eş. 7.4'deki istatistik yoluyla elde edilir. Eşitsizlik 7.5'i sağlayan $\beta_{k(0)}$ değerleri asimptotik yapay-skor güven aralığını oluşturur;

$$X^2(\beta_{k(0)}) = \sum_i \frac{(\hat{Y}_i - \hat{Y}_{i0})^2}{var(\hat{Y}_{i0})} < \chi^2_{1,\alpha}. \quad (7.5)$$

7.5. Parametrik Bootstrap Güven Aralığı Metodu

Bootstrap yaklaşımı ilk olarak Efron (1979) tarafından bilimsel literatüre kazandırılmıştır. Chernick (2008) kitabında yeni geliştirilen bootstrap tekniklerini ve uygulamalarını çok geniş tarihsel notlarla birlikte geçmişten bugüne yapılan birçok bootstrap ile ilgili çalışmanın referanslarını bilimsel literatüre sunmuştur.

Bilinen bir F dağılımının, bilinmeyen parametrelerinin en çok olabilirlik tahminleri gözlenen verilerden hesaplanabilir. Bu tahminlere bağlı olarak üretilen bootstrap setlerinden elde edilen F 'in tahmini dağılımı deneysel (empirical) dağılımı olarak bilinir. Bu tezde, şekli bilinen fakat parametreleri bilinmeyen normal dağılıma dayalı parametrik bootstrap (PBS) üzerinde durulacaktır. Burada, F bilinmeyen ortalamalı ve varyanslı normal dağılım olarak varsayılın, F 'in deneysel dağılımı örnekten elde edilmiş ortalamanın etrafında konumlanmış ve yine örnekten elde edilmiş en çok olabilirlik tahmin edicisine eşit olan varyansa sahip normal dağılım olacaktır. Parametrik bootstrap durumunda, bootstrap (BS) örneklemi deneysel dağılımdan elde edilmiş rastgele örneklem olacaktır. X^* bir bootstrap örneği olarak gösterilsin. T ilgilenilen parameter ve onun EÇO veya KEÇO tahmini t olsun. Her bir bootstrap örneği için, elde edilen bootstrap tahminine t^* denilsin. Çok büyük sayıda bootstrap örnekleri alınarak ve bootstrap

örneklerinden hesaplanmış t^* 'nin gözlenen değerlerinin belirlenen aralığa düşenlerinin oranının olasılığı bulunarak, t^* istatistiğinin örnekleme dağılımı oluşturulur. t^* 'nin bootstrap dağılımı ilgilenilen istatistik t 'nin örnekleme dağılımının tahmini dağılımı olarak kullanılabilir.

PBS metodunun standard analitik metodlarda (en çok olabilirlik tabanlı metodlar gibi) olduğu gibi DKE modellerinde kullanırken rasgele terimlerin normal dağılıma sahip olması varsayımı gibi bir kısıtlaması vardır. Bootstrap aralık yaklaşımı altında güven sınırları, t veya z dağılımının kesim noktalarına (cutoffs) değil bootstrap dağılımının kesim noktalarına bağlı olacaktır. PBS metodunun aralık oluşturmaya yönelik uygulamaları DKE modellerinde yaygın değildir. Carpenter, Goldstein ve Rasbash (2003), PBS ve parametrik olmayan artık bootstrap metodlarını normal olmayan hata terimleri ve normal olmayan rasgele etkilere sahip rasgele katsayı modeli altında bootstrap örnekleri üreterek Efron yüzdelik aralıklarını sabit etkiler için oluşturmuşlardır. Chatterjee ve diğerleri (2008) c sabit değerli bir vektör olmak üzere $T = c'(X\beta + Zb)$ formunda bir karma etkinin güven aralığını kurmak için PBS metodu geliştirmiştir. Chatterjee ve diğerlerinin (2008) karma etki için kullandığı parametrik bootstrap-t metodunu modifiye eden Staggs (2009) yapay pivotu DKE modelinin sabit etki parametresi için uyarlamıştır.

DKE modellerinde parametrik bootstrap metodunun algoritması aşağıdaki gibidir (Staggs, 2009);

- b ve e 'nin normal dağıldığı varsayımı altında ayrıca X ve Z sabit değerli tasarım matrisleri kullanılarak parametrik bootstrap gerçekleştirilir. b ve e rasgele etki ve hata terimlerini üretebilmek için orjinal veri setinden varyans ve kovaryans parametrelerinin (EÇO veya KEÇO) tahminleri elde edilir.
- b 'nin bootstrap rasgele etki vektörü b^* grup veya deney birimi sayısı (n adet) kadar sıfır ortalamalı ve tahmini kovaryans matrisli normal dağılımdan üretilir.
- Her grup veya deney birimi için ayrıca m 'er (gruptaki gözlem sayısı veya ölçüm sayısı) adet rasgele hata vektörü sıfır ortalamalı ve tahmini kovaryans matrisli normal dağılımdan üretilir.
- e^* , $m \times n$ boyutlu rasgele hata vektörü; üretilen her bir deney birimi için bootstrap hata vektörleri bir araya getirilerek oluşturulur.

- $\hat{\beta}_{DEIDYTE}$ sabit etkinin orijinal veriden elde edilen KEÇO tahmini olmak üzere böylece bootstrap gözlem vektörü Y^* , $Y^* = X\hat{\beta}_{DEIDYTE} + Zb^* + e^*$ eşitliği kullanılarak oluşturulur.
- Bu bootstrap prosedürü $N = n \cdot m$ (dengeli veri seti durumu için) boyutlu istenen sayıda Y^* bootstrap gözlem vektörü oluşturulana kadar yukarıdaki adımlar tekrar edilir.
- Her bootstrap veri seti için orijinal veri üzerinde kullanılan parametre tahmin yöntemi kullanılarak sabit etki parametre tahminleri elde edilir.
- Bootstrap örneklerinden elde edilen sabit etki tahminlerinin örnekleme dağılımı bootstrap dağılımından elde edilmiş olur. Böylece tezde ilgilenilen sabit etki parametresi için bootstrap parametre tahmin vektörü, orijinal veriden elde edilmiş sabit etki tahmini ve bootstrap sabit etki tahminlerinin standart hataları kullanılarak yapay-pivot oluşturulur. Yapay pivotun dağılımında ilgili yüzdelikler kullanılarak güven aralığı kurulumu gerçekleştirilebilir.

Literatürde Efron'un yüzdelik aralığı (Efron ve Tibshirani, 1993), Hall'ün yüzdelik aralığı (1992), yanlış tahminler için yanlış düzeltilmiş aralık ve bootstrap-t aralığı gibi bootstrap yöntemleri mevcuttur. İncelenen literatürde genellikle bootstrap-t'nin daha iyi sonuç verdiği görülmüştür. Bu nedenle tezde parametrik bootstrap metodu olarak bootstrap-t metodu ele alınmıştır.

7.5.1. Bootstrap-t güven aralığı

Bootstrap-t (veya yüzdelik-t) aralığı DKE modellerinde sabit etki parametresi için $T = c'(X\beta)$ ve t KEÇO/ EÇO tahmini olmak üzere $(t - T)/s_t$ yapay-pivotunun örnekleme dağılımının BS tahminine göre oluşturulur. Bootstrap tahmini t^* ve t^* 'nin standart hata tahmini s_t^* için $q^* = (t^* - T)/s_t^*$ BS pivot yaklaşımı olarak gösterilsin. Eş. 7.6'daki olasılık bu şekilde varsayılırsa,

$$1 - \alpha \approx P(q_{\alpha/2}^* < (t^* - T)/s_t^* < q_{1-\alpha/2}^*) \approx P(q_{\alpha/2}^* < (t - T)/s_t < q_{1-\alpha/2}^*) \quad (7.6)$$

Eş. 7.6.'daki olasılık dikkate alınarak T için yaklaşık $1 - \alpha$ güven aralığı,

$$1 - \alpha \approx (t - q_{1-\alpha/2}^* s_t^*, t - q_{\alpha/2}^* s_t^*) \text{ olur.}$$

Burada, $q^* = (t^* - T) / s_t^*$ yapay pivotunun bootstrap dağılımı için α kesim noktası q_α^* olarak gösterilir. Bootstrap-t aralığı, $(t^* - T) / s_t^*$ 'yi pivot olarak kullandığı için diğer parametrik bootstrap güven aralığı yöntemlerinden daha iyi performans gösterir. Çünkü $(t^* - T) / s_t^*$ parametre tahminlerine daha az bağlı olduğu için pivot gibi davranır.

DKE modellerinde β_k sabit etki parametresi için DEİDYTE, KEÇO tahmini $\hat{\beta}_k$ ve standart hata $SH(\hat{\beta}_k)$ olmak üzere bootstrap-t metodunun yapay-pivotu $q^* = (\hat{\beta}_k^* - \hat{\beta}_k) / SH(\hat{\beta}_k^*)$ olarak yazılır. Tahmin edicinin standart hatası olarak Wald güven aralığı yönteminde bahsedilen gözlenen bilgi matrisinden (Hessian matrisi yoluyla elde edilmiştir) elde edilen standart hata ile bölüm 4.6.2'de verilen Kenward-Roger (1997) standart hatası kullanılmıştır. Böylece geleneksel (naive) yöntem ve Kenward-Roger yöntemi simülasyon çalışmasında karşılaştırılmıştır. Aslında bootstrap-tabanlı metodlardan biri $SH(\hat{\beta}_k)$ 'nin tahmininde kullanabilir ama bu ekstra bir hesaplama zamanı gerektirir.

Yapay-pivotun paydasında kullanılan $SH(\hat{\beta}_k^*)$ 'nin tahminin kalitesi bootstrap-t yaklaşımının performansını belirler. β_k sabit etki parametresi için yaklaşık $1 - \alpha$ güven aralığı aşağıdaki gibi verilir,

$$1 - \alpha \approx \left(\hat{\beta}_k - q_{1-\alpha/2}^* SH(\hat{\beta}_k), \hat{\beta}_k - q_{\alpha/2}^* SH(\hat{\beta}_k) \right).$$

8. SİMULASYON ÇALIŞMALARI VE VERİ ANALİZLERİ

Bu bölümde DKE modellerinin bazı uygulamalarına ve veri analizlerine yer verilmiştir. Uzun vadeli bir çalışmadan toplanmış veri seti SAS ve R programları ile analiz edilmiştir. Yine tekrarlı ölçümlerden oluşmuş depresyon hastalığı ile ilgili bir veri seti R programı ile analizi gerçekleştirilerek 7. bölümde ele alınan güven aralığı metotları modeldeki sabit etkiler için oluşturulmuştur. Ayrıca iki farklı DKE modeli ele alınarak Bölüm 7’de verilen güven aralığı metotları bu modellerin sabit etki parametrelerini kapsama olasılığı bakımından performansları incelenmiştir. Bölüm 3’de tanıtılan iç içe hata regresyon modeli ile rasgele eğim modelinin sabit etki parametreleri ele alınmıştır.

8.1. Uygulama: 6 Şehir Çalışması

Model parametrelerinin yorumlarına eğilmek için bu bölümde gerçek bir veri setinin analizine yer verilmiştir. Sağlık ve hava kirliliği kurumu yetişkinlerde ve çocuklarda akciğer fonksiyonundaki değişimlere göre akciğer gelişimini karakterize etmek ve onu etkileyen faktörleri belirlemek için uzun vadeli bir çalışma tasarlanmıştır (Dockery ve diğerleri, 1993). 1967’den sonra doğmuş Amerika Birleşik Devletleri’ndeki 6 farklı eyaletin şehirlerinden 13379 çocuk çalışmaya kaydedilerek izlenmiştir. Birçok çocuk çalışmaya ilkökul birinci ve ikinci sınıfta (6 ve 7 yaşları arasında) kaydedilmiştir ve ölçümler liseden mezun olana kadar veya izleme kaybedilene kadar senelik olarak elde edilmiştir. Her senelik incelemede, akciğer fonksiyonu ölçümü spirometri testi yapıp çocuğun resmi ailesi veya koruyucusuna bir solunum sağlık anketi doldurtulmuştur. Basit spirometrideki temel manevra kapalı bir konteynıra olabildiğince hızlı nefes vermeyi izleyen maksimum nefes almadır. Birçok farklı ölçüm çeşidi olsa da bu çalışmada manevradaki ilk saniyede verilen nefesin hacmi ölçüm olarak alınmıştır (FEV1).

Bu bölümde 6 şehirden toplanan verilerin bir bölümü analiz edilecektir. Veri seti Topeka, Kansas’da yaşayan rasgele seçilen kız çocuklarının uzunluk, yaş, FEV1 değerlerini içermektedir (Fitzmaurice ve diğerleri, 2004: 210-216). Rasgele örnek, 300 kız çocuğunun zaman içinde en az bir en fazla on iki gözleminden oluşmuştur. Sadece bir gözleme sahip olan kız çocuklarının uzun vadede değişime veya zaman içinde birey içi değişime doğrudan bir katkısı olmasa da bu ölçümler bireyler arası etkilere ve varyanslara katkısı

olduğu için çalışmaya alınmıştır. Veri seti, Prof. Dr. Garrett Fitzmaurice'in web sitesinden alınmıştır (<https://www.hsph.harvard.edu/fitzmaur/ala/>).

Zaman içinde bu veri seti dengeli değildir her bireyden farklı sayıda ölçüm vardır. Ayrıca çalışmaya her çocuk farklı yaşlarda girmiştir. Yaş zaman için değişim ölçüsü olarak kullanıldığında, yaş ve FEV1 arasında iki bilgi ortaya çıkacaktır. Birincisi farklı yaşlarda çalışmaya alınmış çocuklardan kaynaklanan bilgi veya kesitsel bilgidir. İkincisi çocuklardan zaman içinde tekrarlı ölçümlerinden ortaya çıkan deney birimi içinde farklı zamanlardaki FEV1 ölçümleridir. FEV1 ve yaş arasında iki potansiyel çakışan bilgi kaynağı olduğu için, kesitsel ve uzun vadeli yaştan FEV1 üzerindeki etkilerin farklı tahminlerine olanak sağlayacak biçimde veriyi modellemek önemlidir.

6 şehir çalışması 6 ve 18 yaşındaki çocukların akciğer fonksiyonlarının büyümesini karakterize edebilmek için tasarlanmıştır. Fitzmaurice ve diğerleri (2004)'nin analizine benzer olarak veri setinin analizi SAS yazılımında gerçekleştirilmiştir. Zaman içinde alınan FEV1 ölçümü ile belirlenen akciğer fonksiyonu gelişimi yaş ve boy uzunluğuna bağlı olarak nasıl değiştiği SAS yazılımının DKE modeli için kullanılan Proc Mixed fonksiyonu yardımı ile analiz edilmiştir. FEV1 ölçümünün logaritmik dönüşümünün yaş ve uzunluğun logaritmik dönüşümü ile doğrusal bir bağlantısı olduğu daha önceki çalışmalarda ifade edilmiştir. Bu veri seti dengeli olmadığından rasgele etki yapısı ile aynı çocuktan alınan tekrarlı ölçümlerin kovaryansı modellenenmiştir. Analizde ayrıca yaş değişkeninin kesim ve eğiminin çocuktan çocuğa rasgele değişmesine olanak sağlanmıştır. $\log(FEV1)$ bağımlı değişkenine aşağıdaki model uydurulmuştur,

$$E(Y_{ij} | b_i) = \beta_1 + \beta_2 Yaş_{ij} + \beta_3 \log(boy)_{ij} + \beta_4 Yaş_{i1} + \beta_5 \log(Uz)_{i1} + b_{1i} + b_{2i} Yaş_{ij}.$$

Bu modelde Y_{ij} i . çocuğun j . ölçümdeki $\log(FEV1)$ değeridir. $Yaş_{i1}$ ve $\log(boy)_{i1}$ i . çocuğun çalışmaya alındığındaki yaş ve boy uzunluğunun logaritmik değeridir. Alınan rasgele 300 kız çocuğunun birinin tek bir ölçümü olması nedeniyle veri setinden çıkarılmıştır. Yani analize 299 deney birimi (kız çocukları) ve toplamda 1993 ölçüm alınmıştır. Çizelge 8.1'de SAS Proc mixed ve R lme4 paketi ile analizinden modeldeki sabit etki parametrelerinin KEÇO tahminleri, standart hataları ve z ve t değerleri verilmiştir.

Çizelge 8.1. Sabit etkilerin KEÇO tahminleri ve standart hataları

	Değişken	Tahmin	Standart Hata	Z-değ
SAS	Kesim	-0,2883	0,0387	-7,45
	Yaş	0,0235	0,0014	16,86
	Log(boy)	2,2372	0,0435	51,39
	İlk Yaş	-0,0165	0,0075	-2,21
	İlk log(Boy)	0,2182	0,1455	1,50
R				t-değ
	Kesim	-0.2690	0.0424	-6.34
	Yaş	0.0235	0.0014	16.79
	Log(boy)	2.2400	0.0437	51.26
	İlk Yaş	-0.0238	0.0081	-2.93
	İlk log(boy)	0.3714	0.1585	2.34

β_2 ve β_3 yaş ve boy uzunluğunun “uzun vadedeki” etkileridir. $(\beta_2 + \beta_4)$ ve $(\beta_3 + \beta_5)$ yaş ve boy uzunluğunun “kesitsel” etkileridir. β_4 ve β_5 log(boy) ve yaşı kesitsel ve uzun vadeli etkileri arasındaki farkı temsil eder. β_4 ve β_5 ’in tahminlerinin standart hatalarına ilişkin göreceli büyüklükleri yaşı uzun vadeli ve kesitsel etkisi arasında önemli bir fark olduğunu göstermiştir (0.0235 uzun vadeli etkisi, 0.007 kesitsel etkisi). log(boy)’un kesitsel ve uzun vadeli etkileri benzerdir (2.46 ve 2.24). Yaşı log(FEV1) üzerinde uzun vadeli ve kesitsel etkileri ($e^{0.024} \approx 1.025$ ve $e^{0.007} \approx 1.007$) oldukça farklıdır. Bunun nedeni çalışma boyunca yaşı değişimine göre çalışmaya alınma yaşı benzer olmasıdır. Böylece yaşı kesitsel etkisi tam olarak tahmin edilememiştir. Uzun vadeli yaş ve log(boy) değişkenlerinin log(FEV1) değişimi üzerinde etkisi vardır.

Deney birimine özel rasgele etkilerin dağılımı üzerinde ortalananmış ortalamanın modeli aşağıdaki gibidir;

$$E(Y_{ij}) = \beta_1 + \beta_2 Yaş_{ij} + \beta_3 \log(boy)_{ij} + \beta_4 Yaş_{i1} + \beta_5 \log(boy)_{i1}.$$

Ayrıca bu model iki model bakımından yeniden ifade edilebilir, kesitsel ve uzun vadeli model. İlki aşağıdaki gibidir;

$$\begin{aligned} E(Y_{i1}) &= \beta_1 + \beta_2 Yaş_{i1} + \beta_3 \log(boy)_{i1} + \beta_4 Yaş_{i1} + \beta_5 \log(boy)_{i1}, \\ &= \beta_1 + (\beta_2 + \beta_4) Yaş_{i1} + (\beta_3 + \beta_5) \log(boy)_{i1} \end{aligned}$$

ve uzun vadeli model aşağıdaki gibidir;

$$\begin{aligned}
E(Y_{ij} - Y_{i1}) &= \beta_1 + \beta_2 Ya_{ij} + \beta_3 \log(boy)_{ij} + \beta_4 Ya_{i1} + \beta_5 \log(boy)_{i1} \\
&\quad - \{\beta_1 + \beta_2 Ya_{i1} + \beta_3 \log(boy)_{i1} + \beta_4 Ya_{i1} + \beta_5 \log(boy)_{i1}\} \\
&= \beta_2 (Ya_{ij} - Ya_{i1}) + \beta_3 (\log(boy)_{ij} - \log(boy)_{i1})
\end{aligned}$$

Belirli bir yaş aralığında değişim (örneğin iki yıllık zaman aralığı) için $\log(boy)$ 'taki tek birimlik artışın $\log(FEV1)$ ortalamasındaki değişimi β_3 uzun vadeli $\log(boy)$ etkisi ifade eder. Benzer şekilde, $\log(boy)$ 'daki belirli bir değişim için yaştaki bir yıllık artışın $\log(FEV1)$ ortalamasındaki değişimini β_2 uzun vadeli yaş etkisi ifade eder. $\log(boy)$ için katsayı 2.24 doğrudan yorumlanması yanlış olur çünkü $\log(boy)$ 'ta olan tek birimlik artış boydaki neredeyse 3 kat artışa denk gelir ($e^{1.0} \approx 2.7$). Bunun yerine boydaki %10'luk artışın etkisi bakımından yorum yapmak daha doğru olur. Logaritmik ölçekte, $\log(boy)$ için bu 0.1'lik artışa denk gelir, böylece $e^{0.1} \approx 1.1$. Sonuç olarak, boydaki %10'luk artış ($\log(boy)$ 'ta 0.1'lik artışa denk gelir) $\log(FEV1)$ 'de 0.224 artışla ilgilidir. $\log(FEV1)$ 'deki 0.224'lük artış $FEV1$ 'de %25'lik artışa denk gelir ($e^{0.224} \approx 1.25$). Bir diğer taraftan, yaştan katsayısı, 0.024, doğrudan yorumlanabilir. Yaşın uzun vadeli etkisinin tahmini boydaki sabit bir değişim için yaştaki bir yıllık artış $\log(FEV1)$ 'de 0.024'lük artışa denk gelir ya da $FEV1$ 'de yaklaşık %2.5'luk bir artışa denk gelir ($e^{0.024} \approx 1.025$).

Aşağıdaki tabloda 7-18 yaş arasında çocukların $\log(FEV1)$ tekrarlı ölçümleri arasında tahmini marjinal korelasyonlarını gösteren Çizelge 8.2'de 12x12 boyutlu korelasyon matrisi verilmiştir.

Çizelge 8.2. $\log(FEV1)$ tekrarlanan ölçümlerinin korelasyon matrisi

1,00	0,70	0,69	0,68	0,67	0,66	0,64	0,62	0,60	0,58	0,56	0,54
0,70	1,00	0,7	0,69	0,68	0,67	0,66	0,65	0,63	0,61	0,60	0,58
0,69	0,70	1,00	0,70	0,70	0,69	0,68	0,67	0,66	0,64	0,63	0,61
0,68	0,69	0,70	1,00	0,70	0,70	0,70	0,69	0,68	0,67	0,66	0,64
0,67	0,68	0,70	0,70	1,00	0,71	0,71	0,70	0,70	0,69	0,68	0,67
0,66	0,67	0,69	0,70	0,71	1,00	0,72	0,72	0,71	0,71	0,70	0,70
0,64	0,66	0,68	0,70	0,71	0,72	1,00	0,73	0,73	0,73	0,72	0,72
0,62	0,65	0,67	0,69	0,70	0,72	0,73	1,00	0,74	0,74	0,74	0,74
0,60	0,63	0,66	0,68	0,70	0,71	0,73	0,74	1,00	0,75	0,75	0,75
0,58	0,61	0,64	0,67	0,69	0,71	0,73	0,74	0,75	1,00	0,76	0,76
0,56	0,60	0,63	0,66	0,68	0,70	0,72	0,74	0,75	0,76	1,00	0,77
0,54	0,58	0,61	0,64	0,67	0,70	0,72	0,74	0,75	0,76	0,77	1,00

Çizelge 8.2'deki korelasyon matrisi $\log(FEV1)$ tekrarlı ölçümleri arasında güçlü bir korelasyon olduğunu ve ölçümler arası zaman aralığı arttıkça korelasyonun az bir şekilde düştüğünü göstermiştir (Fitzmaurice ve diğerleri, 2004: 215).

Çizelge 8.3. Rasgele etkilerin varyans, kovaryans ve standart hata tahminleri

Değişken	Tahmin	Standart Hata (x100)	Z
$Var(b_{1i})$	1,2207	0,1924	6,34
$Kov(b_{1i}, b_{2i})$	-0,0435	0,0122	-3,55
$Var(b_{2i})$	0,0050	0,0010	5,11
$Var(e_i) = \sigma^2$	0,3629	0,0133	27,21

8.2. İç-içe Hata Regresyon Modeli için Simülasyon Çalışması

Tek faktörlü rastgele etki modelinden $Y_{ij} = \mu + \alpha_i + e_{ij}$ üretilen veri setleri oluşturulduktan sonra $Y_{ij} = \beta_0 + \beta_1 x_{ij} + b_i + e_{ij}$ modeline uydurulmuştur. Modelin parametre vektörü $\theta = (\beta_0, \beta_1, \sigma_b^2, \sigma_e^2, \rho)$ şeklindedir. Bu simülasyon deneyinde dikkate alınan parametre β_1 sabit etki parametresi olacaktır. % 95'lik eşit kuyruklu güven aralıkları sabit etki parametresi için oluşturulacaktır. Bu sabit etki parametresinin gerçek değerini kapsama olasılıkları bakımından Wald-tipi güven aralığı metodları, profil olabilirlik, Bootstrap-t ve Kenward–Roger standart hatasına dayalı Bootstrap-t güven aralığı metodları karşılaştırılacaktır. İlgilenilen parametre için kapsama özelliğini etkileyen bir takım parametreler olacağı açıktır. Bu nedenle modeldeki parametrelerin ve örnek çaplarının farklı kombinasyonları dikkate alınmıştır.

Grup içi korelasyon terimi $\rho = 0.25$ ve $\rho = 0.5$ dikkate alınmıştır. $\rho = 0.25$ için modeldeki rastgele etki terimi ve hata terimi sırasıyla $\alpha \sim N(0,1)$ ve $e \sim N(0,3)$ olarak üretilmiştir. $\rho = 0.5$ için ise modeldeki rastgele etki terimi ve hata terimi $\alpha \sim N(0,3)$ ve $e \sim N(0,3)$ parametrelili normal dağılımdan üretilmiştir. Deney birimi (sınıf/ grup) sayısı $n = 2, 4, 6, 10, 12, 16, 20$ olarak belirlenmiştir. Her grup için tekerrür (grup içindeki gözlem sayısı) $m = 3, 9, 15$ olarak seçilmiştir. β_0 değerleri $\beta_0 \sim Tekdüze(-10,10)$ parametrelili tek

düze dağılımdan üretilmiştir. Sabit etki terimi β_1 için $x_{ij} \sim N(0,100)$ değerleri normal dağılımdan üretilmiştir. Her üretilmiş veri seti $Y_{ij} = \beta_0 + \beta_1 x_{ij} + b_i + e_{ij}$ iç içe hata regresyon modeline uygulanmıştır. x_{ij} 'ler tek faktörlü rastgele etki modelinden bağımsız olarak üretilmiştir. Simülasyon çalışmasında dikkate alınan kombinasyonlar Çizelge 8.4'de özetlenmiştir.

Çizelge 8.4. Simülasyon şeması

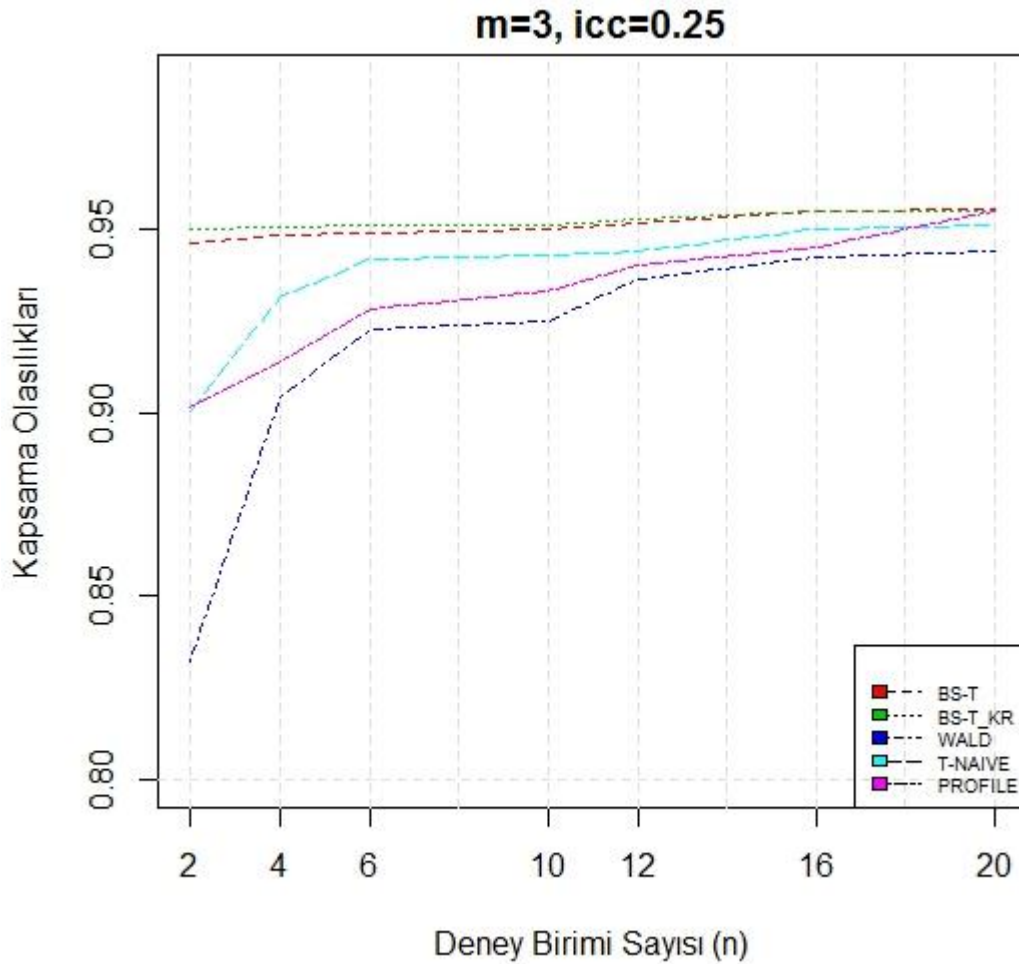
$n = 2, 4, 6, 10, 12, 16, 20$ ve $m = 3, 9, 15$		
$\beta_0 \sim \text{Tekdüze}(-10,10)$ $x_{ij} \sim N(0,100)$	$\gamma = \sigma_\alpha^2 / \sigma_e^2 = 0,333$	$\gamma = \sigma_\alpha^2 / \sigma_e^2 = 1$
	$\rho = \sigma_\alpha^2 / (\sigma_\alpha^2 + \sigma_e^2) = 0,25$	$\rho = \sigma_\alpha^2 / (\sigma_\alpha^2 + \sigma_e^2) = 0,5$
$\left(\begin{array}{l} \alpha \sim N(0, \sigma_\alpha^2 = 1) \\ \alpha \sim N(0, \sigma_\alpha^2 = 3) \end{array} \right)$	$\sigma_\alpha^2 = 1$	$\sigma_\alpha^2 = 3$
$(e \sim N(0, \sigma_e^2 = 3))$	$\sigma_e^2 = 3$	$\sigma_e^2 = 3$

R istatistik yazılımının bir paketi olan lme4 paketi (Bates ve diğerleri, 2014) karışık etkili modellerin tahminleri ve testleri için oluşturulmuş bir R kütüphanesidir. Bu kütüphane gözlenen bilgi matrisinden elde ettiği standart hatayı kullanarak Wald güven aralığını oluşturmuştur. Bu tezde doğrusal karışık-etkili model parametrelerinin EÇO ve KEÇO tahminleri ve Wald güven aralığını elde etmek için bu paket kullanılmıştır. Ayrıca sabit etki teriminin Kenward-Roger standart hatası pbkrtest paketinden (Halekoh ve Højsgaard (2014)) elde edilerek bootstrap-t-KR güven aralığını oluşturmak için kullanılmıştır. Profil olabilirlik, bootstrap-t güven aralıkları güven aralığı metotları için kodlar fonksiyon biçiminde R programı kullanılarak oluşturulmuştur. Aşağıda simülasyon çalışmasının algoritması verilmiştir;

1. Her (n, m, ρ) kombinasyonu için tek faktörlü rastgele etki modeline uygun veri setleri üretilir.
2. Profil olabilirlik, Wald, t-naive güven aralıkları Bölüm 7.1, 7.2 ve 7.3'de anlatıldığı gibi oluşturulur.
3. Bootstrap-t ve Bootstrap-t-Kenward Roger güven aralıklarını elde etmek için parametrik bootstrap bölümünde anlatılan algoritmaya göre 2000 adet bootstrap örneği üretilir. İlgili standart hata değerleri ve parametre tahminleri elde edilerek 2000 adet

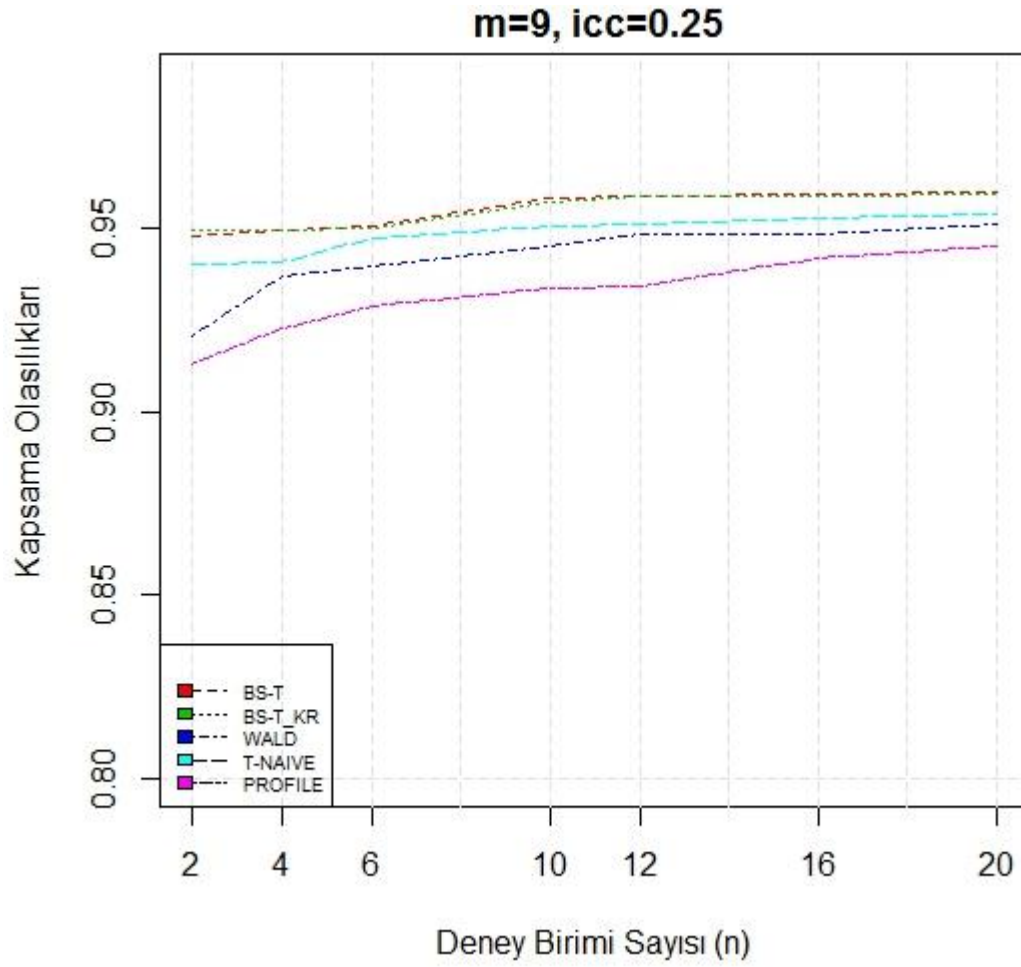
yapay-pivot oluşturulur. Yapay-pivotların bootstrap örnekleme dağılımları oluşturularak istenen yüzdeliklerdeki değerler elde edilir. Bu yüzdelik değerleri ve ilgili standart hata değerleri kullanılarak güven aralıkları oluşturulur.

4. 2. ve 3. adımda elde edilen güven aralıkları sabit etki parametresinin gerçek değeri olan 0'ı içeriyorsa belirlenen sayaç değişkeni 1 değerini alır içermiyorsa sayaç değişkenine sıfır değeri verilir.
5. 1'den 4'e kadar olan adımlar 2000 kez tekrar edilerek toplam simülasyon sayısının toplam sayaç değişkenine oranı hesaplanarak kapsama olasılıkları her bir metot için elde edilir.



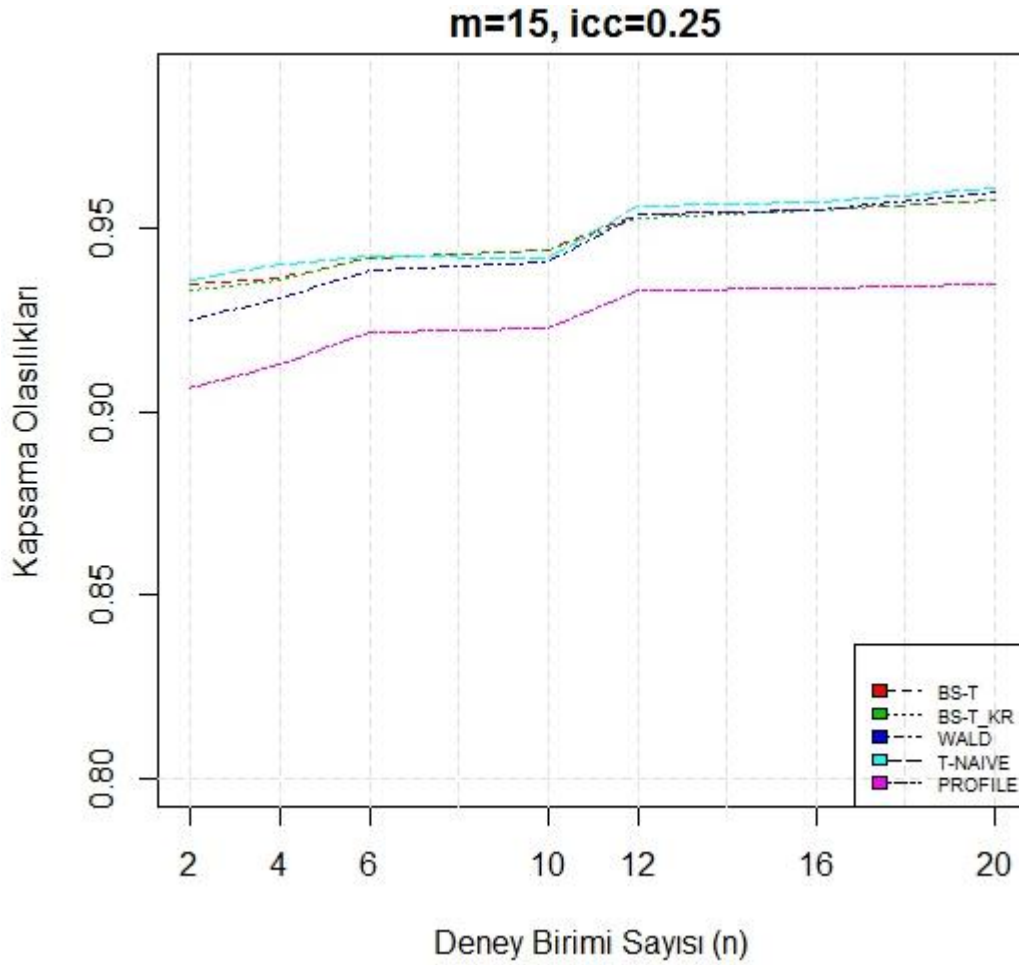
Şekil 8.1. $m = 3$ ve $\rho = 0.25$ kombinasyonu için kapsama olasılığı

Grup içi korelasyon katsayısı $\rho = 0.25$, her bir grup için tekerrür sayısı $m = 3$ durumunda Kenward Roger (BS-t-KR) standart hata terimine parametrik bootstrap metodu diğer metotlara göre en iyi kapsama olasılığına sahiptir.



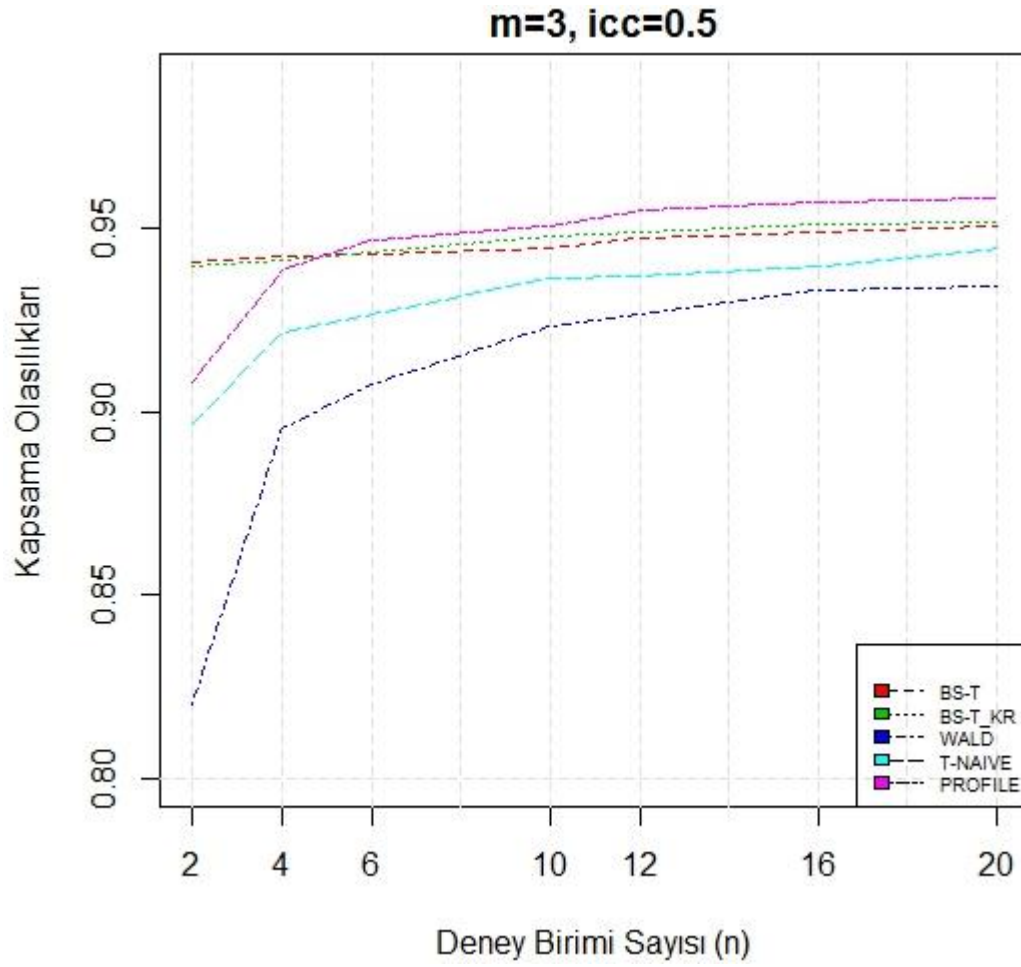
Şekil 8.2. $m = 9$ ve $\rho = 0.25$ kombinasyonu için kapsama olasılığı

Grup içi korelasyon katsayısı $\rho = 0.25$, her bir grup için tekerrür sayısı $m = 9$ durumunda BS-t-KR ve BS-t metotları birbirine yakın kapsama olasılıklarına sahiptir ve diğer metotlara göre daha yüksek kapsama olasılıkları sağlamışlardır.



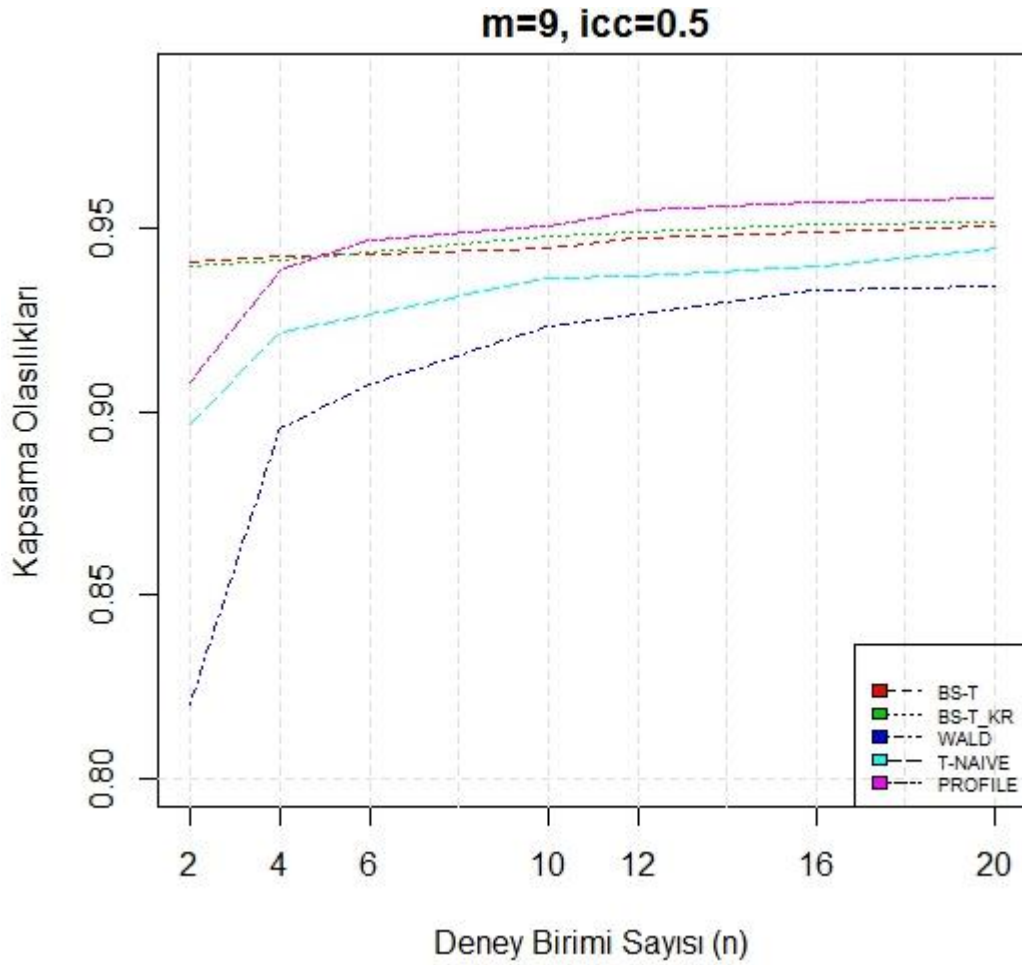
Şekil 8.3. $m=15$ ve $\rho=0.25$ kombinasyonu için kapsama olasılığı

Grup içi korelasyon katsayısı $\rho=0.25$, her bir grup için tekerrür sayısı $m=15$ durumunda küçük örnek çaplarında BS-t, BS-t-KR ve t-naive metodları birbirine yakın kapsama olasılıklarına sahiptir. Örnek çapı arttıkça $n=12,16,20$ kombinasyonları için t-dağılımının kesim noktalarına dayalı olan analitik metot (t-naive) diğerleri arasında daha yüksek kapsama olasılığını sağlamıştır.



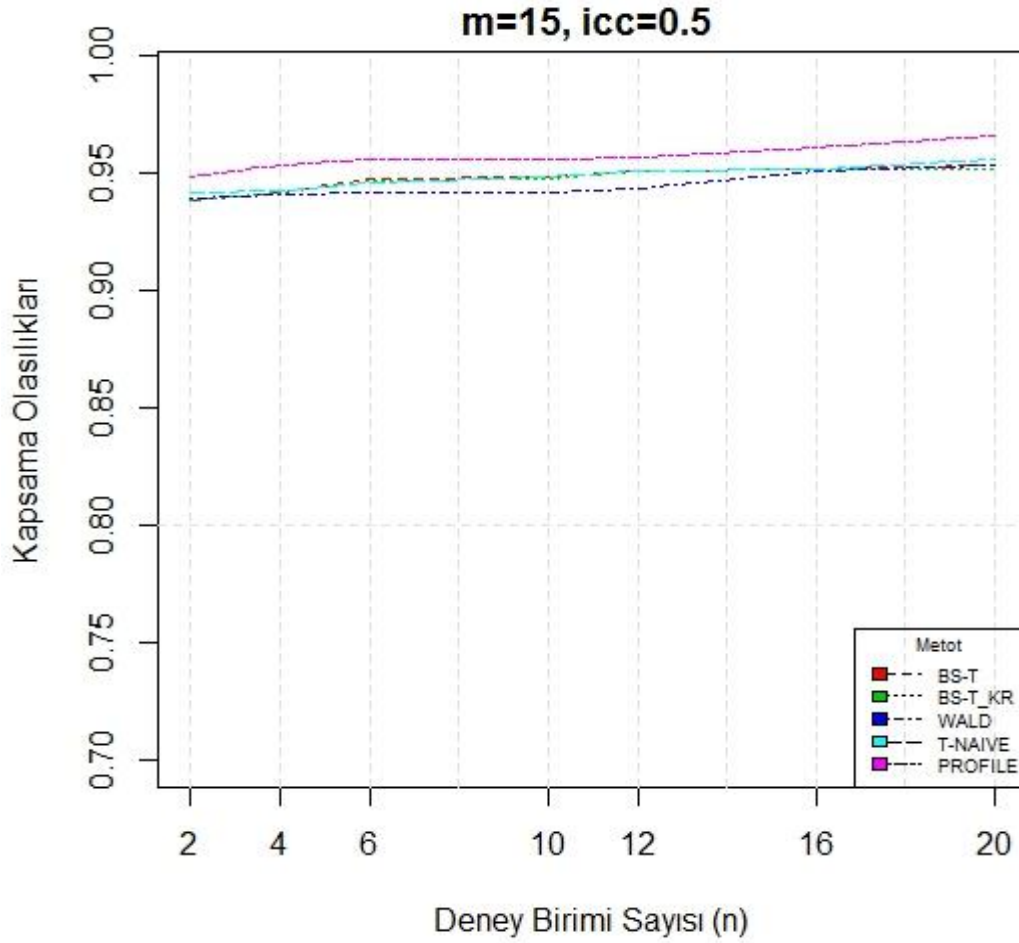
Şekil 8.4. $m=3$ ve $\rho=0.5$ kombinasyonu için kapsama olasılığı

Grup içi korelasyon katsayısı $\rho=0.5$, her bir grup için tekerrür sayısı $m=3$ durumunda $n=2, 4$ için BS-t-KR ve BS-t metotları birbirine yakın ve diğer metotlara göre daha yüksek kapsama olasılıklarına sahiptir. Orta ve büyük örnek çaplarında POGA metodunun kapsama olasılığı diğer metotlara göre daha yüksek olmuştur.



Şekil 8.5. $m=9$ ve $\rho=0.5$ kombinasyonu için kapsama olasılığı

Grup içi korelasyon katsayısı $\rho=0.5$, her bir grup için tekerrür sayısı $m=9$ durumunda $n=2, 4$ için BS-t-KR ve BS-t metotları birbirine yakın ve diğer metotlara göre daha yüksek kapsama olasılıklarına sahiptir. Orta ve büyük örnek çaplarında POGA metodunun kapsama olasılığı diğer metotlara göre daha yüksek olmuştur.



Şekil 8.6. $m=15$ ve $\rho=0.5$ kombinasyonu için kapsama olasılığı

Grup içi korelasyon katsayısı $\rho=0.5$, her bir grup için tekerrür sayısı $m=15$ durumunda tüm n değerleri için POGA metodunun kapsama olasılığı diğer metotlara göre daha yüksektir. Grup içi tekerrür sayısı $m=15$ olduğunda Wald-tipi metodların kapsama olasılıkları yükselmiştir.

Güven aralığı, hipotez testi ilişkisi

İki yanlı, $H_0: \beta_1=0$ hipotezinin $\alpha=0.05$ önem düzeyinde testi yapmak istenirse güven aralıklarının sıfırı içermemesine göre yokluk hipotezi red edilerek hipotez test edilebilir. Çünkü bu çalışmada $\beta_1=0$ alınarak veri setleri oluşturulmuştur. Güven aralığı tabanlı hipotez testinin 1. Tip hatası sıfırı içermeyen aralıklarının sayısının toplam güven aralığı sayısına oranı olarak elde edilecektir. Ya da gözlemlenmiş kapsama olasılığının değeri 1'den çıkarılarak da aynı 1. Tip hata değeri elde edilebilir.

8.3. Rasgele Eğim Modeli için Simülasyon Çalışması

Bölüm 7’de verilen analitik güven aralıklarının performansları bağımlı değişkenin farklı kovaryans yapılarına sahip olduğu durumlar için simülasyon çalışmasında geniş olarak incelenecektir. Simülasyon çalışmasında, aşağıda verilen rasgele kesim ve rasgele eğim terimi içeren DKE modeli ele alınmıştır;

$$Y_{ij} = \beta_0 + \beta_1 x_{ij} + \beta_2 t_{ij} + b_{1i} + b_{2i} t_{ij} + e_{ij} \quad i = 1, \dots, n; \quad j = 1, \dots, m. \quad (8.1)$$

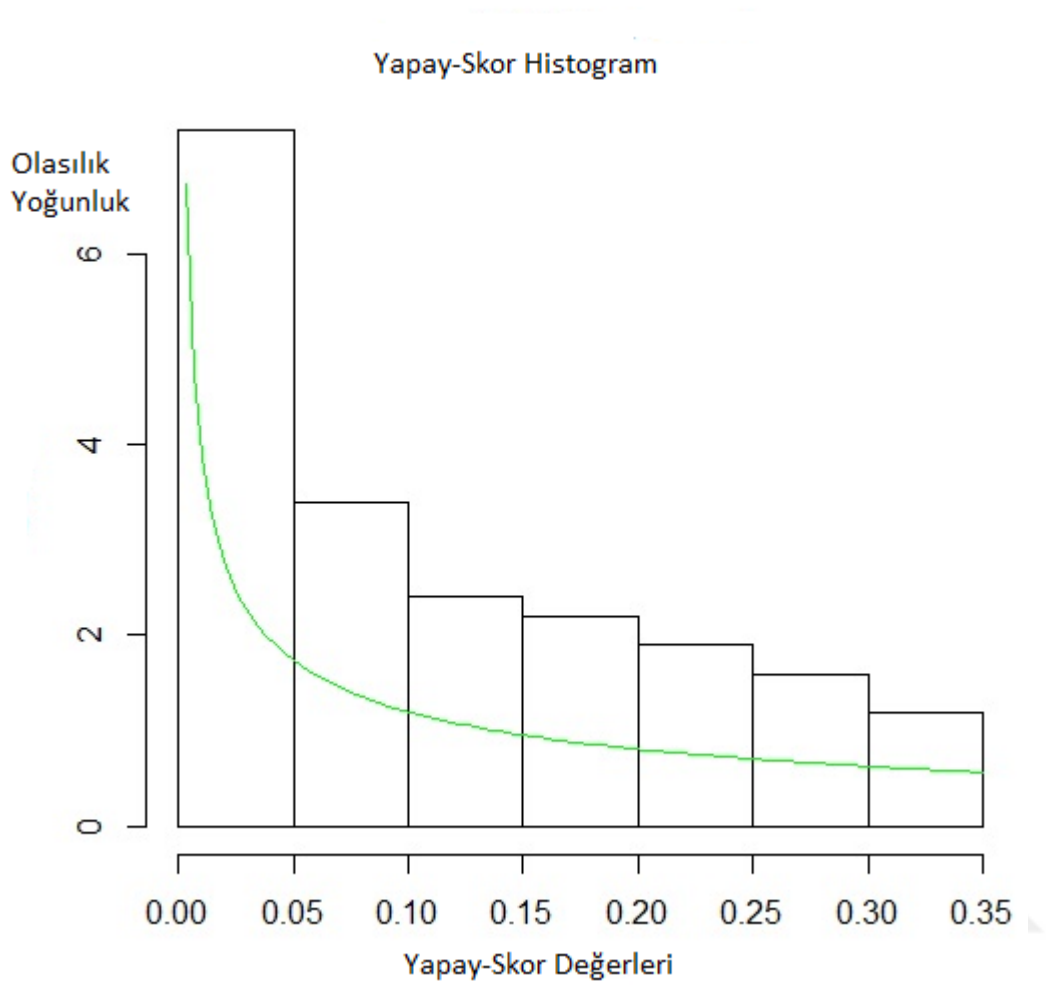
Eş. 8.1’de verilen DKE modelinde sabit etki parametreleri $\beta_0, \beta_1, \beta_2$ olarak ifade edilir. Sabit etki parametreleri kitleye özgü parametrelerdir. Rasgele etki parametreleri verideki her deney birimine özgü b_{1i} ve b_{2i} $i = 1, \dots, n$ parametreleridir. DKE modellerinde bilindiği üzere iki farklı örnek çapı kavramı bulunmaktadır. n çalışmaya alınan deney birimi/ grup/ küme sayısını ifade eder. m deney birimlerinden alınan gözlemler veya ölçümlerin sayısını ifade eder. Simülasyona alınan modelde iki rasgele etki olduğu için ilgili kovaryans matrisi 2×2 boyutlu ve bütün simülasyon durumlarında yapısız

$$G = \text{Var}(b_i) = \begin{bmatrix} \sigma_{b_1}^2 & \sigma_{b_1, b_2} \\ \sigma_{b_2, b_1} & \sigma_{b_2}^2 \end{bmatrix} \text{ olarak ele alınacaktır. Aynı ortalama modeli ve aynı rasgele}$$

etki kovaryans yapısı altında, hata terimlerinin kovaryans yapıları R_i değiştirilerek bağımlı değişkenin kovaryans yapısı $V(Y_i)$ için birçok farklı kovaryans modeli oluşturulmuştur.

β_1 sabit etki parametresi için simülasyon çalışmasında güven aralıkları oluşturulacaktır. Wald, t-naive ve yapay-skor güven aralığı metodlarında β_1 sabit etki parametresi için KEÇO tahmini kullanılacaktır. POGA metodu için EÇO tahmin edicisi kullanılmıştır. β_1 sabit etki parametresi için EÇO ve KEÇO tahminleri, Wald güven aralıkları, bağımlı değişkenin klasik varyans-kovaryans tahmini lme4 R paketiyle elde edilmiştir. Kenward ve Roger yaklaşımına dayalı bağımlı değişkenin varyans-kovaryans tahmini pbkrtest R paketiyle elde edilmiştir.

Ayrıca yapay-skor değerlerinin bazı simülasyon durumlarındaki asimptotik dağılımları incelenmiştir. Örnek olarak $n=5, m=20$ için bantlı kovaryans yapısında elde edilen yapay-skor değerlerinin histogramı verilmiştir. Yapay-skor değerlerinin örnekleme dağılımının 1 serbestlik dereceli Ki-Kare dağılımına (eğri) uyduğu görülebilir;



Şekil 8.7. $n=5$, $m=20$ için bantlı kovaryans yapısında yapay-skor değerleri histogramı

Tüm simülasyon durumlarında kullanılan parametre yapısı Çizelge 8.5’de özetlenmiştir.

Çizelge 8.5. Simülasyon şeması

$n=3,5,10$	$m=5,10,20$	
$\beta_0 \sim \text{Tekdüze}(3,5)$		
$\beta_1 = 3$		
$\beta_2 = 0.84$		
$x_{ij} \sim N(1,10)$		
		$G = \text{Var}(b_i) = \begin{bmatrix} 4.55 & 0.8183 \\ 0.8183 & 1.0238 \end{bmatrix}$
		$b_i \sim N(0, G) \quad e_i \sim N(0, R_i)$

Çizelge 8.5’deki parametre değerleri, Eş. 8.1’deki model yapısı ve varsayımlarına uygun olarak veri setleri üretilmiştir. β_1 parametresi için %95 güven düzeyinde farklı simülasyon kombinasyonlarında her metod için 5000 güven aralığı oluşturulmuştur. Güven

aralıkları metodlarının β_1 parametresinin gerçek değerini kapsama olasılığı (KO) ve güven aralıklarının ortalama uzunlukları (OU) hesaplanmıştır. Simülasyon çalışmasında alınan hata terimlerinin farklı varyans-kovaryans R_i yapıları aşağıda listelenmiştir.

Durum 1. R_i 'nin homojen varyanslı $\sigma_\varepsilon^2 = 5.84$ olduğu ve hata terimlerinin deney birimleri içinde bağımsız olduğu varsayılmıştır.

Durum 2. R_i 'nin birinci mertebeden otoregresif $AR(1)$ ve $\sigma_\varepsilon^2 = 5.84$ olduğu varsayılmıştır. Otokorelasyon katsayılarının $\rho = 0.25, 0.5, 0.75$ değerleri simülasyon çalışmasında incelenecektir. Bu yapının kullanılması için deney birimleri içinde yapılan ölçümler eşit aralıklarla yapılmalıdır bu nedenle Eş. 8.1'de kullanılan zaman vektörü örneğin $m = 20$ için,

$t_{ij} = [1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, 19, 20]$ olarak alınmıştır.

Durum 3. R_i 'nin birleşik simetrik yapıda ve $\sigma_\varepsilon^2 = 5.84$ olduğu varsayılmıştır. Deney birimi içinde hata terimlerinin zaman içinde korelasyonlarının değişmediği düşünülür.

Durum 4. R_i 'nin üstel kovaryansa sahip ve $\sigma_\varepsilon^2 = 5.84$ olduğu varsayılmıştır. Üstel yapı otoregresif yapının zaman içinde ölçümler eşit aralıklı zaman aralığında yapılmadığında genelleştirilmiş halidir. Tekrarlı ölçümler arasında zaman aralığı arttıkça korelasyon üstel olarak azalır. $\{t_{i1}, \dots, t_{im}\}$, i . deney birimi için yapılan ölçümlerin zamanlarını temsil etsin. Deney birimleri içinde hata terimlerinin ilgili korelasyonları ve kovaryansları sabit varyans varsayımı altında $\theta = -\log(\rho)$ olmak üzere ;

$$korr(e_{ij}, e_{ik}) = \rho^{|t_{ij} - t_{ik}|}, \text{ ve } kov(e_{ij}, e_{ik}) = \sigma^2 \rho^{|t_{ij} - t_{ik}|} = \sigma^2 \exp(-\theta |t_{ij} - t_{ik}|),$$

Eş. 8.1'de kullanılan zaman vektörü örneğin $m = 20$ için,

$t_{ij} = [2, 3, 5, 8, 9, 10, 13, 15, 19, 20, 24, 26, 29, 30, 33, 35, 38, 40, 41, 44]$

olarak alınmıştır.

Durum 5. R_i 'nin bantlı kovaryans yapısına sahip olduğu varsayılmıştır. Bu yapı şimdiye kadar bahsedilen kovaryans yapılarına uygulanabilir. Bu çalışmada yapısız kovaryans

yapısı seçilerek bantlı yapı uygulanmıştır. Bantlı yapıyı elde etmek için FinCovRegularization R paketi kullanılmıştır. $m=5$ and bant parametresi $c=1$ olmak üzere yapısız bantlı kovaryans matrisi örneği aşağıdaki gibidir,

$$R_i = \begin{bmatrix} 2.8 & 0.1 & 0 & 0 & 0 \\ 0.1 & 0.3 & 0.2 & 0 & 0 \\ 0 & 0.2 & 0.4 & 0.1 & 0 \\ 0 & 0 & 0.1 & 0.5 & 0.23 \\ 0 & 0 & 0 & 0.23 & 0.6 \end{bmatrix}.$$

Durum 6. R_i 'nin $AR(1)$ ve bantlı yapısız kovaryans matrislerinden oluşmuş hibrid yapıda olduğu varsayılmıştır. Farklı kovaryans modelleri bir araya getirilerek uzun vadeli çalışmalarda ortaya çıkan birçok istenmeyen durum hibrid modellerle üstesinden gelinebilir.

Durum 7. R_i 'nin üstel ve birleşik simetrik kovaryans matrislerinden oluşmuş hibrid yapıda olduğu varsayılmıştır.

lme4 R paketi POGA metodu için bir fonksiyon içerse de birçok durumda sonuç sağlamamıştır. Dolayısıyla yeni bir POGA fonksiyonu sayısal algoritma yoluyla oluşturulmuştur. Buna rağmen $n=3$ durumunda oluşturulan algoritma sonuç vermemiştir. Bunun sebebi küçük örnek çaplarında profil olabilirlik fonksiyonunun yapısı dolayısıyla algoritmada verilen alt ve üst sınır değerlerinin profil olabilirlik fonksiyonu sınırları dışında olmasıdır. Bu nedenle simülasyon sonuç tablolarında bu alanlar boş bırakılmıştır. Ayrıca yapay-skor güven aralığı için sayısal algoritmaya dayalı bir fonksiyon oluşturulmuştur.

Yukarıda ifade edilen tüm durumlar için güven aralığı metodlarının performansları %95 güven düzeyinde aşağıdaki çizelgelerde özetlenmiştir.

Çizelge 8.6. Durum 1’de β_1 ’e ilişkin KO ve OA sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8775		0,9515	0,9178	3,0141		4,2046	3,6184
3	10	0,9278		0,9575	0,9361	2,0313		2,3432	2,1213
3	20	0,9442		0,9567	0,9335	1,4083		1,4975	1,3600
5	5	0,9020	0,9820	0,9495	0,9235	2,3631	3,5884	2,8453	2,5859
5	10	0,9330	0,9780	0,9464	0,9264	1,5857	2,0412	1,6877	1,5440
5	20	0,9392	0,9647	0,9490	0,9210	1,0943	1,2837	1,1300	1,0159
10	5	0,9300	0,9792	0,9540	0,9325	1,6761	2,1986	1,8380	1,6907
10	10	0,9380	0,9622	0,9465	0,9230	1,1272	1,2862	1,1704	1,0559
10	20	0,9499	0,9588	0,9555	0,9198	0,7757	0,8209	0,7864	0,6816

$n=3,5,10$ ve $m=5$ durumları için yapay-skor metodunun kapsama olasılığı Wald metodundan yüksektir. $n=3,5,10$ ve $m=20$ durumları için yapay-skor metodunun ortalama güven aralığı uzunluğu diğer metodlardan düşüktür.

Çizelge 8.7. Durum 2’de $\rho = 0.25$ için β_1 ’e ilişkin KO ve OU sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8713		0,9018	0,8930	0,6916		0,7688	0,7526
3	10	0,9270		0,9370	0,9318	0,4751		0,4974	0,4920
3	20	0,9378		0,9423	0,9395	0,3339		0,3412	0,3391
5	5	0,9103	0,9785	0,9285	0,9283	0,5420	0,7933	0,5734	0,5716
5	10	0,9428	0,9810	0,9478	0,9473	0,3739	0,4859	0,3838	0,3834
5	20	0,9450	0,9785	0,9480	0,9480	0,2601	0,3176	0,2633	0,2633
10	5	0,9290	0,9743	0,9358	0,9355	0,3917	0,4986	0,4020	0,4020
10	10	0,9375	0,9688	0,9420	0,9420	0,2668	0,3132	0,2702	0,2702
10	20	0,9525	0,9698	0,9540	0,9540	0,1844	0,2075	0,1856	0,1856

Çizelge 8.8. Durum 2’de $\rho = 0.5$ için β_1 ’e ilişkin KO ve OU sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8703		0,9008	0,8945	0,6256		0,6954	0,6846
3	10	0,9240		0,9353	0,9300	0,4435		0,4643	0,4597

Çizelge 8.8. (devam) Durum 2’de $\rho = 0.5$ için β_1 ’e ilişkin KO ve OU sonuçları

3	20	0,9430		0,9480	0,9453	0,3207		0,3276	0,3261
5	5	0,9095	0,9745	0,9268	0,9258	0,4962	0,7234	0,5251	0,5239
5	10	0,9416	0,9774	0,9471	0,9463	0,3487	0,4566	0,3579	0,3576
5	20	0,9465	0,9818	0,9488	0,9488	0,2498	0,3089	0,2529	0,2529
10	5	0,9345	0,9748	0,9425	0,9425	0,3573	0,4552	0,3667	0,3667
10	10	0,9433	0,9733	0,9453	0,9453	0,2491	0,2950	0,2522	0,2522
10	20	0,9475	0,9705	0,9490	0,9490	0,1772	0,2017	0,1783	0,1783

Çizelge 8.9. Durum 2’de $\rho = 0.75$ için β_1 ’e ilişkin KO ve OU sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8580		0,9000	0,8880	0,5425		0,6031	0,5918
3	10	0,9344		0,9457	0,9447	0,3885		0,4067	0,4043
3	20	0,9341		0,9381	0,9361	0,2891		0,2953	0,2949
5	5	0,8948	0,9724	0,9129	0,9122	0,4315	0,6249	0,4566	0,4563
5	10	0,9471	0,9844	0,9507	0,9495	0,3052	0,3994	0,3133	0,3132
5	20	0,9475	0,9795	0,9503	0,9503	0,2260	0,2831	0,2289	0,2289
10	5	0,9285	0,9715	0,9333	0,9333	0,3113	0,3935	0,3195	0,3195
10	10	0,9396	0,9691	0,9440	0,9440	0,2163	0,2576	0,2191	0,2191
10	20	0,9479	0,9717	0,9495	0,9495	0,1603	0,1862	0,1613	0,1613

Durum 2’de otoregresif yapının farklı otokorelasyon katsayılarının güven aralığı metodları üzerindeki etkisi incelenmiştir. $n = 3$ ve $\rho = 0.25, 0.5, 0.75$ için yapay-skor metodu Wald metodundan yüksek ve t-naive metoduna çok yakın kapsama olasılığına sahiptir. $n = 5, 10$ ve $\rho = 0.25, 0.5, 0.75$ kombinasyonunda yapay-skor ve t-naive metodu birbirine yakın sonuçlar vermiştir ama genel olarak yapay-skor metodunun ortalama güven aralığı uzunluğu daha düşüktür. $n = 5, 10$ için POGA metodu genellikle yüksek kapsama olasılığına sahiptir ama diğer metodlara göre ortalama uzunluğu daha yüksektir.

Çizelge 8.10. Durum 3’de $\rho = 0.25$ için β_1 ’e ilişkin KO ve OU sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8755		0,9503	0,9135	2,6319		3,6544	3,1337
3	10	0,9158		0,9513	0,9213	1,7721		2,0326	1,8325
3	20	0,9415		0,9568	0,9268	1,2223		1,2957	1,1644
5	5	0,9075	0,9768	0,9513	0,9280	2,0434	3,0434	2,4552	2,2189
5	10	0,9358	0,9713	0,9470	0,9273	1,3816	1,7544	1,4872	1,3488
5	20	0,9455	0,9675	0,9528	0,9188	0,9504	1,1161	0,9769	0,8646
10	5	0,9260	0,9755	0,9493	0,9220	1,4686	1,8914	1,6057	1,4644
10	10	0,9420	0,9618	0,9500	0,9190	0,9864	1,1189	1,0198	0,9071
10	20	0,9463	0,9558	0,9498	0,9002	0,6746	0,7202	0,6824	0,5782

t-naive metodu genellikle güven düzeyine yakın sonuçlar sağlamıştır. $n = 3$ ve $m = 5, 10$; $n = 5$ ve $m = 5$, yapay-skor metodu Wald metodundan daha iyi kapsama olasılığına sahiptir. $n = 5, 10$ için genellikle yapay-skor metodu en düşük ortalama uzunlukları sağlamıştır. $\rho = 0.5, 0.75$ durumlarının tabloları $\rho = 0.25$ sonuçlarına benzer olduğu için dahil edilmemiştir.

Çizelge 8.11. Durum 4’de $\rho = 0.5$ için β_1 ’e ilişkin KO ve OU sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8708		0,9425	0,8863	2,4704		3,3551	2,7717
3	10	0,9230		0,9523	0,9153	1,8838		2,1481	1,8944
3	20	0,9395		0,9498	0,9205	1,3582		1,4356	1,2852
5	5	0,9130	0,9703	0,9500	0,9208	1,9719	2,7479	2,3087	2,0548
5	10	0,9320	0,9655	0,9487	0,9153	1,4786	1,8153	1,5719	1,4218
5	20	0,9483	0,9698	0,9525	0,9243	1,0594	1,2194	1,0872	0,9711
10	5	0,9330	0,9718	0,9490	0,9195	1,4320	1,7345	1,5348	1,3940
10	10	0,9383	0,9528	0,9453	0,9178	1,0547	1,1657	1,0831	0,9694
10	20	0,9470	0,9500	0,9483	0,9053	0,7501	0,7868	0,7581	0,6534

Durum 4’de $\rho = 0.25, 0.75$ sonuçları $\rho = 0.5$ sonuçlarına benzer olduğu için tablolar dahil edilmemiştir. $n = 3, 5$, $m = 5$ olduğunda yapay-skor metodu Wald metodundan daha iyi

kapsama olasılığı sağlamıştır. Yapay-skor metodu genellikle daha düşük ortalama uzunluklara sahiptir.

Çizelge 8.12. Durum 5 için β_1 'e ilişkin KO ve OU sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8798		0,9553	0,9075	3,6074		5,0193	4,1657
3	10	0,9285		0,9588	0,9293	2,9875		3,4114	3,0570
3	20	0,9423		0,9558	0,9308	2,9521		3,1465	2,8945
5	5	0,9040	0,9625	0,9483	0,9183	1,7879	2,3691	2,1210	1,8835
5	10	0,9239	0,9607	0,9386	0,9157	1,6337	1,9682	1,7576	1,5971
5	20	0,9455	0,9734	0,9530	0,9270	1,5223	1,8056	1,5736	1,4396
10	5	0,9343	0,9570	0,9538	0,9303	1,3427	1,5281	1,4592	1,3207
10	10	0,9436	0,9551	0,9495	0,9239	1,1493	1,2302	1,1867	1,0708
10	20	0,9404	0,9561	0,9438	0,9165	1,0804	1,1494	1,0961	0,9875

POGA ve t-naive metodu genellikle yüksek kapsama olasılığı ve yüksek ortalama uzunluklara sahiptir. $n=5,10$ için en düşük ortalama uzunlukları yapay-skor metodu sağlamıştır.

Çizelge 8.13. Durum 6 için β_1 'e ilişkin KO ve OU sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8738		0,9485	0,8758	3,8194		5,3476	4,1835
3	10	0,9318		0,9603	0,9205	2,6615		3,0548	2,6631
3	20	0,9445		0,9558	0,9255	1,8484		1,9622	1,7780
5	5	0,9085	0,9773	0,9525	0,9090	3,0282	4,4449	3,6414	3,1786
5	10	0,9342	0,9796	0,9522	0,9272	2,0844	2,6728	2,2404	2,0487
5	20	0,9317	0,9712	0,9353	0,9245	1,4443	1,7283	1,4873	1,3631
10	5	0,9258	0,9695	0,9483	0,9250	2,1703	2,7462	2,3677	2,1794
10	10	0,9384	0,9635	0,9450	0,9215	1,4826	1,6939	1,5320	1,4077
10	20	0,9483	0,9397	0,9483	0,9310	1,0252	1,1137	1,0374	0,9310

$m=5$ ve $n=3,5$ için yapay skor metodu Wald metodundan yüksek kapsama olasılığına sahiptir. $n=5,10$ için yapay-skor metodunun ortalama uzunlukları daha düşüktür.

Çizelge 8.14. Durum 7 için β_1 ’e ilişkin KO ve OU sonuçları

Örnek Çapları		Kapsama Olasılıkları				Ortalama Uzunluklar			
n	m	Wald	Profile	T-Naive	Yapay-Skor	Wald	Profile	T-Naive	Yapay-Skor
3	5	0,8748		0,9075	0,8860	1,5932		1,7711	1,6424
3	10	0,9233		0,9350	0,9210	1,3699		1,4341	1,3910
3	20	0,9343		0,9383	0,9218	1,2592		1,2865	1,2472
5	5	0,9092	0,9810	0,9224	0,9188	1,2160	1,8820	1,2875	1,2777
5	10	0,9318	0,9870	0,9375	0,9330	1,0729	1,5001	1,1012	1,0914
5	20	0,9383	0,9851	0,9415	0,9389	0,9762	1,3156	0,9885	0,9792
10	5	0,9330	0,9833	0,9385	0,9383	0,8956	1,2030	0,9193	0,9184
10	10	0,9418	0,9825	0,9433	0,9430	0,7553	0,9606	0,7648	0,7639
10	20	0,9455	0,9778	0,9473	0,9465	0,6883	0,8486	0,6926	0,6918

Tüm örnek çapı durumlarında yapay-skor metodu genellikle Wald metodundan daha iyi kapsama olasılıkları üretmiştir. Örnek çapı arttıkça yapay-skor metodunun kapsama olasılıkları t-naive metoduna yaklaşmıştır.

8.4. Uygulama

7. Bölümde ele alınan analitik güven aralıkları gerçek veri setine uydurulmuş modeldeki sabit etki parametreleri için oluşturulacaktır. Bilişsel davranış terapisi (cognitive behavioural therapy), endişe bozukluğu ve depresyon tedavisinde iyi sonuçlar vermeye başlamış yeni bir tedavi yöntemidir. Proudfoot ve diğerleri (2003) interaktif multimedya programı geliştirerek Beating the Blues (BTB) tedavisi adı altında bilişsel davranış terapi yöntemi olarak depresyon ve endişe bozukluğu hastalarına uygulamışlardır. Yeni tedavi yöntemi BTB’nin daha iyi sonuç verdiğini göstermek amacıyla henüz hiçbir tedaviye başlanmamış depresyon ve endişe bozukluğu olan 18-75 yaş arası 167 yetişkin kontrollü rasgele tasarlanmış bir çalışmaya alınmıştır. Hastalar rasgele olarak BTB tedavisi ve geleneksel tedavi yönteminin uygulandığı iki gruba ayrılmıştır. Tedavi uygulanmadan önce tüm hastalardan Beck Depresyon Ölçeği’ne (Beck, Steer ve Brown, 1996) göre depresyon skorları ölçülmüştür. Elde edilen ilk ölçüm açıklayıcı değişken olarak modele alınmıştır. Ayrıca hastaların ne kadar süredir hasta olduğuna göre iki kategoriye ayrılmışlardır. 6 aydan az süredir hasta olanlar süre=0 olarak kodlanmış ve 6 aydan fazla süredir hasta olanlar süre=1 olarak kodlanmıştır. Bağımlı değişken olarak ölçülen depresyon skoru 5

farklı ölçüm noktasında alınmıştır. Tedaviden önce alınan ilk ölçümden 2 ay sonra sonra 2. ölçüm alınmıştır. 2. ölçümden 1 ay sonra 3. ölçüm alınmıştır, 3.'den 2 ay sonra 4. ölçüm alınmıştır. 4. ölçümden 3 ay sonra 5. son ölçüm alınmıştır. Analizden önce tezde dengeli veri seti incelendiğinden ölçümü eksik olan hastaların verileri orijinal veri setinden çıkarılmıştır. Böylece analize tüm zaman noktalarında ölçümü tam olan 45 hasta dahil edilmiştir. Veri setine uydurulan rasgele kesim ve eğim terimine uydurulan her i . hasta $i = 1, \dots, 45$; $j = 1, \dots, 4$, için model aşağıdadır,

$$Y_{ij} = \beta_0 + \beta_1 \text{Tedavi} + \beta_2 \text{Sure} + \beta_3 \text{Ilk} + \beta_4 t_j + b_{1i} + b_{2i} t_{ij} + e_{ij} \quad i = 1, \dots, 45 ; j = 1, \dots, 4 .$$

Veri setinin belli bir bölümü Çizelge 8.15'de verilmiştir.

Çizelge 8.15. Uzun vadeli depresyon çalışması verisinin bir alt seti

Deney Birimi	Süre	Tedavi Türü	İlk Ölçüm	Ay	Depresyon Skoru
2	1	1	32	2	16
2	1	1	32	3	24
2	1	1	32	5	17
2	1	1	32	8	20
4	1	1	21	2	17
4	1	1	21	3	16
4	1	1	21	5	10
4	1	1	21	8	9
7	0	0	17	2	7
7	0	0	17	3	7
7	0	0	17	5	3
7	0	0	17	8	7

Çizelge 8.15'de görüldüğü gibi depresyon skoru ölçümleri eşit aralıklı zamanlarla alınmamıştır. Dolayısıyla bağımlı değişkenin deney birimi içinde ölçüm korelasyonunun üstel kovaryans yapısına uygun olması beklenir. Çizelge 8.16'da sabit etki parametrelerinin tahminleri ve güven aralıkları özetlenmiştir.

Çizelge 8.16. Parametre ve güven aralığı tahminleri

Parametre	Tahmin KEÇO	Standart Hata	t-değeri	Wald	Profile	t-naïve	Yapay-Skor
β_0	10,499	3,442	3,050	(3,752; 17,245)	(3,854; 17,106)	(3,566; 17,431)	(3,742; 16,924)
β_1	-0,617	0,172	-3,577	(-0,955; -0,279)	(-0,959; -0,275)	(-0,958; -0,276)	(-0,905; -0,329)
β_2	-7,319	2,125	-3,444	(-11,485; 3,153)	(-11,611; -3,027)	(-11,615; -3,023)	(-11,536; -3,103)
β_3	5,383	2,295	2,345	(0,884; 9,882)	(0,918; 9,848)	(0,744; 10,022)	(0,826; 9,940)
β_4	0,304	0,118	2,574	(0,072; 0,535)	(0,077; 0,531)	(0,065; 0,543)	(0,098; 0,491)

Çizelge 8.16’da sabit etkilerin KEÇO tahminleri, standart hataları, t-değerleri ve 95% güven düzeyinde güven aralığı tahminleri verilmiştir. Lme4 paketinin geliştiricileri parametre tahminlerinin dağılımlarının güvenilir olmadığını düşündüğü için p-değerlerini rapor etmekten kaçınmışlardır.



9. SONUÇ VE ÖNERİLER

Tez çalışması genel olarak henüz çok çalışılmayan bir problemi ele almayı amaçlamıştır. Doğrusal karışık etki modellerinin küçük örnek çapı ve bağımlı değişkenin farklı kovaryans yapısı durumlarında sabit etki tahminlerinin ve ilgili güven aralığı metotlarının özelliklerini ortaya çıkarmak için yola çıkılmıştır.

Doğrusal karışık etkili model literatüründe var olan analitik metotların özellikle küçük örnek çaplarında yeterli olmadığını düşünen araştırmacılar (Staggs, 2009; Liu, 2013; Akdur ve diğerleri, 2016) son yıllarda parametrik bootstrap yöntemini kullanmaya yönelmişlerdir. Parametrik bootstrap metodunun küçük örnek çaplarında iyi sonuç verdiği görülse de oldukça zaman ve maliyet gerektiren bir yöntemdir. Bu nedenle daha hızlı sonuç alınabilen güvenilir bir analitik güven aralığı metodunun aranması bu çalışmanın önceliği olmuştur.

Birinci simülasyon çalışmasında literatürde var olan analitik yöntemlerle parametrik bootstrap güven aralığı yöntemi DKE modellerinin özel bir formu olan iç içe hata regresyon modelinin sabit etki parametresini kapsamı bakımından simülasyon çalışmasında karşılaştırılmıştır.

Güven aralığı yöntemlerinin karşılaştırılmasının yanı sıra sabit etki tahmini için iki farklı HKO yaklaşımı bootstrap yönteminde içerilerek karşılaştırılmasına olanak sağlanmıştır. Aynı grup içinde gözlemler arası korelasyonun zayıf olduğu ve örnek çapının küçük olduğu durumda Kenward-Roger (1997) HKO yaklaşımına dayalı parametrik bootstrap metodunun en iyi kapsama olasılıklarını sağladığı görülmüştür.

İlk simülasyon çalışmasında güven aralıklarının ortalama uzunlukları dikkate alınmamıştır yorumlar yalnızca kapsama olasılıklarına göre yapılmıştır. Örnek çapı ve korelasyon arttığında profil olabilirlik metodunun kapsama olasılığı diğer metotlara göre daha iyi sonuç vermiştir.

İkinci simülasyon çalışmasında, DKE modellerinin bağımlı değişkeninin kovaryans yapısı iki farklı yapının bir araya getirilmesiyle oluşturulmuştur: birinci yapı rasgele etkilerin kovaryans yapısı diğeri de deney birimi / grup içinde rasgele hata terimlerinin kovaryans

yapısıdır. Böylece bağımlı değişkenin kovaryans yapısı hibrid bir yapı olarak elde edilmiştir. Ayrıca deney birimi içinde rasgele hata terimlerinin kovaryans yapısı da simülasyon çalışmasında iki farklı hibrid durumdan oluşturulmuştur. t-naïve metodu orta ve büyük örnek çaplarında genellikle %95 güven düzeyine yakın sonuçlar sağlamıştır. Wald metodu küçük ve orta örnek çaplarında genellikle liberal kapsama olasılığına sahiptir. Profil olabilirlik metodu genellikle %95 güven düzeyinden fazla kapsama olasılıklarına sahiptir ve ortalama güven aralığı uzunlukları diğer metodlara göre yüksektir. Yapay-skor metodunun sonuçları örnek çapının küçük olduğu durumlarda Wald metoduna ve örnek çapı arttıkça t-naïve metoduna yakın sonuçlar vermiştir.

Yapay-skor metodu küçük örnek çaplarında Wald metodundan daha iyi kapsama olasılıklarına sahiptir ve ortalama güven aralığı uzunlukları t-naïve ve profil olabilirlik metodundan düşüktür. Wald ve t-naïve metodu hata terimlerinin güçlü korelasyonlu otoregresif yapısında diğer kovaryans yapılarına göre en düşük kapsama olasılıklarına sahip olmuştur. Profil olabilirlik metodu en düşük kapsama olasılıklarını hata terimlerinin yapısız bantlı kovaryansa sahip olduğunda üretmiştir. Yapay-skor metodu en düşük kapsama olasılıklarını hata terimlerinin kovaryansı üstel olduğu durumda vermiştir. Hata terimlerinin kovaryans matrisi otoregresif olduğu durumda en iyi güven aralığı metodunun yapay-skor olduğu hem kapsama olasılığı bakımından hem de ortalama uzunluk bakımından söylenebilir.

Yapay-skor metodu, profil olabilirlik ve t-naïve metodları tarafından sağlanan uzun güven aralıkları dezavantajını ve ayrıca Wald metodunun sağladığı düşük kapsama olasılığı dezavantajını bir çok durumda düzeltmiştir. Böylece var olan metodların iki dezavantajını dengelemiştir.

Bu tez çalışmasında yapay-skor metodu Kenward-Roger varyans tahminiyle birlikte kullanılarak yeni bir güven aralığı metodu ortaya çıkarılmıştır. Simülasyon çalışması bu yeni metodun kapsama olasılığı ve ortalama güven aralığı uzunlukları bakımından birçok durumda özellikle küçük örnek çaplarında iyi bir alternatif olabileceğini göstermiştir.

İleriki çalışmalarda farklı DKE modelleri ele alınarak güven aralığı metodlarının performansları incelenebilir. Beklenti ve maksimizasyon algoritması farklı kovaryans yapılarında parameter tahminleri için genişletilebilir. Ayrıca rasgele etkiler için güven

aralığı metotları ele alınabilir. Belirlenen güven düzeyine ulaşmak için her güven aralığı metodu için gerekli örnek çapı büyüklükleri simülasyon çalışmalarında incelenebilir.





KAYNAKLAR

- Agresti, A., and Ryu, E. (2010). Pseudo-score confidence intervals for parameters in discrete statistical models. *Biometrika*, 97(1), 215-222.
- Aitken, A. C. (1936). IV.—On Least Squares and Linear Combination of Observations. *Proceedings of the Royal Society of Edinburgh*, 55, 42-48.
- Akdur, H. T. K., Özönur, D., and Bayrak, H. (2016). A Comparison of Confidence Interval Methods of Fixed Effect in Nested Error Regression Model. *Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, 20(2).
- Alnosaier, W. S. (2007). *Kenward-Roger approximate F test for fixed effects in mixed linear models*, Doctoral dissertation, Oregon State University, Corvallis, Oregon.
- Bain, L. J., and Engelhardt, M. (1992). *Introduction to Probability and Mathematical Statistics*. USA: Brooks/ Cole, 360- 364.
- Bates, D., Maechler, M., Bolker, B., and Walker, S. (2014). lme4: Linear mixed-effects models using Eigen and S4. *R package version*, 1(7).
- Bazzoli, C., Letué, F., and Martinez, M. J. (2014). Modelling finger force produced from different tasks using linear mixed models with lme R function. *Case Studies In Business, Industry And Government Statistics*, 6(1), 16-36.
- Beck, A. T., Steer, R. A., and Brown, G. K. (1996). Beck depression inventory-II. *San Antonio*, 78(2), 490-8.
- Birkes, D., and Wulff, S. S. (2003). Existence of maximum likelihood estimates in normal variance-components models. *Journal of Statistical Planning and Inference*, 113(1), 35-47.
- Brown, H., and Prescott, R. (2014). *Applied mixed models in medicine*. NY, USA: John Wiley and Sons.
- Carpenter, J. R., Goldstein, H., and Rasbash, J. (2003). A novel bootstrap procedure for assessing the relationship between class size and achievement. *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, 52(4), 431-443.
- Catellier, D. J., and Muller, K. E. (2000). Tests for Gaussian repeated measures with missing data in small samples. *Statistics in Medicine*, 19(8), 1101-1114.
- Chatterjee, S., Lahiri, P., and Li, H. (2008). Parametric bootstrap approximation to the distribution of EBLUP and related prediction intervals in linear mixed models. *The Annals of Statistics*, 1221-1245.
- Chemick, M. R. (2008). *Bootstrap methods: A guide for practitioners and researchers*. Hoboken, NJ, USA: John Wiley & Sons.

- Das, S., and Krishen, A. (1999). Some bootstrap methods in nonlinear mixed-effect models. *Journal of statistical planning and inference*, 75(2), 237-245.
- Das, K., Jiang, J., and Rao, J. N. K. (2004). Mean squared error of empirical predictor. *The Annals of Statistics*, 32(2), 818-840.
- Datta, G. S., and Lahiri, P. (2000). A unified measure of uncertainty of estimated best linear unbiased predictors in small area estimation problems. *Statistica Sinica*, 613-627.
- Dawson, K. S., Gennings, C., and Carter, W. H. (1997). Two graphical techniques useful in detecting correlation structure in repeated measures data. *The American Statistician*, 51(3), 275-283.
- Demidenko, E. (2004). *Mixed Models: Theory and Applications*. Hoboken, NJ, USA: John Wiley & Sons, 2-5, 131.
- Dempster, A. P., Laird, N. M., and Rubin, D. B. (1977). Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society. Series B (methodological)*, 1-38.
- DiCiccio, T. J., and Efron, B. (1996). Bootstrap confidence intervals. *Statistical Science*, 189-212.
- Diggle, P. (2002). *Analysis of longitudinal data*. Oxford, UK: Oxford University Press.
- Dockery, D. W., Pope, C. A., Xu, X., Spengler, J. D., Ware, J. H., Fay, M. E., ... and Speizer, F. E. (1993). An association between air pollution and mortality in six US cities. *New England Journal of Medicine*, 329(24), 1753-1759.
- Efron, B. (1979). Computers and the theory of statistics: thinking the unthinkable. *SIAM Review*, 21(4), 460-480.
- Efron, B., & Tibshirani, R. J. (1993). *An Introduction to the Bootstrap: Monographs on Statistics and Applied Probability*, 57. New York and London: Chapman and Hall/CRC.
- Faraway, J. J. (2016). *Extending the linear model with R: generalized linear, mixed effects and nonparametric regression models*, 124. Baco Raton, FL, USA: CRC press.
- Fay III, R. E., and Herriot, R. A. (1979). Estimates of income for small places: an application of James-Stein procedures to census data. *Journal of the American Statistical Association*, 74(366a), 269-277.
- Fei, Y., Pan, Y., Chen, Y., and Pan, J. (2016). Modeling of covariance structures of random effects and random errors in linear mixed models. *Communications in Statistics-Theory and Methods*, 45(9), 2748-2769.
- Fitzmaurice, G. M., Laird, N. M., and Ware, J. H. (2004). *Applied longitudinal analysis*. Hoboken, NJ, USA: Wiley-Interscience, 173, 210-216.

- Giesbrecht, F. G., and Burns, J. C. (1985). Two-stage analysis based on a mixed model: large-sample asymptotic theory and small-sample simulation results. *Biometrics*, 477-486.
- Goldberger, A. S. (1962). Best linear unbiased prediction in the generalized linear regression model. *Journal of the American Statistical Association*, 57(298), 369-375.
- Gurka, M. J. (2006). Selecting the best linear mixed model under REML. *The American Statistician*, 60(1), 19-26.
- Halekoh, U., and Højsgaard, S. (2014). A kenward-roger approximation and parametric bootstrap methods for tests in linear mixed models—the R package pbkrtest. *Journal of Statistical Software*, 59(9), 1-30.
- Hall, P. (1992). On bootstrap confidence intervals in nonparametric regression. *The Annals of Statistics*, 695-711.
- Hartley, H. O., and Rao, J. N. (1967). Maximum-likelihood estimation for the mixed analysis of variance model. *Biometrika*, 54(1-2), 93-108.
- Harville, D. A., and Carriquiry, A. L. (1992). Classical and Bayesian prediction as applied to an unbalanced mixed linear model. *Biometrics*, 987-1003.
- Harville, D. A., and Jeske, D. R. (1992). Mean squared error of estimation or prediction under a general linear model. *Journal of the American Statistical Association*, 87(419), 724-731.
- Harville, D. A. (1976). Confidence intervals and sets for linear combinations of fixed and random effects. *Biometrics*, 403-407.
- Henderson, C. R. (1953). Estimation of variance and covariance components. *Biometrics*, 9(2), 226-252.
- Hox, J. J., Moerbeek, M., and van de Schoot, R. (2010). *Multilevel analysis: Techniques and Applications*. Routledge.
- Hulting, F. L., and Harville, D. A. (1991). Some Bayesian and non-Bayesian procedures for the analysis of comparative experiments and for small-area estimation: computational aspects, frequentist properties, and relationships. *Journal of the American Statistical Association*, 86(415), 557-568.
- Jennrich, R. I., and Schluchter, M. D. (1986). Unbalanced repeated-measures models with structured covariance matrices. *Biometrics*, 805-820.
- Jiang, J. (1998). Asymptotic properties of the empirical BLUP and BLUE in mixed linear models. *Statistica Sinica*, 861-885.
- Jiang, J. (2007). *Linear and generalized linear mixed models and their applications*. NY, USA: Springer Science and Business Media, 71.

- Kackar, R. N., and Harville, D. A. (1984). Approximations for standard errors of estimators of fixed and random effects in mixed linear models. *Journal of the American Statistical Association*, 79(388), 853-862.
- Kenward, M. G., and Roger, J. H. (1997). Small sample inference for fixed effects from restricted maximum likelihood. *Biometrics*, 983-997.
- Laird, N. M., and Ware, J. H. (1982). Random-effects models for longitudinal data. *Biometrics*, 963-974.
- Lang, J. B. (2008). Score and profile likelihood confidence intervals for contingency table parameters. *Statistics in Medicine*, 27(28), 5975-5990.
- Lahiri, P., and Rao, J. N. K. (1995). Robust estimation of mean squared error of small area estimators. *Journal of the American Statistical Association*, 90(430), 758-766.
- Lee, V. E., and Bryk, A. S. (1989). A multilevel model of the social distribution of high school achievement. *Sociology of education*, 172-192.
- Liang, K. Y., and Zeger, S. L. (1986). Longitudinal data analysis using generalized linear models. *Biometrika*, 13-22.
- Littell, R. C. (2002). Analysis of unbalanced mixed model data: a case study comparison of ANOVA versus REML/GLS. *Journal of Agricultural, Biological, and Environmental Statistics*, 7(4), 472-490.
- Little, R. J., and Rubin, D. B. (1987). *Statistical Analysis with Missing Data*. NY, USA: John Wiley & Sons, 131, 150.
- Littell, R. C., Pendergast, J., and Natarajan, R. (2000). Tutorial in biostatistics: modelling covariance structure in the analysis of repeated measures data. *Statistics in medicine*, 19(13), 1793-1819.
- Littell, R. C., Stroup, W. W., Milliken, G. A., Wolfinger, R. D., and Schabenberger, O. (2006). *SAS for mixed models*. NC, USA: SAS Institute.
- Liu, J. (2013). *Statistical inference for functions of the parameters of a linear mixed model*, Doctoral dissertation, Iowa State University, Ames, Iowa.
- Lovison, G. (2005). On Rao score and PearsonX2 statistics in generalized linear models. *Statistical Papers*, 46(4), 555-574.
- McCulloch, C. E., and Neuhaus, J. M. (2001). *Generalized linear mixed models*. NJ, USA: John Wiley and Sons.
- Miller, J. J. (1977). Asymptotic properties of maximum likelihood estimates in the mixed model of the analysis of variance. *The Annals of Statistics*, 746-762.

- McLean, R. A., and Sanders, W. L. (1988). Approximating degrees of freedom for standard errors in mixed linear models. In *Proceedings of the Statistical Computing Section, American Statistical Association* (pp. 50-59).
- Molenberghs, G., and Verbeke, G. (2000). *Linear mixed models for longitudinal data*. NY, USA: Springer-Verlag.
- Munafo, M. R., Clark, T. G., Moore, L. R., Payne, E., Walton, R., and Flint, J. (2003). Genetic polymorphisms and personality in healthy adults: a systematic review and meta-analysis. *Molecular Psychiatry*, 8(5), 471-484.
- Onyango, N. O. (2014). On the linear mixed effects regression (lmer) r function for nested animal breeding data. *Case Studies in Business, Industry and Government Statistics*, 4(1), 44-58.
- Orchard, T., and Woodbury, M. A. (1972). A missing information principle: theory and applications. In *Proceedings of the 6th Berkeley Symposium on mathematical statistics and probability* (Vol. 1, pp. 697-715). Berkeley, CA: University of California Press.
- Patterson, H. D., and Thompson, R. (1971). Recovery of inter-block information when block sizes are unequal. *Biometrika*, 545-554.
- Prasad, N. G. N., and Rao, J. N. (1990). The estimation of the mean squared error of small-area estimators. *Journal of the American Statistical Association*, 85(409), 163-171.
- Proudfoot, J., Goldberg, D., Mann, A., Everitt, B., Marks, I., and Gray, J. A. (2003). Computerized, interactive, multimedia cognitive-behavioural program for anxiety and depression in general practice. *Psychological Medicine*, 33(02), 217-227.
- Raudenbush, S. W., and Bryk, A. S. (2002). *Hierarchical linear models: Applications and data analysis methods* (Vol. 1). Sage.
- Satterthwaite, F. E. (1941). Synthesis of variance. *Psychometrika*, 6(5), 309-316.
- Satterthwaite, F. E. (1946). An approximate distribution of estimates of variance components. *Biometrics Bulletin*, 2(6), 110-114.
- Schaalje, G. B., McBride, J. B., and Fellingham, G. W. (2002). Adequacy of approximations to distributions of test statistics in complex mixed linear models. *Journal of Agricultural, Biological, and Environmental Statistics*, 7(4), 512-524.
- Scheffe, H. (1956). A "mixed model" for the analysis of variance. *The Annals of Mathematical Statistics*, 23-36.
- Searle, S. R. (1971). A Biometrics Invited Paper. Topics in Variance Component Estimation. *Biometrics*, 27(1), 1-76.

- Searle, S. R., Casella, G., and McCulloch, C. E. (1992). *Variance components*. NY, USA: John Wiley and Sons, 34, 91, 298.
- Snedecor, G. W., and Cochran, W. G. (1967). *Statistical methods*. Ames, IO, USA: Iowa State University Press, 12, 349-352.
- Staggs, V. (2009). *Parametric Bootstrap Interval Approach to Inference for Fixed Effects in the Mixed Linear Model*, Doctoral dissertation, University of Kansas, Lawrence, Kansas.
- Terracciano, A., McCrae, R. R., Brant, L. J., and Costa Jr, P. T. (2005). Hierarchical linear modeling analyses of the NEO-PI-R scales in the Baltimore Longitudinal Study of Aging. *Psychology and Aging*, 20(3), 493.
- Verbeke, G., and Molenberghs, G. (1997). *Linear mixed models in practice: a SAS-oriented approach*. NY, USA: Springer-Verlag, 126.
- West, B. T., Galecki, A. T., and Welch, K. B. (2014). *Linear mixed models*. Baco Raton, FL, USA: CRC Press.

ÖZGEÇMİŞ

Kişisel Bilgiler

Soyadı, adı : AKDUR, Hatice Tül Kübra
 Uyruğu : T.C.
 Doğum tarihi ve yeri : 09.08.1985, Isparta
 Medeni hali : Evli
 Telefon : 0 (312) 202 14 30
 Faks : 0 (312) 202 37 10
 e-mail : hatice_senol@wsu.edu



Eğitim

Derece	Eğitim Birimi	Mezuniyet tarihi
Doktora	Gazi Üniversitesi /İstatistik Bölümü	2017
Yüksek Lisans	Washington Eyalet Üniversitesi /İstatistik Bölümü	2011
Lisans	Karadeniz Teknik Üniversitesi / İstatistik ve Bilgisayar Bölümü	2007
Lise	Isparta Anadolu Lisesi	2003

İş Deneyimi

Yıl	Yer	Görev
2011-Halen	Gazi Üniversitesi	Araştırma Görevlisi
2009-2011	Washington Eyalet Üniversitesi	Öğretim Asistanı

Yabancı Dil

İngilizce

Yayınlar ve Konferanslar

Akdur, H. T. K., Özönur, D. and Bayrak, H. (2016). A Comparison of Confidence Interval Methods of Fixed Effect in Nested Error Regression Model. *Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, 20(2).

Özönur, D., Akdur, H. T. K., and Bayrak, H. (2017). Comparisons of tests of distributional assumption in Poisson regression model. *Communications in Statistics-Simulation*

and Computation, 1-11.

- Akdur, H. T. K., Gökpınar, F., Bayrak, H. and Gökpınar, E. (2016). A Modified Jonckheere Test Statistic for Ordered Alternatives in Repeated Measures Design. *Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, 20(3), 391-398.
- Akdur, H. T. K., Gül, H. H., Özönur, D. and Bayrak, H. (2017). *Rank-Based Nonparametric Tests for Ordered Alternatives with Randomized Complete Block Design*. 18th International Symposium on Econometrics, Operations Research and Statistics, Konya.
- Özönur, D., Akdur, H. T. K. and Bayrak, H. (2016). *Gamma Dağılımına sahip Yanıt Değişkenli Model İçin Smooth Testler*. 17th International Symposium on Econometrics, Operations Research and Statistics, Sivas.
- Özönur, D., Akdur, H. T. K. and Bayrak, H. (2016). *Smooth Tests for Inverse Gaussian Response Model*. 2nd International Researchers-Statisticians and Young Statisticians Congress, ANKARA.
- Akdur, H. T. K., Özönur, D. and Bayrak, H. (2016). *Confidence Interval Methods for Fixed Effect Parameter of Random Intercept Model*. 2nd International Researchers-Statisticians and Young Statisticians Congress, ANKARA.
- Özönur, D., Akdur, H. T. K. and Bayrak, H. (2015). *A Comparison Study on Distributional Assumption Tests for Poisson regression Model*. Joint Statistical Meetings Seattle, WA, USA.
- Akdur, H. T. K., Özönur, D. and Bayrak, H. (2015). *Confidence Interval Methods of Fixed Effects in Mixed Models: A Comparison Study*. Joint Statistical Meetings Seattle, WA, USA.
- Akdur, H. T.K., Özönur, D. and Bayrak, H. (2015). *A Comparison of Confidence Interval Methods for Fixed Effects in Linear Mixed Models*. XVI. International Symposium On Econometrics, Operations Research And Statistics, Trakya University, Edirne.
- Özönur, D., Akdur, H. T. K. and Bayrak, H. (2015). *Comparisons of Goodness of Fit tests for Poisson Regression Model*. XVI. International Symposium On Econometrics, Operations Research And Statistics, Trakya University, Edirne.
- Akdur, H. T. K., Gül, H. H., Özönur, D. and Gökpınar, F. (2014). *Permutation Approach to Tests of Ordered Alternatives in Incomplete Block Design*. 3rd International Eurasian Conference on Mathematical Sciences and Applications, Vienna.
- Akdur, H. T. K., Gökpınar, F., Bayrak, H. and Özönur, D. (2014). *Comparison of Two Distribution-Free Test for Longitudinal Data of Randomized Blocks with a Prior Knowledge of Alternative Hypothesis*. 3rd International Eurasian Conference on Mathematical Sciences and Applications, Vienna.

Burslar

2214/A Yurt Dışı Doktora Sırası Araştırma Bursu, TÜBİTAK, Calgary Üniversitesi, Ağustos 2015-Aralık 2015.

Yüksek Lisans Bursu, Milli Eğitim Bakanlığı, Washington Eyalet Üniversitesi, İstatistik Bölümü, Ağustos 2009- Mayıs 2011.

İngilizce Dil Öğrenimi Bursu, Milli Eğitim Bakanlığı, Queens University of Charlotte, Nisan 2008- Ocak 2009.

Hobiler

Seyahat etmek, müzik dinlemek, yüzmek, şiir okumak.





GAZİ GELECEKTİR...