

T.C.
AKDENİZ ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ

ÇOK SEVİYELİ REGRESYON MODELLERİ ve UYGULAMALARI

Burçin ŞİMŞEK

YÜKSEK LİSANS TEZİ
ZOOTEKNİ ANABİLİM DALI

2010

ÇOK SEVİYELİ REGRESYON MODELLERİ ve UYGULAMALARI

Burçin ŞİMŞEK

YÜKSEK LİSANS TEZİ

ZOOTEKNİ ANABİLİM DALI

2010

T.C:
AKDENİZ ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ

ÇOK SEVİYELİ REGRESYON MODELLERİ ve UYGULAMALARI

Burçin ŞİMŞEK

YÜKSEK LİSANS TEZİ

ZOOTEKNİ ANABİLİM DALI

Bu tez .../.../2010 tarihinde aşağıdaki jüri tarafından (.....) not takdir edilerek oy birliği/oy çokluğu ile kabul edilmiştir.

Prof. Dr. Mehmet Ziya FIRAT (Danışman).....

Doç. Dr. M. Soner BALCIOĞLU

Yrd. Doç. Dr. Mehmet MERT.....

ÖZET

ÇOK SEVİYELİ REGRESYON MODELLERİ ve UYGULAMALARI

Burçin ŞİMŞEK

Yüksek Lisans Tezi, Zootekni Anabilim Dalı

Danışman: Prof. Dr. Mehmet Ziya FIRAT

Haziran 2010, 84 Sayfa

Son yirmi yıl civarında geliştirilmiş olan çok seviyeli regresyon modelleri, hiyerarşik ya da iç içe veri yapısına sahip veri setlerinin analizinde oldukça kullanışlı olduğundan birçok farklı alanda kabul görmüş bir yöntemdir. Veri setinin yapısından kaynaklı olarak ortaya çıkan gözlemler arası korelasyon ve veri setinde yer alan seviyeler bu modellerce kolayca ifade edilebilmektedir. Veri setindeki iç içe yapının ve gözlemler arası bağımlılıkların göz ardı edilerek geleneksel istatistiksel yöntemlerle hiyerarşik verilerin analiz edilmesi ile yapılan analiz sonuçlarında yanlış parametre ve yanlış hata terimi tahminleri elde edilmesine sebep olur. Bu da araştırmacılar tarafından hiyerarşik veri yapısına sahip veri setlerinin analizinde çok seviyeli regresyon modellerinin kullanılmasını bir nevi zorunlu kılmaktadır.

Çok seviyeli regresyon modellerinin zirai alandaki kullanımı çok yaygın olmadığı için bu çalışmanın amacı İngiliz Siyah Alaca ırkı sığırlarına süt verimlerini veri setinde çok seviyeli modelleme yöntemini uygulayarak bu alana bu konuyu tanıtmaktır. İlk olarak ortalama süt verimleri bağımlı değişken olarak alınıp regresyon modeli ve çok seviyeli regresyon modelleri kullanılarak analizler yapılmıştır. Bunun yanı sıra kontrol günü süt verimleri için çok seviyeli modelleme yönteminin tekrarlanan ölçümler veri

setine nasıl uygulanacağı gösterilmiştir. Bu çalışmanın sonucunda, veri setindeki seviyelerin göz önüne alınması gerektiği ve çok seviyeli regresyon modellerin veri seti yapısına en uygun model olduğuna karar verilmiştir.

ANAHTAR KELİMELEER: Çok seviyeli regresyon modelleri, grup içi korelasyon, süt verimi, kontrol günü kayıtları

JÜRİ: Prof. Dr. Mehmet Ziya FIRAT (Danışman)

Doç. Dr. M. Soner BALCIOĞLU

Yrd. Doç. Dr. Mehmet MERT

ABSTRACT

MULTILEVEL REGRESSION MODELS and APPLICATION of MULTILEVEL REGRESSION ANALYSIS

Burçin ŞİMŞEK

M. Sc. Thesis in Department of Animal Science

Advisor: Prof. Dr. Mehmet Ziya FIRAT

June 2010, 84 Pages

Multilevel regression models, discovered over the past decade, have widely accepted in many different fields since it is considerably useful way to analyze nested or hierarchical data sets. This method can easily deal with the levels of the data and the possible correlation problems which usually occur as a result of the hierarchical data sets. Ignoring nested structured and correlation in traditional statistical analysis lead to biased parameter estimates and incorrect standard error estimation. Thus, using multilevel regression models for hierarchical data structure is kind of obligation for researchers.

In agriculture field, the using of the multilevel regression model is not common; therefore, the aim of this study is to introduce this method to this area with applying multilevel modeling technique to milk yield data which obtained from England Holstein Friesian cattle. The average milk yield is considered as a dependent variable, and ordinal least square regression models (single level model) and multilevel regression models were used as a modeling method. Besides test day milk yield records is considered and a multilevel approach for repeated measurements were used. As a result

of this study, the level structure of the data had to be considered and multilevel regression model is better fit the data.

KEY WORDS: Multilevel regression models, intra-class corelation, milk yield, test day records

COMMITTEE: Prof. Dr. Mehmet Ziya FIRAT (Advisor)

Associate Professor M. Soner BALCIOĞLU

Assistant Professor Mehmet MERT

ÖNSÖZ

Son yirmi yılda çok seviyeli regresyon modelleri artan bir oranla birçok farklı alanda kullanılmaktadır. Zirai alandaki kullanımı ise diğer alanlara göre nispetten daha azdır. Bu çalışmada, çok seviyeli regresyon modeli kullanılarak süt verimi veri seti analiz edilmiştir. Yapılan bu çalışma ile zirai alanda elde edilen hiyerarşik bir yapı gösteren veri setlerinin analizinde çok seviyeli modellemenin kullanılmasını yaygınlaştırmasını ve bu alana katkısı olmasını dilerim.

Bana bu konuda çalışma olanağı veren, çalışmanın planlanmasında, yürütülmesinde ve sonuçlandırılmasında her türlü yardımını gördüğüm danışmanım Sayın Prof. Dr. Mehmet Ziya FIRAT' a, görüşleriyle bu çalışmaya katkıda bulunmuş Arş. Gör. Emre KARAMAN ve Arş. Gör. Ebru KAYA' ya çok teşekkür ederim. Ayrıca manevi desteklerini çalışma boyunca gördüğüm aileme ve arkadaşlarıma teşekkürlerimi sunarım.

İÇİNDEKİLER

ÖZET.....	i
ABSTRACT.....	iii
ÖNSÖZ.....	v
İÇİNDEKİLER.....	vi
SİMGELER VE KISALTMALAR DİZİNİ.....	viii
ŞEKİLLER DİZİNİ.....	x
ÇİZELGELER DİZİNİ.....	xi
1. GİRİŞ.....	1
2. KURAMSAL BİLGİLER VE KAYNAK TARAMALARI.....	6
2.1. Tarihçe.....	6
2.1.1. Kullanım alanları.....	10
2.1.1.1. Eğitim ve sosyoloji alanındaki kullanımı.....	11
2.1.1.2. Tıp ve sağlık alanındaki kullanımı.....	11
2.1.1.3. Zirai alanındaki kullanımı.....	12
2.1.1.4. Anket çalışmalarındaki kullanımı.....	12
2.1.1.5. Kategorik veri analizindeki kullanımı.....	13
2.1.1.6. Tekrarlanan ölçümler veri setindeki kullanım.....	14
2.1.2. Geliştirilen paket programlar.....	15
2.2. Kuramsal Bilgiler.....	17
2.2.1. Varyans analizi.....	17
2.2.1.1. Tek faktörlü varyans analizi.....	17

2.2.1.2. İki faktörlü varyans analizi.....	19
2.2.1.3. İç içe sınıflandırma.....	21
2.2.2. Çok seviyeli modellerin teorisi.....	23
2.2.2.1. İki seviyeli regresyon modeli.....	24
2.2.2.2. İki seviyeli modelin adımsal olarak elde edilişi.....	27
2.2.2.3. İki seviyeli modelin matrissel ifade edilişi.....	29
2.2.2.4. Birden fazla açıklayıcı değişkenin modellenmesi.....	31
2.2.2.5. Açıklayıcı değişken seçimi.....	34
2.2.2.6 Merkezlenme.....	37
2.2.2.7. Model parametrelerinin tahmininde kullanılan yöntemler.....	39
2.2.2.7.1. ML tahmini.....	40
2.2.2.7.2. REML tahmini.....	45
2.2.3. Tekrarlanan ölçümler verisine uygulanması.....	46
3. MATERYAL VE METOT.....	47
3.1. Materyal.....	47
3.2. Metot.....	50
4. BULGULAR VE TARTIŞMA.....	53
4.1. Toplam Süt Verimi için Elde Edilen Sonuçlar.....	53
4.2. Kontrol Günü Süt Verimleri İçin Elde Edilmiş Sonuçlar.....	64
5. SONUÇ.....	73
6. KAYNAKLAR.....	78

ÖZGEÇMİŞ

SİMGELER VE KISALTMALAR DİZİNİ

Simgeler

Y_{ij} , Y_{ijk}	Bağımlı değişken
μ	Genel ortalama
α_i , β_j , $(\alpha\beta)_{ij}$	A ile B muamele etkilerini ve bunların interaksyonu (sırasıyla)
e_{ij} , e_{ijk}	Hata terimi
β_{0j} , β_{1j}	Kendi regresyon denklemlerine sahip olan sabit ve eğim katsayısı
γ_{00} , γ_{10} , γ_{01} , γ_{11}	Regresyon katsayıları
X_{ij}	En düşük seviyeye ait açıklayıcı değişken
Z_j , W_j	En yüksek seviyeye ait açıklayıcı değişkenler
ε_{ij}	En düşük seviyeye ait hata terimleri
u_{0j} , u_{1j}	En yüksek seviyeye ait hata terimleri
σ_{00} , σ_{10} , σ_{11} , σ_ε^2	Hata terimlerine ilişkin varyans
ρ , r	Grup içi korelasyon katsayısı
n_c	c-inci gruptaki gözlem sayısı

Kısaltmalar

N	Normal dağılım
Var , V	Varyans
iid	Independent and Identically Distributed (birbirinden bağımsız ve aynı dağılımlı)
ICC	Intra-Class Correlation (Grup içi korelasyon)
DE	Design Effect (deneme etkisi)
ML	Maximum Likelihood (en çok olabilirlik)
REML	Restrictd Maximum Likelihood (kısıtlanmış en çok olabilirlik)
l	Likelihood Function (olabilirlik fonksiyonu)

ŞEKİLLER DİZİNİ

Şekil 2.1. Anket çalışmalarında örneklem seçiminde izlenen yol.....	12
Şekil 2.2. İç içe etkili deneme desenine ilişkin veri yapısı.....	22
Şekil 2.3. İki seviyeli hiyerarşik veri yapısı.....	24
Şekil 3.1. Zaman içindeki süt verimindeki değişim.....	49
Şekil 4.1. Alt seviyeye indirgenmiş verinin analizi sonucunda elde edilen süt verimi tahminleri ile hatalara ilişkin grafik.....	54
Şekil 4.2. Üst seviyeye kaydırılmış verinin analizi sonucunda elde edilen süt verimi tahminleri ile hatalara ilişkin grafik.....	56
Şekil 4.3. Holystein yüzdesindeki değişimin toplam süt verimine etkisi.....	57
Şekil 4.4. Hiyerarşi yapısı göz ardı edilmesi ile yapılan analiz sonucunda elde edilen süt verimi tahminleri ile hatalara ilişkin grafik.....	59
Şekil 4.5. Hiyerarşi yapısı göz önüne alınarak yapılan analiz sonucunda elde edilen süt verimi tahminleri ile hatalara ilişkin grafik.....	62
Şekil 4.6. Kontrol günü süt verimi tahminleri ile hatalara ilişkin grafik.....	67
Şekil 4.7. Rastgele etkili model için kontrol günü süt verimi tahminleri ile hatalara ilişkin grafik.....	71

ÇİZELGELER DİZİNİ

Çizelge 2.1. Tek faktörlü deneme deseninin veri seti.....	17
Çizelge 2.2. İki faktörlü deneme deseninin veri seti.....	19
Çizelge 2.3. Tek ve çok seviyeli regresyon modelleri.....	27
Çizelge 2.4. Çok seviyeli regresyon modelindeki parametre sayısı.....	32
Çizelge 3.1. Veri dosyası formatı.....	47
Çizelge 3.2. Yaşlı ve genç sürüdeki boğalara ilişkin frekans tablosu.....	48
Çizelge 3.3. Süt verimi kaydına ait betimsel istatistikler.....	48
Çizelge 4.1. Alt seviyeye indirgenmiş veri için varyans analizi tablosu	53
Çizelge 4.2. Alt seviyeye indirgenmiş veri için parametre tahminleri.....	53
Çizelge 4.3. Seviye-2'ye kaydırılmış veri için varyans analizi tablosu	55
Çizelge 4.4. Üst seviyeye kaydırılmış veri için parametre tahminleri.....	55
Çizelge 4.5. Parametre tahminleri (tüm açıklayıcı değişkenlerin yer aldığı fakat hiyerarşinin göz ardı edildiği model).....	58
Çizelge 4.6. Parametre tahminleri (sadece sabitin yer aldığı model).....	60
Çizelge 4.7. Kovaryans parametrelerinin tahminleri (sadece sabitin yer aldığı model).....	60
Çizelge 4.8. Parametre tahminleri (tüm açıklayıcı değişkenlerin yer aldığı model).....	61
Çizelge 4.9. Kovaryans parametrelerinin tahminleri (tüm açıklayıcı değişkenlerin yer aldığı model).....	62
Çizelge 4.10. Model uygunluk test değerleri.....	63

Çizelge 4.11. Parametre tahminleri (sadece sabitin yer aldığı model).....	64
Çizelge 4.12. Kovaryans parametrelerinin tahminleri (sadece sabitin yer aldığı model).....	64
Çizelge 4.13. Parametre tahminleri (zaman değişkeninin yer aldığı model).....	65
Çizelge 4.14. Kovaryans parametrelerinin tahminleri (zaman değişkeninin yer aldığı model).....	65
Çizelge 4.15. Parametre tahminleri (zaman ve diğer açıklayıcı değişkenlerin yer aldığı model).....	66
Çizelge 4.16. Kovaryans parametrelerinin tahminleri (zaman ve diğer açıklayıcı değişkenlerin yer aldığı model).....	67
Çizelge 4.17. Parametre tahminleri (zamanın yer aldığı model).....	68
Çizelge 4.18. Kovaryans parametrelerinin tahminleri (zamanın yer aldığı model).....	68
Çizelge 4.19. Parametre tahminleri (zaman ve diğer açıklayıcı değişkenlerin yer aldığı model).....	69
Çizelge 4.20. Kovaryans parametrelerinin tahminleri (zaman ve diğer açıklayıcı değişkenlerin yer aldığı model).....	70
Çizelge 4.21. Model uygunluk test değerleri	72

1. GİRİŞ

Bilim, olguların gözlenmesi ve sınıflandırılması ile uğraşan bir çalışmadır ve bilim adamları herhangi bir tasarımın sonucu olarak ortaya çıkan bir olayı veya olaylar topluluğunu gözlemleyebilmelidir. Bilimsel çalışmalarda çok önemli bir yere sahip olan gözlemlerin elde edilmesi eylemi ise deney olarak adlandırılır (Çömlekçi 1988). Deneyler, anlamlı veri elde edebilmek amacıyla, genellenebilir belirli bazı sonuçların elde edilmesinde etkili olabilecek değişkenlerin var olan çözümleme tekniklerinin uygulanabilmesi için gerekli olan varsayımlara uymasını sağlamak ve böylece anlamlı önsel sınamaları yapabilmek için araştırmacılar tarafından düzenlenen süreçlerdir (Çömlekçi 1988, Yıldız 1995).

Araştırmacılar, karşılaştıkları problemlere güvenilir sonuçlar bulmaya çabalarlar. Bu problemlere doğru sonuçlar bulunabilmesi için doğru kararların alınması gereklidir. Doğru kararlar elde edebilmek ise bu bilgilerin doğru bir şekilde kullanılmasına bağlıdır. Bundan dolayı bilginin gerçek kaynağının ne olduğu hakkında bir fikre sahip olmak ve uygulamada karar alınırken hangi bilgi kaynaklarından yararlanılacağına bilincinde olunmalıdır (Serper ve Gürsakal 1989, Yıldız 1995).

Yapılan istatistiksel çalışmalarda doğru kararların alınması ve özellikle veri analizinde kullanılacak istatistiksel yöntemlerin seçimi çok önemlidir. Çünkü bu veriler kullanılarak genelleme yapılacaktır (Yıldız 1995). Ayrıca sonuç çıkarımında yapılacak hataları azaltmak amacıyla birçok farklı tipte veri setleri elde edilmektedir. Hiyerarşik (çok seviyeli) veri seti de bunlardan birisidir.

Birçok alanda yapılan çalışmalarda hiyerarşik veriler yaygın bir şekilde elde edilmektedir. Bu tarz veri setlerinde her bir seviye birbiri içine yuvalanmış haldedir. Bu yüzden hiyerarşik verilerin analizi için uygun deneme planlarının seçimi sanıldığı kadar

kolay değildir. Çünkü bu veri setlerinde seviyelerin modellenmesi karmaşıktır. Ayrıca, hiyerarşik veride aynı grupta bulunan bireyler arasında bir ilişki söz konusudur. Çünkü bireyler arasında bir ilişki yoksa bunlar aynı grup içine yuvalanamazlar. Bundan dolayı hiyerarşik veri yapısında gözlemler arasında ilişkiler söz konusudur. Fakat bu ilişki bizim istemediğimiz bazı problemlere sebep olmaktadır (Leeuw ve Meijer 2006).

Özellikle sosyal ve sağlık alanındaki araştırmalarda hiyerarşik veri setleriyle çok yaygın bir şekilde çalışılmaktadır. Bu alandaki deneyler ve gözlemlere dayalı çalışmalar çok büyük paralar, zaman ve emek harcanarak yapıldığı için deneme planının seçimi büyük önem taşımaktadır. Çünkü deneme planının yanlış seçimi ya da en iyi olmaması bu çabaların hepsinin boşa gitmesine sebep olur (Hsieh 1988, Feng ve Grizzle 1992, Donner 1998, Moerbeek ve ark. 2001a). Ayrıca araştırmalarda hatalı modellenmenin kullanılması ya da deneme planının yanlış seçiminden dolayı ölçüm ve sınıflandırma hatası ortaya çıkmakta ve bu hataları göz ardı etmek sonuç çıkarımında sorunlara neden olmaktadır (Hsieh 1988, Moerbeek ve ark. 2001b).

Bu sebeplerden dolayı hiyerarşik veri yapısına sahip verilerin analizinde kullanılacak olan uygun bir model ya da modellere ihtiyaç duyulmuştur. Bu tarz verilerin standart modeller kullanılarak analiz edilmesi yanlış çıkarımlarda bulunulmasına sebep olacağından çok sayıda seviyesi bulunan çalışmalar için en iyi deneme planları araştırılmıştır. Son yıllarda bu sorunla baş edebilecek uygun ve etkili modeller bulunmuş ve bu karmaşık modellerin gerçek dünyayı açıklamada çok daha etkili olduğu belirtilmiştir (Leeuw ve Meijer 2006).

Çok seviyeli regresyon modelleri son 20 yıl civarında geliştirilmiştir ve özellikle son yirmi yılda eğitim ve sağlık alanını kapsayan birçok farklı alanda çok seviyeli modeller yaygın bir şekilde kullanılmaktadır (Draper 1995, Goldstein ve Spiegelhalter 1996). Bu modeller farklı seviyelerden elde edilmiş değişkenlerin analizi için planlanmış olup, birden fazla bağımlılık içeren modellerdir (Hox 1995). Bu modeller

hiyerarşik veri setlerindeki gözlemler arasındaki ilişkiyi göz önüne alan ve veri setindeki her seviyeye ait açıklayıcı değişkenleri aynı ya da ortak bir modelde yer verebilen modellerdir. Ayrıca hiyerarşik veri setlerinin birçok farklı alanda yaygın bir şekilde elde edilmesinden dolayı bu yöntem artan bir oranla kullanılmaktadır. Hox (1995) bu modelleme yöntemin amacını şu şekilde özetlemiştir;

- Birey ve grup seviyesindeki açıklayıcı değişkenlerin direk etkisinin tespiti
- Seviye-2'deki (Grup seviyesi) değişkenlerin seviye-1 (birey seviyesi) değişkenleriyle ilişkili olarak değerlendirip değerlendirilemeyeceğinin araştırılmasıdır.

Ayrıca bu modeller oluşturulurken modelde yer alacak açıklayıcı değişkenlerin seçimi ve model için belirlenen varsayımlar önem kazanmaktadır. Çok fazla açıklayıcı değişkenin modele alınması ile modelin yapısı karmaşıklaştığı için modellenecek seviye-1'e ve seviye-2'ye ait açıklayıcı değişkenlerin seçimi çok önemlidir. Dahası model oluşturulurken hataların beklenen değerlerinin 0 olması, kümeler arası kovaryans matrisinin homojen olması ve bu matrise ait bazı istenen özelliklerinin belirlenmesi ve değişkenlere ait çok değişkenli dağılımların seçimine ilişkin varsayımlar belirlenmektedir (Hox 1995).

Bu varsayımların yanı sıra, Moerbeek ve ark. (2006) çok seviyeli modelleme yönteminde küme sayısının seçimi ve yapılan analizler sonucunda elde edilen tahminlerin gücü ve modelin dirençlilik (robustness) ilişkin 4 önemli konudan bahsetmişlerdir ve bu noktalar aşağıda verilmiştir.

1. Çok seviyeli modellerde en iyi birimlerin paylaşımı çok önemlidir. Diğer bir deyişle, çok seviyeli veri seti yapısındaki her bir seviyenin örneklem genişliğinin belirlenmesi çok önemlidir.

Aslında çalışmada bulunan kümelerin sayısı, popülasyonda bulunan küme sayısından çok olamayacağından en iyi örneklem sayısı popülasyonda bulunan küme sayısı ile sınırlıdır. Benzer şekilde her bir kümedeki birey sayısı da popülasyondaki küme genişliğinden büyük olamayacağından popülasyondaki küme genişliği ile sınırlıdır.

2. Diğer konu ise belli parametrelerin tahmini, testin gücü ve 1. tip hata oranı tahmini için gerekli bütçeye sahip olmadır ya da kısaca; belli parametre testlerinin gücünü belirleyebilecek bütçeye sahip olmaktır.
3. En iyi deneme planının dirençliliği de çok önemli bir konudur.

Model parametrelerinin değerlerine ilişkin önsel tanımlamalar, özel olarak grup içi (intra-class) korelasyon katsayısı, optimal örneklem genişliğinin belirlenmesi için verilmeli ve dahası bu en iyi deneme deseninin hatalı ya da yanlış belirlenmiş önsel tanımlamalar için dirençli olup olmadığı araştırılmalıdır.

4. Son olarak, birey seviyesine ait tesadüfi dağıtım işlemine karşı yapılan küme rastgeleleştirilmesinin etkinliği göz önüne alınmalıdır.

Muamele etkisinin istatistiksel testinde birey seviyesinde tesadüfi dağıtım yapıldığında daha güçlü sonuçlar vermesine rağmen genellikle küme seviyesinde rastgeleştirme yapılmaktadır. Bu durumda daha etkin sonuçlar elde edilmesine engel olunup olmadığı merak edilebilir ve bu konu göz önüne alınabilir.

Bilindiği gibi zirai alanında yapılan birçok çalışmada, özellikle ıslah çalışmalarında hiyerarşik veri setleri elde edilmektedir. Bu yüzden bu çalışmadaki amaç, çok seviyeli regresyon modellerinin zirai alanında elde edilen veri setlerinin analizinde nasıl kullanılabileceğini uygulamalı olarak büyük bir veri seti ile göstermektir. Verilerin analizinde SAS istatistiksel paket programı kullanılacaktır. Özel olarak, Siyah Alaca ırkı sığır popülasyonuna ilişkin süt verimini etkileyen bazı açıklayıcı değişkenlerin süt verimi üzerine etkisi çok seviyeli regresyon modelleri kullanılarak analiz edilecektir. Burada her bir boğanın yavrusunun kendisi içine yuvalandığı düşünülerek, bir seviyeyi boğaların oluşturduğu ve diğer seviyeyi de bunların kızlarının oluşturduğu hiyerarşik yapı göz önüne alınmıştır. Çok seviyeli modelleme kullanılması ile bu iki seviyeye ait açıklayıcı değişkenlerin bağımlı değişken (süt verimi) üzerine etkisinin araştırılmasına izin vermesinden dolayı çok seviyeli regresyon modellerinin hayvan ıslahı çalışmalarında kullanılmasının gerekliliği açıktır.

Çalışma beş alt bölümden oluşmaktadır. Konunun önemi ve çalışmanın amacı birinci bölümde açıklanmıştır.

İkinci bölümde, çok seviyeli modellemenin tarihçesine, bu yöntem için geliştirilmiş paket programlara ve bu yöntemin hangi alanlarda kullanıldığına değinilmiştir. Ayrıca kuramsal bilgiler kısmında ilk olarak çok seviyeli regresyon modelleme ile yakından ilişkili olan varyans analizi ve iç içe sınıflandırmadan bahsedildikten sonra çok seviyeli modellemenin teorisi detaylı bir şekilde verilmiştir. Burada iki seviyeli regresyon modellerinin nasıl elde edildiği ve bu modellerin matrissel olarak ifade edilmesi verilmiştir. Bunların yanı sıra modelde birden fazla açıklayıcı değişkenin nasıl yer alacağı, bu açıklayıcı değişkenleri seçerken nasıl bir yol takip edileceği anlatıldıktan sonra model parametrelerin tahmin yöntemlerinden bahsedilmiştir.

Üçüncü bölümde ise bu çalışmada kullanılacak materyal ve kullanılan modellerden detaylı bir şekilde bahsedilmiştir. Dördüncü bölümde SAS paket programı kullanılarak yapılan analiz sonuçları ve bu sonuçlara ilişkin açıklamalar verilmiştir. Araştırmadan elde edilen sonuçlar ve bazı önerilerden ise son bölüm olan altıncı bölümde yer verilmiştir.

2. KURAMSAL BİLGİLER VE KAYNAK TARAMALARI

2.1. Tarihçe

Çok seviyeli modelleme yöntemi iç içe (kümelenmiş) veriler için uygun bir yöntem olarak değerlendirilmektedir. Aslında Fisher (1923) faktörlerin farklı seviyeler içinde yer alması fikrini 1920’li yıllarda ortaya atmıştır. Bu modeller geliştirilmeden önce, hiyerarşik veri yapısına sahip verilerin analizinde iç içe olmayan veri setleri için bulunmuş en iyi deneme planları kullanılmıştır (Cochran 1983). Geçmişte yaygın olarak çok seviyeli regresyon modellerin yerine çoklu regresyon modelleri kullanılmaktaydı (Burstein, 1980, Van den Eeden ve Hüttner 1982).

1960 yıllarının sonlarında eğitim ve sosyolojinin birçok alanındaki araştırmacılar ölçümlerin alındığı bireylere ait bilgileri (seviye-1’e ait bilgileri), bireylerin ait olduğu grup bilgileriyle (seviye-2’ye ait bilgilerle) birleştiren istatistiksel yöntemleri inceleyerek başladılar. Bu araştırmacılar ele aldıkları konularda her bir öğrencinin sınıf içine, her bir sınıfın okul içine, her bir okulun... şeklinde iç içe yuvalandığını düşünmüştürler. Bu iç içe seviyelerin hepsinin kendisine ait tahmin edicileri vardır ve zor olan ise bu tahmin edicilerin hepsinin aynı istatistiksel analiz yönteminde yer almasıdır, yani özel bir istatistiksel yöntemde bunların yer almasıdır (Coleman ve ark. 1966).

Galtung (1969) başlarda bu problemi çözmek için yukarı seviyedeki birimlerin aşağı seviyeye kaydırılması (aggregation) veya aşağıdaki seviyelerdeki birimlerin yukarı seviyeye taşınması (disaggregation) yöntemlerini kullanmıştır. Bu yöntemlerle, hiyerarşik veri yapısına sahip veriler, tek seviyeli veri setleri haline dönüştürülmüş ve bunlara bilinen istatistiksel yöntemler uygulanarak analizleri yapılmıştır. Hox’un (1995) da bahsettiği gibi, değişik seviyelerdeki değişkenleri bir seviyeye taşıyarak yapılan analizlerde iki sorunla karşılaşilmektedir.

- Bu sorunların birincisi istatistiksel bir sorundur. Eğer veri aşağı seviyeye indirgenmişse yani yukarı seviyedeki birimler aşağı seviyeye kaydırılmışsa, yüksek seviyelerdeki birimler alt seviyelerde düşük değerler almaktadır. Dahası bazı bilgiler kaybolmakta ve istatistiksel analizler gücünü kaybetmektedir. Diğer taraftan, eğer alt seviyelerdeki birimler yukarı seviyeye kaydırılmışsa, alt birimlerde çok sayıda olan değerler üst birimlerde az sayıdadır. Aslında normal istatistiksel testler aşağı taşınmış veriyi, asıl örneklemden bağımsız olarak değerlendirdiği için p değeri değişmekte ve anlamlılık testlerinin bazılarının red edilmesi ortaya çıkmaktadır.
- İkincisi ise kavramsalıdır. Eğer araştırmacı dikkatli değil ise yanlış bir seviye için yanlış bir değerlendirme yapabilir. Yani farklı bir seviyeyi analiz ederken, başka bir seviye için sonuç çıkarımında bulunabilir. Buna örnek olarak Robinson' un (1950) yaptığı çalışmasında, 1930 yılında 9 bölgedeki zenci oranı ile okuma yazma seviyesi/durumu arasındaki ilişkiyi araştırmıştır. Analizlerini birey seviyesine indirgenmiş veri ile gerçekleştirmiş ve elde ettiği sonuçları yorumlarken yanlışlıklar yapmıştır.

Cochran (1983) tek seviyeli modellerin çok seviyeli modeller yerine kullanımının uygun olmadığını belirtmiştir. Çünkü tek seviyeli modellerde kullanılan formülasyon sadece toplam bireyleri göz önüne alır. Bu formülasyon küme sayılarını ya da her bir bireyin sayısını belirlemez. Bu sebepten dolayı tek seviyeli modellerin çok seviyeli modellerin yerine kullanılması uygun değildir. Bunun sonucunda da hiyerarşik veri setlerinde kullanımı elverişli olan ve yapılan tahminler sonucunda yansız ve tutarlı sonuçlar veren istatistiksel yöntem arayışına devam edilmiştir.

Boyd ve Iversen (1979) ve Moerbeek ve ark. (2006) çalışmalarında hiyerarşik veri seti yapısında aynı grupta bulunan bireyler arasında korelasyon bulunduğunu ve grup içi korelasyon olarak (intra-class correlation (*ICC*)) bilinen bu korelasyonun varlığının veri setlerinin analizlerini zorlaştırdığını belirtmişlerdir. Grup içi korelasyonu göz ardı edip geleneksel analiz yöntemlerini kullanmak yanlış tahminler yapmamıza, hatalı “standart

hatalar” ve yanlış sonuçlar elde etmemize sebep olabilir. Bu yüzden aynı gruptaki bireyler arasında bulunan korelasyonu da göz önüne alabilecek istatistiksel yöntem arayışında bulunulması gerektiğini belirtmişlerdir.

1970’lerde Coleman ve ark. (1966) ve Jencks ve ark. (1972) yaptığı çalışmalarda o zamanlarda popüler olmuş okul etkililiği araştırmalarını ele almışlardır. Bu araştırmacılar grupların yapısını yani verinin hiyerarşik yapısını hesaba katmışlardır. Fakat bu hiyerarşik yapıdan dolayı bireylere ait gözlemler arasında bağımlılıklar çıktığını fark etmişlerdir. Benzer şekilde daha sonraki yıllarda yapılan zirai ve ıslah çalışmalarında, ekonomistler ve bio-istatistikçiler de bunun farkına varmışlar ve varyans analizi için varyans-kovaryans unsurları modelini planlamışlardır (Leeuw ve Meijer 2006).

Bilindiği gibi sosyal bilimlerdeki araştırmaların genel yapısında da bireyler ile bunların içinde bulundukları toplum etkileşim içindedir. Yani bireyler içinde bulundukları sosyal grup tarafından etkilenir ve ayrıca sosyal grupları oluşturan bireyler de bu gruplardan etkilenir. Bu yüzden bu alanda yapılan çalışmalarda da varyans-kovaryans yapısı modelde yer almıştır (Burstein 1980).

Fakat okul etkinliği çalışmalarında bu ilişkili yapıyı açıklamak için farklı bir yol izlenmiştir (Jencks ve ark. 1972, Leeuw ve Meijer 2006). Anlatılacak olan okul öğrenci örneği incelenmiştir. Her okul için farklı regresyon analizi yapmak tatmin edici sonuçlar vermez çünkü

- okullara ait örneklemeler küçük ve regresyon katsayıları sabit değildir,
- aynı zamanda, bu ayrı ayrı yapılan analizlerde her okulun aynı okul sistemi/ eğitim sisteminden olduğu gerçeği göz ardı edilir ve bunun sonucu olarak regresyon katsayılarının benzer olduğunu düşünmek çok normaldir,

- dahası, birçok okulun bulunduğu çalışmalarda, 100'lerce okul ve bunlara ait regresyon katsayıları vardır ve burada veri azaltılması (data reduction) yeterince yapılamaz.

Diğer taraftan, her okuldaki regresyon katsayılarının aynı olduğunu şart koşturmayı sınırlandırıcı bir faktördür. Çünkü okul içi yapılan regresyonun farklı olmasını sağlayacak birçok sebep vardır. Örneğin bazı okullarda puanlama önemli iken, bazılarında sosyo-ekonomik durum önemli olabilir. Bu sebeplerden dolayı regresyon katsayılarının aynı olduğunu varsaymak bazı hatalara sebep olabilir. Bu katsayıların aynı olması varsayımı ile veri azaltılması yapılabilmekte fakat bunun sonucunda yanlış ve anlamsız regresyon katsayıları elde edilebilmektedir. Bu yüzden tek tür regresyon katsayıları tahmini elde etmeyecek ve ayrıca her okul için ayrı ayrı regresyon katsayısı hesaplamayacak bir regresyon modeline ihtiyaç duyulmuştur. Bu da doğrudan doğruya araştırmacıları rastgele etkilerden oluşmuş modeller fikrine sürüklemiştir. Fakat burada da farklı seviyedeki tahmin edicilerin aynı modele yazılmasında sorunlar yaşanmıştır (Burstein ve ark. 1978).

1980'lerin başında Burstein (1980) ilk seviyeye ait (seviye-1'e ait) regresyon katsayılarını, okul seviyesi regresyon modeline (seviye-2 regresyon modeli) bağımlı değişken olarak dahil edilmesi fikrini ortaya atmışlardır. Fakat burada seviye-2'deki gözlemlerin bağımsız olduğunu varsayan standart regresyon modeli kullanılmamıştır, çünkü bu yöntemin kullanılması ile etkisiz regresyon katsayıları tahminleri ve yanlış standart hata tahminleri elde edilmektedir. Dahası, her okuldaki tahmin ediciler farklı dağılım fonksiyonuna sahip olduğundan, seviye-1 regresyon katsayıları farklı standart hatalara sahip bulunmuştur. Burada okul içi regresyon katsayılarının dağılımına karar verirken okul büyüklüğü ve okul içi tahmin ediciler arasındaki kovaryans yapısı göz önüne alınmıştır. Langbein (1977), Burstein ve ark. (1978) ve Burstein (1980) bu seviyelerle ilgili çalışmalar yapmışlar ve ağırlıklandırılmış en küçük kareler metodu ile seviye-2 regresyon katsayılarını tahmin etme girişimlerinde bulunmuşlardır (Boyd ve Iversen 1979, Tate ve Wongbundhit 1983).

Yukarıda bahsedilen bu girişimler tamamen başarılı olmamıştır. Çünkü iki aşamalı tekniklere istatistiksel açıdan yaklaşım o dönemdeki metot farklılıklarından dolayı kolay olmamıştır. Bu durum 1980'lerin ortalarında daha açık bir hal almıştır. Farklı isimler altında ve bazı farklılıklarla 1980'lerde çok seviyeli modeller kullanılmaya başlanmıştır. İlk zamanlarda çok seviyeli modeller, karışık doğrusal model (Hartley ve Rao 1967) ya da hiyerarşik doğrusal modeller (Lindley ve Smith 1972) olarak adlandırılmıştır. Contextual analizdeki problem buradaki doğrusal model taslağında da ortaya çıkınca bu modellere çok seviyeli regresyon modelleri olarak adlandırılmıştır. Bu yüzden çok seviyeli modelin analizi, geleneksel istatistiksel karma model teorisiyle contextual analizin evliliği olarak da tanımlanabilir (Boyd ve Iversen 1979, Leeuw ve Meijer 2006).

Çok seviyeli modeller, birçok farklı seviyesi bulunan iç içe yapıdaki veriyle çalışmaya izin verdiği için, bu seviyelerdeki bir ya da birden fazla açıklayıcı değişkenlerin modelde yer almasına izin vermesi ve değişkenlerin dağılım varsayımı konusunda esnek olmasından dolayı çok yaygın olarak kullanılmaya başlanmıştır (Draper 1995). Günümüzde çok seviyeli regresyon modelleri; rastgele etkili model (de Leeuw ve Kreft 1986, Longford 1989, Kreft ve Leeuw 1993), varyans unsurları modeli (Aitkin ve Longford 1986) ve hiyerarşik doğrusal modeller (Raudenbush ve Bryk 1986) olarak da bilinir.

2.1.1. Kullanım alanları

Hiyerarşik veri yapısına sahip verilerin yaygın bir şekilde birçok farklı alanda elde edilmesinden dolayı çok seviyeli modellerin farklı alanlarda kullanımı oldukça yaygındır. Leeden (1998) çalışmasında belirttiği gibi bu modeller pek çok farklı araştırma alanında uygulama olanağı bulmuştur. Dahası araştırmacılar bu modelleri hiyerarşik veri yapısına sahip olmayan farklı yapıdaki veri setlerinin analizinde de kullanarak bu yöntemin kullanım alanlarını daha da genişletmişlerdir.

Goldstien (2003) tarafından uygulanan çok deęişkenli verileri çok seviyeli model taslaęına uydurmak da yaygın bir yaklaşımdır. n gözlem ve m deęişken varsa; bunları n grup içine yuvalanmış ve her grupta m tane gözlem olan veri seti olarak düşünmek olasıdır. Bu yaklaşım $n \times m$ şeklinde $n \cdot m$ tane elemanı olan bir matris elde etmemizi sağlar. Dahası bu deneme planı, kayıp gözlemlerin bulunduğu veri setlerine de uygulanabilmektedir. Çünkü burada kayıp gözlemlere sahip olmanın anlamı basit olarak o grupta daha az gözlem olması olarak yorumlanabilir. Çok seviyeli modellerin kullanım alanlarını aşağıdaki başlıklarda özetlemek mümkündür.

2.1.1.1. Eğitim ve sosyoloji alanındaki kullanımı

Eğitim ve sosyoloji alanında 1980'lerden başlayarak çok seviyeli model analizi kullanılmaya başlanmıştır. Hiyerarşik veri yapısında bireylere ait bilgiler ya da diğer tip konulara ait bilgiler belli bir grubun içinde yuvalandığı için hiyerarşik veri yapısı pek çok çalışmada ortaya çıkmaktadır (Leeuw ve Meijer 2006). Bununla birlikte çok seviyeli modeller hızlı bir şekilde birçok farklı alanda kullanılmaya başlanmıştır. Örnek olarak, den Brok ve ark.'nın, 2006 yılında yaptığı eğitim alanı ile ilgili çalışmayı ve Papp'ın 2004 yılında davranış bilimleri alanında yaptığı çalışmaları verebiliriz.

2.1.1.2. Tıp ve sağlık alanındaki kullanımı

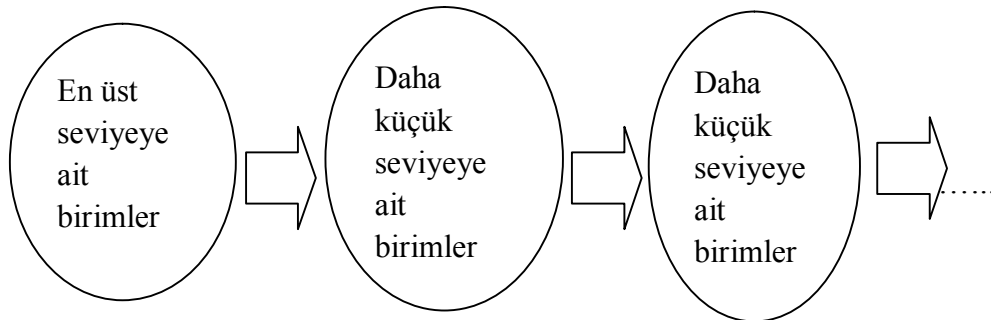
Özellikle son 20 yıl civarında yaygınlaşan çok seviyeli regresyon modeli tıp ve sağlık alanındaki çalışmalarda artan bir oranla kullanılmaktadır. Bu çalışmalara örnek olarak Bardenheier ve ark.'nın 2005'te, Burnside ve ark.'nın 2007'de, Tan ve ark.'nın 2007'de sağlık hizmetleri konusunda çalışmaları verilebilir. Bu çalışmalarda çok seviyeli modeller bilinen diğer istatistiksel yöntemlerle karşılaştırılmış ve çok seviyeli modellerin kullanılmasının avantajlı olduğu noktalar belirtilmiştir.

2.1.1.3. Zirai alandaki kullanımı

Genetik de, bitki ve hayvan ıslahı çalışmalarında hiyerarşik veri yapısı sıkça elde edilmektedir ve bu yüzden iç içe düzenlenmiş deneme planlarının kullanılması çok yaygındır (Fırat 2000). Çok seviyeli modellerin hiyerarşik veri analizinde kullanımı çok elverişli olduğu için zirai alanlarda da bu yöntemin kullanımı mümkündür. Örnek olarak Overmars ve Verburg (2006) , Sheng ve Chih' in (2008) tarla çalışmalarına ilişkin çalışmaları verebiliriz.

2.1.1.4. Anket çalışmalarında kullanımı

Anket çalışmalarını modellemede de çok seviyeli yöntemler kullanılmıştır. Bilindiği gibi birçok anket çalışmasında homojen bir popülasyondan elde edilmiş basit rastgele örneklem kullanılmamaktadır. Genellikle anket çalışmalarındaki örneklem, heterojen alt gruplardan iç içe örneklem yöntemiyle elde edilmektedir. Büyük anket çalışmaları genellikle çok seviyeli iç içe etkiler içerir. Örneklemeler elde edilirken genellikle aşağıdaki yol takip edilir.



Şekil 2.1. Anket çalışmalarında örneklem seçiminde izlenen yol

Bu şekilde ankete ait veri setindeki seviyeler oluşturulur. Bazen her seviyedeki birimlerin tamamı analize katılır ve bu da anketörlerin karmaşık bir veri yapısı elde etmesine neden olur. Birimlerin bazı açılardan farklı olduğu varsayımından dolayı anket

alışmalarında bu tarz karmaşık iç içe yapı kullanılmaktadır. Bu yüzden çok seviyeli modellerle gruplar arası heterojenliği modellemek çok normaldir, çünkü çok seviyeli modeller karmaşık veri yapısını modelleyebilecek ve analizler sonucunda tatminkâr sonuçlar verebilecek bir özelliğe sahiptir (Leeuw ve Kreft 2006).

2.1.1.5. Kategorik veri analizinde kullanımı

Çok seviyeli modellerin kategorik veriler için kullanımı istatistiksel alışmalarda aktif olarak kullanılmaktadır. Bilindiği gibi iki seviyeli (dichotomous) yanıt deęişkeni içeren veri setleri için, rastgele etki paylarının tahmini için lojistik ve probit regresyon modelleri ya da deęişik metotlar kullanılmıştır (Stiratelli ve ark. 1984, Anderson ve Aitkin 1985, Conaway 1989, D'Agostino ve ark. 1990, Drukker 2006, Wong ve Mason 1985). Snijders ve Bosker (1999) çok seviyeli regresyon modeli ile parametre tahminini kısa bir şekilde özetlemiş ve iki seviyeli (dichotomous) yanıt deęişkeni içeren veri setlerinin analizinde çok seviyeli lojistik regresyon modelinin yaygın bir şekilde kullanıldığını göstermişlerdir.

Daniels ve Gatsonis (1997), Revelt ve Train (1998), Skrondal ve rabe-Hesketh (2004) ve Goldstein (2003) çok seviyeli sınıflayıcı veri setleri (multilevel nominal veri) için daha genel regresyon modellerini göz önüne almışlardır. Çok seviyeli sıralayıcı ve çok sınıflı lojistik regresyon modeller (ordinal ve multinominal lojistik regresyon modelleri) kategorik veri analizi için tanımlanmıştır. Hem çok sınıflı verisi için geliştirilen modeller hem de sıralayıcı ve çok sınıflı lojistik regresyon modelleri, çok seviyeli regresyon modellerinin geliştirilmiş halidir. Cevap deęişkeni 2 ya da daha fazla kategori içerdiğinde bu modeller çok kullanışlıdır.

2.1.1.6. Tekrarlanan ölçümler veri setlerinde kullanımı

Tekrarlanan ölçümler içeren veri setlerinin analizinde de çok seviyeli regresyon modelleri kullanılmıştır. Bu tarz veri setlerinde değişkenler arasında bağımlılık, gözlemlerin yer değiştirmemesi gibi sorunlarla sıklıkla karşılaşılabilir (Leeuw ve Meijer 2006). Tekrarlanan ölçümler elde etmek amacıyla düzenlenmiş bir deneme planında, bireyler içinde yuvalanmış tekrarlı ölçümler en düşük seviye (seviye-1) olarak tanımlanabilir (Goldstein 1986). Yani seviye-1'i bireylere ait tekrarlı ölçümler oluştururken, seviye-2'yi bireyler oluşturmaktadır (Omar ve ark. 1999).

Tekrarlanan ölçümler aynı bireylerden belli zaman noktalarında ölçümler alınarak elde edilir. Genellikle burada tek bağımlı değişken vardır. Fakat bunların çoklu çıktı değişkenleri haline getirilmesi oldukça kolaydır. Ayrıca her bireyin aynı zaman diliminde ölçülmesi gibi bir zorunluluk yoktur yani her bir bireyden elde edilen ölçümler farklı zaman dilimlerinde de elde edilmiş olabilir. Bunun yanı sıra, veri setinde kayıp veriler de bulunabilir.

Bu veri setlerinde birden fazla seviyede çoklu cevaplar olduğu zaman problem ortaya çıkmaktadır. Bu tarz problemin tahmini için yapılan çok seviyeli model ve çoklu tekrarlanan doğru bir şekilde uygulanması önemlidir. Örneğin, büyüme verileri gibi tahmin etmeye çalışılan birey seviyesinde ve grup seviyesindeki verilerin tekrarlı ölçümler şeklinde elde edildiği çalışmalarda bu sorunla karşılaşılabilir.

Sonuç olarak bu tarz verilerin analizinde de çok seviyeli regresyon modellerini kullanmak mümkündür, çünkü bu veri setlerinde yer alan gözlemler arasında bağımlılık sorunu ile sıklıkla karşılaşılmaktadır.

2.1.2. Geliştirilen paket programlar

Çok seviyeli modellerin kullanımının yaygınlaşması ile bunların analizi için özel paket programları geliştirilmiş ve geliştirilmeye devam edilmektedir. Bu geliştirilen paket programlarının kısaca tarihçesine bakılacak olursa eğer, çok seviyeli modellerin teorisinin geliştirilmesine öncü olmuş kişiler tarafından çok seviyeli modelleme ile ilgili kitaplar yazılmış ve bunlara ait paket programlar geliştirilmiştir.

Leeuw ve Kreft (2001) tarafından “(and The state of affair) ca. 2000” geliştirilmiştir. Bu programda kullanılan bakış açısı hala geçerli olmasına rağmen, programın yazılımı üzerinde bazı oynamalar yapılmıştır. Aslında çok seviyeli model analizi için yazılmış olan iki program bu alanda yaygın bir şekilde kullanılmaktadır. Bu programlar HLM (Raudenbush ve ark. 2004) ve MLwiN (Rasbash ve ark. 2005)’dir. Bu programlar doğrusal ve doğrusal olmayan çok seviyeli modellerin analizi için geliştirilmiş ve geniş bir uygulama olanağı sağlamaktadır. Zaman içinde bu programların algoritmasında bazı değişiklikler yapılmıştır. Bunun yanı sıra, daha ileri seçeneklere (options) sahiptirler ve çok sık kullanılan model tanımlamalarında da bazı farklılıklar vardır. Fakat bu paket programlarında kullanılan farklı model tanımlamaları, bunların birini diğerinden daha iyi yapma gibi bir sonuca neden olmamaktadır.

VARCL önemli bir paket programıdır (Longford 1990). Fakat bu programın geliştirilmesi ile ilgili çalışmalara son verilmiştir.

Aslında çok seviyeli modellerin çok özel durumlarıyla ilgilenen, bunların farklı seçimlerine sahip olan programlar ya da farklı bakış açısıyla bu modellere yaklaşan birçok paket program vardır. Örneğin MLA (Busing ve ark. 2005), yeniden örneklem seçimi (resampling) konusunda çok kullanışlı bir programdır. PINT (Bosker ve ark. 1999) , güç hesaplama (power calculation) konusunda kullanılmaktadır. MIXFOO

(Hedeker ve Gibbons 1996, Hedeker ve Gibbons 1997) aslında genel çok seviyeli paket program olmasına rağmen, yinede bu kategoride yer alır.

BUGS ve BUGS'ın başka bir türü olan WinBUGS (Spiegelhalter ve ark. 2003), Bayesian veri analizi için kullanılmaktadır. Bayesian Çok Seviyeli Analizi için de çok kullanışlı bir programdır. Dahası, bu programlar doğrusal olmayan çok seviyeli modellerin tahminleri için kullanışlıdır (Spiegelhalter 2001).

SAS, SPSS gibi analiz aşamasında çok yaygın bir şekilde kullanılan istatistiksel paket programlar da çok seviyeli regresyon analizi yapabilecek seçeneğine (multilevel option) sahiptir. Örneğin SAS (2004) "PROC MIXED ve PROC NL MIXED", SPSS (2006), "MIXED" komutları ile çok seviyeli istatistiksel analizler de kullanılmaktadır.

Benzer şekilde Stata (2005) , glamm program (Rabe-Hesketh ve ark. 2004) ve R (2006) programlarında da çok seviyeli regresyon analizi yapabilecek seçimler vardır (Bates ve Sarkar 2006).

Çok seviyeli modellerinin yapısal eşitlik modellerinde kullanılmaya başlanması ile son zamanlarda geliştirilen LISREL (du Toit ve du Toit 2002), EQS (Bentler 2006) ve Mplus (Muhten ve Muhten 1998) paket programları da "multilevel" seçeneğine sahiptir. Fakat bu paket programlar kullanılarak yapılabilecek analizler diğerlerinden farklıdır. Bunlar doğrusal olmayan modeller, üç ya da daha fazla seviyesi olan modeller için daha az seçime sahiptir. Fakat bu programların çok değişkenli modeller ve ölçüm hatası ile latent değişkene sahip olan modellerin analizinde kullanılması daha uygundur. Ölçüm hatası ile latent değişkene sahip olan modellere örnek olarak çok seviyeli yapısal model verilebilir.

2.2. Kuramsal Bilgiler

Çok seviyeli regresyon modelleri, çok değişkenli regresyon modellerinin çok seviyeli şeklidir (Hox 1995). Cohen (1998) ve diğer araştırmacıların gösterdiği gibi çoklu regresyon modelleri değişken yapıdadır. Kategorik değişkenler için yapay (dummy) değişken kodlaması kullanılarak, kategorik değişkenlerin analizinde ANOVA modeli ve çoklu regresyon modelleri kullanılabilir. Bunun sonucunda da çok seviyeli regresyon modelleri oldukça yaygın bir şekilde birçok değişik araştırma problemlerinin analizinde kullanılmaktadır.

Çok seviyeli regresyon modellerinden bahsedilmeden önce kısaca tek değişkenli ve çok değişkenli varyans ve iç içe etkenli varyans analizinden bahsedilecektir.

2.2.1. Varyans analizi

2.2.1.1. Tek faktörlü varyans analizi

Yanıt değişkenini/bağımlı değişkeni etkileyen tek bir faktörün/muamelenin düzeylerinin arasında önemli bir fark olup olmadığı araştırılmak istensin (Erbaş ve Olmuş 2006). Bu durumda veri seti Çizelge 2.1’ deki gibidir.

Çizelge 2.1. Tek faktörlü deneme deseninin veri seti

A_1	A_2	...	A_j
Y_{11}	Y_{12}	...	Y_{1j}
Y_{21}	Y_{22}	...	Y_{2j}
$\vdots$	$\vdots$	...	$\vdots$
Y_{i1}	Y_{i2}	...	Y_{ij}

Albright (2007), modeli (2.1)'ki şekilde vermiştir ve bu modelin genel gösterimidir.

$$y_{ij} = \mu + \alpha_i + e_{ij} \quad i=1,2,\dots,a \quad j=1,2,\dots,n$$

(2.1)

y_{ij} : Bağımlı değişken

μ : Genel ortalama (deneyin ortak) etkisi

α_i : A faktörünün i-inci birimin seviyesinin etkisi

e_{ij} : Hata terimi (j-inci muamelenin i-inci gözlemi y_{ij} 'ye (yanıt değişkenine) ilişkin rastgele hata)

Eğer model sabit etkili model ise varsayımlar;

- α_i 'ler için $\sum_i \alpha_i = 0$
- Hata terimleri $e_{ij} \stackrel{iid}{\sim} N(0, \sigma_e^2)$ ve burada hata terimleri birbirinden bağımsız aynı dağılımlı (*iid* olarak da kısaltılabilir) normal dağılım göstermektedir.

Sabit etkili model için hipotez;

$$\begin{cases} H_0 : \alpha_1 = \alpha_2 = \dots = \alpha_i \\ H_1 : \text{enaz bir } \alpha_i \text{ farklı} \end{cases} \quad \text{veya} \quad \begin{cases} H_0 : \mu_1 = \mu_2 = \dots = \mu_i \\ H_1 : \text{enaz bir } \mu_i \text{ farklı} \end{cases}$$

şeklindedir.

Eğer model rastgele etkili model ise varsayımlar;

- α_i 'ler için $\alpha_i \stackrel{iid}{\sim} N(0, \sigma_\alpha^2)$
- Hata terimleri için $e_{ij} \stackrel{iid}{\sim} N(0, \sigma_e^2)$

Burada α_i terimleri ile e_{ij} terimleri birbirinden bağımsızdır.

Sabit etkili model için hipotez $\begin{cases} H_0 : \sigma_\alpha^2 = 0 \\ H_1 : \sigma_\alpha^2 > 0 \end{cases}$ şeklindedir.

2.2.1.2. İki faktörlü varyans analizi

Çok değişkenli varyans analizi, iki ya da daha fazla bağımsız ve bağımlı gruplarda çok değişkenli normal dağılıma dayanan hipotezleri test etmek üzere geliştirilmiş bir yöntemdir (Özdamar 2004). Veri seti Çizelge 2.2’de verilmiştir.

Çizelge 2.2. İki faktörlü deneme deseninin veri seti

	B_1	B_2	...	B_b
A_1	$Y_{111}, Y_{112}, \dots, Y_{11n}$	$Y_{121}, Y_{122}, \dots, Y_{12n}$	...	$Y_{1b1}, Y_{1b2}, \dots, Y_{1bn}$
A_2	$Y_{211}, Y_{212}, \dots, Y_{21n}$	$Y_{221}, Y_{222}, \dots, Y_{22n}$	...	$Y_{2b1}, Y_{2b2}, \dots, Y_{2bn}$
$\vdots$	$\vdots$	$\vdots$	...	$\vdots$
A_a	$Y_{a11}, Y_{a12}, \dots, Y_{a1n}$	$Y_{a21}, Y_{a22}, \dots, Y_{a2n}$	...	$Y_{ab1}, Y_{ab2}, \dots, Y_{abn}$

Model genel olarak (2.2) eşitliğindeki gibidir (Albright 2007).

$$y_{ijk} = \mu + \alpha_i + \beta_j + e_{ijk} \quad i = 1, 2, \dots, a, \quad j = 1, 2, \dots, b \text{ ve } k = 1, 2, \dots, n \quad (2.2)$$

Burada;

y_{ijk} : Bağımlı değişken

μ : Genel ortalama/deneyin ortak etkisi

α_i : A faktörünün/muamelesinin i-inci seviyesinin etkisi

β_j : B faktörünün/muamelesinin j-inci seviyesinin etkisi

e_{ijk} : Hata terimidir.

Eğer model sabit etkili model ise varsayımları;

- α_i 'ler için $\sum_i \alpha_i = 0$
- β_j 'ler için $\sum_j \beta_j = 0$
- Hata terimleri için $e_{ij} \stackrel{iid}{\sim} N(0, \sigma_e^2)$

Sabit etkili model için hipotez;

$$\begin{cases} H_0 : \alpha_1 = \alpha_2 = \dots = \alpha_i \\ H_1 : \text{enaz bir } \alpha_i \text{ farklı} \end{cases} \quad \text{veya} \quad \begin{cases} H_0 : \beta_1 = \beta_2 = \dots = \beta_b \\ H_1 : \text{enaz bir } \beta \text{ farklı} \end{cases}$$

şeklindedir.

Eğer model rastgele etkili model ise varsayımlar;

- α_i 'ler için $\alpha_i \stackrel{iid}{\sim} N(0, \sigma_\alpha^2)$
- β_j 'ler için $\beta_j \stackrel{iid}{\sim} N(0, \sigma_\beta^2)$
- Hata terimleri için $e_{ij} \stackrel{iid}{\sim} N(0, \sigma_e^2)$

Burada α_i ve β_j terimleri ile e_{ij} terimleri birbirinden bağımsızdır.

Rastgele etkili model için hipotez $\begin{cases} H_0 : \sigma_\alpha^2 = 0 \\ H_1 : \sigma_\alpha^2 > 0 \end{cases}$ şeklindedir.

İki faktörlü etkileşimli (interaksiyonlu) modeller yazılabilir. Örneğin iki faktörlü/muameleli karışık etkili interaksiyonlu model şu şekildedir;

$$y_{ijk} = \mu + \alpha_i + \beta_j + (\alpha\beta)_{ij} + e_{ijk} \quad (2.3)$$

Bu modelde etkenlerden birisi sabit birisi rastgele etkilidir ve

α_i : Sabit etkiyi temsil etmektedir.

β_j : Rastgele etkiyi temsil etmektedir.

$(\alpha\beta)_{ij}$: Sabit etki ile rastgele etki arasındaki etkileşim terimidir.

Hipotezler ;

$$\left\{ \begin{array}{l} H_0 : \alpha_1 = \alpha_2 = \dots = \alpha_i \\ H_1 : \text{enaz bir } \alpha_i \text{ farklı} \end{array} \right. \quad \left\{ \begin{array}{l} H_0 : \sigma_\beta^2 = 0 \\ H_1 : \sigma_\beta^2 > 0 \end{array} \right. \quad \left\{ \begin{array}{l} H_0 : \sigma_{\alpha\beta}^2 = 0 \\ H_1 : \sigma_{\alpha\beta}^2 > 0 \end{array} \right.$$

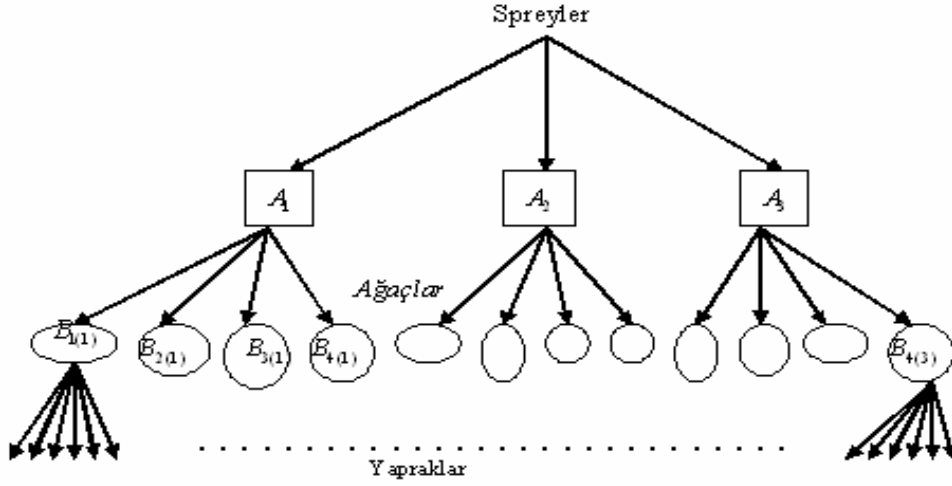
olarak tanımlanır.

2.2.1.3. İç içe sınıflandırma

Etkileri arasında bir hiyerarşi söz konusu olan tasarımlara İç İçe Etkili Deney Tasarımı (Experimental Design with Nested Factors) denir. Çeşitli alt koşullarda yapılmış denemeler ana koşul altında birleştirilerek iç içe gruplar oluşturulur ve bulgular İç içe Gruplar Varyans Analizi ile değerlendirilir.

Veri setinin hiyerarşik yapısını gösterebilmek amacıyla ziraat ile ilgili bir örnek inceleyebiliriz. Sprey şeklindeki 3 çeşit ilacın her biri rastgele seçilen 4'er ağaca uygulansın. Bir hafta sonra 12 ağaçtan rastgele koparılmış 6'şar yaprağın azot değişimi ölçülsün. Burada ilaç düzeyleri özel seçimli, ağaç etkisinin düzeyleri ise rastgele seçimlidir. İlaçların uygulandığı her ağaçta bulunan çok sayıda yaprak arasından rastgele seçilen 6'şar yaprak birer alt örneklem oluşturmaktadır. Bu deneyde 12 ağaca ilaç spreyleri uygulanmış ve toplam 72 yaprağın azot değişimi gözlenmiştir. Azot

derişimi ölçümleri Y_{ijk} olmak üzere i indisi spreyi, j indisi ağacı ve k indisi yaprağı temsil etmektedir. Şekil 2.2’de gözlemlere yer verilmiştir.



Y_{111}	Y_{121}	Y_{131}	Y_{141}	Y_{211}	Y_{221}	Y_{231}	Y_{241}	Y_{311}	Y_{321}	Y_{331}	Y_{341}
Y_{112}	Y_{122}	Y_{132}	Y_{142}	Y_{212}	Y_{222}	Y_{232}	Y_{242}	Y_{312}	Y_{322}	Y_{332}	Y_{342}
...	...	...	...	...	...	...	...	...	...	...	...
Y_{116}	Y_{126}	Y_{136}	Y_{146}	Y_{216}	Y_{226}	Y_{236}	Y_{246}	Y_{316}	Y_{326}	Y_{336}	Y_{346}

Şekil 2.2. İç içe etkili deneme desenine ilişkin veri yapısı

İlgili model;

$$Y_{ijk} = \mu + \alpha_i + \beta_{j(i)} + \varepsilon_{(ij)k} \quad i=1,2,3 \quad j=1,2,3,4 \quad k=1,2,\dots,6 \quad (2.4)$$

biçiminde ifade edilir. Burada

$$\sum_{i=1}^3 \alpha_i = 0 \quad \beta_{j(i)} \sim N(0, \sigma_b^2) \quad \varepsilon_{(ij)k} \sim N(0, \sigma_\varepsilon^2) \quad \text{ve } \beta_{j(i)} \text{ ile } \varepsilon_{(ij)k} \text{ 'lar birbirinden}$$

bağımsızdır. Burada azot derişimi dört parçaya ayrılmış oldu. Bunların ilki genel ortalama μ ’dür, μ üç sprej’in tümü için ortalama azot derişimidir. İkincisi α_i , özel bir spreje ait parçadır. α_i , pozitif ise i -inci sprej ortalamadan daha yüksek azot

derişimi üretir. Üçüncü parça $\beta_{j(i)}$, i -inci spreyn kullanıldığı j -inci ağaca aittir. İlâçlama yaprak-yaprak değil de, ağaç-ağaç yapıldığı için ağaç etkisinin düzeyleri ilâç etkisinin düzeyleri içinde yuvalanmıştır. Bu yuvalanma $\beta_{j(i)}$ terimindeki indisler vasıtasıyla ifade edilmiştir. Yaprakların da ağaçlardan koparıldığı düşünülürse hata terimi $\varepsilon_{(ij)k}$ biçiminde gösterilebilir. Ayrıca burada gözlemler hiyerarşi yapısının en son seviyesinde ortaya çıkmaktadır.

İç içe etkenli modellerin çok seviyeli modellerden farkı; çok seviyeli modellerde her seviyenin kendisine ait özel açıklayıcı değişkenleri vardır ve bunlarla bağımlı değişkenler açıklanmaktadır. İç içe etkili modellerde ise hiyerarşi yapısında yer alan seviyelerde açıklayıcı değişkenler bulunmamaktadır.

2.2.2. Çok seviyeli modellerin teorisi

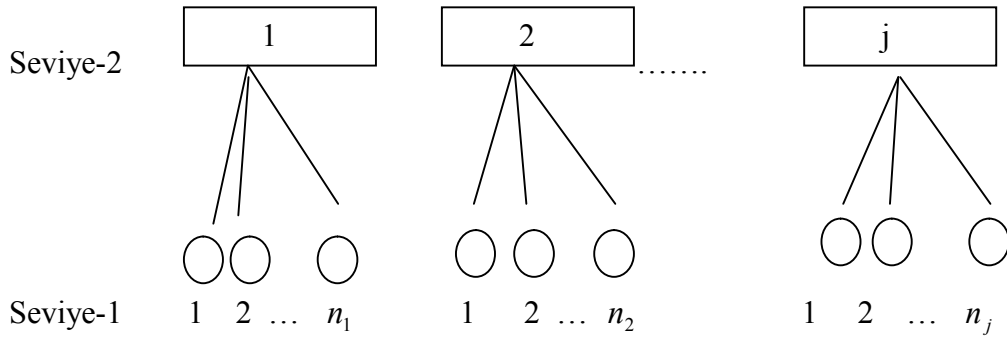
Regresyon analizinde olduğu gibi, gözlenmiş verilerle çok seviyeli regresyon modeline ait regresyon katsayıları ve varyans bileşenleri (yani kısaca model parametrelerini) tahmin edilmek istenmektedir.

Çok seviyeli modellerin teorisi ilerlemiş yöntemlere göre daha az gelişmiş durumdadır (Eeden, 1993). Şimdiye kadar çok seviyeli modellerin kullanılabilmesi için, grup kriterlerinin açık olmasının gerekli olduğu ve ele alınan değişkenlerin ait oldukları seviyenin açık bir şekilde kendi seviyelerine atanması gerektiğinden bahsedildi. Gerçekte, grup sınırları bazen açık değildir ve keyfi olabilir. Bu yüzden değişkenlerin seviyelere atanması görüldüğü kadar basit ve kolay değildir. Ayrıca çok seviyeli problemlerin grup üyeliğine ilişkin kararlar ve gözlemlere göre bir kavrama tanımlama süreci geniş bir teorik varsayım gerektirmektedir (Blalock 1990). Bunların sonucu olarak da çok seviyeli modellerin teorisi karmaşıklığını korumaktadır.

Teorik olarak bakıldığında, çok değişkenli modellerin, veri yapısına ilişkin seviye yapısının da modelde ele alınması ile çok seviyeli regresyon modelleri elde edilir. Çok seviyeli regresyon modelinde, en düşük seviyede bir bağımlı değişkenin olduğu ve bulunan diğer seviyelerin her birinde de açıklayıcı değişkenlerin bulunduğu hiyerarşik veri yapısının olduğu varsayılır. Çok seviyeli regresyon modeli regresyon eşitliklerinin bir hiyerarşik versiyonu olarak da değerlendirilebilir. Ayrıca çok seviyeli modeller, popülasyonun hiyerarşik yapıya sahip olduğu ve seviye-1 ve seviye-2'ye ait örnekleminin rasgele olarak yapıldığını varsaymaktadır. (Hox ve Kreft 1994).

2.2.2.1. İki seviyeli regresyon modeli

Regresyon analizinde olduğu gibi iki seviyeli regresyon modelinde bağımlı değişken ile bağımsız değişken/değişkenler arasındaki ilişki matematiksel olarak modellenir (Özdamar 2004). Sullivan ve ark. (1999) iki seviyeli veri yapısını aşağıdaki şekilde göstermişlerdir.



Şekil 2.3. İki seviyeli hiyerarşik veri yapısı

Seviye-1'de bağımlı değişken Y'nin ve açıklayıcı değişken X'in bulunduğu, seviye-2'de de açıklayıcı değişken Z 'nin bulunduğu iki seviyeli regresyon modelinin elde edilişi aşağıdaki gibi gösterilmiştir. J tane birimin olduğu (seviye-2'deki birim sayısı) durum için model;

$$Y_{ij} = \beta_{0j} + \beta_{1j}X_{ij} + \varepsilon_{ij} \quad i=1,2,\dots, n_j \quad j=1,2,\dots,J \quad \varepsilon_{ij} \sim N(0, \sigma_\varepsilon^2) \quad (2.5)$$

İlk bakışta bu model ile sıradan regresyon modeli arasında bir fark yok gibi gözükmektedir. Fakat çok seviyeli regresyon modelindeki sabit regresyon katsayısı olarak bilinen β_{0j} ve eğim katsayısı β_{1j} 'nın seviye-2'deki her bir birimi için farklı regresyon eşitlikleri vardır. Diğer bir deyişle, sabit ve eğim katsayısı seviye-2 birimlerinin her biri arasında değişim göstermektedir, çünkü her birimin kendisine ait katsayı değerleri vardır. Dolayısıyla bu katsayılar genellikle rasgele katsayılar olarak ele alınmaktadır. β_{0j} ve β_{1j} katsayılarının kendilerine ait regresyon denklemleri (2.6) eşitliğinde verilmiştir.

$$\begin{aligned} \beta_{0j} &= \gamma_{00} + \gamma_{01}Z_j + u_{0j} \\ \beta_{1j} &= \gamma_{10} + \gamma_{11}Z_j + u_{1j} \end{aligned} \quad (2.6)$$

(2.6)'daki eşitlikler ile seviye-2'de bulunan açıklayıcı değişken ya da değişkenlerle regresyon katsayısı olan β_j 'ler kestirilir ve bunların arasındaki varyasyon tahmin edilir. (2.6) eşitliğinde yer alan u_{0j} ve u_{1j} 'ler seviye-2'deki rastgele hata terimlerini göstermektedir. u_j 'lerin ortalamasının 0 olduğu ve ε_{ij} hata terimlerinden bağımsız olduğu varsayılmaktadır. Ayrıca u_{0j} 'nin varyansı σ_{00} ve u_{1j} 'nin varyansı σ_{11} olarak tanımlanmıştır ve u_{0j} ile u_{1j} arasındaki kovaryansı temsil eden σ_{12} 'nin özel olarak 0 olduğu varsayılmıştır.

Not olarak; eşitlik (2.6)'teki regresyon katsayısı γ 'ların seviye-2'deki birimler arasında değişimi göstermediği varsayılın. Bunun sonucu olarak j indisine artık gerek duyulmaz çünkü j indisi ile bu değerlerin hangi birime ait olduğu gösterilmektedir. Bu durumda γ 'lar sabit katsayı olarak adlandırılır ve β katsayıları seviye-2'deki açıklayıcı değişken Z_j 'ler ile tahmin edilir ve hata terimi varyansı da hata terimleri u_j 'lerin içinde yer alır. Bu yüzden u_j 'ler j indisine sahiptir çünkü j indisi sayesinde bu hataların hangi seviye-2 birimine ait olduğu belirtilmektedir.

Eşitlik (2.6) , (2.5)'te yerine yazılırsa eşitlik (2.7)'de yer alan çok seviyeli model elde edilir.

$$Y_{ij} = \gamma_{00} + \gamma_{10}X_{ij} + \gamma_{01}Z_j + \gamma_{11}Z_jX_{ij} + u_{1j}X_{ij} + u_{0j} + \varepsilon_{ij} \quad (2.7)$$

Z_jX_{ij} terimi etkileşim terimidir ve bu terim seviye-2'ye ait Z_j değişkeni ile seviye-1'e ait X_{ij} değişkenine ait değişik β_{ij} regresyon eğim katsayısını gösteren ve modelleme yapısının bir sonucu olarak modelde ortaya çıkan terimdir. Yani veri yapısının çok seviyeli olmasından dolayı böyle bir terim modelde ortaya çıkmaktadır. Çok seviyeli regresyon analizindeki etkileşim terimlerini yorumlamak biraz karmaşıktır (Hox 1995).

Ayrıca çok seviyeli veri setlerinde grup içinde yer alan gözlemler bağımlılıktan dolayı grup içi korelasyon (*ICC*) katsayısı ortaya çıkmaktadır. Grup yapısı ile elde edilmiş, populasyonun varyans tahmini grup içi korelasyonu ρ 'dur. Grup içi korelasyonu grup seviyesine ait varyansı, tahmin edilen toplam varyansa oranlanmasının tahmini olarak ifade edilebilir. $ICC = \rho = \frac{\sigma_{00}}{\sigma_{00} + \sigma^2}$ formülü yardımı ile *ICC* kestirilir. Aynı zamanda *ICC* popülasyondaki açıklanan varyans (explained variance) oranının tahmini olduğu belirtilmelidir.

Çizelge (2.3)'de çok seviyeli regresyon modeli ile tek seviyeli regresyon modeli arasındaki fark görülebilsin diye verilmiştir. Açıkça görüldüğü üzere çok seviyeli modelde yer alan β katsayılarının hepsinin kendisine ait regresyon eşitlikleri vardır ve bu da bu modelleri birbirinden farklı kılmaktadır.

Çizelge 2.3. Tek ve çok seviyeli regresyon modelleri

Analiz Türü	Model
Tek Seviyeli Regresyon Analiz	$Y_{ij} = \beta_0 + \beta_1 X_i + \varepsilon_{ij}$
Çok Seviyeli Regresyon Analiz	$Y_{ij} = \beta_{0j} + \beta_{1j} X_{ij} + \varepsilon_{ij}$ $\beta_{0j} = \gamma_{00} + \gamma_{01} Z_j + u_{0j}$ $\beta_{1j} = \gamma_{10} + \gamma_{11} Z_j + u_{1j}$

2.2.2.2. İki seviyeli modelin adımsal olarak elde edilişi

Adım 1: Hiçbir seviyede açıklayıcı değişkenin olmadığı varsayalım. Model aşağıdaki gibidir.

$$Y_{ij} = \beta_{0j} + \varepsilon_{ij} \quad (2.8)$$

Y : seviye-2'deki j -inci birim ve seviye-1'deki i -inci birime çıktı değişkeni

ε_{ij} : Seviye-1'e ait hata terimi

$$\beta_{0j} = \gamma_{00} + u_{0j} \quad (2.9)$$

Açıklayıcı değişkenler modellenmediği için (2.8)'de ki modelde β_{1j} katsayısı yoktur. β_{0j} katsayısına ilişkin regresyon denklemi ise eşitlik (2.9)'daki şekilde tanımlanmış olup (2.8) eşitliğinde yerine yazılırsa eşitlik (2.10) elde edilir.

$$Y_{ij} = \gamma_{00} + u_{0j} + \varepsilon_{ij} \quad (2.10)$$

Burada ε_{ij} 'ye ait varyans σ^2 ve u_{0j} 'ye ait varyans σ_{00} 'dur.

eşitlik (2.10)'daki modelden ICC katsayısı doğrudan kestirebilir. $ICC = \rho = \frac{\sigma_{00}}{\sigma_{00} + \sigma^2}$ yardımı ile hesaplanır.

Adım 2: Seviye-2'de X açıklayıcı değişkeni olduğu durumda model;

$$Y_{ij} = \beta_{0j} + \beta_{1j}X_{ij} + \varepsilon_{ij} \quad (2.11)$$

şeklinde tanımlanır. Eğer bu açıklayıcı değişkeni rasgele etkili ise β_{1j} katsayısının kendisine ait bir regresyon eşitliği olacaktır ve bu da eşitlik (2.12) verilmiştir.

$$\beta_{1j} = \gamma_{10} + u_{1j} \quad (2.12)$$

u_{1j} 'ye ait varyans σ_{11} 'dur.

Adım 3: σ_{00} anlamlı iken; seviye-2'de varyansa etkisi olan değişkenler modele alınabilir. Eğer seviye-2'de bulunan iki tane açıklayıcı değişkenin (Z_j ve W_j) bağımlı değişken üzerindeki etkisi araştırılmak istenirse, β_{0j} katsayısına ait regresyon modeli eşitlik (2.13) şekilde olur.

$$\beta_{0j} = \gamma_{00} + \gamma_{01}Z_j + \gamma_{02}W_j + u_{0j} \quad (2.13)$$

Adım 4: Seviye-1'deki X açıklayıcı değişkeni, seviye-2'deki değişkenleri farklı şekilde etkiliyor olabilir. Bu durumu araştırabilmek için β_{1j} regresyon modeli de belirtilmelidir ve aşağıdaki gibi düzenlenir.

$$\beta_{1j} = \gamma_{10} + \gamma_{11}Z_j + \gamma_{12}W_j + u_{1j} \quad (2.14)$$

Adım 5: β_{0j} ve β_{1j} eşitlikleri (2.11)'de yerine yazılırsa; seviye-2'de yer alan açıklayıcı değişkenler de modele alınmış olur ve böylelikle eşitlik (2.15) çok seviyeli regresyon modeli elde edilmiş olur.

$$Y_{ij} = \gamma_{00} + \gamma_{01}Z_j + \gamma_{02}W_j + \gamma_{10}X_{ij} + \gamma_{11}X_{ij}Z_j + \gamma_{12}X_{ij}W_j + u_{0j} + u_{1j}X_{ij} + \varepsilon_{ij} \quad (2.15)$$

Yukarıdaki modelde de görüldüğü gibi çok seviyeli regresyon modelinde her bir seviyeye ait açıklayıcı değişkenler ele alınmıştır. Bu modeldeki hangi ifadelerin rastgele hangi ifadelerin sabit etkili olduğu (2.16)'da verilmiştir.

$$Y_{ij} = \underbrace{\gamma_{00} + \gamma_{01}Z_j + \gamma_{02}W_j}_{\text{sabit etkili}} + \underbrace{\gamma_{10}X_{ij} + \gamma_{11}X_{ij}Z_j + \gamma_{12}X_{ij}W_j}_{\text{sabit etkili}} + \underbrace{u_{0j} + u_{1j}X_{ij}}_{\text{rastgele etkili}} + \varepsilon_{ij} \quad (2.16)$$

Bu model kullanılarak tahminler yapılırsa; X değişkeni için hem sabit hem de rastgele etkili kısma ait tahminler elde edilir. Sabit etkili kısım; seviye-1'deki açıklayıcı değişkenin bağımlı değişken üzerindeki beklenen genel etkisini göstermektedir. Rastgele etkili kısım ise bu etkinin seviye-2'deki birimler arasında farklılık gösterip göstermediğine ilişkin bilgi vermektedir (Hox 1995 ve Albright 2007).

2.2.2.3. İki seviyeli modelin matrissel ile ifade edilişi

Yukarıdaki bölümlerde görüldüğü gibi çok seviyeli regresyon modeli uzun ve karmaşıktır. Adım adım yaklaşımı kullanarak çok seviyeli modellerin gösterimi yerine daha genel bir gösterim elde edebilmek için matris notasyonu da kullanılabilir. Her seviyesinde tek açıklayıcı değişkenin bağımlı değişken üzerine etkisi araştırılıyorsa model (2.17)'deki gibi yazılabilir.

$$\mathbf{Y}_j = \mathbf{X}_j \mathbf{Z}_j \boldsymbol{\gamma} + \mathbf{X}_j \mathbf{u}_j + \boldsymbol{\varepsilon}_j \quad j=1,2,\dots,J \quad (2.17)$$

$\mathbf{Y}$: nx1 büyüklüğündeki yanıt vektörünü

$\mathbf{X}$: nxp büyüklüğündeki sabit etkileri içeren matris

$\mathbf{Z}$: nxq büyüklüğündeki rastgele etkileri içeren matris

$\mathbf{u}$: qx1 büyüklüğündeki rastgele etkenler vektörünü

$\boldsymbol{\varepsilon}$: nx1 büyüklüğündeki hata terimleri vektörlerini

temsil etmektedir.

Sullivan ve ark. (1999) makalelerinde matris notasyonu çok seviyeli modellerin gösterimini aşağıdaki gibi vermişlerdir.

Seviye-1 modeli : $\mathbf{Y}_j = \mathbf{X}_j \boldsymbol{\beta}_j + \boldsymbol{\varepsilon}_j$, $j=1,2,\dots,J$ şeklindedir. Burada,

$$\mathbf{Y}_j = \begin{bmatrix} Y_{1j} \\ Y_{2j} \\ \vdots \\ Y_{nj} \end{bmatrix}, \mathbf{X}_j = \begin{bmatrix} 1 & X_{1j} \\ 1 & X_{2j} \\ \vdots & \vdots \\ 1 & X_{nj} \end{bmatrix}, \boldsymbol{\beta}_j = \begin{bmatrix} \beta_{0j} \\ \beta_{1j} \end{bmatrix}, \boldsymbol{\varepsilon}_j = \begin{bmatrix} \varepsilon_{1j} \\ \varepsilon_{2j} \\ \vdots \\ \varepsilon_{nj} \end{bmatrix}$$

Seviye-2 modeli : $\boldsymbol{\beta}_j = \mathbf{Z}_j \boldsymbol{\gamma} + \mathbf{u}_j$, $j=1,2,\dots,J$ şeklindedir. Burada,

$$\boldsymbol{\beta}_j = \begin{bmatrix} \beta_{0j} \\ \beta_{1j} \end{bmatrix}, \mathbf{Z}_j = \begin{bmatrix} 1 & Z_j & 0 & 0 \\ 0 & 0 & 1 & Z_j \end{bmatrix}, \boldsymbol{\gamma} = \begin{bmatrix} \gamma_{00} \\ \gamma_{01} \\ \gamma_{10} \\ \gamma_{11} \end{bmatrix}, \mathbf{u}_j = \begin{bmatrix} u_{0j} \\ u_{1j} \end{bmatrix}$$

Birleştirilmiş model (yani çok seviyeli regresyon modeli) (2.18)deki şekilde olur.

$$\mathbf{Y}_j = \mathbf{X}_j (\mathbf{Z}_j \boldsymbol{\gamma} + \mathbf{u}_j) + \boldsymbol{\varepsilon}_j \quad \Rightarrow \quad \mathbf{Y}_j = \mathbf{X}_j \mathbf{Z}_j \boldsymbol{\gamma} + \mathbf{X}_j \mathbf{u}_j + \boldsymbol{\varepsilon}_j \quad j=1,2,\dots,J \quad (2.18)$$

Özel olarak SAS paket programında kullanılan çok seviyeli model ise;

$\mathbf{A}_j = \mathbf{X}_j \mathbf{Z}_j$ olmak üzere,

$$\mathbf{Y}_j = \mathbf{A}_j \boldsymbol{\gamma} + \mathbf{X}_j \mathbf{u}_j + \boldsymbol{\varepsilon}_j \quad j=1,2,\dots,J \quad (2.19)$$

olarak verilmiştir. Burada $\mathbf{A}_j = \begin{bmatrix} 1 & W_j & X_{1j} & W_j X_{1j} \\ 1 & W_j & X_{2j} & W_j X_{2j} \\ \vdots & \vdots & \vdots & \vdots \\ 1 & W_j & X_{nj} & W_j X_{nj} \end{bmatrix}$ ve ayrıca

$$\varepsilon_j \sim N(0, \sigma^2 I_{nj}), \mathbf{u}_j \sim N(0, G), \mathbf{G} = \begin{bmatrix} \sigma_{00} & \sigma_{10} \\ \sigma_{01} & \sigma_{11} \end{bmatrix} \text{ dir.}$$

2.2.2.4. Birden fazla açıklayıcı değişkenin modellenmesi

Genellikle en yüksek ve en düşük seviyede birden fazla açıklayıcı değişken bulunur. Eğer düşük seviyede P tane X açıklayıcı değişkenimiz varsa ve $p=1,2,...,P$ ise ve en yüksek seviyede Q tane Z açıklayıcı değişkenimiz varsa ve $q=1,2,...,Q$ ise model eşitlik (2.20)'ki gibidir (Hox 1995).

$$Y_{ij} = \gamma_{00} + \gamma_{p0}X_{pj} + \gamma_{0q}Z_{qj} + \gamma_{pq}Z_{qj}X_{pj} + u_{pj}X_{pj} + u_{0j} + \varepsilon_{ij} \quad (2.20)$$

Parametre sayısı: Çok seviyeli regresyon analizindeki parametre sayısı diğer modellere göre nispeten daha fazladır. Eğer en düşük seviyede (seviye-1'de) P tane açıklayıcı değişken ve en yüksek seviyede de (seviye-2'de) Q tane açıklayıcı değişken varsa çok seviyeli modelde yer alacak parametre sayısı Çizelge 2.4' te verilmiştir. Görüldüğü üzere çok seviyeli modellerde açıklayıcı değişken sayısı çok olduğu zaman, modelde yer alan parametreler çoğalmaktadır. Bu durum modelin karmaşık bir hal almasına sebep olmakta ve dolayısıyla çok seviyeli modellere ait istatistiksel teoremin çok karmaşık bir hal almasına sebep olmaktadır. Bu yüzden modele alınacak açıklayıcı değişkenlerin seçimi çok önemlidir (Hox 1995).

Çizelge 2.4. Çok seviyeli regresyon modelindeki parametre sayısı

Parametreler	Değişken Sayısı
Sabit	1
En düşük seviye tahmin edicileri için eğim	P
En düşük seviye tahmin edicileri için eğim	P
Bu hatalar için en yüksek seviye hata terimleri	P
En yüksek seviyedeki sabitlerin kovaryansı ile en yüksek seviyedeki eğimler arasındaki kovaryans	P(P-1)/2
En yüksek seviye tahmin edicileri için eğim	Q
Bu eğim terimleri için croos-level interactions	PQ

Model varsayımları: Çok seviyeli modellerde sabit etkili değişkenler için hata terimlerinin beklenen değerinin 0 olduğu varsayılmaktadır. Rastgele etkili açıklayıcı değişkenlerden oluşan modelde ise hata terimleri ile açıklayıcı değişkenlerin ilişkisiz olduğu varsayılmaktadır. Bu rastgele etkili model varsayımları, modelde yer alan varyans unsurlarına ilişkin olup aşağıda listelenmiştir (Snijders ve Bosker 1993, Hox 1995).

- $\varepsilon_{ij} \sim N(0, \sigma^2)$ olarak dağıldığı varsayılmaktadır.
- u_{0j} artık hatalarının varyansı, gruplar arasındaki sabit terimlerine ait varyans olup özel olarak σ_{00} olarak ifade edilir.
- u_{pj} artık hatalarının varyansı, gruplar arasındaki eğim terimlerine ait varyans olup özel olarak σ_{pp} olarak ifade edilir.
- u terimleri u_{pj} ve u_{0j} yani en yüksek seviyeye ait hata terimleri olup bunların birey seviyesine ait hata terimleri olan ε_{ij} ’lerden bağımsız olduğu varsayılmaktadır. Ayrıca u_{pj} ve u_{0j} ’ler ortalaması 0 olan normal dağılıma sahiptir.

- Hata terimleri arasındaki kovaryansın genellikle 0 olmadığı varsayılır ve bunlar en yüksek seviyedeki varyans kovaryans matrisi olan Σ 'dan elde edilir.

Regresyon analizinde model varsayımları incelendiği zaman hata terimlerinin dağılımının normal dağılıma uygunluk göstermesi istenen bir durumdur. Çok seviyeli regresyon modellerinde ise böyle bir zorunluluk yoktur, bu konuda bu modeller daha esnektir. Eğer rastgele değişkenler için normallik varsayımı sağlanamıyorsa, bu değişkenler için farklı dağılım varsayımları kullanılır. Örneğin t-dağılımı (Seltzer ve ark. 1996) ya da normal dağılımların karışımı (Verbeke ve Lesaffre 1996) kullanılabilir. Buradaki varyansların homojenliği varsayımı, seviye-2'nin popülasyondan rasgele örnekleme ile elde edilmesi varsayımı ile nerdeyse aynıdır.

Ayrıca geleneksel istatistik testleri gözlemlerin bağımsızlığı üzerine kuruludur. Eğer bu varsayım sağlanmıyorsa (ki genellikle çok seviyeli veride bu durumla karşı karşıya kalınmaktadır) geleneksel istatistik testleri ile elde edilen standart hatalar gereğinden çok küçük olmakta ve bunun sonucunda da yanıltıcı “anlamlı” sonuçlar elde edilmesine sebep olmaktadır (Hox 1995). Bundan dolayı çok seviyeli örneklemelerin analizi için çok seviyeli modellere ihtiyaç duyulmaktadır. Çünkü çok seviyeli modellemede, grup içi korelasyona yer verilerek değişkenler arasındaki bağımlılıklar modelde incelenebilmektedir.

En iyi tasarımın dirençliliği: Çok seviyeli çalışmaların analizinde kullanılan çok seviyeli rastgele etkili regresyon modeli, bu alanda kullanılan ileri bir deneme planı olduğundan bahsedilmişti. Bu modellerde karşılaşılan sorunları ortadan kaldırabilmek için, parametre tahminlerine ait bazı önsel tanımlamalar yapılmaktadır. Fakat bazen yapılan önsel tanımlamalar yanlış olabilir ve bu yanlış önsel tanımlamalar da bulunulmuş model parametrelerine karşı en iyi tasarımların dirençliliğinin (robustness) ne kadar olduğu bazı araştırmacılar tarafından hala merak konusu olmuştur (Moerbeek ve ark. 2006). Moerbeek ve ark. (2006) çok seviyeli modeller için dirençliliği en iyi

olan tasarımlar elde etmede farklı yaklaşımlar uygulandığını belirtmişlerdir. Bu yaklaşımlar aşağıda listelenmiştir.

- İlk yaklaşım; örneklem genişliğinin tekrardan tahmin edilmesidir.

Burada en iyi örneklem genişliği, kişinin bilgisine ya da genel bilgiye dayanarak elde edilmiş parametrelerin önsel tahminlerine bağlı olarak hesaplanır (bu önsel tahminleri elde ederken konu ile ilgili genel bilgilerden yararlanılır). Daha sonra küme içine yuvalanmış birey sayısı yeniden tanımlanarak bunlara ait veriler toplanır ve model parametreleri tahmin edilir. Sonra en iyi örneklem genişliği tekrar tahmin edilir ve verinin geri kalanı elde edilir. Tüm veriler kullanılarak son analiz yapılır. Bu yöntem güç ve çalışmanın maliyetini kontrol altında tutmak içinde etkili bir yöntemdir.

- Diğer bir yaklaşım; Bayesian en iyi tasarımıdır (Bayesian optimal design) (Spiegelhalter 2001, Turner ve ark. 2004).

Bu yöntem ile model parametrelerinin önsel dağılımları tanımlanarak model parametrelerinin belirsizliğinin ele alınması sağlanır. Burada model parametrelerinin kendilerine ait önsel dağılımlarından defalarca örnekleme yapılır ve model parametrelerinin istatistiksel testlerinin gücü hesaplanır. Elde edilen güç dağılımı ile model parametrelerinin belirsizliği açıklanır. Ayrıca WINBUGS paket programı kullanılarak Bayesian en iyi deneme planları test edilebilir (Spiegelhalter ve ark. 2003).

- Diğer bir yaklaşım maksimum en iyi tasarımıdır.

2.2.2.5. Açıklayıcı değişkenlerin seçimi

Çok seviyeli modelde yer alan açıklayıcı değişken sayısı makul sayıda olsa bile model yapısı karmaşık bir hal almaktadır. Bu yüzden araştırmacılar tam modeli tahmin etmezler çünkü karmaşık bir modelden çıkarım yapmak çok zordur. Genellikle araştırmacılar modeller oluşturulurken, önceki çalışmalarda önemli bulunmuş parametreler ile ilgilenirler ya da teorik problemin bakış açısına göre ilginç olan parametrelerle ilgilenirler (Hox 1995).

Eğer araştırmacılar güçlü bir “çok seviyeli model teorisine” sahip değillerse, modeldeki açıklayıcı değişkenlere karar verirken Hox’un (1995) çalışmasında bahsettiği açıklayıcı değişken seçimi yöntemini kullanabilirler. Burada işlemlere başlanırken en basit modelle işe başlamak yaygındır. Bu hiçbir açıklayıcı değişkenin bulunmadığı modele sadece sabit katsayının yer aldığı (intercept-only model) model denir ve ilk adımda bu model kullanılır. Daha sonra her bir adımda yeni bir parametre modele eklenerek her bir parametrenin anlamlı olup olmadığı araştırılır ve iki farklı seviyede ne kadar artık hatası oluştuğuna bakılır. Aşağıda seçme işlemi ile ilgili adımlar verilmiştir. Buradaki her bir adımda hangi regresyon katsayısının ve (ko)varyansın modelde kalacağına karar verileceği anlatılacaktır. Tabi ki bunlara karar verirken, modellere ait anlamlılık testleri, sapmadaki değişim ve varyans unsurlarındaki değişim göz önüne alınarak karar verilmektedir.

Adım 1: Hiçbir açıklayıcı değişkenin bulunmadığı sadece sabit katsayı içeren model (intercept-only model) (2.21)’de verilmiştir.

$$Y_{ij} = \gamma_{00} + u_{0j} + e_{ij} \quad (2.21)$$

Aynı gruptaki gözlemlerin, farklı gruptaki gözlemlerden daha çok birbirine benzemesinden dolayı bağımlılık durumu ortaya çıkmaktadır. Bu bağımlılık grup içi korelasyon (*ICC*) olarak adlandırılmakta ve (2.21) modeli yardımı ile grup içi korelasyon tahminini için kullanıldığından dolayı çok kullanışlı bir modeldir.

$$ICC = \rho = \frac{\sigma_{00}}{\sigma_{00} + \sigma^2}$$

Aynı zamanda bu model sapma değerinin tahmininde de kullanılmaktadır (McCullagh ve Nelger 1989).

Adım 2: En düşük seviyede bulunan açıklayıcı değişkenlerin eklendiği model analiz edilir.

$$Y_{ij} = \gamma_{00} + \gamma_{p0}X_{pij} + u_{0j} + e_{ij} \quad (2.22)$$

Bu aşamada her bir açıklayıcı değişkenin katkısı değerlendirilir. Eğer en çok olabilirlik tahmini (maximum likelihood-ML) tahmin metodunu kullanılıyorsa, (2.21) ile (2.22) modellerinin hata terimlerinin varyansları arasındaki fark hesaplanarak son modeldeki iyileşme/gelişme tespit edilir. Bu fark, serbestlik derecesi iki modelin parametre sayılarının farkına eşit olup yaklaşık olarak ki-kare dağılışı gösterir. Aslında bu aşamada ki-kare dağılımının serbestlik derecesi adım 2’de modele eklenen açıklayıcı değişkenlerin sayısı kadardır.

Adım 3: Burada herhangi bir açıklayıcı değişkenin eğiminin anlamlı varyans unsuruna sahip olup olmadığı değerlendirilir.

$$Y_{ij} = \gamma_{00} + \gamma_{p0}X_{pij} + u_{pj}X_{pij} + u_{0j} + e_{ij} \quad (2.23)$$

Rasgele etkili eğim varyasyonunu test etmenin en iyi yolu tek tek bunları ele almaktır. Bir önceki adımda modele alınmamış bazı değişkenler bu aşamada tekrardan analiz edilebilir. Adım 2’de regresyon eğim katsayısı değeri istatistiksel olarak anlamlı çıkmamış olsa da ona ait varyans unsuru anlamlı olabilir. Gruplar arası hangi varyansın anlamlı olduğuna karar verdikten sonra adım 2’deki modelin mi yoksa adım3’teki modelin mi daha iyi olduğunu karar vermek için sapmaları temel alan ki-kare testi uygulanır.

Adım 4: Yüksek seviyedeki açıklayıcı değişkenler modele eklenir.

$$Y_{ij} = \gamma_{00} + \gamma_{p0}X_{pij} + \gamma_{0q}Z_{qj} + u_{pj}X_{pij} + u_{0j} + e_{ij} \quad (2.24)$$

Bu model, en üst seviyeye ait açıklayıcı değişkenlerin bağımlı değişkenlerdeki gruplar arasındaki varyasyonu açıklayıp açıklayamayacağının araştırılmasına izin verir. Burada da eğer ML kullanılıyorsa, modeldeki iyileşmenin testi için ki-kare testi kullanılır.

Adım 5: Adım 3’te elde edilmiş eğim katsayısının varyansı (yani $\text{Var}(u_{0j})$) anlamlı olan birey seviyesindeki açıklayıcı değişkenler ile grup seviyesindeki açıklayıcı değişkenlere ait karşılıklı etkileşim modele etkilenir.

$$Y_{ij} = \gamma_{00} + \gamma_{p0}X_{pij} + \gamma_{0q}Z_{qj} + \gamma_{pq}Z_{qj}X_{pij} + u_{pj}X_{pij} + u_{0j} + e_{ij} \quad (2.25)$$

Bu model için ML kullanıyorsak, modeldeki iyileşmenin testi için ki-kare testi kullanılır.

Özel olarak; eğer en son kullanılan açıklayıcı değişkenler model 2’de modele alınmışsa, en düşük seviyeye ait σ^2 ’nin düşmesi beklenir. Bu yüzden grup seviyesindeki açıklayıcı değişken hem grup varyansının bir kısmını hem de birey seviyesine ait varyansın bir kısmını açıklamaktadır. Ayrıca Adım 4’te eklenen yüksek seviyeye ait açıklayıcı değişkenler sadece grup seviyesine ilişkin varyansı açıklar. Eğer her seviyede ne kadar varyansın açıklandığı incelenmek isteniyorsa çoklu korelasyon katsayısı hesaplanabilir (Raudenbush ve Bryk 1986). Fakat bu çoklu korelasyon katsayısı yaklaşık bir değer olarak elde edilebilir ve modele açıklayıcı değişken eklendikçe azalması mümkündür. Ayrıca “en iyi” modele ulaşmak için açıklayıcı değişken seçme yöntemi kullanılıyor ise buradaki bazı kararların sadece şansa bağlı olması mümkündür.

2.2.2.6. Merkezlenme

Model parametre tahminlerinin nasıl yapıldığına değinilmeden önce, çok seviyeli modellerde genellikle kullanılan (fakat kullanımı zorunlu olmayan) merkezlenmeden (centering) bahsedilmesi gerekir. Çünkü çok seviyeli modellerde açıklayıcı değişkenler yer alırken merkezlenme işlemini uyguladıkları için bu konu oldukça önemlidir.

Regresyon analizi ve çok deęişkenli istatistiksel analiz metotlarında genellikle 0 etrafında merkezlenir. Yani deęişkenler sıfır merkezli olarak ele alınır çünkü sıfır merkezli olmak parametrelerden sonuç çıkarımını kolaylaştırmakta (özellikle de birbiriyle etkileşim içinde olan parametrelerin) ve bazı hesaplamalarda kolaylıklar sağlamaktadır (de Leeuw ve Kreft 2006).

Çok seviyeli model analizinde deęişkenlerin sıfır merkezli olması uygulanmış fakat doğrusal regresyon modelinde olduğu kadar etkili sonuçlar vermemiştir. Bu yüzden çok seviyeli modellerin analiz yöntemlerinde iki tane merkezlenme (odaklanma) yöntemi vardır.

- Genel ortalama etrafında merkezlenme (grand mean centering): Tüm gözlemlere ait ortalama elde edilir ve bunu etrafında odaklanır. Genel ortalama etrafında deęişkenler merkezlenirse; deęişkenler $X_{ij} - \bar{X}$ olarak modelde yer alır.
- Grup ortalaması etrafında merkezlenme (group mean centering/withing group centering): Her grubun kendisine ait ortalaması elde edilerek bunun etrafında odaklanır. Genel ortalama etrafında deęişkenler merkezlenirse; deęişkenler $X_{ij} - \bar{X}_j$ olarak modelde yer alır (de Leeuw ve Kreft 2006 ve Peugh 2010).

Hem sürekli hem de kategorik veriye bu iki yöntem uygulanabilir (Enders ve Tofighi 2007). Seviye-2’de açıklayıcı deęişken/deęişkenlerin bulunduğu durumda, seviye-1 deęişkenleri için genel ortalama etrafında odaklanma en uygun yöntemdir. Diğer taraftan seviye-2’deki deęişkenler için sadece genel ortalama etrafında odaklanma işlemi uygulanabilir çünkü seviye-2 deęişkenleri, seviye-2’de yer alan her ünite için aynıdır. Yani okul- öğrenci örneğini ele alacak olursak; seviye-2’deki açıklayıcı deęişkenler her okulda aynıdır (Peugh 2010).

Çok seviyeli modelin deęişkenlere merkezleme işlemi uygulandıktan sonra nasıl olduğunu gösterebilmek için seviye-2’deki açıklayıcı deęişkenin Z’nin bulunduğu ve

seviye-1'deki açıklayıcı değişken X'in bulunduğu modelde merkezlenme işlemi aşağıdaki şekilde gerçekleştirilir.

Eğer sadece seviye-1de X değişkeni yer alıyorsa ve X değişkenin grup ortalaması etrafında merkezlendiği durum inceleniyorsa (2.26)'daki model elde edilir.

$$Y_{ij} = \beta_{0j} + \beta_{1j}(X_{ij} - \bar{X}_j) + \varepsilon_{ij} \quad (2.26)$$

β_{0j} ve β_{1j} ait regresyon denklemi; $\beta_{0j} = \gamma_{00} + u_{0j}$ ve $\beta_{1j} = \gamma_{10} + u_{1j}$ şeklindedir. Bunları (2.26)'da yerine yazılırsa (2.27) elde edilir.

$$Y_{ij} = \gamma_{00} + \gamma_{10}(X_{ij} - \bar{X}_j) + (\gamma_{10} + u_{1j})(X_{ij} - \bar{X}_j) + u_{0j} + \varepsilon_{ij} \quad (2.27)$$

Seviye-2'de Z değişkenin bulunduğu durumda; Z değişkenin genel ortalama etrafında merkezlenecektir ve model (2.28) modeli ortaya çıkacaktır (Peugh 2009).

$$\begin{aligned} \beta_{0j} &= \gamma_{00} + \gamma_{01}(Z_j - \bar{Z}) + u_{0j} \\ \beta_{1j} &= \gamma_{10} + \gamma_{11}(Z_j - \bar{Z}) + u_{1j} \end{aligned} \quad (2.28)$$

Eşitlik (2.28), (2.26)'daki modelde yerine yazılırsa (2.29) modeline ulaşılır.

$$\begin{aligned} Y_{ij} &= \gamma_{00} + \gamma_{01}(Z_j - \bar{Z}) + \gamma_{10}(X_{ij} - \bar{X}_{.j}) + \gamma_{11}(X_{ij} - \bar{X}_{.j})(Z_j - \bar{Z}) \\ &+ (\gamma_{10} + u_{1j})(X_{ij} - \bar{X}_{.j}) + u_{0j} + \varepsilon_{ij} \end{aligned} \quad (2.29)$$

2.2.2.7. Model parametrelerinin tahmininde kullanılan yöntemler

Çok seviyeli modeli oluşturulduktan sonra bu modelde yer alan parametrelerin tahmini için farklı yöntemler kullanılmaktadır ve doğrusal olan modeller için yapılan tahminlerde genellikle en çok olabilirlik metoduna bağlı olarak elde edilir (Henderson 1953). Bu yöntemler en çok olabilirlik yöntemi(Maximum Likelihood -ML) ve kısıtlanmış en çok olabilirlik yöntemidir (Restricted Maximum Likelihood-REML).

2.2.2.7.1. ML (En çok olabilirlik) tahmini

ML Tarihçesi: Fisher (1921) tarafından en çok olabilirlik tahmin edicileri yöntemi daha öncelerde kullanılmaktaydı. 1970'lere kadar, varyans unsurlarının tahmini, Fisher tarafından bulunan ve ANOVA tekniğinde kullanılan yöntemler kullanılarak tahmin edilmekteydi (Henderson 1953). Bu yöntemle ilgili en kapsamlı makale Hartley ve J.N.K Rao (1967) tarafından yazılmıştır. Bu makalede tüm rasgele ve karma etkili modeller ve bunlar için dengeli veya dengeli olmayan veri setleri incelenerek ML yöntemi geliştirilmiştir.

Crump (1951) dengeli ve dengeli olmayan tek yönlü sınıflama ile ilgilenmiş ve eşitliklerde iterasyon işlemini kullanmıştır. Herbach (1959) belli modellere ait dengeli veri setlerinden ML tahmin edicilerini bulmuş ve bu tahmin edicilerin negatif olmama zorunluluğunu getirmiştir. Çünkü ML yöntemi bir maksimizasyon işlemi olduğu için parametrelerin ve varyans unsurlarının tahmin edicilerinin negatif olmaması gerekmektedir. Miller (1977) da hem dengeli hem de dengeli olmayan veri setleri üzerinde çalışmış, 2 etkenli rasgele etkili hem etkileşimli hem de etkileşimsiz modelleri kurmuş ve bunlara ait ML eşitliklerini yazmış fakat bu eşitlikleri çözmemiştir. Searle (1970) dengeli olmayan veri seti için ML tahmin edicilerine ait matrisleri çözümlenmiştir.

1970'lerden önce paket programların kullanılamamasından dolayı en çok olabilirlik tahmin yönteminin karışık modeller ve rasgele etkili modellere uygulanması mümkün değildi. Çünkü bu modellerde olabilirlik eşitlikleri ancak iteratif işlemlerle çözülebiliyordu. Fakat paket programların geliştirilmesiyle, en çok olabilirlik tahmin yönteminin kullanımı oldukça yaygın bir şekilde model parametre tahmininde kullanılmaya başlanmıştır Giesbrecht (1983) ve Thompson (1980) tarafından bu eşitlikleri hesaplayacak paket programlar geliştirilmiştir.

Çok seviyeli regresyon modelleri için de en çok olabilirlik tahminleri kullanılmaktadır. En çok olabilirlik tahmin yöntemi ile modele ait parametreler tahmin edilirken olabilirlik fonksiyonlarına maksimum değeri verecek popülasyon değerleri belirlenmeye çalışılır. Basit olarak ifade edecek olursak, en çok olabilirlik tahmin edicileri, olabilirlik fonksiyonunu maksimum yapacak parametre tahminleri belirlemektedir.

Ayrıca en çok olabilirlik tahmin yönteminin kullanılması sonucunda bazı sapma değerleri elde edilmektedir. Bu sapmalar modelin veri setine ne kadar uygunluk olduğunu gösterir. Genel olarak sapması az olan model, sapması çok olan modele göre veri setine daha çok uygunluk gösterir.

En çok olabilirlik tahmin yönteminin ANOVA yöntemine göre bazı avantajları vardır,

- Varyanslar için negatif tahminler elde edilmez. ANOVA kullanılarak varyans unsurlarını tahmin etmedeki temel prensip, kareler ortalamalarının beklenen değerleri eşitlendikten sonra elde edilen doğrusal eşitlik sistemlerini çözmekten ibarettir (Fırat 1997). ANOVA'nın kullanılması ile bazen negatif varyans unsuru tahminleri elde edilmesinden dolayı sorunlar yaşanmaktadır. En çok olabilirlik yöntemi ile elde edilen tahminlerde ise bu negatif tahmin elde etme sorunu ile karşılaşılmamaktadır (Fırat 2000).
- Ayrıca bu yöntemde rasgele ve karışık etkili modeldeki rasgele değişkene ait dağılımı belirtmeliyiz. Böyle bir zorunluluk ANOVA yönteminde bulunmamaktadır. Önceki çalışmalarda ML ile tahmin edilen varyans unsurlarının dağılımı normal olduğu varsayılmıştır. Ayrıca ML yöntemi ile çok yönlü sınıflamaya da bu uyarlanmıştır.

ML Teorisi: ML teorisinin temelinde yatan iki temel varsayım vardır. Bunlar;

- Genellikle dağılımı çoklu normal dağılıma uygunluk gösteren hata terimlerin birbirinden bağımsız aynı dağılımlı olması gerekir.
- Örneklem büyüklüğü yeterince büyük olmalıdır.

Fakat pratikte bu varsayımlar tam olarak sağlanamamaktadır ve bunun sonucu olarak da yanlış tahminler ve doğru olmayan hata terimleri elde edilmektedir (Sullivan ve ark. 1999). En çok olabilirlik tahmin yönteminin teorisi kısaca aşağıdaki şekildedir.

Olabilirlik fonksiyonu; parametre tahmini yapabilmek için gözlenmiş örnekleme ait verilerin olasılıklarını veren bir fonksiyon olup aşağıdaki şekilde ifade edilir.

$$\mathbf{V} = \text{var}(\mathbf{Y}) = \mathbf{XGX}^T + \mathbf{r}, \quad \mathbf{r} = \mathbf{Y} - \mathbf{A}(\mathbf{A}^T \mathbf{V}^{-1} \mathbf{A})^T \mathbf{A}^T \mathbf{V}^{-1} \mathbf{Y}, \quad \mathbf{G} = \begin{bmatrix} \sigma_{00} & \sigma_{10} \\ \sigma_{01} & \sigma_{11} \end{bmatrix}$$

ve $A_j = X_j Z_j$ olmak üzere logaritması alınmış olabilirlik fonksiyonu;

$$l_{ML} = -\frac{1}{2} \log |\mathbf{V}| - \frac{N}{2} \log \mathbf{r}^T \mathbf{V}^{-1} \mathbf{r} - \frac{N}{2} \left[1 + \log \frac{2\pi}{N} \right] \quad (2.30)$$

şeklinde tanımlanır. Bu fonksiyon yardımı ile parametre tahminleri yapılır. En çok olabilirlik tahmin ediciler yöntemi kullanılarak yapılan tahminler sonucunda bazı standart hatalar ortaya çıkmaktadır. Bu standart hatalar anlamlılık testlerinin formülünde kullanılmaktadır (Sullivan ve ark. 1999).

$$\text{Test istatistiği; } Z = \frac{\text{parametre}}{\text{parametreye ait standart hata}}$$

Bu test Wald Testi olarak bilinir (Wald, 1943) ve buradaki yokluk hipotez, parametrenin 0 olduğunu test eder ve buna ait p-değeri bulunur. Z test istatistiğinin $Z \sim N(0,1)$ olduğu varsayılmaktadır. Fakat Raudenbush ve Bryk'un (198) çalışmalarından sabit etkililer için Z test istatistiğinin serbestlik derecesi J-q-1 olan t dağılımına sahip olması gerektiğini tartışmışlardır. Ayrıca Z test istatistiğinin varyanslar için kullanılmasının uygun olmadığını ve bunun yerine ki-kare testinin artıklar için kullanılması gerektiğini de tartışmışlardır

Eğer iki model iç içe ise; iki model arasındaki sapmanın dağılımı ki-kare olup, serbestlik derecesi 2 modeldeki parametre sayısının farkı olarak elde edilir. Bu aynı zamanda modelin analizinde ki-kare testi kullanılarak elde edilen model uyumlarının basit modelden anlamlı bir şekilde daha iyi olup olmadığının test edilmesinde de kullanılır. Sapmalar için kullanılan ki-kare testi aynı zamanda rasgele etkilerin önemini test etmede de kullanılır. Eğer karşılaştırılan modeller iç içe değilse, modelin prensibi olabildiğince basit olmalıdır (Hox 1995).

Çok seviyeli modellemede sadece en çok olabilirlik fonksiyonu kullanılması bazı sorunlara neden olmaktadır. Draper (2006) bunları aşağıdaki gibi özetlemiştir.

- En çok olabilirlik tahminleri ve bunlara ait standart hatalar Gaussian çok seviyeli modellerdeki parametreler için kolayca bulunabilir. Fakat seviye-2'deki örneklem küçük ise bunlara ait güven aralıklarını bulmada sorunlar ortaya çıkmaktadır.
- Bağımlı değişken Gaussian modelden farklı olarak, rastgele etkili lojistik regresyondaki gibi ise durum daha karmaşık olmaktadır.

Yukarıdaki sorunlara bir çözüm bulabilmek ve modeldeki değişkenlerin önsel dağılımlarını tahmin aşamasında göz önüne alabilmek için çok seviyeli regresyon modellerine Bayesian yaklaşımı kullanılmıştır. Markov Chain Monte Carlo (MCMC)'nin değiştirilmesi ve bilgisayarların geliştirilmesi ile 1990'lı yıllardan sonra çok seviyeli modeller için Bayesian modeller artan bir oranla kullanılmaktadır.

- Ayrıca iterasyonların sonlanmaması sorunu ile de karşılaşılır.

En çok olabilirlik tahmin edicilerinin hesaplanmasında iteratif yöntem uygulanır. Paket programlar çeşitli parametreler için uygun bir başlangıç değeri atayarak işlemlere başlar. Bu başlangıç değeri belirlenirken regresyon analizinde başlangıç değeri belirlenmesi temel alınarak çok seviyeli regresyon analizi için başlangıç değeri belirlenmektedir. Yapılan iterasyonlarla bu başlangıç değeri iyileştirilmeye çalışılmaktadır. Bu iyileştirme adımları birçok kere tekrarlanır. Her bir iterasyon sonrası

elde edilen deęerler bir öncekiyle karşılaştırılır ve tahminler arasındaki fark yeterince küçük olduğunda bilgisayar tahmin işleminin bir sonucuna yakınsadığına karar verir ve iterasyon işlemini bitirir. Bazen programlar iteratif ML yönteminin durmasına garanti verilemediğı için bazı sorunlar yaşanır çünkü program iterasyon işlemlerine son verememektedir. Bu yüzden birçok programın yazılımında belli sayıda iterasyon sayısı sınırı konulmuştur. Bu belirlenen iterasyon sayısına ulaşıldıktan sonra hala belirli bir yakınsama elde edilmemişse iterasyon sayısı arttırılarak yeniden program yürütülebilir. Buna rağmen (yani çok sayıda iterasyon yapılmasına rağmen) bir yakınsama elde edilemiyorsa, araştırmacılar yakınsama elde edilemeyeceğini düşünürler. Burada da yakınsama elde edilemeyen bir modelle nasıl sonuç çıkarımında bulunacağı problemi karşımıza çıkmaktadır.

Yakınsama elde edilemeyen modeller için “iyi bir model değildir” yorumunda bulunulur. Yeterli özelliklere sahip olmayan bu modelden tahmin yapıp yapılamayacağı da hala bir tartışma konusudur (Hox 1995). Aslında buradaki problemin sebebi veri seti de olabilir. Örneğin kurulan model doğru olsa bile, küçük bir örneklem kullanılması sonucunda tahmin yöntemindeki bu sorunlar ortaya çıkıyor olabilir.

Fakat yapılan çalışmalar gösteriyor ki, eğer yeterli örneklem genişliğine/büyüklüğüne sahip bir model yakınsama göstermiyorsa, bunun sebebi genellikle yanlış model olarak görülmektedir. Çok seviyeli regresyon modelinde, yakınsama problemi genellikle değeri 0’a yakın olan varyans unsuru tahmin edilmeye çalışıldığında ortaya çıkmaktadır. Bu sorunu ortadan kaldırabilmek için genellikle bazı varyans unsurları modelden çıkarılır.

2.2.2.7.2. REML tahmini

1970'lerden sonra FML ve REML yöntemleri tercih edilen yöntemler olmaya başlamıştır. Özellikle paket programların kullanımı bu iki en çok olabilirlik tahmin yönteminin popüler olarak kullanılmasına sebep olmuştur. Çok seviyeli regresyon modellerinin analizi için geliştirilen paket programlarda son zamanlarda kullanılmaya başlanması sonucunda REML yönteminin kullanımı yaygınlaşmıştır (Hox 1995, Albright 2007).

REML yönteminde sadece varyans unsurları olabilirlik fonksiyonunda bulunur.

$$\mathbf{V} = \text{var}(\mathbf{Y}) = \mathbf{XGX}^T + \mathbf{r}, \quad \mathbf{r} = \mathbf{Y} - \mathbf{A}(\mathbf{A}^T \mathbf{V}^{-1} \mathbf{A})^{-1} \mathbf{A}^T \mathbf{V}^{-1} \mathbf{Y}, \quad \mathbf{G} = \begin{bmatrix} \sigma_{00} & \sigma_{10} \\ \sigma_{01} & \sigma_{11} \end{bmatrix},$$

$$\mathbf{A}_j = \mathbf{X}_j \mathbf{Z}_j \text{ ve } p = \text{rank}(\mathbf{A})$$

olmak üzere logaritması alınmış olabilirlik fonksiyonu;

$$l_{REML} = -\frac{1}{2} \log |\mathbf{V}| - \frac{1}{2} \log |\mathbf{A}^T \mathbf{V}^{-1} \mathbf{A}| - \frac{N-p}{2} \log \mathbf{r}^T \mathbf{V}^{-1} \mathbf{r} - \frac{N-p}{2} \left[1 + \log \frac{2\pi}{N-p} \right]$$

şeklinde tanımlıdır (Sullivan ve ark. 1999). Görüldüğü gibi fonksiyon karmaşık bir yapıdadır. Bu yüzden, bu tarz fonksiyonların maksimize edilmesi çok karmaşıktır (Searle 1970).

ML yönteminin REML yöntemine göre 2 avantajı vardır (Hox 1995). Bunlar;

- Hesaplamalar daha kolaydır.
- ML yöntemi regresyon katsayılarını da olabilirlik fonksiyonunda bulundurduğundan, bu yöntem sadece sabit etkili kısmı farklı olan iki model arasındaki farkın testi için kullanılır. Ayrıca modelde yer alan sabit etkili kısımda ki-kare testi kullanılarak test edilir. Diğer taraftan REML yöntemi ile

sadece varyans unsurlarına ait kısım yani rasgele etkili kısım ki-kare testi ile test edilir.

2.2.3. Tekrarlanan ölçümler verisine uygulanması

Bir önceki bölümde de bahsedildiği gibi tekrarlanan ölçümler verisine çok seviyeli modellemeyi kullanılarak analizler yapılması mümkündür. Burada tekrarlanan ölçümler, bireyler içine yuvalanmış olan veri setleri olarak ele alınırlar (Peugh 2010). Bilindiği gibi iç içe yapıdaki veri setlerinde gözlemlerin bağımsızlığı varsayımı bozulmaktadır. Bu bağımsızlığın sağlanmaması durumu, çok seviyeli modellerin kullanımını gerekli kılmaktadır. Geleneksel istatistiksel metotların kullanımı 1. tip hataların oluşmasına ve yanlış parametre tahmini elde edilmesine neden olmaktadır. Tekrarlanan ölçümlerde iç içe veri setlerine dönüştüğünde benzer sorunlar çıktığından dolayı çok seviyeli modelleme kullanılmalıdır.

Burada kullanılan çok seviyeli modeller de bir farklılık yoktur. Ama özel olarak zaman değişkenini modele katarak, zamanın açıklanan değişken üzerindeki etkisi araştırılmak istenebilir. Böyle bir durumda model;

$$Y_{ij} = \beta_{0j} + \beta_{1j}(Time_{ij}) + \varepsilon_{ij} \quad (2.31)$$

$$\beta_{0j} = \gamma_{00} + \gamma_{01}Z_j + u_{0j} \quad \text{ve} \quad \beta_{1j} = \gamma_{10} + \gamma_{11}Z_j + u_{1j}$$

şeklinde tanımlanmış olup çok seviyeli modellerle tekrarlanan ölçümler verisinin analizine zaman etkisi katılarak bunun etkisi araştırılabilir.

3. MATERYAL VE METOT

3.1. Materyal

Bu çalışmada kullanılan hayvan materyalini İngiliz Siyah Alaca süresindeki boğaların dişi yavrularına ait süt verimi kayıtları oluşturmaktadır. Bu hayvanların ilk kontrolleri Kasım 1998 ve Ekim 1999 tarihleri arasında yapılmış olup 10 adet verim kaydı bulunmaktadır. Verim kayıtları yaklaşık olarak birer aylık aralıklarla alınmıştır. Her bir kontrol günü süt verimi 24 saatlik süre içerisinde alınan bireysel günlük süt verimlerinin toplamıdır ve kilogram (kg) cinsinden kayıt edilmiştir.

Veri seti 40 denenmiş ve 649 denenmemiş boğanın 23,873 dişi yavrusuna ait süt verimi kayıtlarından oluşmaktadır. Ayrıca denenmemiş boğa başına düşen dişi yavru sayısı 1 ile 31 arasında değişmekte iken denenmiş boğa başına düşen yavru sayısı 187 ile 1,371 arasında değişim göstermektedir.

Çizelge 3.1. Veri dosyasının formatı

Boğa	Sym	fx	ps	yas	iskgs	hy	L1	L2	...	L10	LT
4	1	1	2	29.5	4	0.19	23.80	24.60	...	14.6	5793.5
4	1	1	2	29.6	4	0.19	17.80	16.2	...	11.8	5195.2
4	2	1	2	30.6	23	0.19	16.8	18.8	...	14.2	4776.9
4	2	1	2	31	20	0.19	22.8	21	...	16.8	5476.5
6	3	1	2	30.4	15	0.25	22.6	23.8	...	26.4	7578.2
469	3	2	1	28.9	15	0.50	17.6	21.2	...	12	4458.3

Veri girişi yapılırken Çizelge 3.1’deki format kullanılmıştır. Burada ‘obs’ gözlem sayısını, ‘boğa’ erkek sayısını, ‘sym’ sürü-yıl-mevsimi etkisini, ‘fx’ boğaların yaşlı ve genç sürülerden hangisine ait olduğunu, ‘ps’ ebeveyn kaydı tutulup

tutulmadığını, ‘yas’ yaşı, ‘iskgs’ doğumdan ilk kontrole kadar geçen süreyi, ‘hy’ Holystein yüzdesini, ‘L1, L2,...,L10’ Kasım 1998 ile ekim 1999 arasında tutulan 10 adet kontrol günü süt verim kaydını ve ‘LT’ toplam süt verimini temsil etmektedir. Ayrıca fx 1 değerini alması boğanın yaşlı sürüden geldiğini ve 2 değerini alıyorsa boğanın genç sürüden geldiğini göstermektedir.

Çizelge 3.2. Yaşlı ve genç sürüdeki boğalara ilişkin frekans tablosu

fx	Frekans	Yüzdelik	Birikimli Frekans	Birikimli Yüzdelik
1	18975	79.48	18975	79.48
2	4898	20.52	23873	100.00

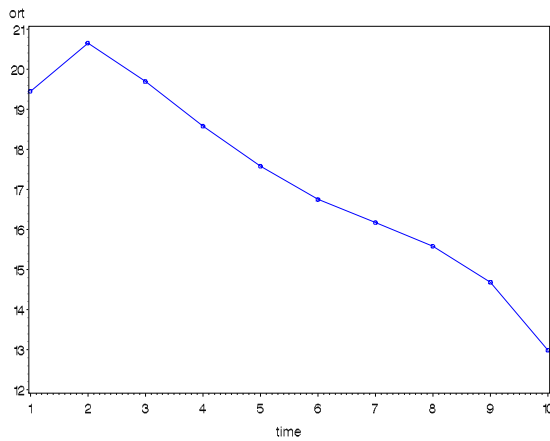
Yaşlı ve genç sürüdeki boğa sayısı Çizelge 3.2.’ de verilmiştir. Bu frekans tablosundan görüldüğü gibi verideki boğaların % 79.48’i yaşlı boğa sürüsüne aittir ve geri kalan %20.52’si ise genç boğa sürüsündendir. Ayrıca yaşlı sürüde bulunan boğaların yavrularının süt verimi ortalaması 5268.82 kg ve genç sürünün yavrularının süt verimi ortalaması 5258.32 kg’ dır.

Çizelge 3.3. Süt verimi kaydına ait tanımlayıcı istatistikler

Değişken	N	Ortalama	Standart Hata	Standart Sapma
Yaş	23873	30.2491681	0.0295031	4.5584986
iskgs	23873	18.3304151	0.0575655	8.8943844
Hy	23873	0.3670527	0.0025607	0.3956586
L1	23873	19.4498304	0.0260207	4.0204351
L2	23873	20.6566791	0.0252186	3.8964948
L3	23873	19.6954467	0.0253203	3.9122041
L4	23873	18.5825200	0.0247413	3.8227526
L5	23873	17.5815147	0.0242947	3.7537502
L6	23873	16.7524442	0.0240150	3.7105389
L7	23873	16.1753780	0.0241651	3.7337220
L8	23873	15.5848532	0.0242155	3.7415102
L9	23873	14.6880870	0.0242814	3.7516966
L10	23873	12.9894986	0.0251094	3.8796248
LT	23873	5266.66	6.0075603	928.2215367

Çizelge 3.3. Süt verimi dosyasında bulunan yaş (age), doğumdan ilk kontrole kadar geçen süre (dac), Holystein yüzdesi (holp), toplam verime (LT) ve 10 kontrol günü süt verimine ilişkin ortalama, standart hata ve standart sapma değerleri verilmiştir. Görüldüğü üzere, 23873 tane gözlem bulunmaktadır. Bu veri setinde yer alan ineklerin ortalama yaşlarının 30.25 ve ortalama Holystein yüzdeleri 0.367 olduğu görülmektedir. İlk süt sağımı yapılarına kadar ortalama 18.33 gün geçmiştir. Ayrıca inek başına düşen ortalama toplam süt verimi de 5266.66 bulunmuştur.

Bunun yanı sıra kontrol günü süt verimleri ortalamalarında bir azalış görülmektedir çünkü ilk kontrol günü (L1) süt verimi 19.45 kg iken son kontrol günü (L10) süt verimi ortalaması 12.99 kg' dır. Yani %33'lük $\left(\frac{19.45 - 12.99}{19.45} \cong 0.33\right)$ bir azalış görülmektedir. 10 farklı zaman noktasındaki süt verimlerinin 10 farklı zamanda nasıl bir değişim grafiksel olarak şekil 3.1' de incelenebilir. Bu grafikte açıkça görüldüğü üzere ilk 2 aydan sonra süt veriminde düzenli bir düşüş görülmektedir. İlk kontrol günündeki süt verimi ortalama 19,5 kg iken 10. kontrol günündeki süt verimi ortalama 13 kg' dır. Ayrıca gözlemler alınan inekler arasında süt verimleri değişimleri yaklaşık olarak aynıdır.



Şekil 3.1. Zaman içindeki süt verimindeki değişim

3.2. Metot

Bu çalışmada yer alan İngiliz Siyah Alaca ineklerine ilişkin ortalama toplam süt verimleri ve 10 adet kontrol günü süt verimleri çok seviyeli regresyon modeli aracılığı ile analiz edilmiştir.

a)Toplam süt verimi olan LT' ler için yapılan analizlerde kullanılan modeller;

1. İlk olarak regresyon modeli kullanılmıştır.

$$y_{ij} = \beta_0 + \beta_1(yaş)_i + \beta_2(hy)_i + \beta_3(iskgs)_i + \beta_4(sym)_i + \varepsilon_{ij} \quad (3.1)$$

Y_{ij} ; i . boğanın j . yavrusuna ilişkin süt verimini ve ε_{ij} ise inek seviyesine ait hata terimini göstermektedir. Ayrıca β regresyon katsayılarını temsil etmektedir.

2. Seviyelerin hiçbirindeki açıklayıcı değişkenin bulunmadığı ve sadece sabit terimden oluşan çok seviyeli model (intercept-only model)

$$Y_{ij} = \beta_{0j} + \varepsilon_{ij} \quad (3.2)$$

Y_{ij} : i_inci boğanın j_inci yavrusuna ilişkin süt verimi

ε_{ij} : inek seviyesine ait hata terimi

Modelde açıklayıcı değişken olmadığı için modelde β_{1j} katsayısı yoktur. β_{0j} katsayısına ait regresyon denklemi ise (3.3) 'te olduğu gibi tanımlanmıştır.

$$\beta_{0j} = \gamma_{00} + u_{0j} \quad (3.3)$$

(3.3) eşitliği (3.2)'de yerine yazılırsa (3.4) elde edilir.

$$Y_{ij} = \gamma_{00} + u_{0j} + \varepsilon_{ij} \quad (3.4)$$

şeklinde elde edilir.

3. İnek ve boğa seviyesindeki açıklayıcı değişkenlerin modele alınması ile elde edilen çok seviyeli regresyon modeli;

$$y_{ij} = \beta_{0j} + \beta_1(yas)_i + \beta_{2j}(hy)_i + \beta_{3j}(iskgs)_i + \varepsilon_{ij} \quad (3.5)$$

$$\beta_{0j} = \gamma_{00} + \gamma_{01}(sym)_j + u_{0j}$$

$$\beta_{2j} = \gamma_{20} + \gamma_{21}(sym)_j + u_{2j} \text{ olan regresyon denklemleri tanımlanmıştır bu denklemler}$$

$$\beta_{3j} = \gamma_{30} + \gamma_{31}(sym)_j + u_{3j}$$

(3.6)'da verilen çok seviyeli regresyon modeli yazılır.

$$\begin{aligned} Y_{ij} = & \gamma_{00} + \gamma_{01}(sym)_j + \gamma_{10}(yas)_{ij} + \gamma_{20}(hy)_{ij} + \gamma_{30}(iskgs)_{ij} \\ & + \gamma_{21}(sym_j * hy_{ij}) + \gamma_{31}(sym_j * iskgs_{ij}) + u_{0j} \\ & + u_{2j}(hy)_{ij} + u_{3j}(iskgs)_{ij} + \varepsilon_{ij} \end{aligned} \quad (3.6)$$

Burada ilk seviye yani inek seviyesi için ele alınan açıklayıcı değişkenler ineğin yaşı (yas), Holystein yüzdesi (hy) ve doğumdan ilk kontrole kadar geçen süre (iskgs) olarak alınmıştır. Ayrıca ineğin yaşı modelde kovaryete olarak incelenmiştir. Diğer taraftan ikinci seviye yani boğa seviyesi için sürü-yıl-mevsim (sym) etkisi modele alınarak (3.5)' teki model elde edilmiştir.

b) On adet kontrol günü süt verimleri için kullanılan modeller:

Burada veri seti tekrarlanan ölçümler veri seti şeklinde olduğu için tekrarlanan ölçümler veri setinde kullanımı uygun olan çok seviyeli model analiz aşamasında kullanılmıştır. Ayrıca veri tekrarlanan ölçümler veri setine dönüştüğü için buradaki seviye-1'i tekrarlanan ölçümler şeklindeki kontrol günü süt verimleri ve seviye-2'yi inekler oluşturmaktadır.

1. Sadece zaman değişkenin bulunduğu çok seviyeli model;

(Burada zaman değişkeni 10 kontrol gününü temsil etmektedir)

$$y_{ij} = \beta_{0j} + \beta_1(zaman)_i + \varepsilon_{ij} \quad (3.7)$$

İnek seviyesine ait açıklayıcı değişkenlerinde modele alınması ile elde edilen model;

$$\begin{aligned} \beta_{0j} &= \gamma_{00} + \gamma_{01}(hy)_j + \gamma_{02}(sym)_j + \gamma_{01}(yaş)_j + u_{0j} \\ \beta_{1j} &= \gamma_{10} + \gamma_{11}(hy)_j + \gamma_{12}(sym)_j + \gamma_{11}(yaş)_j + u_{1j} \end{aligned} \quad (3.8)$$

olmak üzere çok seviyeli model;

$$\begin{aligned} Y_{ij} &= \gamma_{00} + \gamma_{01}(hy)_j + \gamma_{01}(sym)_j + \gamma_{01}(yaş)_j \\ &+ \gamma_{10}(zaman)_{ij} + \gamma_{11}(hy_j * zaman_{ij}) + \gamma_{31}(yaş_j * zaman_{ij}) \\ &+ u_{0j} + u_{1j}(yaş)_{ij} + \varepsilon_{ij} \end{aligned} \quad (3.9)$$

Y_{ij} ; i . inek için j . zamandaki süt verimidir.

Bu modeller yardımı ile süt verimi üzerinde etkisi olan değişkenler tespit edilmeye çalışılmıştır.

4. BULGULAR VE TARTIŞMA

4.1. Toplam Süt Verimi İçin Elde Edilen Sonuçlar

Burada ilk olarak veri setindeki hiyerarşi yapı göz ardı edilerek veri setinin alt seviyeye indirgenerek geleneksel regresyon analizi ile veriler modellenip analiz uygulanmış ve elde edilen sonuçlar varyans analiz tablosu Çizelge 4.1.'de sunulmuştur.

Çizelge 4.1. Alt seviyeye indirgenmiş veri için varyans analiz tablosu

Değişimin Kaynağı	s.d.	KT	KO	F	p değeri
Model	4	960634431	240158608	292.34	<0.0001
Hata	23868	19607366687	821492		
Düzeltilmiş Toplam	23872	20568001119			

Çizelge 4.2. Alt seviyeye indirgenmiş veri için parametre tahminleri

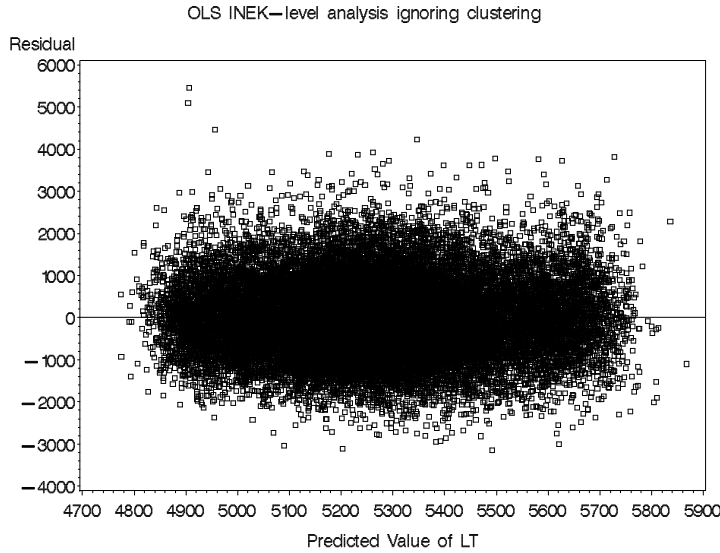
Değişken	sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	1	4420.92966	43.71776	101.12	<.0001
Yaş	1	24.37712	1.29335	18.85	<.0001
Hy	1	412.15876	14.89621	27.67	<.0001
İskgs	1	2.47485	0.65993	3.75	0.0002
Sym	1	-0.02461	0.00291	-8.47	<.0001

Yapılan regresyon analizi sonucunda kareler toplamı ve kareler ortalamaları Çizelge 4.1.'de verilmiştir ve bunlar kullanılarak F test istatistiği hesaplanmıştır. Burada

p değeri 0.05 yanılma payından küçük olduğu için model parametrelerin sıfıra eşitliği hipotezi red edilir. Dolayısıyla modelin anlamlı olduğuna karar verilir. Çizelge 4.2.de ise model parametrelerinin anlamlı olup olmadığını araştırmak için yapılar t testi sonuçları verilmiştir. Yapılan analizler sonucunda parametrelerin hepsi anlamlı bulunmuştur. Dolayısıyla açıklayıcı değişkenlerin hepsinin toplam süt verimi üzerinde bir etkisi vardır. Modelde tahmin edilmiş olan parametre değerleri yerine yazıldığında tahmin eşitliği (4.1)'deki gibi elde edilir.

$$\hat{y}_{ij} = 4420.93 + 24.38(yaş)_i + 412.16(hy)_i + 2.48(iskgs)_i - 0.025(sym)_i \quad (4.1)$$

Yani yaş değişkenindeki 1 birimlik değişim toplam süt verimini 24.38kg. etkilemektedir. Benzer şekilde Holstein yüzdesi değişkenindeki değişim toplam süt verimini 412.16kg, doğumdan ilk kontrole kadar geçen süre değişkenindeki değişim toplam süt verimini 2.48kg etkilemektedir. Diğer taraftan sürü-yıl-mevsim etkisindeki değişim ortalama süt veriminin 0.025kg düşmesine sebep olmaktadır. Ayrıca R^2 değeri 0.0467 bulunmuştur. Bunun anlamı modeldeki açıklayıcı değişkenler ile çıktı değişkeninin %4'ü açıklanabilmektedir.



Şekil 4.1. Alt seviyeye indirgenmiş verinin analizi sonucunda elde edilen süt verimi tahminleri ile hatalara ilişkin grafik

Şekil 4.1.’de kestirilen regresyon denklemi soncunda yapılan LT tahminlerinin (süt verimi) hatalara karşı grafiği görülmektedir. Bu grafikte hataların homojen bir yayılıma sahip olduğu görülmektedir.

Daha sonraki aşamada ise veri setindeki hiyerarşi yapısı göz ardı edilip veriler üst seviyeye kaydırılmış ve veri tek seviyeli veri setine dönüştürülmüştür ve bu tek seviyeli verinin analizden kullanılan regresyon sonuçları Çizelge 4.3. ve Çizelge 4.4.’te verilmiştir.

Çizelge 4.3. Seviye-2’ ye kaydırılmış veri için varyans analiz tablosu

Değişimin Kaynağı	s.d.	KT	KO	F	p değeri
Model	4	23015276	5753819	17.41	<0.0001
Hata	684	226012448	330428		
Düzeltilmiş Toplam	688	249027725			

Çizelge 4.4. Üst seviyeye kaydırılmış veri için parametre tahmini

Değişken	sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	1	4248.16693	246.98945	17.20	<.0001
Yaş	1	25.71154	7.17653	3.58	0.0004
Hy	1	585.59980	73.40198	7.98	<.0001
İskgs	1	2.36874	4.69803	0.50	0.6143
Sym	1	-0.02009	0.01251	-1.61	0.1087

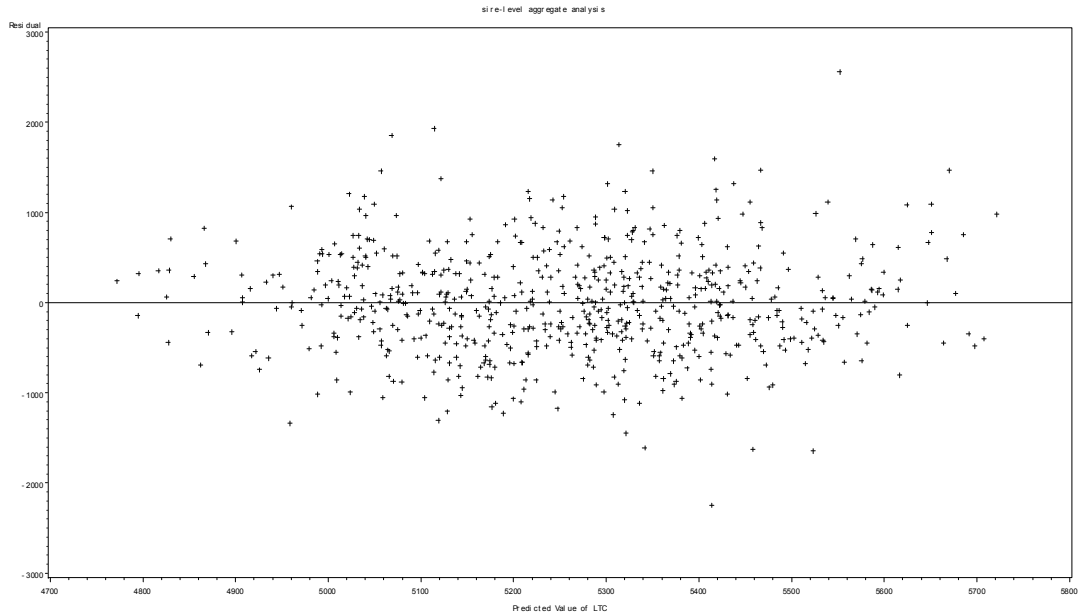
Çizelge 4.1.’ de olduğu gibi Çizelge 4.3.’ te de kareler toplamı ve kareler ortalamaları verilmiştir ve bunlar kullanılarak hesaplanan F test istatistiği verilmiştir.

Burada p değeri 0.05 yanılma payından küçük olduğu için model parametrelerinin sıfıra eşitliği hipotezi red edilir, yani model anlamlıdır. Bu durumda regresyon tahmin eşitliği;

$$\hat{y}_{ij} = 4248.17 + 25.71(yaş)_i + 585.6(hy)_i + 2.37(iskgs)_i - 0.02(sym)_i \quad (4.2)$$

şeklinde elde edilir. Çizelge 4.4. ise model parametrelerinin her biri için yapılmış olan t testi sonuçlarını içermektedir. Burada iskgs ve sym değişkenlerinin katsayılarının sıfıra eşit olduğu hipotezi red edilememiştir. Sonuç olarak bu iki değişkenin ortalama süt verimi üzerinde önemli bir etkisi yoktur. Diğer taraftan yaş ve Hoystein yüzdesi (hy) değişkenlerin katsayılarının sıfıra eşitliği hipotezi red edildiği için bu değişkenlerin ortalama süt verimi üzerinde anlamlı bir etkisi olduğu sonucuna varılır. Ayrıca yaş değişkenindeki değişim ortalama süt veriminde 25.711 kg.'lık bir değişikliği neden olurken, hy değişkenindeki değişim 585.6 kg.'lık bir değişime neden olmaktadır.

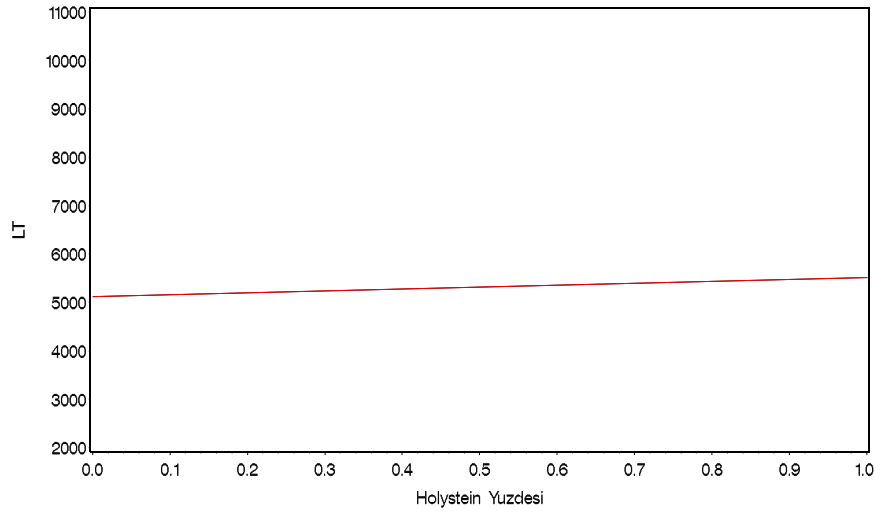
R^2 değeri 0.0924 olarak kestirilmiştir. Bunun anlamı modeldeki açıklayıcı değişkenler ile çıktı değişkenin %9'u açıklanabilmektedir.



Şekil 4.2. Üst seviyeye kaydırılmış verinin analizi sonucunda elde edilen süt verimi tahminleri ile hatalara ilişkin grafik

Üst seviyeye kaydırılmış veri seti için elde edilen LT tahminlerinin (süt verimi) hatalara karşı grafiği Şekil 4.2.'de verilmiştir. Burada veri seti üst seviyeye kaydırıldığı için her bir boğanın yavrularının süt verimi ortalamaları alınarak analizler yapılmıştır. Bu yüzden analizdeki örneklem hacmi aşağı seviyeye yuvarlanarak elde edilen veri setinin örnekleminden daha küçüktür. Dolayısıyla şekil 4.2'de daha az LT tahmini ve hata terimi bulunmaktadır.

Şekil 4.3.'de ise her iki regresyon modelinde anlamlı bulunmuş ve ortalama süt verimi üzerinde çok önemli bir etkiye sahip olan Holystein yüzdesi değişkeninin süt verimini nasıl etkilediği grafiksel olarak verilmiştir. Açık bir şekilde Holystein yüzdesi %100'e yaklaştıkça (ya da Holystein oranı 1'e yaklaştıkça) ortalama süt veriminde bir artış gösterdiği söylenebilir.



Şekil 4.3. Holystein yüzdesindeki değişimin toplam süt verimine etkisi

Şu ana kadar veri setindeki seviyeler göz ardı edilerek geleneksel regresyon yöntemi ile analizler yapılmıştır. Bu işlemlerin yapılmasındaki amaç hiyerarşi yapısının göz ardı edilmesi yapılan tahminlerde hatalı kestirimlerde bulunulup bulunulmadığı ya da yanlışlığa neden olup olmadığını araştırmaktır. Bundan sonraki analiz aşamasında ise

çalışmadaki asıl amaç olan çok seviyeli regresyon modeller analiz yöntemi olarak kullanılmıştır.

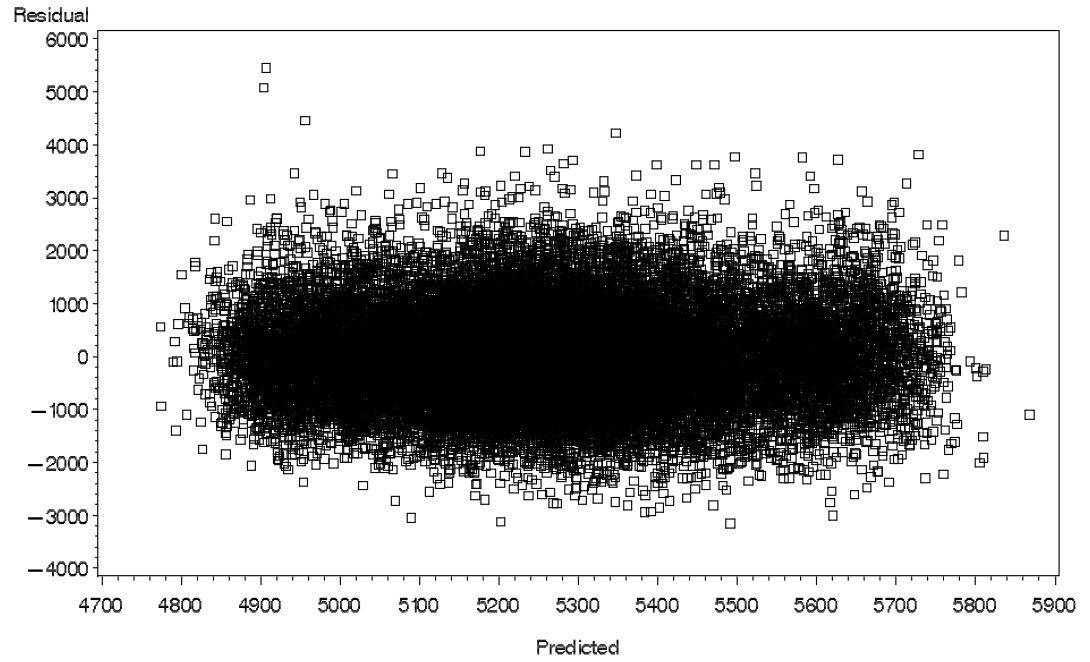
İlk olarak veri setindeki hiyerarşi yapısı göz ardı edilerek çok seviyeli regresyon analizi ile veriler modellenip analiz edilmiştir. Seviyeler göz ardı edildiği için bu analizde kullanılan modelin regresyon analizinde kullanılan modelden bir farkı yoktur. Seviyelerin göz ardı edilmesindeki amaç ise seviyeler göz önüne alınmadığı sürece çok seviyeli regresyon yönteminin kullanılmasının hatalı sonuçlar elde etmemizi sağlayacağını göstermektir. Çizelge 4.5.' de modelde yer alan açıklayıcı değişkenlerin katsayılarının tahmin değerleri, bu tahminlerde yapılan standart hatalar ve bunların her birine ait t test istatistiği yer almaktadır. Buradaki her bir katsayının sıfıra eşitliği hipotezi 0.05 yanılma payı ile red edildiği için tüm katsayılar anlamlı bulunmuştur ve model;

$$\hat{y}_{ij} = 4420.93 + 24.38(yaş)_i + 412.16(hy)_i + 2.48(iskgs)_i - 0.025(sym)_i \quad (4.3)$$

olarak kestirilmiştir. Ayrıca hataların varyansı 821320 olarak tahmin edilmiştir.

Çizelge 4.5. Parametre tahminleri (tüm açıklayıcı değişkenlerin yer aldığı fakat hiyerarşinin göz ardı edildiği model)

Değişken	Sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	23873	4420.93	43.7132	101.13	<.0001
Yaş	23873	24.3771	1.2932	18.85	<.0001
Hy	23873	412.16	14.8946	27.67	<.0001
İskgs	23873	2.4748	0.6599	3.75	0.0002
Sym	23873	-0.02461	0.002906	-8.47	<.0001



Şekil 4.4. Hiyerarşi yapısı göz ardı edilmesi ile yapılan analiz sonucunda elde edilen süt verimi tahminleri ile hatalara ilişkin grafik

Şekil 4.4.'de analiz sonucunda tahmin edilen LT'ler (süt verimi) ve bunların her birine ait hatalara ilişkin grafiği görülmektedir. Bu grafikte çok fazla sapan değerin olmadığı görülmektedir.

Seviye-2'nin göz önüne alındığı ve seviyelerin hiçbirinde açıklayıcı değişkenin bulunmadığı modele ML uygulanması sonucunda elde edilen tahminler Çizelge 4.6.'da verilmiştir. Bu modelde sadece rasgele etki olarak ele alınmış sabit katsayı yer almaktadır. Ayrıca bu model grup içi korelasyon katsayısının (*ICC*) hesaplanmasına da izin verir çünkü grup içi korelasyon katsayısının hesaplama formülünde yer alan seviye-1 hata varyansı (level-1 residual variance) olan σ^2 bu model sayesinde tahmin edilir ve buna ilişkin bulgular Çizelge 4.6.'da yer verilmiştir. Ayrıca Çizelge 4.7.'de yer alan tahminler yardımı ile *ICC* hesaplanabilir. Elde edilen tahmin değerlerine göre $\hat{\sigma}^2 = 770532$ ve $\hat{\sigma}_{00} = 186639$ bulunmuştur. Buna göre

$$ICC = \frac{186639}{186639 + 760532} = 0.19499 \quad (4.4)$$

olarak tahmin edilir. Bu katsayı açıklanamayan varyansının %20'si seviye-2'den dolayı ortaya çıktığının göstermektedir.

Çizelge 4.7'de $\hat{\sigma}^2 = 760532$ ve $\hat{\sigma}_{00} = 186639$ olduğundan ortalama süt verimine ilişkin %95 güven aralığı: $5256.69 \pm 1.96\sqrt{186639} : (4409.94 ; 6103.34)$ olarak bulunmuştur. Yani ortalama süt verimi %95 güvenle 4.5 ton ile 6 ton arasında bir değer alması beklenmektedir.

Çizelge 4.6. Parametre tahmini (sadece sabitin yer aldığı model)

Değişken	Sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	688	5256.69	20.9259	251.20	<.0001

Çizelge 4.7. Kovaryans parametrelerinin tahminleri (sadece sabitin yer aldığı model)

Kovaryans Parametresi	Tahmin Değeri	Standart Hata	Z	p değeri
Boğa	186639	16386	11.39	<.0001
Hata	760532	7065.21	107.64	<.0001

Ayrıca bu model yardımı ile deneme etkisi (design effect) hesaplanarak da veri setindeki hiyerarşi yapısının göz önüne alınarak çok seviyeli regresyon analizinin yapılmasının gerekli olup olmadığına karar verilebilir. Bunun için $DE = 1 + (n_c - 1)ICC$ eşitliğinden yararlanılır (Peugh 2010). DE “design effect”in kısaltılmış hali olup ICC ve boğa başına düşen ortalama yavru sayısı (n_c)’ye bağlı olarak hesaplanmaktadır.

$$n_c = 23873/689 = 34.65$$

$$DE = 1 + (34.65 - 1)(0.20) = 7.73 \quad (4.5)$$

olarak hesaplanır. Burada n_c değeri ve DE yeterince büyük olduğu için çok seviyeli modellenmenin kullanılması gerektiği sonucuna varılır.

Seviye-2'nin göz önüne alındığı ve açıklayıcı değişkenlerin modele alınması ile elde edilen tahminler Çizelge 4.8.'de yer almıştır. Bu sonuçlara göre (4.6)'daki çok seviyeli regresyon modeli elde edilir.

$$\hat{Y}_{ij} = 4134.27 - 0.018(sym)_j + 28.93(yas)_{ij} + 584.61(hy)_{ij} + 3.1671(iskgs)_{ij} - 0.0017(sym_j * hy_{ij}) - 0.00013(sym_j * iskgs_{ij}) \quad (4.6)$$

Çizelge 4.8. Parametre tahminleri (tüm açıklayıcı değişkenlerin yer aldığı model)

Değişken	Sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	688	4134.27	59.2534	69.77	<.0001
Yaş	23873	28.9300	1.3063	22.15	<.0001
Hy	23873	584.61	69.0932	8.46	0.0125
İskgs	23873	3.1671	1.3123	2.41	<.0001
Sym	23873	-0.01750	0.007009	-2.50	0.0158
hy*sym	23873	-0.00172	0.007632	-0.23	0.8216
iskgs*sym	23873	-0.00013	0.000315	-0.40	0.6881

Burada $hy*sym$ ve $iskgs*sym$ etkisinin 0.05 yanılma payı ile anlamlı değil iken modelde yer alan diğer değişkenlerin anlamlı bir etkiye sahip olduğu sonucu varılır. Diğer modellerde olduğu gibi burada da sym etkisi ortalama süt veriminin düşmesini sağlamak ve bu değer 0.018 kg. olarak tahmin edilmiştir. Yaş değişkenindeki değişim süt veriminin 28.93kg. değişmesine neden olurken Holystein yüzdesindeki değişim süt veriminde 584.61 kg.'lık bir değişikliğe sebep olmaktadır.

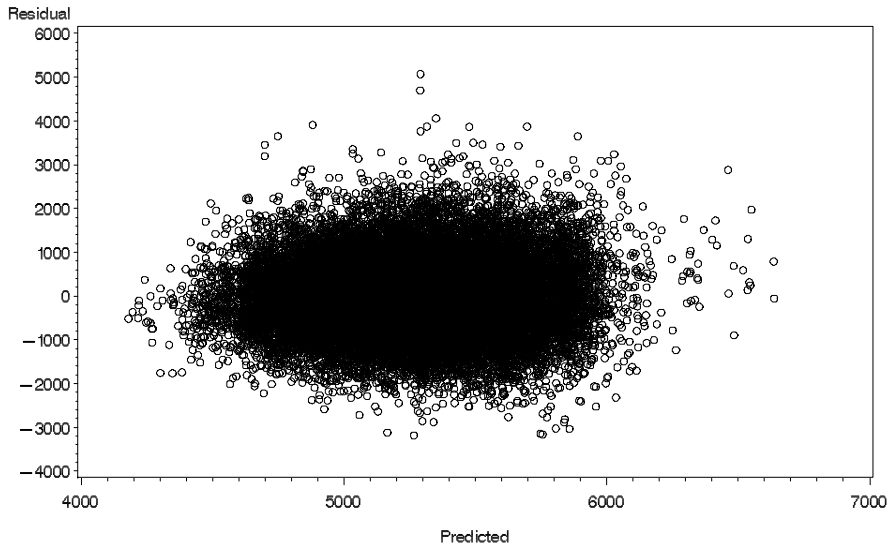
Çizelge 4.9. da ise hatalara ve seviye-2'ye ilişkin varyans tahminleri yer almaktadır. Bu değerler ile veri setindeki varyasyonun yüzde kaçının modelde yer alan açıklayıcı değişkenler tarafından açıklanamadığı araştırılabilir.

$$\hat{r} = \frac{\hat{\sigma}_a^2}{\hat{\sigma}_a^2 + \hat{\sigma}} = \frac{160534}{160534 + 742711} = 0.216 \quad (4.7)$$

Veri setindeki varyasyonun yaklaşık olarak %22' si modeldeki açıklayıcı değişkenler tarafından açıklanamamaktadır sonucuna varılır.

Çizelge 4.9. Kovaryans parametrelerinin tahmini (tüm açıklayıcı değişkenlerin yer aldığı model)

Kovaryans Parametresi	Tahmin Değeri	Standart Hata	Z	p değeri
UN(1,1)	160534	14804	10.84	<.0001
Hata	742711	6900.62	107.63	<.0001



Şekil 4.5. Hiyerarşi yapısı göz önüne alınarak yapılan analiz sonucunda elde edilen süt verimi tahminleri ile hatalara ilişkin grafik

Şekil 4.5.'de hiyerarşik yapının modellenmesi sonucunda yapılan analizden elde edilen LT tahminlerinin (süt verimi) hatalara karşı grafiği çizdirilmiştir. Bu grafikte hataların homojen bir yayılıma sahip olduğu ve çok fazla sapan değerin olmadığı görülmektedir. Ayrıca Şekil 4.4.'e kıyasla, bu grafikteki hata terimleri birbirine daha yakın değerler aldığı görsel olarak söylenebilir.

Yapılan analizler sonucunda hangi modelin veri setinin analizi için daha uygun olduğuna karar vermek için model uygunluk testleri yapılması gereklidir. Bu testlerin içinde en küçük olan veri setinin analizi için en uygun model olarak kabul edilmektedir. İncelenilen çalışmalarda çok seviyeli modeller için yapılan model uygunluk testlerinden olabilirlik oranı (Likelihood Ratio -LR) test sonuçlarına göre model uygunluğa karar verildiğinden bu çalışmada da bu test istatistiği değerine göre karar verilecektir.

Çizelge 4.10.'da seviyeler göz ardı edilmesi ile yapılan çok seviyeli regresyon modeli sonuçları ile seviyelerin göz önüne alınması sonucunda elde edilen modellere ilişkin model uygunluk testlerinin değerleri sunulmuştur.

Çizelge 4.10. Model uygunluk test değerleri

Seviyelerin göz ardı edildiği çok seviyeli regresyon analizi		Seviye-2'nin Göz önüne alındığı çok seviyeli regresyon analizi	
-2 Log Likelihood	392867.1	-2 Log Likelihood	391231.3
AIC	392879.1	AIC	391249.3
AICC	392879.1	AICC	91249.3
BIC	392927.6	BIC	391290.1

Buna göre hiyerarşi yapısının göz önüne alınması ile yapılan çok seviyeli regresyon modeli için elde edilen -2LL değeri daha küçük bulunmuştur. Test istatistiği değeri $\chi^2(6) = 392867.1 - 391231.3 = 1635.8$ bulunmuş ve bu değer istatistiksel olarak anlamlı olduğu için seviyelerin göz önüne alınarak açıklayıcı değişkenlerin

modellendiği çok seviyeli regresyon modeli veri seti analizine en uygun modeldir sonucuna varılır.

4.2. Kontrol Günü Süt Verimleri İçin Elde Edilen Sonuçlar

Bu bölümde veri setinde kontrol günü süt verimleri göz önüne alınmış ve 10 farklı kontrol günündeki süt verimleri incelenmiştir. İlk olarak hiçbir açıklayıcı değişkenin yer almadığı sadece sabit katsayının bulunduğu model analiz edilmiştir. ML yönteminin uygulanması sonucunda elde edilen tahminler Çizelge 4.11.'de verilmiştir. Ayrıca Çizelge 4.12.'de de varyans tahminine yer verilmiştir.

Çizelge 4.11. Parametre tahminleri (sadece sabitin yer aldığı model)

Değişken	sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	688	17.1857	0.07237	237.48	<.0001

Çizelge 4.12. Kovaryans parametrelerinin tahminleri (sadece sabitin yer aldığı model)

Kovaryans Parametresi	Tahmin Değeri	Standart Hata	Z	p değeri
Sabit	3.2645	0.1987	16.43	<.0001
Hata	18.6627	0.05410	344.97	<.0001

Burada $\hat{\sigma}^2 = 18.6627$ ve $\hat{\sigma}_{00} = 3.2645$ bulunmuştur. Ayrıca kontrol günü süt verimine ait %95 güven aralığı: $17.1857 \pm 1.96\sqrt{3.27} : (10.7765 ; 27.9622)$ olarak bulunmuştur. Yani kontrol günü süt verimlerinin %95 güvenle 10.78kg ile 27.96kg arasında olması beklenmektedir.

$$ICC = \frac{3.2645}{3.2645 + 18.6627} = 0.148879 \text{ olarak tahmin edilir. Yani açıklanamayan}$$

varyansının %15'si seviye-2'den dolayı ortaya çıkmaktadır.

Sadece zaman değişkeninin modele alınması ile elde edilen tahminler Çizelge 4.13. ve Çizelge 4.14.'te yer almıştır. Buna göre kestirilen çok seviyeli regresyon modeli;

$$\hat{y}_{ij} = 21.47 - 0.78(\text{zaman})_i \quad (4.8)$$

Burada sabit ve zaman değişkenleri 0.05 anlam düzeyinde önemli bulunmuştur. Zaman değişkeni süt veriminde düşmeye neden olmaktadır. Yani zamanla zamandaki 1 birimlik değişimin sonucunda süt veriminde 0.78kg.'lık bir azalma görülmektedir.

Çizelge 4.13. Parametre tahminleri (zaman değişkeninin yer aldığı model)

Değişken	Sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	688	21.4705	0.07429	288.99	<.0001
Zaman	23873	-0.7789	0.002632	-295.97	<.0001

Çizelge 4.14. Kovaryans parametrelerinin tahminleri (zaman değişkeninin yer aldığı model)

Kovaryans Parametresi	Tahmin Değeri	Standart Hata	Z	p değeri
U(1,1)	3.4022	0.2006	16.96	<.0001
Hata	13.6419	0.03954	344.97	<.0001

Ayrıca zaman değişkeninin veri setindeki varyasyonun yüzde kaçını açıklayabildiği Çizelge 4.14. yardımı ile hesaplanabilir.

$$ICC = \hat{r} = \frac{\hat{\sigma}_\alpha^2}{\hat{\sigma}_\alpha^2 + \hat{\sigma}^2} = \frac{3.4020}{3.4020 + 13.6419} = 0.1996 \quad (4.9)$$

Veri setindeki varyasyonun yaklaşık olarak %20' si modeldeki açıklayıcı değişkenler tarafından açıklanamamaktadır sonucuna varılır.

Zaman değişkeni ve etkisi araştırılmak istenen diğer değişkenlerin modele alınması ile yapılmış olan analiz sonucu Çizelge 4.15. ve Çizelge 4.16.'te yer almıştır. Buna göre kestirilen çok seviyeli regresyon modeli eşitlik 4.10'daki şekildedir.

$$\begin{aligned} \hat{Y}_{ij} = & 17.93 + 2.32(hy)_j - 0.00013(sym)_j + 0.099(yaş)_j \\ & - 0.77(zaman)_{ij} - 0.073(hy_j * zaman_{ij}) - 0.0006(yaş_j * zaman_{ij}) \\ & - 0.000011(sym_j * zaman_{ij}) \end{aligned} \quad (4.10)$$

Modelde yer alan zaman*yaş değişkeninin etkisinin 0.05 anlam düzeyinde süt verimi üzerinde önemli bir etkisi yoktur. Fakat diğer değişkenler anlamlı bulunmuştur ve yaş dışındaki diğer değişkenler süt veriminde az da olsa azalmaya sebep olmaktadır.

Çizelge 4.15. Parametre tahminleri (zaman ve diğer açıklayıcı değişkenlerin yer aldığı model)

Değişken	Sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	688	17.9272	0.1658	108.14	<.0001
zaman	23874	-0.7742	0.01884	-41.08	<.0001
Hy	23874	2.3197	0.23	10.09	<.0001
Sym	23874	-0.00013	0.00829	-15.47	<.0001
Yaş	23874	0.09875	0.003637	27.16	<.0001
zaman*hy	23874	-0.07248	0.006637	-10.92	<.0001
zaman*sym	23874	0.000011	1.295	8.62	<.0001
zaman*yaş	23874	-0.00060	0.000576	-1.04	0.2963

Çizelge 4.16. da ise hatalara ve seviye-2'ye ilişkin varyans tahminleri yer almaktadır. Bu değerler veri setindeki varyasyonun yüzde kaçının modelde yer alan açıklayıcı değişkenler tarafında açıklanamadığı araştırılabilir.

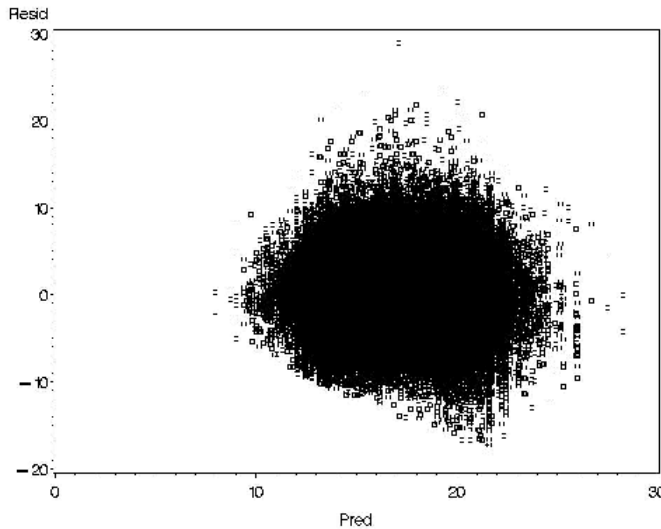
Çizelge 4.16. Kovaryans parametrelerinin tahminleri (zaman ve diğer açıklayıcı değişkenlerin yer aldığı model)

Kovaryans Parametresi	Tahmin Değeri	Standart Hata	Z	p değeri
U(1,1)	3.0765	0.1827	16.84	<.0001
Hata	13.4525	0.039	344.97	<.0001

Buna göre;

$$\hat{r} = \frac{\hat{\sigma}_\alpha^2}{\hat{\sigma}_\alpha^2 + \hat{\sigma}^2} = \frac{3.0765}{3.0765 + 13.4525} = 0.186 \quad (4.11)$$

Veri setindeki varyasyonun yaklaşık olarak %19' ü modeldeki açıklayıcı değişkenler tarafından açıklanamamaktadır sonucuna varılır.



Şekil 4.6. Kontrol günü süt verimi tahminleri ile hatalara ilişkin grafik

Şekil 4.6.'de sabit etkili zaman değişkeni ve diğer değişkenlerin modellenmesi sonucunda yapılan analizden elde edilen kontrol günü süt verimleri tahminlerinin (süt verimi) hatalar terimlerine karşı grafiği çizdirilmiştir. Görsel olarak yapılan tahminlerin yaklaşık 10 kg ile 25 kg arasında değişim gösterdiği söylenebilir. Ayrıca çok fazla sapan değerin olmadığı görülmektedir.

Zaman değişkeninin rasgele etkili olarak modellenmesi;

Modelde sadece zaman değişkeninin bulunduğu analiz sonucunda elde edilen tahminler Çizelge 4.17. ve Çizelge 4.18.'te yer almıştır. Buna göre kestirilen çok seviyeli regresyon modeli (4.12)'deki gibidir.

$$\hat{y}_{ij} = 21.443 - 0.77(\text{zaman})_i \quad (4.12)$$

Elde edilen sonuçlar eşitlik (4.8)'de yer alan modeldeki kestirim değerlerine yakın bulunmuştur. Burada da sabit ve zaman değişkenleri 0.05 anlam düzeyinde önemli bulunmuştur. Zamandaki 1 birimlik değişimin sonucunda süt veriminde 0.77kg.'lık bir azalmaya neden olmaktadır.

Çizelge 4.17. Parametre tahminleri (zamanın yer aldığı model)

Değişken	Sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	688	21.4430	0.08991	238.49	<.0001
Zaman	688	-0.7739	0.009084	-85.20	<.0001

Çizelge 4.18. Kovaryans parametrelerinin tahminleri (zamanın yer aldığı model)

Kovaryans Parametresi	Tahmin Değeri	Standart Hata	Z	p değeri
UN(1,1)	4.5824	0.3124	14.67	<.0001
UN(2,1)	-0.1974	0.02616	-7.55	<.0001
UN(2,2)	0.03272	0.003131	10.45	<.0001
Hata	13.5255	0.03927	344.45	<.0001

Modelde hem rasgele etkili zaman değişkeninin hem de diğer değişkenlerin bulunduğu analiz sonucunda elde edilen tahminler Çizelge 4.19. ve Çizelge 4.20.'te yer almıştır. Buna göre kestirilen çok seviyeli regresyon modeli;

$$\begin{aligned}\hat{Y}_{ij} = & 17.93 + 2.08(hy)_j - 0.00012(sym)_j + 0.099(yaş)_j \\ & - 0.77(zaman)_{ij} - 0.029(hy_j * zaman_{ij}) - 0.0006(yaş_j * zaman_{ij}) \\ & + 9.648(sym_j * zaman_{ij})\end{aligned}\quad (4.13)$$

Çizelge 4.19. Parametre tahmini (zaman ve diğer açıklayıcı değişkenlerin yer aldığı model)

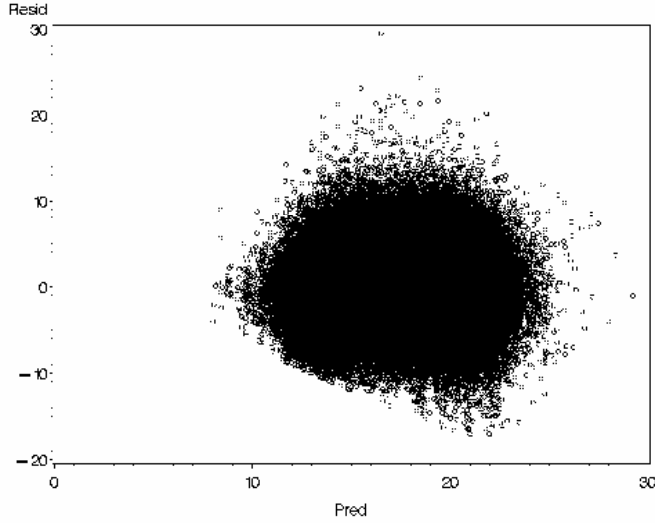
Değişken	Sd	Tahmin Değeri	Standart Hata	t değeri	p değeri
Sabit	688	17.9343	0.1878	95.48	<.0001
zaman	23874	-0.7755	0.02477	-31.31	<.0001
Hy	23874	2.0785	0.2772	7.50	<.0001
Sym	23874	-0.00012	8.88	-13.51	<.0001
Yaş	23874	0.09867	0.003791	25.02	<.0001
zaman*hy	23874	-0.02862	0.02888	-0.99	0.3217
zaman*sym	23874	9.648	1.42	6.80	<.0001
zaman*yaş	23874	-0.00059	0.000609	-0.96	0.3359

Modelde yer alan zaman*yaş ve zaman*hy değişkeninin etkisinin 0.05 anlam düzeyinde süt verimi üzerinde önemli bir etkisi yoktur. Fakat diğer değişkenler anlamlı bulunmuştur. Holystein yüzdesindeki değişim süt veriminde 2.08'lik bir artışa sebep olmakta, yaş değişkenindeki değişim az da olsa (0.1 kg) süt veriminde bir artışa sebep olmakta ve sürü-yıl-mevsim ile zaman etkileşimindeki değişim süt veriminde 9.65kg.'lık bir artışa sebep olmaktadır. Diğer taraftan sürü-yıl-mevsim süt veriminde 0.00012 kg.'lık azalışa ve zaman değişkeni süt veriminde 0.77 kg.'lık azalışa neden olmaktadır.

Çizelge 4.20. Kovaryans parametrelerinin tahmini (zaman ve diğer açıklayıcı değişkenlerinin yer aldığı model)

Kovaryans Parametresi	Tahmin Değeri	Standart Hata	Z	p değeri
σ_{00}	4.0381	0.2820	14.32	<.0001
σ_{10}, σ_{01}	-0.1743	0.02445	-7.13	<.0001
σ_{11}	0.03162	0.003056	10.35	<.0001
σ^2	13.3454	0.03874	344.45	<.0001

Sadece zaman değişkenin bulunduğu modelde tahmin edilen varyans unsurları Çizelge 4.18’da seviye-2 açıklayıcı değişkenlerinin modellenmediği sadece zaman değişkeninin modellenmesi ile tahmin edilen varyans unsurları verilmiş ve $\hat{\sigma}_{00} = 4.58$, $\hat{\sigma}_{10} = \hat{\sigma}_{01} = -0.2$, $\hat{\sigma}_{11} = 0.033$ ve $\hat{\sigma}^2 = 13.53$ bulunmuştur. Çizelge 4.20.’te ise seviye-2’ye ilişkin açıklayıcı değişkenlerinde modele alınması sonucunda tahmin edilen varyans unsurları tahminlerinin $\hat{\sigma}_{00} = 4.04$, $\hat{\sigma}_{10} = \hat{\sigma}_{01} = -0.17$, $\hat{\sigma}_{11} = 0.032$ ve $\hat{\sigma}^2 = 13.35$ olduğu belirtilmiştir. Buna göre seviye-2 sabit varyans (level-2 intercept variance) $\hat{\sigma}_{00}$ modele açıklayıcı değişken alınması ile yani $\left(\frac{[4.58 - 4.04]}{4.58} \cong 0.12 \right) \%12$ azalmıştır. Benzer şekilde seviye-2 eğim varyans (level-2 slope variance) $\hat{\sigma}_{11}$ modele açıklayıcı değişken alınması ile $\left(\frac{[0.033 - 0.032]}{0.033} \cong 0.03 \right) \%3$ azalmıştır. Benzer şekilde bu açıklayıcı değişkenlerin modele alınması seviye-1 hata varyansı olan $\hat{\sigma}$ ’da da azalmaya sebep olmuştur. Bu azalma $\left(\frac{[13.53 - 13.35]}{13.53} \cong 75.167 \right)$ olarak bulunur.



Şekil 4.7. Rasgele etkili model için kontrol günü süt verimi tahminleri ile hatalara ilişkin grafik

Şekil 4.7.'de rasgele etkili zaman değişkeni ve diğer değişkenlerin modellenmesi sonucunda yapılan analizden elde edilen kontrol günü süt verimleri tahminlerinin (süt verimi) hatalar terimlerine karşı grafiği çizdirilmiştir. Bu grafik görsel olarak Şekil 4.5' teki ile birbirine benzerlik göstermektedir. Burada görsel olarak yapılan tahminlerin yaklaşık 10 kg ile 22 kg arasında değişim gösterdiği söylenebilir. Ayrıca çok fazla sapan değer olmadığı ve hataların birbirine yakın değerler aldığı görülmektedir.

Bu bölümde de kullanılan modellerin hangisinin veri setinin analizi için daha uygun olduğuna karar vermek için -2 Likelihood Ratio (LR) test sonuçları kıyaslanır. Çizelge 4.21.'de kontrol günü süt verimleri veri setinin analizinde kullanılmış olan modellerin her biri için yapılan model uygunluk testleri sonuçlarına yer verilmiştir.

Burada -2LL ilişkin değerler baz alınarak model uygunluğa karar verilecektir. Buna göre rasgele etkili zaman etkisi ve diğer açıklayıcı değişkenlerin modellenmesi ile elde edilmiş çok seviyeli regresyon modeli için elde edilen -2LL değeri daha küçük bulunmuştur. Bu değer istatistiksel olarak anlamlı olduğu için bu model kontrol günü

süt verimi analizinde kullanılan veri setinin analizine en uygun modeldir sonucuna varılır. Buna göre rasgele etkili zaman değişkeni ve diğer açıklayıcı değişkenlerin modeli veri seti analizine en uygun modeldir sonucuna varılmıştır.

Çizelge 4.21. Model uygunluk test değerleri

Model	Model Uygunluk Testleri ve Sonuçları	
Sadece sabit etkinin modellenmesiyle yapılan çok seviyeli regresyon analizi sonuçları	-2 LL AIC AICC BIC	1378015 1378021 1378021 1378035
Sabit etkili zaman değişkeninde modele alınması ile yapılan çok seviyeli regresyon analizi sonuçları	-2 LL AIC AICC BIC	1303425 1303433 1303433 1303452
Sabit etkili zaman değişkeni ve diğer açıklayıcı değişkenlerin modele alınması ile yapılan çok seviyeli regresyon analizi sonuçları	-2 LL AIC AICC BIC	1300033 1300053 1300053 1300099
Rasgele etkili zaman değişkeninde modele alınması ile yapılan çok seviyeli regresyon analizi sonuçları	-2LL AIC AICC BIC	1302118 1302130 1302130 1302158
Rasgele etkili zaman değişkeni ve diğer açıklayıcı değişkenlerin modele alınması ile yapılan çok seviyeli regresyon analizi sonuçları	-2 LL AIC AICC BIC	1298856 1298880 1298880 1298935

5. SONUÇ

Birçok farklı alanda yaygın bir kullanıma sahip olan çok seviyeli modelleme yöntemi zirai alandaki özellikle de hayvansal veri setlerindeki kullanımı pek yaygın değildir. Veri setindeki hiyerarşik yapıyı göz önüne alınarak modelleme yapılmış ve modelde yer alan regresyon katsayılarının her birinin kendine ait farklı regresyon denklemi oluşturularak Siyah Alaca boğaları ve onların yavrularının ortalama süt verimi kayıtları ve 10 farklı zamanda ölçülmüş kontrol günü süt verimlerine ilişkin analizler yapılmıştır.

Veride setinin analizinde ilk olarak ortalama süt verimi göz önüne alınarak analizler yapılmıştır. Burada her ineğin boğaların içine yuvalandığı (her bir boğanın yavrusu kendi içine yuvalandığı) düşünülerek seviye-1'i ineklerin ve seviye-2'yi boğaların oluşturduğu hiyerarşi yapısı göz önüne alınmıştır. Daha sonra 10 kontrol günü süt verimi üzerinde analiz yapılmak istenmiş ve veri seti tekrarlanan ölçümler verisi haline dönüşmüştür. Tekrarlanan ölçümler verisinin çok seviyeli modelleme ile analizinde tekrarlanan ölçümlerin bireyler içine yuvalandığı ve seviye-1'i oluşturduğu seviye-2'yi de bireylerin oluşturduğu düşünülerek çok seviyeli modelleme yöntemi uygulanmaktaydı. Burada ise kontrol günü süt verimi seviye-1'i ve ineklerinde seviye-2'yi oluşturduğu düşünülerek tekrarlanan ölçümler veri setine çok seviyeli modelleme uygulanmıştır.

Ortalama süt verimi için yapılan çok seviyeli regresyon yönteminden önce veri setindeki hiyerarşi yapısı göz önüne alınmadan regresyon analizi uygulanmıştır. Hem üst seviyedeki birimlerin aşağı seviyeye kaydırılması ile yapılan regresyon analizi yapılmıştır. Alt seviyeye indirgenmiş veri setine yapılan regresyon analizi sonucunda hataların varyansı 821492, üst seviyeye kaydırılmış veri setine yapılan regresyon analizi sonucunda hataların varyansı 330428 bulunmuştur. Hata varyansının çok küçük bulunması sonucunda yanıltıcı sonuç çıkarımında bulunacağı belirtilmişti. Alt seviyeye indirgenmiş modelde yer alan tüm açıklayıcı değişkenlerin ortalama süt verimine

anlamli etkisi olduđu sonucuna varilirken üst seviyeye kaydırılmış modelde iskgs (ilk sađım yapıłana kadar geen süre) ve sym (sürü-yıl-mevsim) deđiřkenlerinin ortalama süt verimi üzerinde anlamli etkisi yoktur sonucuna varılmıştır. Görüldüğü üzere yapılan regresyon analizi sonuçları birbirinden farklıdır. Bu da gösteriyor ki veri setindeki hiyerarřı yapısı göz önüne alınarak analizler yapılmalıdır.

Benzer řekilde seviyeler göz önüne alınmadan yapılan ok seviyeli regresyon analiz sonucunda elde edilen sonuçlar regresyon analizi sonucu tahminlerine oldukça yakın bulunmuřtur. Modelde yer alan tüm açıklayıcı deđiřkenlerin ortalama süt verimine anlamli etkisi olduđu sonucuna varılır. Sadece sürü-yıl-mevsim deđiřkeni ortalama süt verimi üzerinde negatif etkiye sahip bulunmuřtur. Yani sadece bu deđiřken ortalama süt veriminde azalışa sebep olmaktadır sonucuna varılmıştır. Dahası hataların varyansı 821320 olarak tahmin edilmiştir.

Seviye-2'nin göz önüne alınması sonucunda yapılan ok seviyeli regresyon analizinde ise modelde yer alan $hy*sym$ ve $iskgs*sym$ etkileřim terimlerinin ortalama süt verimi zerinde önemli bir etkiye sahip deđildir (0.05 yanılma payı ile). Ayrıca sym, $hy*sym$ ve $iskgs*sym$ deđiřkenlerine ait katsayılar negatif bulunduđu için ortalama süt veriminde azalışa sebep olduđu söylenebilir. Hata terimlerinin varyansının tahmini 742711 (izelge 4.9'dan) olarak kestirilmiştir.

ok seviyeli regresyon analizi ile test edilmiş iki modelin birincisinde hata terimi varyansı 821320 ikincisinde ise hata terimi varyansı 742711 bulunmuřtur. Görüldüğü üzere seviyelerin göz önüne alınması ile hata terimlerinin varyansında %10'luk $\left(\frac{[821320 - 742711]}{821320} \cong 0.0957 \right)$ azalış görülmüřtür. Burada da veri setindeki hiyerarřı yapısı göz önüne alınarak veri setine ok seviyeli modelleme uygulanmalıdır sonucuna varılmıştır.

Deneme etkisi hesaplanarak da veri setindeki hiyerarşi yapısının göz önüne alınarak çok seviyeli regresyon analizinin yapılmasının gerekli olup olmadığına karar verilebilir. *ICC* ve boğa başına düşen ortalama yavru sayısı (n_c)'ye bağlı olarak hesaplandığı belirtilmişti. n_c 34.65 ve deneme etkisi 7.73 bulunmuştur. n_c ve DE değeri yeterince büyük olduğundan çok seviyeli modellemenin kullanılması gerektiği sonucuna varılır.

Ayrıca hiyerarşi yapısının göz önüne alındığı ve sadece sabit terimin bulunduğu modele ilişkin hata terimlerinin varyansı 760532 (Çizelge 4.7.'de) olarak bulunmuştur. Açıklayıcı değişkenlerin de yer aldığı model için hata terimlerin varyansının 742711 (Çizelge 4.9.'de) olarak tahmin edilmiştir. Yani seviye-1'e ait açıklayıcı değişkenler modele alındığında bu seviyeye ilişkin hata terimlerinin varyansında % 2 $\left(\frac{[760532 - 742711]}{760532} = 0.0234 \right)$ düşürmüştür. Dolayısıyla her seviyeye ait açıklayıcı değişkenlerin modele alınması ile yapılan analizler sonucunda elde edilen hata değerleri diğer modellere göre daha küçüktür denilebilir.

Hangi modelin veri setine daha uygun olduğuna karar vermek -2LL ilişkin değerler baz alınarak karar verilmiştir. Burada hiyerarşi yapısının göz önüne alınması ile yapılan çok seviyeli regresyon modeli için elde edilen -2LL değeri daha küçük bulunmuştur. Test istatistiği değeri $\chi^2(6) = 392867.1 - 391231.3 = 1635.8$ bulunduğu için bu değer istatistiksel olarak anlamlı olduğu için seviyelerin göz önüne alınarak açıklayıcı değişkenlerin modellendiği çok seviyeli regresyon modeli veri seti analizine en uygun modeldir sonucuna varılır.

Seçilen bu model için *ICC* katsayısı yaklaşık olarak 0.20 ($ICC = r = 0.20$) olduğu ve açıklanamayan varyansının %20'si seviye-2'den dolayı ortaya çıktığı söylenebilir. Bunun yanı sıra $r^2 = 0.4$ olduğu için ortalama süt verimindeki varyasyonun %4'nün modelde yer alan değişkenler tarafından açıklandığı sonucuna varılır.

10 farklı kontrol günündeki süt verimleri incelemek için veri setinde kontrol günü süt verimleri göz önüne alınmıştır. Sadece sabit etkinin modellenmesi ile yapılan analiz sonucunda kontrol günü süt verimi %95 güvenle 10.78kg ile 27.96kg arasında olması beklenmektedir sonucuna varılır. Ayrıca ICC katsayısı yaklaşık olarak 0.15 (yani $r=0.20$) olarak bulunmuştur. Bu da açıklanamayan varyansının %15'ini seviye-2'den dolayı ortaya çıktığı anlamına gelmektedir.

Modelde sadece açıklayıcı değişken olarak zamanın ele alındığı zaman yapılan analiz sonucunda zaman değişkene ilişkin katsayı negatif tahmin edildiği için zamanın süt veriminde düşmesine sebep olduğu söylenebilir. Yani ilk aylardaki süt verimleri son aylardan daha yüksektir.

Ayrıca seviye-1 ve seviye-2'ye ait açıklayıcı değişkenlerin modele alınması ile veri setindeki varyasyonun yaklaşık olarak %19'u modeldeki açıklayıcı değişkenler tarafından açıklanamamaktadır. Dahası yapılan analiz sonucunda zaman, sym (sürü-yıl-mevsim) ve zaman*hy (zaman ile Holstein yüzdesi etkileşimi) etkenlerinin süt verimi düşürmektedir sonucuna varılmıştır. Diğer taraftan zaman değişkeninin rastgele etkili olduğu düşünülerek tüm değişkenlerin de modele alınması ile analizler yapıldığı zaman sym mevsim etkisinin ortalama süt verimi üzerinde etkisi olduğu ve bir azalmaya sebep olduğu bulunmuştur.

Seviye-2'ye ait açıklayıcı değişkenler göz ardı edilerek yapılan analizlerde tahmin edilen varyans unsurları değerleri bu açıklayıcı değişkenlerin göz önüne alınması ile yapılan analizler sonucunda tahmin edilen varyans unsuru değerlerinden daha yüksektir. Yani bu değişkenlerin modellenmesi ile hata varyanslarında bir azalma görülmüştür. Buna göre seviye-2 sabit varyans (level-2 intercept variance) σ_{00} modele açıklayıcı değişken alınması ile %12 azalmıştır. Benzer şekilde seviye-2 eğim varyans (level-2 slope variance) σ_{11} modele açıklayıcı değişken alınması ile %3 azalmıştır. Benzer şekilde bu açıklayıcı değişkenlerin modele alınması seviye-1 hata varyansı olan σ 'da da azalmaya sebep olmuştur. Bunun anlamı modelde eklenen açıklayıcı değişkenler

yardımları ile kestirilen modelde elde edilen hatalar daha düşüktür ve bu da araştırmada istenilen bir durumdur.

Burada da model uygunluk testlerin -2LL sonuçlarına göre hangi modelin veri setinin analizinde kullanılmasının daha avantajlı olduğuna karar verilmiştir. Rasgele etkili zaman değişkeni ve diğer açıklayıcı değişkenlerin modeline ilişkin -2LL tahmini diğer modellere göre daha küçük olduğu için bu modelin veri seti analizine en uygun modeldir sonucuna varılmıştır.

Sonuç olarak bakıldığında hem ortalama süt verimine ilişkin yapılan analizlerde hem de kontrol günü süt verimleri için yapılan analizlerde çok seviyeli modellemeye ihtiyaç duyulmuştur. Bu da gösteriyor ki zirai alanlarda yapılan çalışmalarda hiyerarşi yapısı göz nüne alınarak çok seviyeli modelleme kullanılmalıdır. Bir çok farklı alanda kullanımı çok yaygın olan bu modelleme yönteminin zirai alandaki kullanımı da elverişli olup yaygınlaşmalıdır.

Yapılan bu çalışma hayvansal veri setlerinin analizinde çok seviyeli modellemeyi kullanması bakımından öncü bir çalışma niteliğindedir. Elde edilen sonuçlar gösteriyor ki çok seviyeli regresyon analizi hiyerarşik veri yapısı gösteren hayvansal veri setlerinin analizinde kullanımı çok elverişlidir ve yapılan tahminlerde seviyelerin göz önüne alınması ile elde edilen hataların azaldığı ve daha gerçekçi sonuçlar elde edildiğini göstermektedir. Bu nedenlerden dolayı çok seviyeli modelleme farklı hayvansal veri setlerinde de kullanılarak bu alana çok seviyeli modelleme tanıtılmalıdır.

6. KAYNAK

- Abright, J. J. 2007. Estimating model using SPSS, Stata, and SAS. online. <http://www.indiana.edu/~statmath/stat/all/hlm/hlm.pdf>.
- Anderson, D. A. and Aitkin, M. 1985. Variance component models with binary response: Interviewer variability. *Journal of the Royal Statistical Society, Series B*, 47:203–210.
- Aitkin, M and N. Longford 1986. Statistical modelling issues in school effectiveness studies. *Journal of the Royal Statistical Society*, 149:1-43.
- Bardenheier, B. H., Shefer, A., Lawrence, B., Winston, C. A. and Sionean, K. 2005. Public health application comparing multilevel analysis with logistic regression: Immunization coverage among long-term care facility residents. *Elsevier*, 15: 749-755.
- Bates, D. and Sarkar, D. 2006. The lme4 Package online. [Http://cran.r-project.org](http://cran.r-project.org).
- Bentler, P. M. 2006. EQS 6 Structural Equations Program Manual. Multivariate Software, Encino, CA.
- Blalock, H. M. 1990. Auxiliary measurement theories revisited. In: J.J. Hox & J. De Jong-Gierveld (eds.). *Operationalization and Research Strategy*. Amsterdam: Swets & Zeitlinger.
- den Bork, P., Brekelmas, M. and Wubbels, T. 2006. Multilevel issue in research using students' perceptions of learning environments: The case of the questionnaire on teacher interaction. *Springer*, 9: 199-203.
- Bosker, R. J., Snijders, T. A. B. and Guldemon, H. 1999. *PINT: Estimating Standard Errors of Regression Coefficients in Hierarchical Linear Models for Power Calculations. User's Manual Version 1.6*. University of Twente, Enschede, The Netherlands.
- Boyd, L. H. and Iversen, G. R. 1979. *Contextual Analysis: Concepts and Statistical Techniques*. Wadsworth, Belmont, CA.
- Burnside, G., Pine, C. M. and Williamson, P. R. 2007. The application of multilevel modelling to dental caries data. *InterScience*, 26: 4139-4149.
- Burstein, L., Linn, R. L. and Capell, F. J. 1978. Analyzing multilevel data in the presence of heterogeneous within-class regressions. *Journal of Educational Statistics*, 3:347–383.
- Burstein, L. 1980. The analysis of multilevel data in educational research and evaluation. *Review of Research in Education*, 8:158–233.

- Busing, F. M. T. A., Meijer, E. and Van der Leeden, R. 2005. MLA: Software for MultiLevel Analysis of Data with Two Levels. User's Guide for Version 4.1. Leiden University, Department of Psychology, Leiden.
- Cochran, W. G. 1983. Planning and Analysis of Observational Studies. Wiley, New York.
- Coleman, J. S., Campbell, E. Q., Hobson, C. J., McPartland, J., Mood, A. M., Weinfeld, F. D. and York, R. L. 1966. Equality of Educational Opportunity. U.S. Government Printing Office, Washington, DC.
- Cohen, M. P. 1998. Determining sample sizes for surveys with data analyzed by hierarchical linear models. *Journal of Official Statistics*, 14: 267–275.
- Conaway, M. R. 1989. Analysis of repeated categorical measurements with conditional likelihood methods. *Journal of the American Statistical Association*, 84: 53–61.
- Çömlekçi, N. 1988. Deney Tasarımı ve Çözümlemesi. Eğitim, Sağlık ve Bilimsel Araştırma Çalışmaları Vakfı Yayınları, Eskişehir, 312ss.
- Crump, S. L. 1951. The present status of variance components analysis. *Biometrics* 7, 1-16.
- Daniels, M. J. and Gatsonis, C. 1997. Hierarchical polytomous regression models with applications to health services research. *Statistics in Medicine*, 16: 2311–2325.
- D'Agostino, R. B., Lee, M. L., Belanger, A. J., Cupples, L. A., Anderson, K. and Kannel, W. B. 1990. Relation of pooled logistic regression to time dependent Cox regression analysis: The Framingham Heart Study. *Statistics in Medicine*, 9: 1501–1515.
- Donner, A. 1998. Some aspects of the design and analysis of cluster randomization trials. *Applied Statistics*, 47: 95–113.
- Draper, D. 1995. Inference and hierarchical modeling in the social sciences. *Journal of Educational and Behavioral Statistics*, 20: 115–147, 233–239.
- Draper, D. 2006. Bayesian multilevel analysis and MCMC. In: J. de Leeuw and E. Meijer (Editors), Handbook of Multilevel Analysis. *Springer*, pp. 76-139 Los Angeles, CA.
- Drukker, D. M. 2006. Maximum simulated likelihood: Introduction to a special issue. *The Stata Journal*, 6: 153–155.
- Eeden, P. V. D. 1993. Multilevel theory and the underspecification of multilevel models. *Quality & Quantity*, 26, 307-322.

- Ender, C. K. and Tofighi D. 2007. Centering predictor variable in cross-sectional multilevel model: A new look at an old issue. *Psychological Methods*, 12: 121-138.
- Erbaşı, S. O. ve Olmuş, H. 2006. Deney Düzenleri ve İstatistik Analizler, 1. Baskı. Gazi Kitabevi, Ankara, 412 sayfa.
- Feng, Z. and Grizzle, J. E. 1992. Correlated binomial variates: Properties of estimator of intraclass correlation and its effect on sample size calculation. *Statistics in Medicine*, 11: 1607–1614.
- Fırat, M. Z. 1997. Hayvan ıslahında negatif varyans unsuru tahmini ve tahmin yöntemlerinin incelenmesi. *Ç. Ü. Ziraat Fakültesi Dergisi*, 12: 169-176.
- Fırat, M. Z. 2000. Dengeli iki seviyeli şansa bağlı iç içe düzenlenmiş denemelerde varyans bileşenlerinin tahmini için varyans analizi, maksimum olabilirlik ve kısıtlanmış maksimum olabilirlik etotlarının karşılıklı olarak incelenmesi. *Anadolu Üni. Bilim ve Teknoloji Dergisi*, 1: 105-113.
- Fisher, R. A. 1921. On the mathematical foundations of theoretical statistics. *Philosophical Transactions of the Royal Society of London, Series A*, 222: 309–368.
- Fisher, R. A. 1923. Studies in crop variation: II. The manurial response of different potato varieties. *Journal of Agricultural Science*, 13: 311–320.
- Galtung, J. 1969. *Theory and Methods of Social Research*. New York: Columbia University Press.
- Giesbrecht, F. G. 1983. An efficient procedure for computing MINQUE of variance components and generalized least squares estimates of fixed effects. *Commun. Stat. A : Theory & Methods*, 12: 2169-2177.
- Goldstein, H. 1986. Efficient statistical modelling of longitudinal data. *Annals of Human Biology*, 13: 129-141.
- Goldstein, H. and Spiegelhalter, D. J. 1996. League tables and their limitations: Statistical issues in comparisons of institutional performance. *Journal of the Royal Statistical Society, Series A*, 159: 385–444.
- Goldstein, H. 2003. *Multilevel Statistical Models*, 3rd edition. Edward Arnold, London.
- Hartley, H. O. and Rao J. N. K. 1967. Maximum Likelihood estimation for mixed model analysis of variance model. *Biometrika*, 54: 93-98.
- Hedeker, D. and Gibbons, R. D. 1996. MIXOR: A computer program for mixed-effects ordinal regression analysis. *Computer Methods and Programs in Biomedicine*, 49: 157–176.

- Hedeker, D. and Gibbons, R. D. 1997. MIXREG: A computer program for mixed effects regression analysis with autocorrelated errors. *Computer Methods and Programs in Biomedicine*, 49: 229–252.
- Henderson, C. R. 1953. Estimation of variance and covariance components. *Biometrics*, 9: 226–252.
- Herbach, L. H. 1959. Properties of model I1 type analysis of variance tests, A: optimum nature of the F-test for Model I1 in the balanced case. *Ann. Math. Stat.*, 30: 939–959.
- Hox, J. J. and Kreft, Ita G.G. 1994. Multilevel analysis methods. *Sociological Methods and Research*, 22: 283–299.
- Hox, J. J. 1995. Applied Multilevel Analysis. TT-Publikaties, Amsterdam, 126 ss.
- Hsieh, F. Y. 1988. Sample size formulae for intervention studies with the cluster as unit of randomization. *Statistics in Medicine*, 8: 1195–1201.
- Jencks, C., Smith, M., Acland, H., Bane, M. J., Cohen, D., Gintis, H., Heyns, B. and Michelson, S. 1972. Inequality: A Reassessment of the Effect of Family and Schooling in America. Basic Books, New York.
- Kreft, Ita G.G. and de Leeuw, J. 1993. The gender gap in earnings: A two-way nested multiple regression analysis with random effects. *Sociological Methods & Research*, 22: 319–341.
- Langbein, L. I. 1977. Schools or students: Aggregation problems in the study of student achievement. *Evaluation Studies Review Annual*, 2: 270–298.
- de Leeuw, J. and Kreft, Ita G.G. 1986. Random coefficient models. *Journal of Educational Statistics*, 11(1): 55–85.
- de Leeuw, J. and Kreft, I. G. G. 2001. Software for multilevel analysis. In A. H. Leyland and H. Goldstein, editors, *Multilevel Modelling of Health Statistics*, pages 187–204. Wiley, Chichester, 2001.
- de Leeuw J. and Meijer E. 2006. Handbook of Multilevel Analysis. *Springer*, Los Angeles, CA, 504 pp.
- Leeden R. V. D. 1998 Multilevel analysis of repeated measures data. *Quality & Quantity*, 32: 15–29.
- Lindley, D. V. and Smith, A. F. M. 1972. Bayes estimates for the linear model. *Journal of the Royal Statistical Society, Series B*, 34: 1–41.
- Longford, N.T. 1989. To center or not to center in multilevel analysis. *Multilevel Modelling Newsletter*, 1(3): 7–11.

- Longford, N. T. 1990. *VARCL*. Software for Variance Component Analysis of Data with Nested Random Effects (Maximum Likelihood). Educational Testing Service, Princeton, NJ.
- McCullagh, P. and Nelder, J. A. 1989. *Generalized Linear Models*. London Chapman & Hall.
- Miller, J.J. 1977. Asymptotic properties of maximum likelihood estimates in the mixed model of the analysis of variance. *Ann. Stat.* 5: 746-762.
- Moerbeek, M., Van Breukelen, G. J. P. and Berger, M. P. F. 2001a. Optimal experimental designs for multilevel logistic models. *The Statistician*, 50: 1–14.
- Moerbeek, M., Van Breukelen, G. J. P. and Berger, M. P. F. 2001b. Optimal experimental designs for multilevel models with covariates. *Communications in Statistics, Theory and Methods*, 30: 2683–2697.
- Moerbeek, M., Van Breukelen G. J. P. and Berger, M. P. F. 2006. Optimal designs for multilevel studies. In: J. de Leeuw and E. Meijer (Editors), *Handbook of Multilevel Analysis*. Springer, pp. 177-205 Los Angeles, CA.
- Muthen, L. K. and Muthen, B. O. 1998. *Mplus User's Guide*, 4th edition. Muthen and Muthen, Los Angeles.
- Omar, R. Z., Wright, M. E., Turner, R. M. and Thompson, S. G. 1999. *Statistics in Medicine*, 18: 1587-1603.
- Overmars, K. P. and Verburg, P. H. 2006. Multilevel modelling of land use from field to village level in the Philippines. *Agricultural Systems*, 89: 435-456.
- Özdamar, K. 2004. Paket programları ile istatistiksel veri analizi. Kaan Kitapevi, Eskişehir, 649 ss.
- Papp, L. M. 2004. Capturing the interplay among within- and between- person processes using multilevel modeling techniques. *Applied and Preventive Psychology*, 11, 115-124.
- Peugh, J. L. 2010. A practical guide to multilevel modeling. *Journal of School Psychology*, 48: 85-112.
- R Development Core Team 2006. R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria online. <http://www.r-project.org>.

- Rabe-Hesketh, S., Skrondal, A. and Pickles, A. 2004. GLLAMM manual. Working Paper 160, U.C. Berkeley Division of Biostatistics, Berkeley, CAonline. (Downloadable from <http://www.bepress.com/ucbbiostat/paper160/>).
- Rasbash, J., Steele, F., Browne, W. J. and Prosser, B. 2005. *A User's Guide to MLwiN. Version 2.0*. Centre for Multilevel Modelling, University of Bristol, Bristol, UK.
- Raudenbush, S.W. and Bryk, A.S. 1986. A hierarchical model for studying school effects. *Sociology of Education*, 59, 1-17.
- Raudenbush, S. W., Bryk, A. S., Cheong, Y. F. and Congdon, R. 2004. HLM 6: Hierarchical Linear and Nonlinear Modeling. Scientific Software International, Chicago.
- Revelt, D. and Train, K. 1998. Mixed logit with repeated choices: Household's choices of appliance efficiency level. *Review of Economics and Statistics*, 80: 647–657.
- Robinson, W.S. 1950. Ecological correlations and the behavior of individuals. *American Sociological Review*, 15: 351-357.
- SAS/Stat 2004. SAS/Stat User's Guide, version 9.1. SAS Institute, Cary, NC.
- Searle, S.R. 1970. Large sample variances of maximum likelihood estimators of variance components using unbalanced data. *Biometrics* 26: 505-524.
- Seltzer, M. H., Wong, W. H. and Bryk, A. S. 1996. Bayesian analysis in applications of hierarchical models: Issues and methods. *Journal of Educational and Behavioral Statistics*, 21: 131–167.
- Serper, Ö. ve Gürsakal, N. 1989. Araştırma Yöntemleri. Filiz Kitabevi, İstanbul 191ss.
- Sheng, Y. and Chih, Y. 2008. New approaches to multilevel analysis. *Journal of Urban Health*, 85.
- Skrondal, A. and rabe-Hesketh, S. 2003. Multilevel logistic regression for polytomous data and rankings. *Psychometrika*, 68: 267–287.
- Snijders, T. A. B. and Bosker, R. J. 1999. Multilevel Analysis: An Introduction to Basic and Advanced Multilevel Modeling. Sage, Thousand Oaks, CA.
- Snijders, T. A. B. and Bosker, R. J. 1993. Standard errors and sample sizes for two-level research. *Journal of Educational Statistics*, 18: 237–259.
- Sullivan, L. M., Dukes, K. A. and Losina, E. 1999. Tutorial in Bistatistics an introduction to hierarchical linear modelling. *Statistics in Medicine* 18: 855-888.
- Spiegelhalter, D. J. 2001. Bayesian methods for cluster randomized trials with continuous responses. *Statistics in Medicine*, 20: 435–452.

- Spiegelhalter, D. J., Thomas, A., Best, N. G. and Lunn, D. 2003. WinBUGS User Manual, Version 1.4. MRC Biostatistics Unit, Cambridge, UK.
- SPSS, 2006. *SPSS Advanced Modelstm 15.0 Manual*. SPSS, Chicago.
- Stata Corp 2005. Stata Statistical Software: Release 9. Stata Corporation, College Station, TX.
- Stiratelli, R., Laird, N. M. and Ware, J. H. 1984. Random-effects models for serial observations with binary response. *Biometrics*, 40: 961–971.
- Tan, A., Freeman, J. L. and Freeman Jr., D. H. 2007. Evaluating health care performance: Strengths and limitations of multilevel analysis. *Biometrical Journal*, 49 (5): 707-718.
- Tate, R. L. and Wongbundhit, Y. 1983. Random versus nonrandom coefficient models for multilevel analysis. *Journal of Educational Statistics*, 8: 103–120.
- du Toit, M. and du Toit, S. H. C. 2002. *Interactive LISREL: User's Guide*. Scientific Software International, Chicago.
- Thompson, R. 1980. Maximum likelihood estimation of variance components. *Math. Operationsforsch.*, 11: 545-561.
- Turner, R. M. , Prevost, A. T. and Thompson, S. G. 2004. Allowing for imprecision of the intracluster correlation coefficient in the design of cluster randomized trials. *Statistics in Medicine*, 23: 1195–1214.
- Van den Eeden, P. And Hüttner, H. J. M. 1982. Introduction. *Current Sociology*, 30(3):1-9.
- Verbeke, G. and Lesaffre, E. 1996. A linear mixed-effects model with heterogeneity in the random-effects population. *Journal of the American Statistical Association*, 91: 217–221.
- Wald, A. 1943. Tests of statistical hypotheses concerning several parameters when the number of observations is large. *Transactions of the American Mathematical Society*, 54: 426-482.
- Wong, G. Y. and Mason, W. M. 1985. The hierarchical logistic regression model for multilevel analysis. *Journal of the American Statistical Association*, 80:513–524.
- Yıldız, Z. 1995. Tek etmenli tekrarlanan ölçümlü iki etmenli deneylerde etkin çözümleme yaklaşımı ve aday öğretmenlerin öğrencilere yönelik tutumlarının belirlenmesinde uygulama denemesi. Doktora tezi (yayınlanmamış), Osmangazi Üniversitesi, 104.

ÖZGEÇMİŞ

Burçin ŞİMŞEK 1987 yılında Adana’da doğdu. İlk, orta, lise öğretimlerini Mersin’de tamamladı. 2004 yılında Ankara Üniversitesi Fen Fakültesi İstatistik Bölümü’nde başladığı üniversite eğitiminden 2008 yılında mezun oldu. 2008 Eylül ayında başladığı Akdeniz Üniversitesi Fen Bilimleri Enstitüsü, Zootekni Anabilim Dalı’nda Yüksek Lisans öğretimine devam etmektedir.