

**REGRESYON ANALİZİNDE ÇOKLU BAĞLANTI:
PARAMETRİK VE SEMİPARAMETRİK TAHMİN**

Esra AKDENİZ DURAN

**DOKTORA TEZİ
İSTATİSTİK**

**GAZİ ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ**

**MAYIS 2011
ANKARA**

Esra Akdeniz Duran tarafından hazırlanan “REGRESYON ANALİZİNDE ÇOKLU BAĞLANTI: PARAMETRİK VE SEMİPARAMETRİK TAHMİN” adlı bu tezin Doktora tezi olarak uygun olduğunu onaylarım.

Prof. Dr., Salih Çelebioğlu

.....

Tez Danışmanı, İstatistik

Bu çalışma, jürimiz tarafından oy birliği ile İstatistik Anabilim Dalında Doktora tezi olarak kabul edilmiştir.

Prof. Dr. Reşat Kasap

.....

İstatistik Anabilim Dalı, Gazi Üniversitesi

Prof. Dr. Hamza Gamgam

.....

İstatistik Anabilim Dalı, Gazi Üniversitesi

Prof. Dr. Salih Çelebioğlu

.....

İstatistik Anabilim Dalı, Gazi Üniversitesi

Prof. Dr. Fikri Öztürk

.....

İstatistik Anabilim Dalı, Ankara Üniversitesi

Yrd. Doç. Dr. Güvenç Arslan

.....

Matematik Anabilim Dalı, İzmir Ekonomi Üniversitesi

Tarih: 03 / 05 / 2011

Bu tez ile G.Ü. Fen Bilimleri Enstitüsü Yönetim Kurulu Doktora derecesini onamıştır.

Prof. Dr. Bilal TOKLU

.....

Fen Bilimleri Enstitüsü Müdürü

TEZ BİLDİRİMİ

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

Esra AKDENİZ DURAN

**REGRESYON ANALİZİNDE ÇOKLU BAĞLANTI:
PARAMETRİK VE SEMİPARAMETRİK TAHMİN
(Doktora Tezi)**

Esra AKDENİZ DURAN

**GAZİ ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ
Mayıs 2011**

ÖZET

Çoklu bağlantı çoklu regresyon modellerinde iki veya daha fazla açıklayıcı değişken arasında doğrusal bir ilişkiye yakın bir ilişki olması durumudur. Doğrusal ve semiparametrik regresyon modellerinde açıklayıcı değişkenler arasında çoklu bağlantı olması durumunda en küçük kareler yöntemi ile elde edilen parametre tahminlerinin olumsuz etkilendiği bilinmektedir. Bu duruma çözüm olarak birçok yöntem önerilmiştir. Bu yöntemlerden bir tanesi en küçük kareler tahmin edicisi yerine yanlı tahmin edicileri kullanmaktır. Yanlı tahmin edicilerin çoklu bağlantı durumunda daha etkin sonuçlar verdiği bilinmektedir. Bu çalışmada parametrik ve semiparametrik regresyon modellerinde çoklu bağlantı durumunda kullanılabilecek yeni tahmin ediciler önerilmiştir. Bu tahmin edicilerin üstün olma koşulları verilmiştir. Teorik bulgular uygulamalarla ve simülasyon çalışmaları ile desteklenmiştir.

Bilim Kodu	205.1.066
Anahtar Kelimeler	Çoklu Bağlantı, Liu-tipi tahmin edici, Ridge Regresyon tahmin edicisi, Semiparametrik Regresyon
Sayfa Adedi	123
Tez Yöneticisi	Prof. Dr. Salih Çelebioğlu

MULTICOLLINEARITY IN REGRESSION ANALYSIS: PARAMETRIC AND SEMIPARAMETRIC ESTIMATION

(Ph. D. Thesis)

Esra AKDENİZ DURAN

**GAZİ UNIVERSITY
INSTITUTE OF SCIENCE AND TECHNOLOGY**

May 2011

ABSTRACT

Multicollinearity is a statistical phenomenon in which two or more predictor variables in a multiple regression model have a nearly linear relation. In this situation the coefficient estimates may change erratically in response to small changes in the data. Multicollinearity seriously affects calculations regarding individual predictors. That is, a multiple regression model with correlated predictors may give invalid results about any individual predictor; therefore it is a phenomenon that should be considered carefully. There have been many attempts in literature as a remedy to multicollinearity problem. The main stream approach is using biased estimators in place of ordinary least squares (OLS) estimators. It is well known that biased estimators are more efficient than OLS estimators in case of multicollinearity. In this study, new biased estimators are proposed for parametric or semiparametric regression models that are exposed to multicollinearity problem. The mean squared error matrix (MSEM) superiority conditions are given for each estimator. Theoretical findings are supported with applications and simulation studies.

Science Code

205.1.066

Key Words

Multicollinearity, Liu-type Estimator, Ridge Regression Estimator, Semiparametric Regression

Page Number

123

Adviser

Prof. Dr. Salih Çelebioğlu

TEŞEKKÜR

Çalışmalarım boyunca değerli yardım ve katkılarıyla beni yönlendiren hocam Prof. Dr. Salih Çelebioğlu'na yine kıymetli tecrübelerinden faydalandığım babam Prof. Dr. Fikri Akdeniz'e, ayrıca tez izleme komitelerimde bakış açımı genişleten Prof. Dr. Hamza Gangam ve Yrd. Doç. Dr. Güvenç Arslan'a ve beni koşulsuz destekleyen anneme, eşime, ablama, enişteme ve yiğenlerime teşekkürü bir borç bilirim. Ayrıca çalışmalarında bana destek veren Berlin Humboldt Üniversitesindeki hocam Prof. Dr. Wolfgang Hårdle'ye teşekkürlerimi sunarım. Çalışmalarım boyunca manevi destek veren çalışma arkadaşlarıma minnetarım. Doktora süresince bana yurtiçi araştırma bursunu sağlayan TÜBİTAK'a da ayrıca bu tez vesilesi ile teşekkürlerimi sunarım.

İÇİNDEKİLER

	Sayfa
ÖZET.....	iv
ABSTRACT.....	v
TEŞEKKÜR.....	vi
ŞEKİLLERİN LİSTESİ	xi
SİMGELER VE KISALTMALAR.....	xii
1. GİRİŞ	1
2. TAHMİN EDİCİLERİ KARŞILAŞTIRMA KRİTERLERİ: KAYIP FONKSİYONLARI	6
2.1. Karesel (kuadratik) Kayıp Fonksiyonu	6
2.2. LINEX (Linear Exponential) Kayıp Fonksiyonu	9
2.3. Sınırlı LINEX Kayıp Fonksiyonu (BLINEX Loss Function).....	10
2.4. Dengelenmiş Kayıp Fonksiyonu.....	11
3. YANLI TAHMİN EDİCİLER	13
3.1. Ridge Regresyon Tahmin Edicisi	13
3.2. Liu Tahmin Edicisi	17
4. EŞİTLİK KISITLAMALI YANLI TAHMİN EDİCİLER VE EŞİTLİK KISITLI YENİ LIU-TİPİ TAHMİN EDİCİ	25
4.1. Kısıtlı En Küçük Kareler Tahmin Edicisi	25
4.2. Kısıtlı Ridge Regresyon Tahmin Edicisi	29
4.3. Kısıtlı Liu Tahmin Edicisi	31
4.4. Kısıtlı Liu Tahmin Edicisi ve Kısıtlı EKK Tahmin Edicinin HKOM Ölçütüne göre Karşılaştırılması	40
5. JACKKNİFE LIU-TİPİ TAHMİN EDİCİ	44
5.1. Jackknife Genelleştirilmiş Liu-tipi Tahmin Edici.....	47

Sayfa

5.2. Yeni bir Jackknife Liu-tipi Tahmin Edici (YJL)	50
5.3. YJGL Tahmin Edicisi ile GL Tahmin Edicisinin HKOM Ölçütüne göre Karşılaştırılması	52
5.4. YJGL Tahmin Edicisi ile JGL Tahmin Edicisinin HKOM Ölçütüne göre Karşılaştırılması	55
6. KISMİ DOĞRUSAL MODELLER.....	58
6.1. Semiparametrik Regresyon Modelinde Cezalı EKK ve Speckman Tahmini	60
6.2. Semiparametrik Regresyon Modelinde Liu-tipi Tahmin Edici	63
6.3. Semiparametrik Regresyon Modelinde Eşitlik Kısıtlı Cezalı EKK ve Speckman Tahmin Edicileri	66
6.4. Semiparametrik Regresyon Modelinde Eşitlik Kısıtlı Liu Tahmin Edicisi	67
7. KISMİ DOĞRUSAL MODELLERDE FARKA DAYALI TAHMİN EDİCİLER ...	70
7.1. Model ve Fark Matrisi	71
7.2. Farka-dayalı Ridge Regresyon ve Liu-tipi Tahmin Ediciler.....	75
7.3. Farka-dayalı Liu-tipi Tahmin Edicisinin HKOM Ölçütüne Göre Farka Dayalı Tahmin Edici ile Karşılaştırılması	76
8. SİMÜLASYON ÇALIŞMALARİ VE UYGULAMALAR.....	81
8.1. Jackknife Liu-tipi Tahmin Edici ile İlgili Simülasyon Çalışması.....	81
8.2. Jackknife Liu-tipi Tahmin Edici ile İlgili Uygulama.....	85
8.3. Speckman ve CEKK Tahmin Edicileri ile İlgili Simülasyon Çalışması.....	90
8.4. Speckman ve CEKK Tahmin Edicileri ile İlgili Uygulama.....	94
8.5. Farka Dayalı Liu-tipi Tahmin Edici ile İlgili Simülasyon Çalışması	97
8.6. Farka Dayalı Liu-tipi Tahmin Edici ile İlgili Uygulama	102
9. SONUÇ VE ÖNERİLER	109
KAYNAKLAR	112

	Sayfa
EKLER.....	119
EK-1 Jackknife tahmin edicilerinin elde edilmesi	120
ÖZGEÇMİŞ	122

ÇİZELGELERİN LİSTESİ

Çizelge	Sayfa
Çizelge 7.1. Optimum fark ağırlıkları	72
Çizelge 8.1. EKK, LT, JLT ve YJLT tahmin edicilerinin SHKO'ları, $n=15$	83
Çizelge 8.2. EKK, LT, JLT ve YJLT tahmin edicilerinin SHKO'ları, $n=50$	83
Çizelge 8.3. EKK, LT, JLT ve YJLT tahmin edicilerinin SHKO'ları, $n=100$	83
Çizelge 8.4. LT, JLT ve YJLT tahmin edicilerinin OMS'leri, $n=15$	84
Çizelge 8.5. LT, JLT ve YJLT tahmin edicilerinin OMS'leri, $n=50$	84
Çizelge 8.6. LT, JLT ve YJLT tahmin edicilerinin OMS'leri, $n=100$	84
Çizelge 8.7. Değişkenler arasındaki korelasyon katsayıları	86
Çizelge 8.8. Farklı d değerleri için β vektörü tahminleri	87
Çizelge 8.9. Tahmin edicilerin SHKO'ları	92
Çizelge 8.10. Kısıtlı tahmin edicilerin SHKO'ları	93
Çizelge 8.11. Ev fiyatı örneği için parametre tahminleri ($d=0,628$)	95
Çizelge 8.12. Kısıt altında ev fiyatı örneği için parametre tahminleri, ($d=0,6846$)	97
Çizelge 8.13. $\hat{\beta}_{diff}(l=1)$ ve $\hat{\beta}_{diff}(l)$ tahmin edicilerinin tahmini SHKO değerleri	99
Çizelge 8.14. $\hat{\beta}_{diff}(l=1)$ ve $\hat{\beta}_{diff}(l)$ tahmin edicilerinin, standart hataları, tahmini varyans ve SHKO değerleri	107

ŞEKİLLERİN LİSTESİ

Şekil	Sayfa
Şekil 8.1. Beta tahminlerinin normlarının karesi ve Liu d parametresi	88
Şekil 8.2. Beta tahminlerinin normlarının karesi ve Liu d parametresi, LTE (mavi --),	88
Şekil 8.3. Beta tahminlerinin varyansları ve Liu d parametresi, LTE (mavi --),	89
Şekil 8.4. Beta tahminlerinin SHKO'ları ve Liu d parametresi, LTE (mavi --),	89
Şekil 8.5. $\Delta = \hat{S}HKO(\hat{\beta}_{diff}(l)) - \hat{S}HKO(\hat{\beta}_{diff})$ 'nın $\rho = 0,8$ ve $\sigma = 1$ için l değerlerine karşı grafiği(mavi).....	100
Şekil 8.6. $\Delta = \hat{S}HKO(\hat{\beta}_{diff}(l)) - \hat{S}HKO(\hat{\beta}_{diff})$ 'nın $\rho = 0,9$ ve $\sigma = 1$ için l	101
Şekil 8.7. $\Delta = \hat{S}HKO(\hat{\beta}_{diff}(l)) - \hat{S}HKO(\hat{\beta}_{diff})$ 'nın $\rho = 0,99$ ve $\sigma = 1$ için l değerlerine karşı grafiği(mavi).....	101
Şekil 8.8. Açıklayıcı değişkenlerin yanıt değişkene karşı grafikleri, doğrusal model (yeşil), kernel regresyonu (kırmızı), %95 güven bantları (mavi).....	103
Şekil 8.9. $y_i - x_i \hat{\beta}_{diff}(l = 0,85)$ 'in ms değişkenine karşı grafiği, doğrusal model (yeşil), kernel regresyonu (kırmızı), %95 güven bantları (mavi).....	108

SİMGELER VE KISALTMALAR

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

Simgeler	Açıklama
κ	Koşul sayısı
$ \cdot $	Determinant
$\ \cdot\ $	Norm
I	Birim matris
I_p	$p \times p$ tipinde birim matris
σ^2	Varyans
λ_i	i nci özdeğer
Λ	Özdeğerlerin köşegenel matrisi
T	Özvektörlerin matrisi
P	İzdüşüm matrisi
$b(\hat{\theta})$	$\hat{\theta}$ tahmin edicisinin yanlılık terimi
$M(\hat{\theta})$	$\hat{\theta}$ tahmin edicisinin HKOM si
V	Kovaryans matrisi
$\hat{\beta}$	EKK tahmin edicisi
$\hat{\beta}_R$	Kısıtlı EKK tahmin edicisi
$\hat{\beta}(k)$	Ridge regresyon tahmin edicisi
$\hat{\beta}_R(k)$	Kısıtlı ridge regresyon tahmin edicisi
$\hat{\beta}(K)$	Genelleştirilmiş ridge regresyon tahmin edicisi
$\hat{\beta}_R(K)$	Kısıtlı genelleştirilmiş ridge regresyon tahmin edicisi
$\hat{\beta}(d)$	Liu tahmin edicisi

Simgeler	Açıklama
$\hat{\beta}_R(d)$	Kısıtlı Liu tahmin edicisi
$\hat{\beta}(D)$	Genelleştirilmiş Liu tahmin edicisi
$\hat{\beta}_R(D)$	Kısıtlı genelleştirilmiş Liu tahmin edicisi
$\tilde{\beta}^0$	Hemen Hemen Yansız Genelleştirilmiş Liu tahmin edicisi
$\tilde{\beta}^{00}$	Uygulanabilir Hemen Hemen Yansız Genelleştirilmiş Liu tahmin edicisi
$\tilde{\beta}_{BC}$	Yanlılığı Düzeltilmiş Liu tahmin edicisi
$\tilde{\beta}_J(D)$	Jackknife Genelleştirilmiş Liu-tipi tahmin edici
$\hat{\beta}_{YJ}(D)$	Yeni Jackknife Genelleştirilmiş Liu-tipi tahmin edici
$\hat{\beta}_p$	Cezalı EKK tahmin edicisi
$\tilde{\beta}_s$	Speckman tahmin edicisi
$\hat{b}_p$	Kısıtlı Cezalı EKK tahmin edicisi
$\tilde{b}_s$	Kısıtlı Speckman tahmin edicisi
$\hat{\beta}_p(d)$	Cezalı Liu-tipi tahmin edici
$\tilde{\beta}_s(d)$	Speckman Liu-tipi tahmin edici
$\hat{b}_p(d)$	Kısıtlı Cezalı Liu-tipi tahmin edici
$\tilde{b}_s(d)$	Kısıtlı Speckman Liu-tipi tahmin edici
$\hat{\beta}_{diff}(k)$	Farka dayalı Ridge regresyon tahmin edicisi
$\hat{\beta}_{diff}(l)$	Farka dayalı Liu-tipi tahmin edici

Kısaltmalar	Açıklama
b.a.d.	Bağımsız Aynı Dağılımlı
CEKK	Cezalı En Küçük Kareler (Backfitting)
DKF	Dengelenmiş Kayıp Fonksiyonu

Kısaltmalar	Açıklama
EKK	En Küçük Kareler
GÇG	Genelleştirilmiş Çapraz Geçerlilik
GHHKO	Genelleştirilmiş Hata Kareler Ortalaması
GL	Genelleştirilmiş Liu
GRR	Genelleştirilmiş Ridge Regresyon
HHYGL	Hemen Hemen Yansız Genelleştirilmiş Liu
HHYL	Hemen Hemen Yansız Liu
HHYGRR	Hemen Hemen Yansız GRR
HHYUGL	Hemen Hemen Yansız Uygulanabilir GL
HHYUGR	Hemen Hemen Yansız UGRR
HKO	Hata Kareler Ortalaması
HKOM	Hata Kareler Ortalaması Matrisi
JL	Jackknife Liu
JGL	Jackknife Genelleştirilmiş Liu
KCEKK	Kısıtlı Cezalı En Küçük Kareler
KDM	Kısmi Doğrusal Model
KEKK	Kısıtlı EKK
KL	Kısıtlı Liu
KRR	Kısıtlı Ridge Regresyon
KSPC	Kısıtlı Speckman
L	Liu
LCEKK	Liu-tipi Cezalı En Küçük Kareler
LSPC	Liu-tipi Speckman
maks	Maksimum
OMS	Ortalama Mutlak Sapma
RR	Ridge Regresyon
SHKO	Skaler Hata Kareler Ortalaması
SKL	Stokastik Kısıtlı Liu
SPC	Speckman
TE	Tahmin Edici

Kısaltmalar	Açıklama
UGL	Uygulanabilir Genelleştirilmiş Liu
UGRR	Uygulanabilir Genelleştirilmiş Ridge Regresyon
YJGL	Yeni Genelleştirilmiş Jackknife Liu

1. GİRİŞ

Doğrusal regresyon modeli

$$y_i = z_i^T \gamma + \varepsilon_i, \quad i = 1, 2, \dots, n \quad (1.1)$$

biçiminde ifade edilir. Burada y_i 'ler bağımlı gözlemler, $z_i = (z_{i1}, z_{i2}, \dots, z_{ip})^T$ 'ler $p \times 1$ ($p \leq n$) tipinde gözlenmiş açıklayıcı değişkenler, $\gamma = (\gamma_1, \gamma_2, \dots, \gamma_p)^T$, $p \times 1$ tipinde bilinmeyen parametreler vektörü ve ε_i 'ler ise $E[\varepsilon_i | z_i] = 0$ ve $V[\varepsilon_i | z_i] = \sigma^2$ olan birbirinden bağımsız rasgele hatalardır. Hatalar açıklayıcı değişkenlerden bağımsızdır. Eş. 1.1'deki doğrusal regresyon modelinin matris biçiminde ifadesi

$$y = Z\gamma + \varepsilon, \quad (1.2)$$

şeklinde olur. Burada y , $n \times 1$ tipinde yanıt değişken vektörü, Z , $n \times p$ tipinde stokastik olmayan açıklayıcı değişkenlerin gözlenen matrisi, γ , $p \times 1$ tipinde bilinmeyen parametreler vektörü ve ε , $n \times 1$ tipinde $E[\varepsilon | z] = 0$ ve $V[\varepsilon | z] = \sigma^2 I$ olan birbirinden bağımsız rasgele hataların vektörüdür. Eş. 1.1 veya Eş. 1.2'deki doğrusal regresyon modelinin katsayılarını tahmin etmek için en çok kullanılan yöntem en küçük kareler (EKK)'dir. Çoklu regresyon modelinde açıklayıcı değişkenlerin genellikle ilişkisiz olduğu varsayılır ancak bu her zaman doğru olmayabilir. Açıklayıcı değişkenler ilişkili olduğunda parametre tahminleri yapmak için EKK tahmin edicisine alternatif birçok tahmin edici geliştirilmiştir. Bu çalışmada yanlı bir tahmin edici olan ve Liu (1993) tarafından önerilen Liu-tipi tahmin edicisi anlatılacak ve çeşitli regresyon modelleri için geliştirilecektir. Öncelikle teorik çıkarımların daha kolay yapılabilmesi için Eş. 1.2'deki model kanonik biçimde yazılmalıdır. Bir sonraki bölümde kanonik model verilecektir.

Doğrusal modellerle ilgili çıkarımlarda özellikle Z matrisinin kolonları arasında çoklu bağlantı olduğunda ve dolayısıyla $Z^T Z$ matrisi tekil veya yaklaşık tekil olduğunda genellikle model kanonik forma dönüştürülür [Rao, 2002: 40]. Bunun nedeni parametre tahminlerinin varyanslarını özdeğerler cinsinden ifade edebilmek ve tahmin edicilerin özelliklerini daha kolay incelemektir.

Eş. 1.1 modeli kanonik formda yazılırsa

$$y = X\beta + \varepsilon \quad (1.3)$$

modeli elde edilir. Eş. 1.3 modelinde $X = ZT$ ve $\beta = T^T \gamma$, ve T kolonları $Z^T Z$ matrisinin özvektörlerinden oluşan ortogonal matristir. Bu durumda $X^T X = (ZT)^T ZT = \Lambda$ olur. Burada $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_p)$ dır ve $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p$ ler $Z^T Z$ matrisinin sıralı özdeğerleridir. Kanonik modelde β 'nin EKK tahmin edicisi

$$\hat{\beta} = \Lambda^{-1} X^T y \quad (1.4)$$

ve $\hat{\beta}$ 'nin kovaryans matrisi

$$V(\hat{\beta}) = \sigma^2 \Lambda^{-1} \quad (1.5)$$

olur. $V(\hat{\beta})$ köşegen matris olduğundan parametre tahminleri birbiri ile ilişkisizdir ayrıca her bir $\hat{\beta}_i$ istatistiğinin varyansı $V(\hat{\beta}_i) = \sigma^2 \lambda_i^{-1}$ olur. Çoklu bağlantı olması özdeğerlerden bazılarının çok küçük olması dolayısıyla bazı $V(\hat{\beta}_i)$ 'lerin çok büyük olmasına neden olur.

Tam çoklu bağlantı, Eş. 1.1 modelinde $|Z^T Z| = 0$ ve yeniden parametrelenmiş Eş. 1.3 kanonik modelinde $|X^T X| = |\Lambda| = 0$ olması anlamına gelir. Bunun anlamı en az bir

özdeğerin 0 olduğudur. Güçlü çoklu bağlantı ise $|X^T X| \cong 0$ olması anlamına gelir, bu da özdeğerlerden en az birinin 0'a çok yakın olduğunu ifade eder [Rao ve ark., 2008]. Çalışmamızın izleyen bölümlerinde regresyon modellerinin Eş. 1.3'deki kanonik formları kullanılacaktır.

Çoklu bağlantı, doğrusal regresyon modelindeki açıklayıcı değişkenler arasında tam veya yaklaşık bir doğrusal ilişki olması (veri vektörlerinin ortogonal olmaması) anlamına gelir. Açıklayıcı değişkenler arasında çoklu bağlantı olması durumunda

- Regresyon katsayıları mutlak değerce oldukça büyük olma eğilimindedir [Montgomery ve Peck, 1992],
- Bazı katsayıların yanlış işaretli olması mümkündür [Farrar ve Glauber, 1967],
- Regresyon katsayılarının varyansları büyük olacaktır,
- Korelasyon matrisinin özdeğerlerinden biri veya çoğu çok küçük olacaktır,
- Şiddetli çoklu bağlantı altında parametre tahminleri kararsız olma eğilimi gösterecektir,
- Tahminlerin geçerliliğini görmek için yeni örneklemeler kullanıldığında tahminler şiddetle etkilenecek değişirler [Foucart, 1999].

Doğrusal regresyon modelinde açıklayıcı değişkenler arasında çoklu bağlantı olması durumunda regresyon katsayılarının standart hataları büyük olur ve bu da katsayıların tahminlerinin kesin doğrulukla yapılamayacağı anlamına gelir. Açıklayıcı değişkenler arasındaki korelasyon arttıkça tahminlerin güvenilirliği azalır. Özellikle katsayı tahmini ile ilgileniliyorsa daha güvenilir tahmin ediciler kullanılması gereklidir.

Eş 1.1 modelinde γ parametresinin tahmini genellikle

$$\hat{\gamma} = (Z^T Z)^{-1} Z^T y \quad (1.6)$$

en küçük kareler (EKK) tahmin edicisi ile yapılır. Eş. 1.6'deki EKK tahmin edicisi $\hat{y}$ yansızdır ve tüm yansız tahmin ediciler içinde minimum varyansa sahiptir. EKK tahmin edicisi yukarıda açıklayıcı değişkenler arasında çoklu bağlantı olması durumunda belirtilen nedenlerden (i-vi) iyi bir tahmin edici değildir. Çoklu bağlantı durumunda EKK tahmin edicileri yansızlığını korur ancak bazı parametre tahminlerinin varyansları çok büyük olacağından tahminler güvenilir olmaktan uzaktır. EKK tahmin edicisinin neden güvenilirmez olduğunu açıklayalım. $X^T X$ kötü koşullu olduğunda $(X^T X)^{-1}$ 'nin elemanları büyük olur ve y 'deki küçük değişimleri daha da büyütür bu da EKK tahmin edicisinin güvenilirmez olmasına sebep olur.

Çoklu bağlantının yapısını anlamamanın en iyi yolu $X^T X$ matrisinin özdeğerlerine ve özvektörlerine bakmaktır. Çoklu bağlantının derecesini göstermek için kullanılan ölçütlerden biri X matrisinin (ya da $X^T X$ in) koşul sayısı (κ)dır. $\lambda_1, \dots, \lambda_p$ 'ler $X^T X$ 'in özdeğerleri olmak üzere $\kappa = \lambda_{\max} / \lambda_{\min}$ şeklinde tanımlanır. Belsley ve ark. (1980:100-104) koşul sayısı 5 ve 10 arasında ise zayıf çoklu bağlantı, KS 'nin değeri 30'dan 100'e doğru arttıkça orta dereceli çoklu bağlantıdan güçlü çoklu bağlantıya geçiş eğilimi olduğu yorumunu yapmışlardır. Bazı araştırmacılar X matrisinin tekil değerlerinin ($X^T X$ matrisinin özdeğerlerinin karekök değeri) oranı olan $\kappa = (\lambda_{\max} / \lambda_{\min})^{1/2}$ formülünü tercih etmektedirler [Vinod ve Ullah, 1981:122; Weisberg, 1985; Chatterjee ve Hadi, 1988]. $X^T X$ matrisinin tekil olması durumunda tam çoklu bağlantı vardır, $\lambda_{\min} = 0$ 'dır ve $\kappa \rightarrow \infty$ dir. $X^T X = I$ olması durumunda $\kappa = 1$ dir. Bu durumda κ , $[1, \infty)$ açık aralığında değerler almaktadır.

Çoklu bağlantının tespit edilmesinde sıkça kullanılan bir diğer yöntem çoklu regresyon modelindeki değişkenlerin varyans şişirme katsayılarına (VIF) bakmaktır. VIF çoklu bağlantının regresyon katsayılarının varyanslarını ne kadar arttırdığını göstermektedir. Bu teşhis yönteminin adı EKK tahmin edicisi $\hat{y}$ 'nin varyansının yazılış biçiminden gelmektedir.

$$V(\hat{\gamma}_j) = \sigma^2 (X^T X)^{-1}_{jj}$$

$$= \frac{\sigma^2}{SS_j} VIF_j$$

Belsley, Kuh ve Welsch (1980, Bölüm 3) de çoklu bağlantının tespit edilme yöntemlerini ayrıntılı incelemişlerdir.

Çoklu bağlantı problemini ortadan kaldırmak için literatürde değişik yöntemler önerilmiştir. Literatürde en çok kullanılan yöntemlerden birkaç tanesi geleneksel Ridge regresyon (RR) tahmin edicisi [Hoerl ve Kennard, 1970], Stein-rule (SR) tahmin edicisi [Stein, 1956], Liu tahmin edicisi [Liu, 1993], maksimum entropi tahmin edicisi [Golan ve ark., 1996] ve temel bileşenler regresyon tahmin edicisidir [Massy, 1965], ayrıca son zamanlarda LASSO [Tibshirani, 1996] tahmin edicisi de sıklıkla çalışılmaktadır.

Bu çalışmada öncelikle tahmin edicilerin karşılaştırılması için ölçüt olarak kullanılan kayıp fonksiyonları incelenecektir. Daha sonra çoklu bağlantı probleminin etkilerini ortadan kaldırmak için yaygın olarak kullanılan Ridge regresyon (RR) ve Liu (L) yanlı tahmin edicileri ele alınacaktır. Yanlı tahmin edicilerden, ilk olarak, ridge regresyon ve kısıtlı ridge regresyon (KRR) tahmin edicisi ile ilgili literatür çalışması verilecek ve bu tahmin edicilerin özellikleri incelenecektir. İkinci olarak, Liu ve kısıtlı Liu (KL) tahmin edicisi ile ilgili literatür çalışması verilecek ve bu tahmin edicilerin özellikleri incelenecektir. Ayrıca Gross (2003) ve Hu (2005) makalelerinden yola çıkılarak yeni bir tam doğrusal eşitlik kısıtlı Liu tahmin edicisi önerilecek, özellikleri incelenecek ve bir örnek üzerinde uygulanacaktır.

2. TAHMİN EDİCİLERİ KARŞILAŞTIRMA KRİTERLERİ: KAYIP FONKSİYONLARI

Tahmin edicilerin karşılaştırılması için ölçüt olarak sıklıkla kullanılan kayıp fonksiyonlarını inceleyelim. θ , parametrelerin $p \times 1$ tipinde bir vektörü ve $\tilde{\theta}$ da θ parametresi için bir tahmin edici olsun.

2.1. Karesel (kuadratik) Kayıp Fonksiyonu

A , $p \times p$ tipinde simetrik ve en az negatif olmayan tanımlı matris olmak üzere $\tilde{\theta}$ rasgele değişkeni için karesel kayıp fonksiyonu

$$L^*(\tilde{\theta}, \theta, A) = (\tilde{\theta} - \theta)^T A (\tilde{\theta} - \theta) \quad (2.1)$$

olarak tanımlanır. Açıkça görülür ki bu kayıp fonksiyonu örnekleme bağlıdır. Bu nedenle tüm olabilir örneklemeler üzerinde beklenen kayıp fonksiyonunu düşünmek zorundayız. Bu fonksiyonun beklenen değerine tahmin edicinin karesel risk fonksiyonu veya genelleştirilmiş hata kareler ortalaması (GHKO) denir ve aşağıdaki gibi tanımlanır:

$$\text{GHKO}(\tilde{\theta}, \theta, A) = R(\tilde{\theta}, \theta, A) = E[(\tilde{\theta} - \theta)^T A (\tilde{\theta} - \theta)] \quad (2.2)$$

Özel olarak $A = I$ birim matris olarak alındığında elde edilen

$$L^*(\tilde{\theta}, \theta) = (\tilde{\theta} - \theta)^T (\tilde{\theta} - \theta) \quad (2.3)$$

karesel kayıp fonksiyonunun beklenen değeri

$$\text{SHKO}(\tilde{\theta}) = E[L^*(\tilde{\theta}, \theta)] = E[(\tilde{\theta} - \theta)^T (\tilde{\theta} - \theta)] \quad (2.4)$$

değerine skaler değerli hata kareler ortalaması (SHKO) denir. $SHKO(\tilde{\theta}), A=I$ olması durumunda $GHKO(\tilde{\theta}, \theta, A)$ 'nın özel bir halidir. SHKO tahmin edicinin iyiliğini gösteren bir ölçüdür.

$$L(\tilde{\theta}, \theta) = (\tilde{\theta} - \theta)(\tilde{\theta} - \theta)^T \quad (2.5)$$

matris değerli kayıp fonksiyonunun beklenen değeri

$$HKOM(\tilde{\theta}) = E[L(\tilde{\theta}, \theta)] = E[(\tilde{\theta} - \theta)(\tilde{\theta} - \theta)^T] \quad (2.6)$$

değerine tahmin edici için hata kareler ortalaması matrisi (HKOM) denir. HKOM tahmin edicinin iyiliğini gösteren diğer bir ölçüdür.

$\tilde{\theta}$ tahmin edicisinin kovaryans matrisi

$$V(\tilde{\theta}) = E[(\tilde{\theta} - E(\tilde{\theta}))(\tilde{\theta} - E(\tilde{\theta}))^T] \quad (2.7)$$

dır. $E(\tilde{\theta}) = \theta$ ise $\tilde{\theta}, \theta$ için yansız bir tahmin edicidir denir. $E(\tilde{\theta}) \neq \theta$ ise $\tilde{\theta}$ yanlı bir tahmin edicidir. $E(\tilde{\theta})$ ile θ arasındaki farkı (yanlılık terimi)

$$b(\tilde{\theta}) = E(\tilde{\theta}) - \theta \quad (2.8)$$

ile gösterirsek Eş. 2.6

$$\begin{aligned} HKOM(\tilde{\theta}) &= E[(\tilde{\theta} - E(\tilde{\theta})) + (E(\tilde{\theta}) - \tilde{\theta})][(\tilde{\theta} - E(\tilde{\theta})) + (E(\tilde{\theta}) - \tilde{\theta})]^T \\ &= V(\tilde{\theta}) + [b(\tilde{\theta})][b(\tilde{\theta})]^T \end{aligned} \quad (2.9)$$

şeklinde gösterilebilir. Eş. 2.9'dan HKOM, kovaryans matrisi ve yanlılık teriminin karesinin toplamıdır. HKOM tahmin edicinin varyans ve yanlılık özelliklerini aynı anda yansıttığından farklı tahmin ediciler HKOM ölçütüne göre karşılaştırılabilir.

Ayrıca, HKOM'nin iz'i alınarak

$$\text{SHKO}(\tilde{\theta}) = \text{iz}[\text{HKOM}(\tilde{\theta})] = \text{iz}[V(\tilde{\theta})] + [b(\tilde{\theta})]^T [b(\tilde{\theta})]. \quad (2.10)$$

yazılır. Karesel risk fonksiyonu HKOM'nin skaler değerli biçimidir.

2.1. Teorem

[Theobald, 1974; Trenkler, 1981] $\tilde{\theta}_1$ ve $\tilde{\theta}_2$, θ 'nın iki tahmin edicisi olsun. Bu durumda ancak ve ancak $\text{HKOM}(\tilde{\theta}_1) - \text{HKOM}(\tilde{\theta}_2) \geq 0$ ise $R(\tilde{\theta}_1, A) - R(\tilde{\theta}_2, A) \geq 0$ dir.

2.1. Tanım

(HKOM ölçütüne göre üstünlük) [Rao ve ark., 2008] $\tilde{\theta}_1$ ve $\tilde{\theta}_2$, θ 'nın iki tahmin edicisi olsun. $\text{HKOM}(\tilde{\theta}_1) - \text{HKOM}(\tilde{\theta}_2)$ farkı negatif olmayan tanımlı bir matris, başka bir gösterimle

$$\Delta(\tilde{\theta}_1, \tilde{\theta}_2) = \text{HKOM}(\tilde{\theta}_1) - \text{HKOM}(\tilde{\theta}_2) \geq 0. \quad (2.11)$$

ise Teorem 2.1'e göre $\tilde{\theta}_2$ tahmin edicisi HKOM ölçütüne göre $\tilde{\theta}_1$ tahmin edicisinden üstündür denir.

$A = I$ olması durumunda GHKO, SHKO ya eşit olacağından SHKO ölçütüne göre $\tilde{\theta}_2$ tahmin edicisinin, $\tilde{\theta}_1$ tahmin edicisine göre üstünlüğü HKOM leri karşılaştırarak kolayca incelenebilir.

2.2. Tanım

(SHKO ölçütüne göre üstünlük) [Rao ve ark. 2008: 234] $\tilde{\theta}_1$ ve $\tilde{\theta}_2, \theta$ 'nın iki tahmin edicisi olsun.

$$\text{iz}(\Delta(\tilde{\theta}_1, \tilde{\theta}_2)) = \text{iz}(\text{HKOM}(\tilde{\theta}_1)) - \text{iz}(\text{HKOM}(\tilde{\theta}_2)) \geq 0. \quad (2.12)$$

ise $\tilde{\theta}_2$, SHKO ölçütüne göre $\tilde{\theta}_1$ 'den üstündür denir.

$\tilde{\theta}_2$, HKOM ölçütüne göre $\tilde{\theta}_1$ e göre üstünse, SHKO ölçütüne göre de $\tilde{\theta}_1$ 'e göre üstündür. Ancak tersi her zaman doğru değildir. Bu durumda SHKO ölçütünün HKOM ölçütüne göre daha zayıf bir ölçüt olduğu söylenebilir.

2.3. Tanım

(Risk ölçütüne göre üstünlük) [Rao ve ark., 2008: 235] $\tilde{\theta}_1$ ve $\tilde{\theta}_2, \theta$ nın iki tahmin edicisi olsun.

$$R(\tilde{\theta}_1, A) - R(\tilde{\theta}_2, A) \geq 0 \quad (2.13)$$

ise $\tilde{\theta}_2$ karesel risk ölçütüne göre $\tilde{\theta}_1$ den üstündür denir.

Tahmin edicilerin performanslarını karşılaştırmak amacıyla kullanılan diğer bir ölçüt LINEX kayıp fonksiyonudur.

2.2. LINEX (Linear Exponential) Kayıp Fonksiyonu

Bazı tahmin edicilerin risk fonksiyonlarının karşılaştırılmasında karesel (simetrik) kayıp fonksiyonlarının çok yaygın olarak kullanıldığını biliyoruz. Özellikle yanlı

tahmin edicilerin performans ölçümü üzerindeki hemen hemen tüm çalışmalarda HKO veya eşdeğer olarak karesel kayıp fonksiyonu kullanılmıştır. Pozitif ve negatif hatalar farklı sonuçlar verdiği için simetrik kayıp fonksiyonları uygun olmayabilir.

Varian (1975), Eş. 2.14'deki kayıp fonksiyonunu verdi. Zellner (1986), bu fonksiyonu geniş bir biçimde irdeledi. Bir θ parametresinin $\hat{\theta}$ ile tahmin edilmesi durumunda LINEX kayıp fonksiyonu: $a \neq 0$, $b > 0$ ve $\Delta^* = \hat{\theta} - \theta$ ya da $\Delta = (\hat{\theta} - \theta) / \theta$ olmak üzere

$$L(\Delta) = b[e^{a\Delta} - a\Delta - 1] \quad (2.14)$$

biçiminde gösterilmiştir. Δ relatif tahmin hatasının bir birimi olmadığından, uygulamada çoğu kez Δ kullanılır. Şekil parametresi a 'nın işareti, simetrisinin yönünü, a 'nın büyüklüğü simetrisinin derecesini yansıtır. $\Delta < 0$ iken $L(\theta)$ üstel olarak, $\Delta > 0$ iken hemen hemen doğrusal olarak yükselir. Mutlak değerce küçük a için $L(\Delta) \cong (1/2)ba^2((\hat{\theta} - \theta)/\theta)^2$ karesel hata kayıp fonksiyonu ile orantılıdır. Bu nedenle LINEX kayıp fonksiyonuna karesel hata kayıp fonksiyonunun genelleştirilmesi olarak bakılabilir. Baraj yapımında maksimum su seviyesini olması gerekenden düşük tahmin etmenin (underestimation), yüksek tahmin etmeye göre (overestimation) daha ciddi bir sorun yaratacağı bilinmektedir. Ya da bir uzay mekiğinin bileşenlerinin ortalama ömrünü tahmin ederken yüksek tahmin etme, düşük tahmin etmeden daha ciddi sorunlar yaratır. Bu nedenle bu durumlarda simetrik olmayan LINEX kayıp fonksiyonun karesel kayıp fonksiyonundan daha uygun olacağı belirtilmiştir. LINEX kayıp fonksiyonu son yıllarda literatürde kuramsal olarak geniş ve etkili bir biçimde kullanıldı.

2.3. Sınırlı LINEX Kayıp Fonksiyonu (BLINEX Loss Function)

Bazı yoğunluk fonksiyonları altında beklenen LINEX kayıp fonksiyonu tanımlanamadığından uygulama alanı sınırlıdır. Örneğin, LINEX kayıp fonksiyonu

Stein-Rule (SR) tahmin edicilerinin fonksiyonlarının risk hesaplanmasında kullanılamaz [Ohtani, 1999] . Bu güçlüğü ortadan kaldırmak amacıyla Wen ve Levy (2001) tarafından

$$L^*(\Delta) = L(\Delta)/(1 + \lambda L(\Delta))$$

$$= (1/\lambda)[1 - 1/(1 + \lambda \hat{\beta}(e^{a\Delta} - a\Delta - 1))] \quad (2.15)$$

biçiminde 0 ve $1/\lambda$, ($\lambda > 0$), ile sınırlı ve LINEX kayıp fonksiyonundan çıkarılan BLINEX tanımlanmıştır.

2.4. Dengelenmiş Kayıp Fonksiyonu

Wan (1994) dengelenmiş kayıp fonksiyonu (DKF) kullanarak EKK ve eşitsizlik kısıtlamalı EKK tahmin edicilerinin risklerini karşılaştırdı. Ohtani (1995), Giles ve Giles (1995) regresyon katsayılarının ön (pre)-test tahmin edicisini DKF kullanarak incelediler. Chung ve Kim (1997) DKF kullanarak çok değişkenli normal dağılıma sahip bir vektörün eşanlı tahmin edilmesi problemini incelediler. Dipak, Ghosh ve Strawderman (1999) yaptıkları çalışmada $y = \theta + u$, $\theta \in \Omega \subset \mathbb{R}^n$ genel doğrusal modelinde $\delta(y) = \hat{\theta} + g(y)$, $g(y) \in \Omega$ biçimindeki tahmin edicileri dengelenmiş kayıp fonksiyonu kullanılarak incelenmiştir. Ayrıca θ parametre vektörü için Bayes tahmin edicilere de yer verilmiştir. Wan (2002) dengelenmiş kayıp fonksiyonunu kullanarak uygulanabilir (feasible) genelleştirilmiş ridge (UGR) tahmin edici ve hemen hemen yansız uygulanabilir genelleştirilmiş ridge (HHYUGR) tahmin edicinin risk fonksiyonlarını nümerik olarak da hesapladı. Şiddetli bağlantı olması durumu da dahil olmak üzere EKK, UGR ve HHYUGR tahmin edicilerinin relatif riskleri karşılaştırıldı. Gruber (2004) de Bayes ve deneysel Bayes tahmin ediciler ve yaklaşık minimum HKO'lu tahmin ediciler elde edilmiş ve bu tahmin edicilerin etkinliği Zellner'in DKF si kullanılarak incelenmiştir.

Zellner (1994) hem modelin iyiliğini hem de tahminin yakınlığını (estimation precision) içeren bir dengelenmiş kayıp fonksiyonu önermiştir. Gruber (2004) Zellner'in DKF sinin geliştirilmiş halini vermiştir. Gruber tarafından verilen geliştirilmiş DKF

$$L_w(\tilde{\theta}, \theta) = w(y - X\tilde{\theta})^T (y - X\tilde{\theta}) + (1-w)(\theta - \tilde{\theta})^T (\theta - \tilde{\theta}) \quad (2.16)$$

şeklinde tanımlanmıştır. Burada w modelin iyiliğinin ölçüsüne verilen ağırlıktır. $(1-w)$ ise tahminin yakınlığına verilen ağırlıktır.

Bu bölümde tahmin edicileri karşılaştırmak için sıklıkla kullanılan ölçülerden bahsedildi. Bu çalışmada önerilen tahmin ediciler tezin kapsamı açısından sadece HKOM ölçütüne göre karşılaştırılacaktır. Diğer kriterlere göre karşılaştırmalar ileriki çalışmalarda yapılabilir.

3. YANLI TAHMİN EDİCİLER

Regresyon analizinde açıklayıcı değişkenler arasında çoklu bağlantı olduğunda çoğunlukla yanlı tahmin ediciler kullanılır.

3.1. Ridge Regresyon Tahmin Edicisi

EKK tahmin edicisinde bulunan $X^T X$ matrisi açıklayıcı değişkenler arasında çoklu bağlantı olduğunda kötü koşullu olur. Ridge regresyon (RR) tahmin edicisi, $X^T X$ matrisinin her bir köşegen elemanına sabit bir katsayı ekleyerek $X^T X$ matrisinin kötü koşulluluğunu ortadan kaldırmak amacı ile Hoerl ve Kennard (1970) tarafından önerilmiştir.

Hoerl ve Kennard (1970), daha sonra da Vinod (1978) ridge tahmin edicisini hata kareler ortalaması (HKO) ölçütüne göre incelemiş ve HKO'su EKK tahmin edicisinin HKO'sundan daha küçük olan bir ridge tahmin edicisinin her zaman var olduğunu göstermişlerdir [Vinod ve Ulah, 1981].

Açıklayıcı değişkenler arasında çoklu bağlantı olması durumunda RR yöntemi ile tahmin edilen β regresyon katsayılarının EKK yöntemiyle yapılan tahminlerden daha küçük HKO'ya sahip olduğu Hoerl ve Kennard (1970) tarafından gösterilmiştir. Hoerl ve Kennard (1970) tarafından önerilen RR tahmin edicisi

$$\hat{\beta}(k) = (X^T X + kI)^{-1} X^T y, \quad k \geq 0, \quad (3.1)$$

şeklinde ifade edilir. Burada $k \geq 0$ yanlılık parametresidir. $k = 0$ olması durumunda $\hat{\beta}(0) = (X^T X)^{-1} X^T y = \hat{\beta}$, EKK tahmin edicisi elde edilir. Ridge regresyon tahmin edicisi ile elde edilen tahmin vektörünün uzunluğu EKK tahmin edicisine göre elde edilen tahmin vektörünün uzunluğundan daha azdır ($\|\hat{\beta}(k)\| \leq \|\hat{\beta}\|$). Bu yöntem çoklu bağlantıdan dolayı $\|\hat{\beta}\|$ çok büyük olduğunda bir çare olarak görülebilir.

Ridge regresyon tahmin edicisi $\hat{\beta}(k)$

$$\hat{\beta}(k) = W_k \hat{\beta}, \quad W_k = (I + k\Lambda^{-1})^{-1}, \quad k \geq 0, \quad (3.2)$$

şeklinde yazılabilir. Burada $\hat{\beta}$, EKK tahmin edicisi ve W_k da simetrik ve pozitif tanımlı ridge küçültücü (shrinkage) matris ve $\Lambda = X^T X$ 'dir.

Ridge tahmin edicisinin elde edilmesi için Lagrange çarpanlar yöntemi kullanılır. Hata terimleri minimum yapılırken aynı zamanda parametre vektörünün normuna kısıt konur.

$$\begin{aligned} F(\beta) &= \arg \min_{\beta} \{ (y - X\beta)^T (y - X\beta) + k\beta^T \beta \} \\ &= y^T y - y^T X\beta - \beta^T X^T y + \beta^T X^T X\beta + k\beta^T \beta \end{aligned}$$

$$\begin{aligned} \frac{\partial F}{\partial \beta} &= -2X^T y + 2X^T X\beta + 2k\beta = 0 \\ &= (X^T X + kI)\beta = X^T y. \end{aligned}$$

Buradan

$$\hat{\beta}(k) = (X^T X + kI)^{-1} X^T y = (\Lambda + kI)^{-1} X^T y = (\Lambda + kI)^{-1} \Lambda \hat{\beta} \quad (3.3)$$

bulunur.

β 'nin genelleştirilmiş ridge regresyon (GRR) tahmin edicisi

$$\hat{\beta}(K) = (X^T X + K)^{-1} X^T y = (\Lambda + K)^{-1} X^T y = (\Lambda + K)^{-1} \Lambda \hat{\beta} \quad (3.4)$$

dir. Burada $K = \text{diag}(k_1, k_2, \dots, k_p)$, $k_i \geq 0$ ve K 'nin elemanları yanalılık parametreleridir. $K = kI$, $k \geq 0$ alınırsa, tahmin edici ridge tahmin edici olarak adlandırılır ve

$$\hat{\beta}(k) = (X^T X + kI)^{-1} X^T y = (\Lambda + kI)^{-1} X^T y = (\Lambda + kI)^{-1} \Lambda \hat{\beta}$$

olur. Burada görüldüğü gibi ridge ve genelleştirilmiş ridge tahmin ediciler $\hat{\beta}$ 'nin bir fonksiyonu olarak yazılabilir. $\hat{\beta}$ 'nin önündeki katsayının küçültücü etkisinden dolayı bu tahmin edicilere küçültülmüş (shrinkage) tahmin ediciler denir.

$\hat{\beta}(K)$ 'nin beklenen değeri

$$E(\hat{\beta}(K)) = (\Lambda + K)^{-1} \Lambda \beta \quad (3.5)$$

dir. $\hat{\beta}(K)$ 'nin temel özelliklerinden biri yanlı olmasıdır ve yanalılık terimi

$$b(\hat{\beta}(K)) = E(\hat{\beta}(K)) - \beta = (\Delta - I)\beta \quad (3.6)$$

dır. Burada $\Delta = (\Lambda + K)^{-1} \Lambda$ 'dır. $\hat{\beta}(K)$ 'nin varyansı

$$V(\hat{\beta}(K)) = \sigma^2 (\Lambda + K)^{-1} \Lambda (\Lambda + K)^{-1} \quad (3.7)$$

dir. $\hat{\beta}(K)$ 'nin HKOM'si ve SHKO'su sırasıyla

$$\begin{aligned} \text{HKOM}(\hat{\beta}(K)) &= E((\hat{\beta}(K) - \beta)(\hat{\beta}(K) - \beta)^T) = V(\hat{\beta}(K)) + b(\hat{\beta}(K))b(\hat{\beta}(K))^T \\ &= \sigma^2 (\Lambda + K)^{-1} \Lambda (\Lambda + K)^{-1} + (I - \Delta)\beta\beta^T (I - \Delta). \end{aligned} \quad (3.8)$$

ve

$$\text{SHKO}(\hat{\beta}(K)) = \text{iz}(\text{HKOM}(\hat{\beta}(K))) \quad (3.9)$$

bulunur. $\hat{\beta}(k)_i$, $\hat{\beta}(k)$ vektörünün i -nci elemanı olmak üzere $\hat{\beta}(k)_i$ 'nin beklenen değeri, yanlışlık parametresi ve varyansı sırasıyla Eş. 3.10-3.12 eşitliklerindeki gibidir.

$$\text{E}(\hat{\beta}(k)_i) = \left(\frac{\lambda_i}{\lambda_i + k_i} \right) \beta_i \quad (3.10)$$

$$\text{b}(\hat{\beta}(k)_i) = \left(\frac{\lambda_i}{\lambda_i + k_i} - 1 \right) \beta_i = -\frac{k_i}{\lambda_i + k_i} \beta_i \quad (3.11)$$

$$\text{V}(\hat{\beta}(k)_i) = \sigma^2 \frac{\lambda_i}{(\lambda_i + k_i)^2} \quad (3.12)$$

ve $\hat{\beta}(k)$ 'nin skaler değerli hata kareler ortalaması

$$\text{SHKO}(\hat{\beta}(k)) = \sigma^2 \sum \frac{\lambda_i}{(\lambda_i + k_i)^2} + \sum \frac{k_i^2 \beta_i^2}{(\lambda_i + k_i)^2}. \quad (3.13)$$

olur. $k_1 = k_2 = \dots = k_p = k$ alınırsa $K = kI$ özel durumu için $\hat{\beta}(k)$ 'nin beklenen değeri, yanlışlık parametresi, varyansı, HKOM'si ve SHKO'su sırasıyla Eş.3.14-3.18'deki gibidir.

$$\text{E}(\hat{\beta}(k)) = (\Lambda + kI)^{-1} \Lambda \beta, \quad (3.14)$$

$$\text{b}(\hat{\beta}(k)) = ((\Lambda + kI)^{-1} \Lambda - I) \beta, \quad (3.15)$$

$$\text{V}(\hat{\beta}(k)) = \sigma^2 (\Lambda + kI)^{-1} \Lambda (\Lambda + kI)^{-1}, \quad (3.16)$$

$$\text{HKOM}(\hat{\beta}(k)) = \sigma^2 (\Lambda + kI)^{-1} \Lambda (\Lambda + kI)^{-1} + (I - \Delta) \beta \beta^T (I - \Delta), \quad (3.17)$$

ve

$$\text{SHKO}(\hat{\beta}(k)) = \text{iz}(\text{HKOM}(\hat{\beta}(k))) \quad (3.18)$$

olarak bulunur.

Ohtani (1986), Kadiyala (1984) de önerilen yanlılık terimi düzeltme yaklaşımını ve Singh, Chaubey ve Dwivedi (1986) jackknife yöntemini kullanarak aşağıdaki hemen hemen yansız genelleştirilmiş ridge (HHYGR) tahmin ediciyi elde etmişlerdir.

$\hat{\beta}(K) = (\Lambda + K)^{-1} X^T y$ genelleştirilmiş ridge tahmin edici olmak üzere

$$\hat{\beta}_K^* = [I - ((\Lambda + K)^{-1} K)^{-2}] \hat{\beta} \quad (3.19)$$

ve $K = kI$ özel durumu için

$$\hat{\beta}_k^* = [I - ((\Lambda + kI)^{-1} kI)^{-2}] \hat{\beta} \quad (3.20)$$

yi tanımladılar. Ohtani (1986) kullanılabilir (operational) HHYGR tahmin edicinin HKO özelliklerini incelemiştir.

3.2. Liu Tahmin Edicisi

Çoklu bağlantı problemini ortadan kaldırmak için Liu (1993), Stein-Rule (1956) ve geleneksel RR tahmin edicilerini birleştirmiştir. Liu'nun önerdiği ve her iki tahmin edicinin avantajlarını barındıran tahmin edici Akdeniz ve Kaçiranlar (1995) tarafından Liu (Linear unified) tahmin edicisi olarak adlandırılmıştır. Liu (doğrusal

olarak birleştirilmiş) tahmin edicisi ridge regresyon ve Stein-Rule tahmin edicisinin doğrusal bir kombinasyonu olması nedeniyle bu adı almıştır. Liu tahmin edicisi

$$\begin{aligned}\hat{\beta}(d) &= (X^T X + I)^{-1} (X^T y + d \hat{\beta}), & 0 < d < 1 \\ &= \underbrace{(X^T X + I)^{-1} X^T y}_{\hat{\beta}(k=1)} + \underbrace{(X^T X + I)^{-1} d \hat{\beta}}_{c \hat{\beta}}\end{aligned}\quad (3.21)$$

şeklinde ifade edilir. Liu tahmin edicisi Eş. 3.21’de görüldüğü gibi Ridge regresyon ve Stein-Rule tahmin edicisinin doğrusal bir kombinasyonu şeklinde ifade edilebilir. Liu tahmin edicisi, EKK tahmin edicisinin doğrusal bir fonksiyonu şeklinde

$$\hat{\beta}(d) = F_d \hat{\beta}, \quad F_d = (\Lambda + I)^{-1} (\Lambda + dI) \quad (3.22)$$

Eş. 3.22’deki gibi yazılabilir. $\hat{\beta}(k)$ uygulamada etkilidir fakat k ’nın karmaşık bir fonksiyonudur. $\hat{\beta}(d)$ ise d ’nin doğrusal bir fonksiyonu ve d ’nin seçiminin k ’nın seçimine göre daha kolay olması $\hat{\beta}(d)$ ’nin $\hat{\beta}(k)$ ’ya göre avantajlarıdır. $\hat{\beta}(k)_i$, k ’nın azalan fonksiyonu ; $\hat{\beta}(d)_i$ ise d ’nin artan bir fonksiyonudur.

Liu tahmin edicisi hata kareler toplamına katsayıların şişmemesi için regresyon katsayılarının normunun karesini $((d\hat{\beta} - \beta)^T (d\hat{\beta} - \beta))$ içeren bir ceza fonksiyonu eklenmesi ile elde edilir.

$$F(\beta) = \arg \min_{\beta} \left\{ (y - X\beta)^T (y - X\beta) + (d\hat{\beta} - \beta)^T (d\hat{\beta} - \beta) \right\}$$

$$\frac{\partial F(\beta)}{\partial \beta} = -2X^T y + 2X^T X \beta + 2\beta - 2d\hat{\beta} = 0$$

eşitliğinden

$$\hat{\beta}(d) = (X^T X + I)^{-1} (X^T y + d\hat{\beta}) = (\Lambda + I)^{-1} (\Lambda + dI) \hat{\beta} = F_d \hat{\beta}$$

elde edilir.

Genelleştirilmiş Liu tahmin edicisi d yerine köşegen elemanları farklı olabilen D matrisinin getirilmesiyle

$$\begin{aligned}\hat{\beta}(D) &= (X^T X + I)^{-1}(X^T y + D\hat{\beta}), \quad 0 < d_i < 1 \\ &= (\Lambda + I)^{-1}(\Lambda + D)\hat{\beta} \\ &= [I - (\Lambda + I)^{-1}(I - D)]\hat{\beta}\end{aligned}\tag{3.23}$$

olarak elde edilir. Burada $D = \text{diag}(d_1, d_2, \dots, d_p)$ yanlılık parametrelerinden oluşan köşegenel matristir. Burada özel olarak $d = d_1 = d_2 = \dots = d_p$ alınırsa

$$\begin{aligned}\hat{\beta}(d) &= (X^T X + I)^{-1}(X^T y + d\hat{\beta}), \quad 0 < d < 1 \\ &= (\Lambda + I)^{-1}(\Lambda + dI)\hat{\beta} \\ &= [I - (\Lambda + I)^{-1}(1 - d)]\hat{\beta}\end{aligned}$$

Liu tahmin edicisi elde edilir.

$\hat{\beta}(D)$ tahmin edicisinin beklenen değeri

$$E(\hat{\beta}(D)) = \beta - (\Lambda + I)^{-1}(I - D)\beta\tag{3.24}$$

olduğundan $\hat{\beta}(D)$ 'nin yanlılık terimi

$$b(\hat{\beta}(D)) = -(\Lambda + I)^{-1}(I - D)\beta\tag{3.25}$$

bulunur. $\hat{\beta}(D)$ 'nin varyansı $V(\hat{\beta}(D))$ ile gösterilirse

$$V(\hat{\beta}(D)) = \sigma^2 (\Lambda + I)^{-1}(\Lambda + D)\Lambda^{-1}(\Lambda + D)(\Lambda + I)^{-1}\tag{3.26}$$

ve $\hat{\beta}(D)$ 'nin HKOM'si $\text{HKOM}(\hat{\beta}(D))$ ve SHKO'su $\text{SHKO}(\hat{\beta}(D))$ ile gösterilirse

$$\text{HKOM}(\hat{\beta}(D)) = \sigma^2 (\Lambda + I)^{-1} (\Lambda + D) \Lambda^{-1} (\Lambda + D) (\Lambda + I)^{-1} + (\Lambda + I)^{-1} (I - D) \beta \beta^T (I - D) (\Lambda + I)^{-1} \quad (3.27)$$

$$\text{SHKO}(\hat{\beta}(D)) = \text{iz}(\text{HKOM}(\hat{\beta}(D))). \quad (3.28)$$

$\hat{\beta}(d)_i$, $\hat{\beta}(d)$ vektörünün i .inci elemanı olmak üzere sırasıyla yanlılık parametresi, beklenen değeri ve varyansı

$$E(\hat{\beta}(d)_i) = \beta_i - (\lambda_i + 1)^{-1} (1 - d_i) \beta_i \quad (3.29)$$

$$b(\hat{\beta}(d)_i) = -(\lambda_i + 1)^{-1} (1 - d_i) \beta_i \quad (3.30)$$

$$V(\hat{\beta}(d)_i) = \sigma^2 \frac{(d_i + \lambda_i)^2}{(1 + \lambda_i)^2 \lambda_i} \quad (3.31)$$

Eş. 3.29-3.31'deki gibi bulunur. Özel olarak $d_1 = d_2 = \dots = d_p = d$ alınırsa

$$E(\hat{\beta}(d)) = \beta - (\Lambda + I)^{-1} (1 - d) \beta \quad (3.32)$$

$$b(\hat{\beta}(d)) = -(\Lambda + I)^{-1} (1 - d) \beta \quad (3.33)$$

$$V(\hat{\beta}(d)) = \sigma^2 (\Lambda + I)^{-1} (\Lambda + dI) \Lambda^{-1} (\Lambda + dI) (\Lambda + I)^{-1} \quad (3.34)$$

$$\text{HKOM}(\hat{\beta}(d)) = \sigma^2 (\Lambda + I)^{-1} (\Lambda + dI) \Lambda^{-1} (\Lambda + dI) (\Lambda + I)^{-1} + (\Lambda + I)^{-1} (I - dI) \beta \beta^T (I - dI) (\Lambda + I)^{-1} \quad (3.35)$$

bulunur. Bu tahmin edici için SHKO

$$\text{SHKO}(\hat{\beta}(d)) = \text{iz}(\text{HKOM}(\hat{\beta}(d))) = \sigma^2 \sum_{i=1}^p \frac{(d + \lambda_i)^2}{\lambda_i (1 + \lambda_i)^2} + (1-d)^2 \sum_{i=1}^p \frac{\beta_i^2}{(1 + \lambda_i)^2} \quad (3.36)$$

bulunur. $g(d) = \sigma^2 \sum_{i=1}^p \frac{(d + \lambda_i)^2}{\lambda_i (1 + \lambda_i)^2} + (1-d)^2 \sum_{i=1}^p \frac{\beta_i^2}{(1 + \lambda_i)^2}$ olmak üzere $g(d)$

fonksiyonunun d 'ye göre türevi alınır

$$\frac{\partial g(d)}{\partial d} = \sigma^2 \sum_{i=1}^p \frac{2(d + \lambda_i)}{\lambda_i (1 + \lambda_i)^2} - 2(1-d) \sum_{i=1}^p \frac{\beta_i^2 \lambda_i}{\lambda_i (1 + \lambda_i)^2} = 0$$

eşitliğinden d yanlılık parametresinin optimum değeri

$$d_{opt} = \frac{\sum_{i=1}^p \frac{(\beta_i^2 - \sigma^2)}{(1 + \lambda_i)^2}}{\sum_{i=1}^p \frac{\lambda_i \beta_i^2 + \sigma^2}{\lambda_i (1 + \lambda_i)^2}} \quad (3.37)$$

bulunur. Bu formülde Liu (1993) te β_i^2 'nin yansız tahmin edicisi $\hat{\beta}_i^2 - \frac{\hat{\sigma}^2}{\lambda_i}$ ve σ^2

nin yansız tahmin edicisi $\hat{\sigma}^2$ kullanarak d 'nin tahmini olarak

$$\hat{d}_{mm} = 1 - \hat{\sigma}^2 \left[\frac{\sum_{i=1}^p \frac{1}{\lambda_i (1 + \lambda_i)}}{\sum_{i=1}^p \frac{\hat{\beta}_i^2}{(1 + \lambda_i)^2}} \right] \quad (3.38)$$

yi verdi. Burada $\hat{\sigma}^2 = \frac{1}{n-p} (y - X\hat{\beta})^T (y - X\hat{\beta})$ ve $\hat{\beta} = (X^T X)^{-1} X^T y$ dir.

Akdeniz ve Kaçıranlar (1995), Kadiyala (1984) tarafından önerilen yanlılığı düzeltilmiş genelleştirilmiş Liu tahmin edicisindeki

$(\hat{\beta}_D^* = \hat{\beta}(D) + (\Lambda + I)^{-1}(I + D)\beta)$ β yı Ohtani (1986) ile aynı mantık doğrultusunda

$\hat{\beta}(D)$ ile değiştirerek $\tilde{\beta}_D^* = [I - (\Lambda + I)^{-2}(I - D)^2]\hat{\beta}$ hemen hemen yansız genelleştirilmiş Liu (HHYGL) tahmin ediciyi önerdiler ve uygulanabilir genelleştirilmiş Liu (UGL) tahmin edici için tam HKO ve yanlılık terimini verdiler. Ayrıca HKO'yu kullanarak HHYGL, GL ve EKK tahmin edicilerini karşılaştırdılar.

Akdeniz, Erol ve Öztürk (1999) GL ve EKK tahmin edicilerinin konveks kombinasyonunu olarak geliştirilmiş $\beta^0 = A\hat{\beta} + (I - A)\tilde{\beta}_D$, $A = \text{diag}(a_1, a_2, \dots, a_p)$ 'yi tanımladılar. HKO cinsinden önerilen tahmin edicinin performansını HHYGL ve GL tahmin edicileri ile karşılaştırdılar.

Akdeniz (2001) bir doğrusal modelde regresyon parametrelerini tahmin etmek için yanlı tahmin etme tekniklerini kullandı. Bu tahminleri artıkları hesaplamak için kullandı. Artıkların kareleri ortalaması ve yanlılıkları elde edilerek HKO cinsinden Liu, HHYL ve EKK artıkları karşılaştırıldı.

Liu (2003) çoklu bağlantı problemini çözmek amacıyla Liu-tipi tahmin edici denilen

$$\hat{\beta}(k, d) = (X^T X + kI_p)^{-1}(X^T y - d\tilde{\beta}), \quad k > 0, \quad -\infty < d < \infty \quad (3.39)$$

biçiminde iki yanlılık parametrelili bir tahmin edici önerdi. Burada $\tilde{\beta}$, β 'nin herhangi bir tahmin edicisidir. Liu (2003), $\tilde{\beta}$ yerine EKK ve ridge tahmin edicilerini almış ve bu tahmin edicinin ridge regresyon tahmin edicisinden daha küçük HKO'ya sahip olacak şekilde k ve d yanlılık parametrelerini vermiştir. k yanlılık parametresi $X^T X + kI$ nin koşul sayısını uygun bir seviyeye düşürmek, d yanlılık ise tahmin edicinin istatistiksel özelliklerini geliştirmek için kullanılır. Eş. 3.37'de β_i yerine $\hat{\beta}_i$ ve σ^2 yerine $\hat{\sigma}^2$ alınarak

$$\hat{d}_{opt} = \frac{\sum \frac{(\hat{\beta}_i^2 - \hat{\sigma}^2)}{(1 + \lambda_i)^2}}{\sum \frac{\lambda_i \hat{\beta}_i^2 + \hat{\sigma}^2}{\lambda_i (1 + \lambda_i)^2}} \quad (3.40)$$

eşitliği elde edilir.

Akdeniz ve Erol (2003) hemen hemen yansız genelleştirilmiş ridge regresyon (HHYGRR) tahmin edici ile GL tahmin edicisini HKOM ölçütüne göre karşılaştırdılar. Teorik sonuçları nümerik örneklerle açıkladılar.

Liu (2004), Liu (2003)' deki sonuçları k 'nın rasgele olduğu duruma genellemiştir ve ridge parametresi k Hoerl-Kennard metodu ile seçildiğinde Liu-tipi tahmin edicinin ridge regresyondan üstün olabileceğini göstermiştir. Ayrıca, koşul sayısı çok kötü olduğunda k 'nın seçimi için literatürdeki yöntemlerden farklı olan koşul sayısına dayalı bir yöntem önermiştir.

Akdeniz (2004) simetrik olmayan $L(x) = b[ae^{ax} - ax - 1]$, LINEX kayıp fonksiyonunu kullanarak GL ve HHYGL tahmin edicilerinin risk fonksiyonunu elde etti. GL tahmin edicisinin EKK tahmin edicisine üstün olabilmesi için yeter koşul verdi. Ayrıca bu tahmin edicilerin risk performanslarını inceleyerek relatif etkinliklerini karşılaştırdı.

Akdeniz ve ark. (2005) dengelenmiş kayıp fonksiyonunu ele alarak FGL ve FAUGL tahmin edicilerinin risk performanslarını karşılaştırdılar.

Akdeniz ve Öztürk (2005) normallik varsayımı altında uygulanabilir Liu tahmin edicisinin stokastik yanlılık parametresinin yoğunluk fonksiyonunu elde etmişlerdir.

Akdeniz ve ark. (2006) d yanlılık parametresinin optimum değerinin beklenen değeri için üst sınır elde ettiler. Ayrıca, Liu ve genelleştirilmiş Liu tahmin

edicilerinin stokastik yanlılık parametrelerinin momentleri için genel gösterim elde ettiler.

Yang ve Xu (2008) Liu-tipi tahmin edicinin ridge regresyon tahmin edicisinden [Hoerl ve Kennard, 1970] daha iyi olan modelin uygunluğunu ölçen karakteristikleri (beta katsayıları, çoklu belirleme katsayısı, artık varyansı ve tahminin genelleştirilmiş çapraz geçerlilik sayısı vb.) ele aldılar ve tüm bu karakteristikler için Liu tahmin edicisinin RR tahmin edicisinden daha iyi olduğunu buldular. Ayrıca, Liu-tipi tahmin edicinin [Liu, 2003] yanlılık parametresi d 'nin seçimi için iki metod önerdiler ve ridge parametresi k 'nın matrisin kötü koşulluluğunu gidermek için kullanılabileceğini d 'nin ise modelin daha uygun olması için kullanılabileceğini tespit ettiler.

4. EŞİTLİK KISITLAMALI YANLI TAHMİN EDİCİLER VE EŞİTLİK KISITLI YENİ LIU-TİPİ TAHMİN EDİCİ

Çoklu regresyon modelinin kullanılması ve yorumlanması regresyon katsayılarının tahminine bağlıdır. Eşitlik kısıtlaması şeklindeki önsel bilgiler daha iyi tahmin yapılmasını sağlayabilir. $y = X\beta + \varepsilon$ doğrusal regresyon modelinde β parametre vektörünün

$$R\beta = r \quad (4.1)$$

doğrusal kısıtını sağladığını düşünelim. Burada R , $m \times p$ tipinde, bilinen ve rankı $m < p$ olan bir matris ve r , $m \times 1$ tipinde ve bilinen bir vektör olsun. Burada amaç Eş. 4.1'deki doğrusal kısıt altında β parametre vektörünü tahmin etmektir. $n \times 1$ tipinde her y vektörü için $R\tilde{\beta}(y) = r$ kısıtını sağlayan $\tilde{\beta}(y)$ tahmin edicisi, $R\beta = r$ kısıtına göre β nın kısıtlı tahmin edicisi kabul edilecektir [Gross,2003]. Parametre vektörleri üzerine kısıt konulması çoklu bağlantının giderilmesi için en direk yöntemdir. Ayrıca, $H_0 : R\beta = r$ hipotezindeki kısıtın sağlanıp sağlanmadığı F istatistiği ile test edilebilir [Rao ve ark., 2008 : 51].

4.1. Kısıtlı En Küçük Kareler Tahmin Edicisi

Örnek bilgisini ve önsel bilgiyi aynı anda kullanmak için, $S(\beta) = (y - X\beta)^T (y - X\beta)$ 'yı $R\beta = r$ kısıtı altında minimize etmeliyiz. Bunun için

$$S(\beta^*, \lambda) = (y - X\beta^*)^T (y - X\beta^*) - 2\lambda(R\beta^* - r) \quad (4.2)$$

eşitliğini β ve λ 'ya minimize etmeliyiz. Burada $\lambda(m \times 1)$, Lagrange çarpanları vektörüdür. $S(\beta, \lambda)$ 'nın β ve λ 'ya göre kısmi türevleri alındığında

$$\frac{1}{2} \frac{\partial S(\beta^*, \lambda)}{\partial \beta^*} = -X^T y + X^T X \beta^* - R^T \lambda = 0, \quad (4.3)$$

$$\frac{1}{2} \frac{\partial S(\beta^*, \lambda)}{\partial \lambda} = R\beta^* - r = 0. \quad (4.4)$$

normal denklemleri elde edilir. Buradan kısıtlı EKK (KEKK) tahmin edicisi

$$\beta^* = \hat{\beta}_R = \hat{\beta} + (X^T X)^{-1} R^T \lambda \quad (4.5)$$

elde edilir. Burada $\hat{\beta}_R$ 'deki R indisi β 'nın $R\beta = r$ kısıtı altında tahmin edildiği anlamına gelmektedir. Eş. 4.4'teki kısıt kullanılarak

$$R\hat{\beta}_R = r = R\hat{\beta} + R(X^T X)^{-1} R^T \lambda$$

elde edilir. $(X^T X)^{-1}$ pozitif tanımlı matris olduğundan $R(X^T X)^{-1} R^T$ de pozitif tanımlıdır, ayrıca rankı m 'dir ve tekil değildir. Bu durumda λ 'nın optimum değeri

$$\hat{\lambda} = (R(X^T X)^{-1} R^T)^{-1} (r - R\hat{\beta}) \quad (4.6)$$

bulunur. $\hat{\lambda}$ 'yı Eş.4.5'te yerine koyarsak, $\Lambda = X^T X$ olmak üzere, kısıtlı EKK tahmin edicisini

$$\hat{\beta}_R = \hat{\beta} + \Lambda^{-1} R^T [R\Lambda^{-1} R^T]^{-1} (r - R\hat{\beta}) \quad (4.7)$$

şeklinde elde ederiz. $\hat{\beta}_R$, yukarıda tanımladığımız anlamda $R\beta = r$ kısıtına göre kısıtlı EKK tahmin edicisidir. Eş. 4.7 denkleminde görüldüğü gibi KEKK tahmin edicisi, EKK tahmin edicisi ile Eş. 4.1 kısıtlamasını tam olarak sağlayacak bir

düzeltilme teriminin toplamına eşittir. $\hat{\beta}_R$ 'nin $R\hat{\beta}_R = r$ kısıtını her zaman sağladığı kolaylıkla görülebilir.

$$R\hat{\beta}_R = R\hat{\beta} + R\Lambda^{-1}R^T[R\Lambda^{-1}R^T]^{-1}(r - R\hat{\beta}) = r$$

olduğundan $\hat{\beta}_R$ 'nin $R\hat{\beta}_R = r$ kısıtını sağladığını söyleyebiliriz.

$\hat{\beta}_R$ 'nin beklenen değeri

$$E(\hat{\beta}_R) = \beta + (X^T X)^{-1} R^T [R(X^T X)^{-1} R^T]^{-1} \delta \quad (4.8)$$

olduğundan $\hat{\beta}_R$ tahmin edicisinin yanlılık terimi

$$b(\hat{\beta}_R) = (X^T X)^{-1} R^T [R(X^T X)^{-1} R^T]^{-1} \delta \quad (4.9)$$

dır. Burada $\delta = r - R\beta$ 'dır. $\delta \neq 0$ durumunda $\hat{\beta}_R$ tahmin edicisi her zaman yanlıdır aksi halde yansızdır. Ayrıca

$$V(\hat{\beta}_R) = \sigma^2 (X^T X)^{-1} - \sigma^2 (X^T X)^{-1} R^T [R(X^T X)^{-1} R^T]^{-1} R (X^T X)^{-1} = \sigma^2 M_0 \quad (4.10)$$

Burada

$$\begin{aligned} M_0 &= (X^T X)^{-1} - (X^T X)^{-1} R^T (R(X^T X)^{-1} R^T)^{-1} R (X^T X)^{-1} \\ &= \Lambda^{-1} - \Lambda^{-1} R^T (R\Lambda^{-1}R^T)^{-1} R\Lambda^{-1} \end{aligned} \quad (4.11)$$

dir.

$\delta \neq 0$ olması durumunda $\hat{\beta}_R$ yanlı olur fakat $V(\hat{\beta}_R)$, δ 'ya bağlı olmadığından etkilenmez. Kısıtlar sağlansa da sağlanmasa da $\hat{\beta}_R$ 'nin varyansı $\hat{\beta}$ nin varyansından küçük olur.

$$V(\hat{\beta}) - V(\hat{\beta}_R) = \sigma^2 (X^T X)^{-1} R^T [R(X^T X)^{-1} R^T]^{-1} R (X^T X)^{-1} \geq 0 \quad (4.12)$$

Eş. 4.12'nin sağındaki ifade negatif olmayan tanımlı bir matris olduğundan doğrusal kısıtların kullanılması varyansta azalmayı sağlayarak etkinliği arttırmıştır denilebilir.

δ 'nın $\hat{\beta}_R$ üzerindeki etkisini görmek için HKOM'sine bakmak gereklidir. Kısıtlar sağlanmadığında ($\delta \neq 0$) $\hat{\beta}_R$ 'nin HKOM'si

$$\begin{aligned} \text{HKOM}(\hat{\beta}_R) &= V(\hat{\beta}_R) + b(\hat{\beta}_R)b(\hat{\beta}_R)^T \\ &= \sigma^2 M_0 + (X^T X)^{-1} R^T [R(X^T X)^{-1} R^T]^{-1} \delta \delta^T [R(X^T X)^{-1} R^T]^{-1} R (X^T X)^{-1} \end{aligned} \quad (4.13)$$

ve kısıtlar sağlandığında ($\delta = 0$) $\hat{\beta}_R$ yansız bir tahmin edici olacağından $\hat{\beta}_R$ 'nin HKOM'si

$$\text{HKOM}(\hat{\beta}_R) = V(\hat{\beta}_R) = \sigma^2 M_0 \quad (4.14)$$

bulunur.

$\hat{\beta}_R$ 'nin skaler değerli hata kareler ortalaması kısıtlar sağlanmadığında

$$\begin{aligned} \text{SHKO}(\hat{\beta}_R) &= \text{iz}(\text{HKOM}(\hat{\beta}_R)) \\ &= \sigma^2 \text{iz}(M_0) + \text{iz}(b(\hat{\beta}_R)b(\hat{\beta}_R)^T) \\ &= \sigma^2 \text{iz}(M_0) + \text{iz}(b(\hat{\beta}_R)^T b(\hat{\beta}_R)) \\ &= \sigma^2 \text{iz}(M_0) + \delta^T [R(X^T X)^{-1} R^T]^{-1} R (X^T X)^{-2} R^T [R(X^T X)^{-1} R^T]^{-1} \delta \end{aligned} \quad (4.15)$$

ve kısıtlar sağlandığında $\delta = 0$ olacağından

$$\begin{aligned}
\text{SHKO}(\hat{\beta}_R) &= \text{iz}(\text{HKOM}(\hat{\beta}_R)) \\
&= \sigma^2 \text{iz}(M_0) + \text{iz}(\text{b}(\hat{\beta}_R) \text{b}(\hat{\beta}_R)^T) \\
&= \sigma^2 \text{iz}(M_0)
\end{aligned} \tag{4.16}$$

olur.

X matrisinin çoklu bağlantı problemi olması ve dolayısıyla $X^T X$ matrisinin kötü koşullu olması durumunda kısıtlı EKK yöntemiyle tahmin edilen β parametre vektörü sağlam değildir ve yanıltıcı bilgiler verebilir. Bu problemi ortadan kaldırmak için önerilen ridge ve Liu tahmin edicilerinin EKK tahmin edicisine alternatif olarak yaygın olarak kullanıldığından 3. bölümde bahsetmiştik. Bu tahmin ediciler genellikle kısıt olmaması durumunda incelenmiştir. Son zamanlarda kısıtlı EKK tahmin edicisinin çoklu bağlantı problemini ortadan kaldırmak için ridge ve Liu tahmin edicileri parametre vektörleri üzerinde kısıtlamalar olduğu varsayımı altında incelenmektedir.

4.2. Kısıtlı Ridge Regresyon Tahmin Edicisi

Sarkar (1992), $\hat{\beta}(k) = W_k \hat{\beta}$, $W_k = (I + k(X^T X)^{-1})^{-1}$, $k \geq 0$ ridge tahmin edicisinde $\hat{\beta}$ yerine $\hat{\beta}_R$ kullanarak

$$\hat{\beta}_R^*(k) = W_k \hat{\beta}_R \tag{4.17}$$

kısıtlı ridge regresyon tahmin edicisini önermiştir. $\hat{\beta}_R^*(k)$ tahmin edicisi $\hat{\beta}_R$ 'nin elemanlarını 0'a yakınsatmaktadır öyle ki $\hat{\beta}_R^*(0) = \hat{\beta}_R$ ve $\lim_{k \rightarrow \infty} \hat{\beta}_R^*(k) = 0$ 'dır. Sarkar (1992), ridge regresyonda kullanılan W_k matrisini kısıtlı EKK tahmin edicisi $\hat{\beta}_R$ 'ye uygulayarak varyansı küçültmeyi amaçlamıştır.

$$E(\hat{\beta}_R^*(k)) = W_k \beta + W_k \Lambda^{-1} R^T (R \Lambda^{-1} R^T)^{-1} \delta \tag{4.18}$$

olduğundan $k = 0$ ve $\delta = 0$ olmadığı durumda $\hat{\beta}_R^*(k)$ tahmin edicisi her zaman yanlıdır.

Gross (2003), $\hat{\beta}_R^*(k)$ 'nin y 'nin her değeri için $R\hat{\beta}_R^*(k) = r$ eşitliğini sağlamadığını ve $R\beta = r$ kısıtına göre kısıtlı bir tahmin edici olmadığını belirtmiştir. Ayrıca Sarkar (1992), $R\beta = r$ kısıtının her zaman doğru olması gibi bir varsayımda bulunmadığı için önerdiği tahmin edici $\hat{\beta}_R^*(k)$ 'nın $R\beta = r$ eşitliği ile üretilen altuzayın dışında değer alabileceğini belirtmiştir. Bu düşünceyle Gross (2003), $R\beta = r$ kısıtı altında yeni bir ridge-tipi tahmin edici önermiştir.

Swindel (1976), EKK tahmin edicisini $p \times 1$ tipindeki b vektörüne yaklaştıran

$$\hat{\beta}(k, b) = (X^T X + kI)^{-1} (X^T y + kb), \quad k \geq 0 \quad (4.19)$$

tahmin edicisini önermiştir. b vektörü hakkında hiçbir ön bilginiz yoksa $b = 0$ vektörüne doğru yaklaştırmak ridge regresyon tahmin edicisini elde edebiliriz. $R\beta = r$ doğrusal kısıtını sağlandığı varsayımı altında $b = 0$ vektörüne yaklaştırmak yerine kısıtları sağlayan en kısa vektöre yaklaştırmak daha uygun olur. Doğrusal kısıtı sağlayan en kısa vektör $\beta_0 = R^T (RR^T)^{-1} r$ Eş. 4.19'da b yerine konursa

$$\hat{\beta}(k, \beta_0) = (X^T X + kI)^{-1} (X^T y + kR^T (RR^T)^{-1} r), \quad k \geq 0 \quad (4.20)$$

tahmin edicisi elde edilir fakat bu tahmin edici Eş. 4.1 kısıtına göre yukarıda tanımlanan anlamda kısıtlı bir tahmin edici değildir. Gross (2003) doğrusal regresyon modelinde parametre vektörü üzerinde Eş. 4.1'deki kısıtları sağlayan yeni bir ridge regresyon tahmin edicisi tanımlamıştır.

$$\hat{\beta}_r(k) = \hat{\beta}(k, \beta_0) - \Lambda_k^{-1} R^T [R \Lambda_k^{-1} R^T]^{-1} (R \hat{\beta}(k, \beta_0) - r), \quad \Lambda_k = X^T X + kI, \quad k \geq 0 \quad (4.21)$$

Zhong ve Yang (2007), Gross (2003)'te önerilen Eş. 4.21'deki tahmin edicinin çok kullanılabilir olmadığını düşünmüşler ve doğrusal kısıt ($R\beta = r$) ve küresel kısıt ($\beta^T \beta \leq \phi^2$) altında hata kareler toplamını minimize ederek yeni bir kısıtlı ridge tahmin yöntemi önermişlerdir .

$$\hat{\beta}_R(k) = \hat{\beta}(k) - \Lambda_k^{-1} R^T [R \Lambda_k^{-1} R^T]^{-1} (R \hat{\beta}(k) - r) \quad (4.22)$$

Ayrıca, Zhong ve Yang (2007) önerdikleri kısıtlı ridge tahmin edicisinin HKOM ve SHKO'sunu sırasıyla

$$\text{HKOM}(\hat{\beta}_R(k)) = \sigma^2 M_k X^T X M_k + k^2 M_k \beta \beta^T M_k, \quad M_k = \Lambda_k^{-1} - \Lambda_k^{-1} R^T [R \Lambda_k^{-1} R^T]^{-1} R \Lambda_k^{-1} \quad (4.23)$$

ve

$$\text{SHKO}(\hat{\beta}_R(k)) = \sigma^2 \text{iz}(X M_k^2 X^T) + k^2 \beta^T M_k^2 \beta \quad (4.24)$$

olarak elde etmişlerdir. Ridge yanlılık parametresi k 'nın seçimi için KEKK tahmin edicisinin SHKO'sunu Eş. 4.24'deki eşitlikle karşılaştırarak gerekli ve yeterli koşulları elde etmişlerdir.

4.3. Kısıtlı Liu Tahmin Edicisi

Kaçıranlar ve ark. (1999), Sarkar (1992) ile aynı mantık doğrultusunda kısıtlı Liu tahmin edicisini önerdiler. $\hat{\beta}_R$ kısıtlı EKK tahmin edici olmak üzere

$$\hat{\beta}_R^*(d) = F_d \hat{\beta}_R, \quad F_d = (\Lambda + I)^{-1} (\Lambda + dI), \quad d \in (-\infty, +\infty) \quad (4.25)$$

ile gösterilen kısıtlı Liu (KL) tahmin ediciyi tanımladılar. Burada $\hat{\beta}_R = \hat{\beta} + \Lambda^{-1}R^T[R\Lambda^{-1}R^T]^{-1}(r - R\hat{\beta})$ ve $\Lambda = X^T X$ 'dir. β üzerindeki doğrusal kısıtlamalar $R\beta = r$ biçimindedir. Skaler değerli hata kareleri ortalaması ölçütüne göre kısıtlamalar gerçekten doğru olduğunda kısıtlı Liu tahmin edicinin, kısıtlı EKK ve Liu tahmin edicilerinin her ikisinden de daha iyi olduğunu gösterdiler. Ayrıca kısıtlamalar doğru olmadığında da kısıtlı Liu tahmin edicinin kısıtlı EKK ve Liu tahmin edicilerinin her ikisinden de üstün olabileceği koşulları verdiler.

Torigoe ve Ujiie (2006) kısıtlı EKK tahmin edicisi ile kısıtlı Liu tahmin edicisini Gauss Markov modeli altında karşılaştırdılar. Matrislerin genelleştirilmiş inverslerini kullanarak HKO ölçütüne göre kısıtlı Liu tahmin edicisinin üstünlüğü için bazı eşdeğer koşullar vermişlerdir.

Parametre vektörü ile ilgili stokastik kısıtlama ($R\beta + e = r$, $e \sim (0, \Sigma)$) olması durumunda Theil and Goldberger (1961) β parametre vektörünün tahmin edicisi olarak karma (mixed) tahmin edici $\hat{\beta}_m$ 'yi önerdiler. Hubert ve Wijekoon (2006) bir doğrusal modelde β için $\hat{\beta}_{rsd}$ ile gösterilen stokastik kısıtlı Liu tahmin edici (SKLTE) tanımlayarak HKOM ölçütüne göre etkinliğini düşünmüşlerdir. Ayrıca, SHKO ölçütüne göre yeni önerilen tahmin edicinin Liu (1993) tahmin edicisine ve karma tahmin edici $\hat{\beta}_m$ 'ye üstün olduğunu göstermişlerdir.

Yang ve Xu (2009), geleneksel karma tahmin edicide EKK tahmin edicisi yerine Liu tahmin edici koyarak yeni bir stokastik kısıtlı Liu tahmin edicisi önermişlerdir. Bu tahmin ediciyi Hubert ve Wijekoon (2006) tarafından önerilen tahmin edici, Liu (1993) tahmin edici ve geleneksel karma tahmin edici ile HKOM'ye göre karşılaştırmışlardır.

Yang ve ark. (2009) yeni bir ağırlıklı karma Liu-tipi tahmin edici önerdiler. Ayrıca, ağırlıklandırılmış karma Liu-tipi tahmin edicinin Liu tahmin edicisinden ve Schaffrin

ve Toutenburg (1990) tarafından önerilen ağırlıklı karma tahmin edicisinden HKOM ye göre üstün olması için gerekli ve yeterli koşulları elde ettiler.

Liu tahmin edicisinin elde edilmesinde kullanılan

$$\min_{\beta} \{(y - X\beta)^T (y - X\beta) + (d\hat{\beta} - \beta)^T (d\hat{\beta} - \beta)\}$$

eşitliğine $R\beta = r$ kısıtı eklenerek optimizasyon problemine dönüştürülür ve bu problemin çözümü için Lagrange fonksiyonu kullanılırsa

$$\tilde{L} = (y - X\beta)^T (y - X\beta) + (d\hat{\beta} - \beta)^T (d\hat{\beta} - \beta) + 2\lambda^T (R\beta - r). \quad (4.26)$$

şeklinde yazılabilir. $\tilde{L}$ türevlenebilir bir fonksiyondur ve λ , $(m \times 1)$ tipinde Lagrange çarpanları vektörüdür. $\tilde{L}$ fonksiyonunun β ve λ 'ya göre kısmi türevleri alınırsa

$$\begin{aligned} \frac{\partial \tilde{L}}{\partial \beta} &= -2X^T y + 2X^T X\beta - 2d\hat{\beta} + 2\beta + 2R^T \lambda = 0 \\ \frac{\partial \tilde{L}}{\partial \lambda} &= R\beta - r = 0. \end{aligned} \quad (4.27)$$

Eş. 4.27'den elde edilen sonuç $\hat{\beta}_R(d)$ olarak gösterilsin.

$$\begin{aligned} \hat{\beta}_R(d) &= (X^T X + I)^{-1} (X^T y + d\hat{\beta}) - (X^T X + I)^{-1} R^T \lambda, \\ &= \hat{\beta}(d) - (\Lambda + I)^{-1} R^T \lambda \end{aligned} \quad (4.28)$$

$R\hat{\beta}_R(d) - r = 0$ kısıtını ekleyerek

$$R\hat{\beta}_R(d) = r = R\hat{\beta}(d) - R(\Lambda + I)^{-1} R^T \lambda, \quad (4.29)$$

ve $R(\Lambda + I)^{-1}R^T > 0$, olduğundan λ 'nın optimal değeri aşağıdaki gibi elde edilir.

$$\hat{\lambda} = [R(\Lambda + I)^{-1}R^T]^{-1}(r - R\hat{\beta}(d)). \quad (4.30)$$

Eş. 4.25'te $\hat{\lambda}$ 'nın değeri yerine konularak ve $\Lambda_1 = \Lambda + I$ alınarak

$$\hat{\beta}_R(d) = \hat{\beta}(d) + \Lambda_1^{-1}R^T(R\Lambda_1R^T)^{-1}(r - R\hat{\beta}(d)) \quad (4.31)$$

kısıtlı Liu tahmin edicisi elde edilir. $\hat{\beta}_R(d)$ 'nin $R\hat{\beta}_R(d) = r$ kısıtını sağladığı görülmektedir.

Eş. 4.31 ile Eş. 4.7'yi karşılaştırsak Eş. 4.31'in Eş. 4.7'de $\hat{\beta}$ yerine $\hat{\beta}(d) = \Lambda_1^{-1}(X^T y + d\hat{\beta})$ ve $\Lambda = X^T X$ yerine $\Lambda_1 = X^T X + I$ konulduğunda elde edilebildiğinden KEKK tahmin edicisi ile tutarlı olduğu söylenebilir.

Eş. 4.31'i başka bir biçimde yazarak kısıtlı Liu tahmin edicisinin başka bir tanımını verebiliriz.

$$\hat{\beta}_R(d) = (I - \Lambda_1^{-1}R^T(R\Lambda_1R^T)^{-1}R)\hat{\beta}(d) + \Lambda_1^{-1}R^T(R\Lambda_1R^T)^{-1}r \quad (4.32)$$

$R[\Lambda_1^{-1}R^T(R\Lambda_1R^T)^{-1}]R = R$ olduğundan ve genelleştirilmiş tersin tanımından, $[\Lambda_1^{-1}R^T(R\Lambda_1R^T)^{-1}]$ 'nin R 'nin genelleştirilmiş tersi olduğunu söyleyebiliriz.

$R^- = [\Lambda_1^{-1}R^T(R\Lambda_1R^T)^{-1}]$ olarak tanımlayalım bu durumda Eş. 4.32

$$\hat{\beta}_R(d) = (I - R^-R)\hat{\beta}(d) + R^-r. \quad (4.33)$$

şeklinde yazılabilir. $R\beta = r$ doğrusal denklem sistemi tutarlı ise her bir h vektörü için

$$b^* = (I - R^-R)h + R^-r$$

vektörü bir çözümdür. Burada h , $p \times 1$ tipinde herhangi bir vektör ve R^- , R 'nin herhangi bir genelleştirilmiş tersidir. Bu durumda $\hat{\beta}_R(d)$ 'nin $R\beta = r$ doğrusal denkleminin β 'ya göre çözümünün özel bir hali olduğu görülmektedir.

4.1. Lemma.

[Gross, 2003; Zhong ve Yang, 2007]. $M_1 = \Lambda_1^{-1} - \Lambda_1^{-1}R'(R\Lambda_1^{-1}R')^{-1}R\Lambda_1^{-1}$ ve $\Lambda_1 = X'X + I$ olsun.

i) M_1 matrisi

$$M_1 = (Q\Lambda_1Q)^+ = V \begin{pmatrix} (\Psi + I)^{-1} & 0 \\ 0 & 0 \end{pmatrix} V^T,$$

şeklinde yazılabilir. Burada $M_1 = (Q\Lambda_1Q)^+$, $Q\Lambda_1Q$ matrisinin Moore-Penrose tersini göstermektedir. $Q = I - R^T(RR^T)^{-1}R$ ve Ψ , köşegen elemanları $Q\Lambda_1Q$ matrisinin sıfırdan farklı $(p-m)$ tane özdeğeri olan köşegen matristir ve V ise ortogonal matristir.

$$Q = V \begin{pmatrix} I_{p-m} & 0 \\ 0 & 0 \end{pmatrix} V^T, \quad (Q\Lambda Q)^+ = V \begin{pmatrix} \Psi^{-1} & 0 \\ 0 & 0 \end{pmatrix} V^T,$$

ii) $M_1\Lambda_1M_1 = M_1$ ve $M_0\Lambda M_0 = M_0$, burada $\Lambda = X^T X$ ve $M_0 = \Lambda^{-1} - \Lambda^{-1}R^T(R\Lambda^{-1}R^T)^{-1}R\Lambda^{-1}$.

$\hat{\beta}_R(d)$ 'nin beklenen değeri, yanlılık terimi ve varyans özellikleri ilerleyen teoremlerde verilecektir.

4.1. Teorem

β 'nın Eş. 4.1'de verilen doğrusal kısıtı sağladığı varsayımı altında $\hat{\beta}_R(d)$ 'nin yanlılık terimi

$$b(\hat{\beta}_R(d)) = E(\hat{\beta}_R(d)) - \beta = -(1-d)M_1\beta \quad (4.34)$$

dir. Burada $M_1 = \Lambda_1^{-1} - \Lambda_1^{-1}R^T(R\Lambda_1^{-1}R^T)^{-1}R\Lambda_1^{-1} = (I - R^T R)\Lambda_1^{-1}$. $\hat{\beta}_R(d)$ sadece $d=1$ olması durumunda yansız olur.

İspat

Liu –tipi tahmin edicinin yanlılık terimi aşağıdaki gibi hesaplanır.

$$E(\hat{\beta}(d)) = [I - (1-d)\Lambda_1^{-1}]\beta \quad (4.35)$$

$\beta_0 = R^T(RR^T)^{-1}r$ olsun, β_0 'ın $R\beta = r$ doğrusal kısıtını sağladığı görülmektedir. Bu durumda

$$\begin{aligned} \hat{\beta}_R(d) &= \hat{\beta}(d) + \Lambda_1^{-1}R^T(R\Lambda_1^{-1}R^T)^{-1}(r - R\hat{\beta}(d)) \\ &= \Lambda_1^{-1}\Lambda_1\hat{\beta}(d) + \Lambda_1^{-1}R^T(R\Lambda_1^{-1}R^T)^{-1}r - \Lambda_1^{-1}R^T(R\Lambda_1^{-1}R^T)^{-1}R\Lambda_1^{-1}\Lambda_1\hat{\beta}(d) \\ &= M\Lambda_1\hat{\beta}(d) + \Lambda_1^{-1}R^T(R\Lambda_1^{-1}R^T)^{-1}r \\ &= M\Lambda_1\hat{\beta}(d) + \Lambda_1^{-1}R^T(R\Lambda_1^{-1}R^T)^{-1}R\beta_0 \\ &= M\Lambda_1\hat{\beta}(d) + \Lambda_1^{-1}R^T(R\Lambda_1^{-1}R^T)^{-1}R\Lambda_1^{-1}\Lambda_1\beta_0 + \beta_0 - \Lambda_1^{-1}\Lambda_1\beta_0 \\ &= M\Lambda_1\hat{\beta}(d) + \beta_0 - M\Lambda_1\beta_0 \\ &= M\Lambda_1(\hat{\beta}(d) - \beta_0) + \beta_0 \end{aligned} \quad (4.36)$$

Eş. 4.35 ve Eş. 4.36'dan

$$\begin{aligned} E(\hat{\beta}_R(d)) &= M_1 \Lambda_1 [I - (1-d) \Lambda_1^{-1}] \beta + \beta_0 - M_1 \Lambda_1 \beta_0 \\ &= M_1 \Lambda_1 (\beta - \beta_0) + \beta_0 - (1-d) M_1 \beta \end{aligned} \quad (4.37)$$

.

$R\beta = r$ kısıtının doğru olduğu varsayımı altında ve $R\beta_0 = r$ olduğundan

$$\begin{aligned} M_1 \Lambda_1 (\beta - \beta_0) &= (I - R^- R) (\beta - \beta_0) \\ &= (\beta - \beta_0) - R^- (R\beta - R\beta_0) = \beta - \beta_0 \end{aligned} \quad (4.38)$$

.

Burada $R^- = [\Lambda_1^{-1} R^T (R \Lambda_1^{-1} R^T)^{-1}]$ dir. Bu durumda

$$\begin{aligned} E(\hat{\beta}_R(d)) &= \beta - \beta_0 - (1-d) M_1 \beta + \beta_0, \\ &= \beta - (1-d) M_1 \beta. \end{aligned} \quad (4.39)$$

Buradan $\hat{\beta}_R(d)$ 'nin yanlışlık terimi $-(1-d) M_1 \beta$ bulunur. M_1 , 0 olmayan bir matris olduğundan, $\hat{\beta}_R(d)$ sadece $d = 1$ olduğunda yansız bir tahmin edici olur.

4.2. Teorem

Doğrusal regresyon modeli için $R\beta = r$ kısıtını sağlayan $\hat{\beta}_R(d)$ tahmin edicisinin HKOM'si

$$\text{HKOM}(\hat{\beta}_R(d)) = \sigma^2 M_1 (\Lambda + dI) \Lambda^{-1} (\Lambda + dI) M_1 + (1-d)^2 M_1 \beta \beta^T M_1. \quad (4.40)$$

İspat

Eş. 4.36'daki gösterimden,

$$\begin{aligned}
\hat{\beta}_R(d) &= M_1 \Lambda_1 \hat{\beta}(d) - M_1 \Lambda_1 \beta_0 + \beta_0, \\
&= M_1 \Lambda_1 [\Lambda_1^{-1} (X^T y + d \hat{\beta})] - M_1 \Lambda_1 \beta_0 + \beta_0, \\
&= (M_1 + d M_1 \Lambda^{-1}) X^T y + (I - M_1 \Lambda_1) \beta_0.
\end{aligned} \tag{4.41}$$

yazılır. Bu durumda $\hat{\beta}_R(d)$ 'nin kovaryans matrisi

$$V(\hat{\beta}_R(d)) = V(Ay + \eta) = V(Ay) = \sigma^2 A A^T$$

ve $\hat{\beta}_R(d)$ 'nin yanlılık terimi

$$b(\hat{\beta}_R(d)) = -(1-d) M_1 \beta.$$

dir. Burada $A = (M_1 + d M_1 \Lambda^{-1}) X^T = M_1 (\Lambda + d I) \Lambda^{-1} X^T$ ve $\eta = (I - M_1 \Lambda_1) \beta_0$ 'dır.

$$HKOM(\hat{\beta}_R(d)) = V(\hat{\beta}_R(d)) + b(\hat{\beta}_R(d)) b(\hat{\beta}_R(d))^T,$$

olduğundan $\hat{\beta}_R(d)$ 'nin HKOM'si

$$HKOM(\hat{\beta}_R(d)) = \sigma^2 M_1 (\Lambda + d I) \Lambda^{-1} (\Lambda + d I) M_1 + (1-d)^2 M_1 \beta \beta^T M_1. \tag{4.42}$$

olarak elde edilir.

4.3. Teorem

Doğrusal regresyon modeli için $R\beta = r$ kısıtının doğruluğu varsayımı altında $\hat{\beta}_R$ ve $\hat{\beta}_R(d)$ tahmin edicileri verilmiş olsun. Bu durumda $D = V(\hat{\beta}_R) - V(\hat{\beta}_R(d))$ her zaman negatif olmayan tanımlı bir matristir.

İspat

Kısıtlı geleneksel EKK tahmin edicisinin $R\beta = r$ kısıtı altında çok iyi bilinen matris riski

$$V(\hat{\beta}_R) = \sigma^2 \Lambda^{-1} - \sigma^2 \Lambda^{-1} R^T (R \Lambda^{-1} R^T)^{-1} R \Lambda^{-1} = \sigma^2 M_0 \quad (4.43)$$

şeklinde yazılabilir [Rao ve ark., 2008: 225]. Teorem 4.2'den

$$V(\hat{\beta}_R(d)) = \sigma^2 M_1 (\Lambda + dI) \Lambda^{-1} (\Lambda + dI) M_1. \quad (4.44)$$

Bu durumda $D = V(\hat{\beta}_R) - V(\hat{\beta}_R(d)) = \sigma^2 M_0 - \sigma^2 M_1 (\Lambda + dI) \Lambda^{-1} (\Lambda + dI) M_1$. Lemma 4.1 kullanılarak

$$D = \sigma^2 [M_0 - M_1 (\Lambda + 2dI + d^2 \Lambda^{-1}) M_1]$$

veya

$$M_1 \Lambda M_1 = M_1 (I + \Lambda) M_1 = M_1 \text{ olduğundan}$$

$$D = \sigma^2 [M_0 - M_1 + (1 - 2d) M_1^2 - d^2 M_1 \Lambda^{-1} M_1]$$

.

elde edilir. Buradan aşağıdaki sonuçlar yazılır.

$$M_0 = (Q \Lambda Q)^+ = V \begin{pmatrix} \Psi^{-1} & 0 \\ 0 & 0 \end{pmatrix} V^T \quad \text{and} \quad M_1 = (Q \Lambda_1 Q)^+ = V \begin{pmatrix} (\Psi + I)^{-1} & 0 \\ 0 & 0 \end{pmatrix} V^T,$$

$$M_1^2 = V \begin{pmatrix} (\Psi + I)^{-2} & 0 \\ 0 & 0 \end{pmatrix} V^T.$$

Bu durumda

$$D = \sigma^2 V \begin{pmatrix} \Psi^{-1} - (\Psi + I)^{-1} + (1 - 2d)(\Psi + I)^{-2} - d^2 \Psi^{-1}(\Psi + I)^{-2} & 0 \\ 0 & 0 \end{pmatrix} V^T,$$

$$= \sigma^2 V \begin{pmatrix} \Gamma & 0 \\ 0 & 0 \end{pmatrix} V^T$$

olarak yazılabilir. Burada $\Gamma = \Psi^{-1} - (\Psi + I)^{-1} + (1 - 2d)(\Psi + I)^{-2} - d^2 \Psi^{-1}(\Psi + I)^{-2}$. Γ 'nin köşegen elemanları γ_i 'lerin aşağıdaki gibi yazılabildikleri kolayca görülebilir.

$$\gamma_i = \frac{1}{\psi_i} - \frac{1}{1 + \psi_i} + \frac{1 - 2d}{(1 + \psi_i)^2} - \frac{d^2}{\psi_i(1 + \psi_i)^2}.$$

$$0 < d < 1 \text{ ve } \psi_i > 0 \text{ olduğunda } \gamma_i = \frac{(1 + \psi_i)^2 - (\psi_i + d)^2}{\psi_i(1 + \psi_i)^2} = \frac{(1 - d)(1 + 2\psi_i + d)}{\psi_i(1 + \psi_i)^2} > 0 \text{ 'dır.}$$

Bu durumda $\Gamma > 0$ olur ve $D = V(\hat{\beta}_R) - V(\hat{\beta}_R(d)) > 0$ olur.

4.4. Kısıtlı Liu Tahmin Edicisi ve Kısıtlı EKK Tahmin Edicinin HKOM Ölçütüne göre Karşılaştırılması

Önerilen $\hat{\beta}_R(d)$ tahmin edicisinin matris riskinin hangi koşullar altında $\hat{\beta}_R$ tahmin edicisinin matris riskinden daha küçük olduğunu araştıralım. Bunun için

$$\Delta(\hat{\beta}_R, \hat{\beta}_R(d)) = \text{HKOM}(\hat{\beta}_R) - \text{HKOM}(\hat{\beta}_R(d)). \quad (4.45)$$

farkını inceleyelim. $R\beta = r$ kısıtının doğru olduğu varsayımı altında EKK tahmin edicisi $\hat{\beta}$ 'nin matris riski

$$V(\hat{\beta}) = \sigma^2 \Lambda^{-1} - \sigma^2 \Lambda^{-1} R^T (R \Lambda^{-1} R^T)^{-1} R \Lambda^{-1} = \sigma^2 M_0. \quad (4.46)$$

Eş. 1.1’deki doğrusal regresyon modeli altında, $\hat{\beta}$ ’nin HKOM’si

$$\text{HKOM}(\hat{\beta}_R) = \sigma^2 M_0 + \Lambda^{-1} R^T (R \Lambda^{-1} R^T)^{-1} \delta \delta^T (R \Lambda^{-1} R^T)^{-1} R \Lambda^{-1}, \quad (4.47)$$

burada

$$M_0 = \Lambda^{-1} - \Lambda^{-1} R^T (R \Lambda^{-1} R^T)^{-1} R \Lambda^{-1}, \quad \delta = R\beta - r \quad \text{ve} \quad b(\hat{\beta}_R) = -\Lambda^{-1} R^T (R \Lambda^{-1} R^T)^{-1} \delta \text{’dir.}$$

Teorem 4.2’den, Δ farkı $R\beta = r$ kısıtının sağlandığı varsayımı altında

$$\Delta(\hat{\beta}_R, \hat{\beta}_R(d)) = \sigma^2 M_0 - \sigma^2 M_1 (\Lambda + dI) \Lambda^{-1} (\Lambda + dI) M_1 - (1-d)^2 M_1 \beta \beta^T M_1 \quad (4.48)$$

şeklinde bulunur.

4.2. Lemma.

[Baksalary ve Kala, 1983] A negatif olmayan tanımlı bir matris olsun ve a kolon vektörü olsun. Ancak ve ancak a, A ’nın tanım kümesindeyse ve $a^T A^- a \leq 1$ ise $A - aa^T \geq 0$ olur. Burada A^- , A ’nın herhangi bir g-inversidir, $AA^-A = A$.

4.4. Teorem

$\hat{\beta}_R$ and $\hat{\beta}_R(d)$ tahmin edicileri verilmiş olsun. $0 < d < 1$ ise, bu durumda risk matrisleri farkı, $\Delta(\hat{\beta}_R, \hat{\beta}_R(d)) = \sigma^2 D - (1-d)^2 M_1 \beta \beta^T M_1$ ancak ve ancak

$$\beta^T \left[\frac{2}{1+d} Q + (QX^T XQ)^+ \right] \beta \leq \alpha \sigma^2, \quad (4.49)$$

olması durumunda negatif olmayan tanımlıdır. Burada

$$D = \sigma^2 M_0 - \sigma^2 M_1 (\Lambda + dI) \Lambda^{-1} (\Lambda + dI) M_1, \alpha = \frac{1+d}{1-d} \text{ ve } Q = I_p - R^T (RR^T)^{-1} R^T.$$

İspat

$\varphi = (1-d)M_1\beta$ olarak alınırsa Δ farkı $\Delta = \sigma^2 D - \varphi\varphi^T$ olur, burada D simetrik, negatif olmayan tanımlı matristir ve φ, D ’nin tanım kümesindedir. Lemma 4.1’i uygulayarak, ancak ve ancak $\varphi^T D \varphi \leq \alpha \sigma^2$ olması yani $\beta' M_1' D^+ M_1 \beta \leq \sigma^2 / (1-d)^2$ durumunda Δ negatif olmayan tanımlıdır. Lemma 4.1’ i kullanarak,

$$\varphi = (1-d)M_1\beta = V \begin{pmatrix} (1-d)(\Psi + I_{p-m})^{-1} & 0 \\ 0 & 0 \end{pmatrix} V^T \beta. \quad (4.50)$$

elde edilir. Ayrıca,

$$D^+ = V \begin{pmatrix} \Gamma^{-1} & 0 \\ 0 & 0 \end{pmatrix} V^T,$$

olduğundan

$$\varphi^T D^+ \varphi = \beta^T V \begin{pmatrix} (1-d)^2 (\Psi + I_{p-m})^{-1} \Gamma^{-1} (\Psi + I_{p-m})^{-1} & 0 \\ 0 & 0 \end{pmatrix} V^T \beta = \beta^T V \begin{pmatrix} \Omega & 0 \\ 0 & 0 \end{pmatrix} V^T \beta,$$

olur. Burada

$$\begin{aligned} \Omega &= (1-d)^2 (\Psi + I_{p-m})^{-1} \Gamma^{-1} (\Psi + I_{p-m})^{-1} = (1-d)^2 \left[(\Psi + I_{p-m}) \Gamma (\Psi + I_{p-m}) \right]^{-1} \\ &= (1-d)^2 \left[2(1-d) I_{p-m} + (1-d^2) \Psi^{-1} \right]^{-1} = \frac{1-d}{1+d} \left(\frac{2}{1+d} I_{p-m} + \Psi^{-1} \right)^{-1}. \end{aligned}$$

Buradan

$$\begin{aligned}\varphi^T D^+ \varphi &= \frac{1}{\alpha} \beta^T \left[V \begin{pmatrix} \left(\frac{2}{1+d} I_{p-m} + \Psi^{-1} \right) & 0 \\ 0 & 0 \end{pmatrix} V^T \right]^+ \beta \\ &= \frac{1}{\alpha} \beta^T \left[V \begin{pmatrix} \frac{2}{1+d} I_{p-m} & 0 \\ 0 & 0 \end{pmatrix} V^T + V \begin{pmatrix} \Psi^{-1} & 0 \\ 0 & 0 \end{pmatrix} V^T \right]^+ \beta,\end{aligned}$$

elde edilir. Burada D^+, D matrisinin Moore-Penrose inversini göstermektedir.

Tekrar Lemma 4.1' i kullanarak

$$\varphi^T D^+ \varphi = \frac{1}{\alpha} \beta^T \left[\frac{2}{1+d} Q + (QX^T XQ)^+ \right]^+ \beta \leq \sigma^2. \quad (4.51)$$

elde edilir. Bu sonuçla Teorem 4.4 ispatlanmış olur.

5. JACKKNİFE LİU-TİPİ TAHMİN EDİCİ

Doğrusal regresyon modelinde çoklu bağlantı olması durumunda EKK tahmin edicisine alternatif olarak Ridge veya Liu tahmin edicilerinin kullanılabileceğinden bahsedilmişti. Bu tahmin ediciler yanlı olduğu için bunların yanlılık terimlerinin jackknife tekniği ile küçültülmesine ilişkin birçok çalışma yapılmıştır [Singh ve ark.,1986; Nyquist, 1988; Batah ve ark., 2008]. Burada genelleştirilmiş Liu tahmin edicisine jackknife yöntemi uygulanarak yeni bir jackknife Liu (JL) tahmin edicisi önerilecektir. Jackknife tahmin edicileri genellikle yansız bir tahmin edici bulunamadığında ve parametre tahmininde kullanılabilecek uygun yanlı bir istatistik olduğunda kullanılır. EKK tahmin edicisi yansız olmasına rağmen çoklu bağlantı olması durumunda güvenilir bir tahmin edici olmadığından EKK tahmin edicisi yerine Ridge veya Liu-tipi tahmin edicilerine jackknife yöntemi uygulamak daha makul bir tahmin edici elde edilmesini sağlayacaktır.

$y = X\beta + \varepsilon$ kanonik modelinde β nın EKK ve genelleştirilmiş Liu-tipi tahmin edicisi sırasıyla

$$\hat{\beta} = (X^T X)^{-1} X^T y \quad (5.1)$$

ve

$$\begin{aligned} \hat{\beta}(D) &= (X^T X + I)^{-1} (X^T y + D\hat{\beta}), \quad 0 < d_i < 1 \\ &= (\Lambda + I)^{-1} (\Lambda + D)\hat{\beta} \\ &= [I - (\Lambda + I)^{-1} (I - D)]\hat{\beta} \end{aligned} \quad (5.2)$$

olarak bulunmuştu.

Jackknife tekniği ilk defa Quenouille (1949; 1956) tarafından yanlılık terimini küçültmek amacıyla önerilen parametre dışı bir yöntemdir. Tukey (1958) bu tekniği

parametre dışı varyans tahmin yöntemi olarak geliştirmiştir. Quenouille (1949) serisel korelasyonun yanlışlık terimini küçültmek amacıyla örneği iki kısma ayırma üzerine dayanan bir teknik önermiştir. 1956 daki çalışmasında bu düşüncesini örneği her biri h büyüklüğünde g gruba ayırarak genelleştirmiştir ve kullanım alanlarını incelemiştir.

$Y_1, Y_2, \dots, Y_n$ b.a.d. rasgele değişkenlerin bir örneği olsun. $\hat{\theta}$, θ 'nın n çaplı örnekten elde edilen bir tahmin edicisi olsun. $\hat{\theta}_{-i}$ ise h çaplı i . grubun çıkarılması ile elde edilen $(g-1)h$ çaplı örnekten hesaplanan ve $\hat{\theta}$ 'ya karşı gelen tahmin edici olsun. Bu durumda jackknife tahmin edicisinin hesaplanmasında kullanılan sözde değerler

$$\tilde{\theta}_i = g\hat{\theta} - (g-1)\hat{\theta}_{-i} \quad (i=1, \dots, g). \quad (5.3)$$

dir. Eş. 5.3'teki sözde değerleri kullanarak hesaplanan jackknife tahmin edicisi

$$\tilde{\theta} = \frac{1}{g} \sum_{i=1}^g \tilde{\theta}_i = g\hat{\theta} - (g-1) \frac{1}{g} \sum_{i=1}^g \hat{\theta}_{-i} \quad (5.4)$$

dir. $y = X\beta + \varepsilon$ kanonik modeli için özel olarak $g = n$ ve $h = 1$ alınarak elde edilen sözde değerler

$$P_i = n\hat{\beta} - (n-1)\hat{\beta}_{-i}, \quad (i=1, \dots, n) \quad (5.5)$$

ve i . gözlemin çıkarılması ile elde edilen jackknife tahmin edicisi

$$\tilde{\beta} = \frac{1}{n} \sum P_i \quad (5.6)$$

bulunur. Eş. 5.1'deki EKK tahmin edicisinin i . gözlemi (x_i^T, y_i) çıkarılarak elde edilen β tahmin değerini bulmak için aşağıdaki Lemma 5.1'den yararlanılacaktır.

5.1. Lemma

[Baksalary ve Kala, 1983] A , $n \times n$ tipinde simetrik bir matris ve $a \in R(A)$, $b \in R(A)$ olsun. A^+ , A matrisinin Moore-Penrose inversi olmak üzere $1 + b'A^+a \neq 0$ varsayımı altında

$$(A + ab^T)^+ = A^+ - \frac{A^+ ab^T A^+}{1 + b^T A^+ a} \quad (5.7)$$

eşitliği sağlanır.

X_{-i} , X matrisinden i . satırın çıkarılmasıyla elde edilen matris ve y_{-i} , y vektöründen i . satırın çıkarılmasıyla elde edilen vektör olmak üzere, $X_{-i}^T X_{-i} = X^T X - x_i x_i^T$ ve $X_{-i}^T y_{-i} = X^T y - x_i y_i$ eşitlikleri sağlanır. Lemma 5.1 kullanılarak i . gözlemin (x_i^T, y_i) çıkarılmasıyla elde edilen EKK β tahmin değeri

$$\begin{aligned} \hat{\beta}_{-i} &= \hat{\beta} - \frac{(X^T X)^{-1} x_i (y_i - x_i^T \hat{\beta})}{1 - x_i^T (X^T X)^{-1} x_i} \\ &= \hat{\beta} - \frac{(X^T X)^{-1} x_i r_i}{1 - w_i} \quad (i = 1, \dots, n) \end{aligned} \quad (5.8)$$

bulunur. Burada x_i^T , X matrisinin i . satırıdır ve $r_i = y_i - x_i^T \hat{\beta}$ i . EKK artık terimidir, $w_i = x_i^T (X^T X)^{-1} x_i$ uzaklık faktörüdür. EKK tahmin edicisi için sözde değerler

$$P_i = \hat{\beta} - (n-1)(X^T X)^{-1} x_i r_i (1 - w_i)^{-1} \quad (i = 1, \dots, n) \quad (5.9)$$

ve Eş. 5.9'dan yararlanılarak bulunan jackknife EKK tahmin edicisi

$$\tilde{\beta} = \hat{\beta} + (n-1)n^{-1}(X^T X)^{-1} \sum x_i r_i (1 - w_i)^{-1} \quad (i = 1, \dots, n) \quad (5.10)$$

5.1. Jackknife Genelleştirilmiş Liu-tipi Tahmin Edici

Hinkley (1977) ve Singh ve ark. (1986) makalelerindeki düşüncelerden yola çıkılarak i . gözlem (x_i^T, y_i) çıkarılarak elde edilen Liu-tipi tahmin edici

$$\hat{\beta}(D)_{-i} = (X_{-i}^T X_{-i} + I)^{-1} (X_{-i}^T y_{-i} + D \hat{\beta}_{-i}) \quad (5.11)$$

bulunur. Cebir işlemlerinden sonra Lemma 5.1'den i . gözlemin (x_i^T, y_i) çıkarılmasıyla elde edilen Liu-tipi β tahmin edicisi

$$\begin{aligned} \hat{\beta}(D)_{-i} &= (A - x_i x_i^T)^{-1} (X^T y - x_i y_i) = (A^{-1} + \frac{A^{-1} x_i x_i^T A^{-1}}{1 - x_i^T A^{-1} x_i}) (X^T y - x_i y_i) \\ &= \hat{\beta}(D) - \frac{A^{-1} x_i e_i}{1 - w_i} \end{aligned} \quad (5.12)$$

Eş. 5.12'deki gibi elde edilmiştir. Eş. 5.12'nin elde edilmesi Ek 1'de açık bir biçimde verilmiştir. Burada x_i^T , X matrisinin i . satırıdır ve $e_i = y_i - z_i^T \hat{\beta}(D)$, i . Liu artık terimidir, $w_i = x_i^T A^{-1} x_i$ uzaklık faktörüdür ve $A^{-1} = (\Lambda + I)^{-1} (I + D \Lambda^{-1})$ 'dir. Sıfırdan farklı w_i değeri modeldeki denge eksikliğini ifade etmektedir. Ağırlıklandırılmış sözde (pseudo) değerler

$$Q_i = \hat{\beta}(D) + n(1 - w_i)(\hat{\beta}(D) - \hat{\beta}(D)_{-i}) \quad (5.13)$$

dir [Hinkley, 1977; Nyquist, 1988]. β 'nın ağırlıklandırılmış jackknife genelleştirilmiş Liu-tipi tahmin edicisi

$$\tilde{\beta}_j(D) = \frac{1}{n} \sum_{i=1}^n Q_i = \bar{Q} = \hat{\beta}(D) + A^{-1} \sum_{i=1}^n x_i e_i \quad (5.14)$$

Ayrıca

$$\begin{aligned}
\sum_{i=1}^n x_i e_i &= \sum_{i=1}^n x_i (y_i - x_i^T \hat{\beta}(D)) \\
&= X^T y - X^T X (\Lambda + I)^{-1} (I + D \Lambda^{-1}) X^T y \\
&= (I - \Lambda A^{-1}) X^T y
\end{aligned} \tag{5.15}$$

olduğundan

$$\tilde{\beta}_j(D) = \hat{\beta}(D) + A^{-1} X^T y - A^{-1} \Lambda A^{-1} X^T y = (2I - A^{-1} \Lambda) \hat{\beta}(D). \tag{5.16}$$

elde edilir. Ayrıca $I - A^{-1} \Lambda = (\Lambda + I)^{-1} (I - D)$ olduğundan

$$\tilde{\beta}_j(D) = (I + (\Lambda + I)^{-1} (I - D)) \hat{\beta}(D). \tag{5.17}$$

elde edilir. Eş. 5.2'den

$$\tilde{\beta}_j(D) = (I - (\Lambda + I)^{-2} (I - D)^2) \hat{\beta} \tag{5.18}$$

bulunur.

Eş. 5.2 ve Eş. 5.18'deki parametrelerin bireysel tahmin edicileri

$$\tilde{\beta}_j(D)_i = [g(d_i)] \hat{\beta}_i. \tag{5.19}$$

dir. Burada $0 \leq g(d_i) \leq 1$ dir. Böylece, jackknife Liu-tipi tahmin edici elde edilmiş olur. Akdeniz and Kaçiranlar (1995) makalesinde önerilen hemen hemen yansız Liu-tipi tahmin edici farklı bir yoldan elde edilmesine rağmen jackknife Liu-tipi tahmin edici ile aynıdır.

Jackknife yöntemi yanlışlık terimi daha küçük olan bir tahmin edici bulunmasını sağlar. Jackknife Liu-tipi tahmin edicisinin yanlışlık terimi

$$b(\tilde{\beta}_j(D)) = -(\Lambda + I)^{-2}(I - D)^2 \beta \quad (5.20)$$

olarak bulunur. $\hat{\beta}(D)_i$ tahmin edicisinin yanlılık teriminin mutlak değeri $|b(\hat{\beta}(D)_i)|$ ile $\tilde{\beta}_j(D)_i$ tahmin edicisinin yanlılık teriminin mutlak değeri $|b(\tilde{\beta}_j(D)_i)|$ karşılaştırıldığında

$$\begin{aligned} |b(\hat{\beta}(D)_i)| - |b(\tilde{\beta}_j(D)_i)| &= \frac{(1-d_i)}{1+\lambda_i} |\beta_i| - \frac{(1-d_i)^2}{(1+\lambda_i)^2} |\beta_i| \\ &= \frac{(1-d_i)(\lambda_i + d_i)}{(1+\lambda_i)^2} |\beta_i| \end{aligned} \quad (5.21)$$

farkının pozitif olduğu görülür. Buradan jackknife yönteminin yanlılık terimini azalttığı görülmektedir. $d_1 = d_2 = \dots d_p = d$, $0 < d < 1$ olduğunda β nın geleneksel jackknife Liu-tipi (JL) tahmin edicisi

$$\tilde{\beta}_j(d) = (I - (1-d)^2(\Lambda + I)^{-2})\hat{\beta}. \quad (5.22)$$

şeklinde yazılabilir. $d=1$ durumunda $\tilde{\beta}_j(d)$ nın EKK tahmin edicisini vereceği açıktır. Eş. 5.2 ve Eş. 5.18’de görüldüğü gibi $\hat{\beta}(D)$ ve $\tilde{\beta}_j(D)$ tahmin edicileri her zaman β ’nın yanlı tahmin edicileridir. Akdeniz ve Erol (2003)’den

$$V(\hat{\beta}(D)) = \sigma^2 \left[I - (\Lambda + I)^{-1}(I - D) \right] \Lambda^{-1} \left[I - (\Lambda + I)^{-1}(I - D) \right]^T \quad (5.23)$$

ve

$$V(\tilde{\beta}_j(D)) = \sigma^2 \left[I - (\Lambda + I)^{-2}(I - D)^2 \right] \Lambda^{-1} \left[I - (\Lambda + I)^{-2}(I - D)^2 \right]^T \quad (5.24)$$

bulunur.

$\hat{\beta}(D)$ nin ve $\tilde{\beta}_j(D)$ tahmin edicilerinin HKOM'leri sırasıyla

$$\begin{aligned} \text{HKOM}(\hat{\beta}(D)) &= V(\hat{\beta}(D)) + \left[b(\hat{\beta}(D)) \right] \left[b(\hat{\beta}(D)) \right]^T \\ &= \sigma^2 \left[I - (\Lambda + I)^{-1} (I - D) \right] \Lambda^{-1} \left[I - (\Lambda + I)^{-1} (I - D) \right] + \\ &\quad (\Lambda + I)^{-1} (I - D) \beta \beta^T (I - D) (\Lambda + I)^{-1} \end{aligned} \quad (5.25)$$

ve

$$\begin{aligned} \text{HKOM}(\tilde{\beta}_j(D)) &= V(\tilde{\beta}_j(D)) + \left[b(\tilde{\beta}_j(D)) \right] \left[b(\tilde{\beta}_j(D)) \right]^T \\ &= \sigma^2 \left[I - (\Lambda + I)^{-2} (I - D)^2 \right] \Lambda^{-1} \left[I - (\Lambda + I)^{-2} (I - D)^2 \right] + \\ &\quad (\Lambda + I)^{-2} (I - D)^2 \beta \beta^T (I - D)^2 (\Lambda + I)^{-2} \end{aligned} \quad (5.26)$$

bulunur.

5.2. Yeni bir Jackknife Liu-tipi Tahmin Edici (YJL)

JGL tahmin edicisi yanlılık terimini küçültürken, çoklu bağlantı olması durumunda genelleştirilmiş Liu-tipi tahmin edici örneklem varyansını küçültür. Bu nedenle burada her iki tahmin edicideki avantajları birleştiren ve çoklu bağlantı durumunda kullanılabilecek yeni bir JGL tahmin edicisi önerilmiştir. Eş. 5.18'de EKK tahmin edicisi $\hat{\beta}$ yerine genelleştirilmiş Liu-tipi tahmin edici $\hat{\beta}(D)$ konularak yeni bir jackknife Liu-tipi (YJGL) tahmin edici elde edilmiştir. Yeni tahmin edici

$$\begin{aligned} \hat{\beta}_{YJ}(D) &= \left[I - (\Lambda + I)^{-2} (I - D)^2 \right] \hat{\beta}(D) \\ &= \left[I - (\Lambda + I)^{-2} (I - D)^2 \right] \left[I - (\Lambda + I)^{-1} (I - D) \right] \hat{\beta} \end{aligned} \quad (5.27)$$

olarak bulunur. $d_i = 1, i = 1, 2, 3, \dots, p$ olması durumunda $\hat{\beta}_{YJ}(D)_i = \hat{\beta}_i$ olacağı açıktır. $d_1 = d_2 = \dots = d_p = d, 0 < d < 1$ olması durumunda YJGL, $\hat{\beta}_{YJ}(d)$ ile gösterilen yeni jackknife Liu-tipi (YJL) tahmin edici adını almaktadır.

$$\hat{\beta}_{YJ}(d) = \left[I - (1-d)^2 (\Lambda + I)^{-2} \right] \left[I - (1-d)(\Lambda + I)^{-1} \right] \hat{\beta} \quad (5.28)$$

Eş. 5.19 ve Eş. 5.28'den

$$\begin{aligned} g(d_i) &= \left[1 - \frac{(1-d_i)^2}{(1+\lambda_i)^2} \right] \left[1 - \frac{1-d_i}{1+\lambda_i} \right] \\ &= \frac{(\lambda_i + d_i)^2}{(1+\lambda_i)^3} (2 + \lambda_i - d_i) \leq 1. \end{aligned} \quad (5.29)$$

bulunur. YJGL tahmin edicisinin yanlılık terimi

$$b(\hat{\beta}_{YJ}(D)) = -(\Lambda + I)^{-1} W (I - D) \beta \quad (5.30)$$

dir. Burada $W = I + (\Lambda + I)^{-1} (I - D) - (\Lambda + I)^{-2} (I - D)^2$ 'dir.

$\tilde{\beta}_{YJ}(D)_i$ tahmin edicisinin yanlılık teriminin mutlak değeri $|b(\tilde{\beta}_{YJ}(D)_i)|$ ile $\hat{\beta}_j(D)_i$ tahmin edicisinin yanlılık teriminin mutlak değeri $|b(\hat{\beta}_j(D)_i)|$ karşılaştırıldığında

$$\begin{aligned} |b(\tilde{\beta}_{YJ}(D)_i)| - |b(\hat{\beta}_j(D)_i)| &= \frac{(1-d_i)}{1+\lambda_i} \left(1 + \frac{(1-d_i)}{1+\lambda_i} - \frac{(1-d_i)^2}{(1+\lambda_i)^2} \right) |\beta_i| - \frac{(1-d_i)^2}{(1+\lambda_i)^2} |\beta_i| \\ &= \frac{(1-d_i)}{1+\lambda_i} \left(1 + \frac{(1-d_i)^2}{(1+\lambda_i)^2} \right) |\beta_i| \end{aligned}$$

farkının pozitif olduğu görülür. Buradan yeni jackknife Liu tahmin edicisinin jackknife Liu tahmin edicisinden daha yanlı olduğu görülmektedir. Bunun nedeni yeni jackknife Liu tahmin edicisinde yansız bir tahmin edici olan EKK yerine yanlı bir

tahmin edici olan Liu-tipi tahmin edicinin kullanılmasıdır. İleriki bölümlerde yapacağımız varyans karşılaştırmalarında yeni jackknife Liu tahmin edicisinin varyansının jackknife Liu tahmin edicisinin varyansından daha küçük olduğu gösterilecektir. Buna ek olarak YJGL tahmin edicinin HKOM'sinin JGL tahmin edicisinin HKOM'sinden daha küçük olması için gerekli ve yeterli koşul verilecektir. YJGL tahmin edicisinin varyans ve HKOM si sırasıyla

$$V(\hat{\beta}_{YJ}(D)) = \sigma^2 \Phi \Lambda^{-1} \Phi^T \quad (5.31)$$

burada $\Phi = \left[I - (\Lambda + I)^{-2} (I - D)^2 \right] \left[I - (\Lambda + I)^{-1} (I - D) \right]$ 'dir.

$$\text{HKOM}(\hat{\beta}_{YJ}(D)) = \sigma^2 \Phi \Lambda^{-1} \Phi^T + (\Lambda + I)^{-1} W (I - D) \beta \beta^T (I - D) (\Lambda + I)^{-1} \quad (5.32)$$

bulunur.

5.3. YJGL Tahmin Edicisi ile GL Tahmin Edicisinin HKOM Ölçütüne göre Karşılaştırılması

Bir önceki bölümde görüldüğü gibi YJGL tahmin edicisi yanlı olduğundan bu tahmin edicinin performansını ölçmek için kullanılacak en uygun ölçüt HKOM ölçütüdür. Bu ve sonraki bölümde YJGL tahmin edicisi, JGL tahmin edicisi ve GL tahmin edicisi ile karşılaştırılacaktır.

5.1. Teorem

D , $p \times p$ tipinde pozitif tanımlı köşegenel bir matris olsun. Bu durumda YJGL tahmin edicisinin varyansı GL tahmin edicisinin varyansından daha küçüktür.

İspat

Eş. 5.22 ve Eş. 5.30'dan

$$V(\hat{\beta}(D)) - V(\hat{\beta}_{YJ}(D)) = \sigma^2 H \quad (5.33)$$

burada

$$\begin{aligned} H &= \left[I - (\Lambda + I)^{-1}(I - D) \right] \Lambda^{-1} \left[I - (\Lambda + I)^{-1}(I - D) \right]^T - \Phi \Lambda^{-1} \Phi^T \\ &= \sigma^2 (I - F_D) \left[\Lambda^{-1} - (I - F_D)^2 \Lambda^{-1} (I - F_D)^T \right] (I - F_D)^T \end{aligned} \quad (5.34)$$

burada $F_D = (\Lambda + I)^{-1}(I - D) \cdot H$, i . köşegen elemanı

$$h_{ii} = \frac{\lambda_i(\lambda_i + d)(1 - d)^2}{(1 + \lambda_i)^5} \left[(1 + \lambda_i)^2 + (\lambda_i + d)(2 + \lambda_i - d) \right] \quad (5.35)$$

pozitif olan köşegenel bir matristir. Bu durumda H 'nin pozitif tanımlı bir matris olduğu söylenebilir.

YJGL tahmin edicisinin HKOM ölçütüne göre GL tahmin edicisinden üstünlüğünü göstermek için aşağıdaki lemmadan yararlanılacaktır.

5.2. Lemma

[Farebrother, 1976] B pozitif tanımlı bir matris ve α , $p \times 1$ tipinde bir vektör olsun.

Bu durumda ancak ve ancak $\alpha^T B^{-1} \alpha \leq 1$ koşulu sağlanırsa $B - \alpha \alpha^T$ matrisi negatif olmayan tanımlı matris olur.

5.2. Teorem

D $p \times p$ tipinde pozitif tanımlı köşegenel bir matris olsun. Bu durumda

$$\Delta_1 = \text{HKOM}(\hat{\beta}(D)) - \text{HKOM}(\hat{\beta}_{yt}(D)) \quad (5.36)$$

farkı ancak ve ancak

$$\beta^T \left[L^{-1}(\sigma^2 H + F_D \beta \beta^T F_D^T)(L^T)^{-1} \right]^{-1} \beta \leq 1 \quad (5.37)$$

olması durumunda negatif olmayan tanımlı bir matristir. Burada

$$L = (\Lambda + I)^{-1} \left[I + F_D - F_D^2 \right] (I - D) = (\Lambda + I)^{-1} W (I - D) = F_D W \text{ ve}$$

$$W = I + F_D - F_D^2 \text{ 'dir.}$$

İspat

Eş. 5.25 ve Eş. 5.32'den

$$\begin{aligned} \Delta_1 &= \text{HKOM}(\hat{\beta}(D)) - \text{HKOM}(\hat{\beta}_{yt}(D)) \\ &= \sigma^2 H + F_D \beta \beta^T F_D^T - F_D W \beta \beta^T W^T F_D^T \end{aligned} \quad (5.38)$$

burada $W = I + F_D - F_D^2$ pozitif tanımlı bir matristir. Teorem 5.1'de H matrisinin pozitif tanımlı olduğu bulunmuştu. Bu durumda, Δ_1 farkı ancak ve ancak $L^{-1} \Delta_1 (L^T)^{-1}$ negatif olmayan tanımlı olduğunda negatif olmayan tanımlıdır. $L^{-1} \Delta_1 (L^T)^{-1}$ matrisi

$$L^{-1} \Delta_1 (L^T)^{-1} = L^{-1} (\sigma^2 H + F_D \beta \beta^T F_D^T) (L^T)^{-1} - \beta \beta^T. \quad (5.39)$$

şeklinde yazılabilir.

$\sigma^2 H + F_D \beta \beta^T F_D^T$ matrisi simetrik ve pozitif tanımlı olduğundan Lemma 5.2'den ancak ve ancak

$$\beta^T \left[L^{-1}(\sigma^2 H + F_D \beta \beta^T F_D^T)(L^T)^{-1} \right] \beta \leq 1 \quad (5.40)$$

koşulu sağlandığında $L^{-1}\Delta_1(L^T)^{-1}$ matrisinin negatif olmayan tanımlı bir matris olduğu söylenebilir.

5.4. YJGL Tahmin Edicisi ile JGL Tahmin Edicisinin HKOM Ölçütüne göre Karşılaştırılması

YJGL tahmin edicisi örneklem varyansı bakımından JGL tahmin edicisine göre üstündür.

5.3. Teorem

D , $p \times p$ tipinde pozitif tanımlı köşegenel bir matris olsun. Bu durumda YJGL tahmin edicisinin varyansı JGL tahmin edicisinin varyansından daha küçüktür.

İspat

Eş. 5.24 ve Eş. 5.31'den YJGL ve JGL tahmin edicilerinin varyansları sırasıyla

$$\begin{aligned} V(\tilde{\beta}_J(D)) &= \sigma^2 \left[I - (\Lambda + I)^{-2} (I - D)^2 \right] \Lambda^{-1} \left[I - (\Lambda + I)^{-2} (I - D)^2 \right]^T \\ &= VU \Lambda^{-1} U^T V^T \end{aligned} \quad (5.41)$$

ve

$$V(\hat{\beta}_{YJ}(D)) = \sigma^2 \Phi \Lambda^{-1} \Phi^T = \sigma^2 VUV \Lambda^{-1} V^T U^T V^T \quad (5.42)$$

dir. Buradan YJGL tahmin edicisinin varyansı ve JGL tahmin edicisinin varyansı arasındaki fark

$$V(\tilde{\beta}_J(D)) - V(\hat{\beta}_{YJ}(D)) = \sigma^2 \Sigma \quad (5.43)$$

bulunur. Burada $\Sigma = VU(\Lambda^{-1} - V\Lambda^{-1}V^T)U^TV^T$, $V = I - F_D$ ve $U = I + F_D$. Σ köşegenel bir matristir. Σ 'nın i . köşegen elemanı

$$V(\tilde{\beta}_J(D))_i - V(\hat{\beta}_{YJ}(D))_i = \frac{\sigma^2(\lambda_i + d)^3(2 + \lambda_i - d)^3}{\lambda_i(1 + \lambda_i)^5} > 0 \quad (5.44)$$

olduğundan Σ pozitif tanımlı bir matristir.

Bir sonraki teoremden YJGL tahmin edicisinin JGL tahmin edicisine HKOM ölçütü bakımından üstün olması için gerekli ve yeterli koşul elde edilmiştir. Teoremin ispatı Teorem 5.2 nin ispatı ile benzerlik göstermektedir.

5.4. Teorem

D , $p \times p$ tipinde pozitif tanımlı köşegenel bir matris olsun. Bu durumda

$$\Delta_2 = \text{HKOM}(\tilde{\beta}_J(D)) - \text{HKOM}(\hat{\beta}_{YJ}(D)) \quad (5.45)$$

farkı ancak ve ancak

$$\beta^T \left[L^{-1}(\sigma^2 \Sigma + F_D^2 \beta \beta^T F_D^T)(L^T)^{-1} \right]^{-1} \beta \leq 1 \quad (5.46)$$

olması durumunda negatif olmayan tanımlı bir matristir.

İspat

Eş. 5.24 ve Eş. 5.31'den $\tilde{\beta}_J(D)$ ve $\hat{\beta}_{YJ}(D)$ tahmin edicilerinin varyansları arasındaki fark

$$\begin{aligned}\Delta_2 &= V(\tilde{\beta}_J(D)) - V(\hat{\beta}_{YJ}(D)) + F_D^2 \beta \beta^T (F_D^T)^2 - F_D W \beta \beta^T W^T F_D^T \\ &= \sigma^2 \Sigma + F_D^2 \beta \beta^T (F_D^T)^2 - F_D W \beta \beta^T W^T F_D^T\end{aligned}\quad (5.47)$$

Teorem 5.3'te Σ matrisinin pozitif tanımlı bir matris olduğu bulunmuştu. Bu durumda Δ_2 farkı ancak ve ancak $L^{-1} \Delta_2 (L^T)^{-1}$ negatif olmayan tanımlı olduğunda negatif olmayan tanımlıdır. Buna göre $L^{-1} \Delta_2 (L^T)^{-1}$

$$L^{-1} \Delta_2 (L^T)^{-1} = L^{-1} (\sigma^2 \Sigma + F_D^2 \beta \beta^T (F_D^T)^2) (L^T)^{-1} - \beta \beta^T. \quad (5.48)$$

şeklinde yazılabilir. $(\sigma^2 \Sigma + F_D^2 \beta \beta^T (F_D^T)^2)$ simetrik ve pozitif tanımlı olduğundan Lemma 5.2'den ancak ve ancak

$$\beta^T \left[L^{-1} (\sigma^2 \Sigma + F_D^2 \beta \beta^T (F_D^T)^2) (L^T)^{-1} \right]^{-1} \beta \leq 1 \quad (5.49)$$

koşulu sağlandığında $L^{-1} \Delta_2 (L^T)^{-1}$ matrisinin negatif olmayan tanımlı bir matris olduğu söylenebilir.

Bu bölümde önerilen tahmin edici ile ilgili uygulama ve simulasyon çalışması 8. bölümde verilmiştir.

6. KISMİ DOĞRUSAL MODELLER

Bilindiği gibi klasik regresyon modeli yanıt değişkeni ve açıklayıcı değişkenler arasında doğrusal bir ilişki olduğu varsayımına dayanır. Uygulamada bu ilişkinin şeklinin doğrusal olmadığı veya bilinmediği birçok durumla karşılaşılır. Parametrik olmayan regresyon yanıt değişkeni ve açıklayıcı değişkenler arasındaki ilişkinin şekli bilinmediğinde sıklıkla kullanılan bir yöntemdir. Açıklayıcı değişken sayısı fazla olduğunda parametrik olmayan regresyon kullanmak çok boyutluluk problemine (curse of dimensionality) sebep olur. Bu problem, parametrik olmayan regresyonda açıklayıcı değişken sayısı arttıkça tahminlerin gerçeğe yakınsama hızının düşmesi anlamına gelmektedir. Parametrik olmayan regresyon modeline alternatif olarak, Buja ve ark. (1989) toplamsal regresyon modellerini önermişlerdir. Toplamsal regresyon modellerinde, her bir açıklayıcı değişken ile bağımlı değişken arasında bireysel bir fonksiyonel ilişki vardır. Semiparametrik regresyon modeli, parametrik ve parametrik olmayan regresyon fonksiyonunun toplamsal olarak birleşiminden oluşur. Semiparametrik regresyon modelleri genellikle parametrik olmayan modellerin iyi sonuç vermediği veya araştırmacının parametrik model kullanmak istemesine rağmen bazı değişkenlerin fonksiyonel biçimlerinin bilinmemesi nedeniyle tercih edilir. Literatürde birçok semiparametrik regresyon modeli önerilmiştir. Kısmi doğrusal modelleri, indeks modelleri ve değişken katsayılı (varying coefficient) regresyon modelleri sıklıkla kullanılan semiparametrik regresyon modelleridir. Kısmi doğrusal model (KDM) Eş. 6.1 biçimindeki semiparametrik regresyon modelidir. Bu çalışmada kısmi doğrusal modelinde çoklu bağlantı olması durumunda kullanılabilecek Liu-tipi tahmin ediciler elde edilecektir.

KDM'ler yanıt değişkenin açıklayıcı değişkenlerin bazıları ile doğrusal diğerleri ile parametrik olmayan bir şekilde ilişkili olduğu modellerdir. Bu anlamda KDM'ler klasik regresyon modellerinin genelleştirilmiş halidir.

Basit çoklu kısmi doğrusal modelini düşünelim.

$$y_i = z_i^T \gamma + f(t_i) + \varepsilon_i, \quad i = 1, 2, \dots, n \quad (6.1)$$

Bu modele basit semiparametrik regresyon modeli denilmesinin nedeni yanıt değişkeni ile doğrusal ilişkide olmayan tek bir değişkenin olduğunun varsayılmasıdır. Burada y_i 'ler yanıt gözlemleri, $z_i = (z_{i1}, z_{i2}, \dots, z_{ip})^T$ 'ler $p \times 1$ ($p \leq n$) tipinde gözlenmiş açıklayıcı değişkenler, $\gamma = (\gamma_1, \gamma_2, \dots, \gamma_p)^T$, $p \times 1$ tipinde bilinmeyen parametreler vektörü, t_i 'ler y_i 'ler ile doğrusal bir ilişki içinde olmayan değişkenlerdir, $f(\cdot)$ bilinmeyen düzgün bir fonksiyon, ve ε_i 'ler ise birbirinden bağımsız sıfır ortalamalı, σ^2 varyanslı rasgele hatalardır. Engle ve ark. (1986) tarafından ortaya atılan bu model kısmi doğrusal model olarak da adlandırılır.

Eş. 6.1'deki semiparametrik regresyon modelinin matris biçiminde ifadesi

$$y = Z\gamma + f + \varepsilon, \quad \varepsilon \sim (0, \sigma^2 I) \quad (6.2)$$

şeklinde gösterilir. Burada $y = (y_1, y_2, \dots, y_n)^T$, $Z = (z_1, z_2, \dots, z_n)^T$, $f = (f(t_1), f(t_2), \dots, f(t_n))^T$, $\varepsilon = (\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n)^T$ olmak üzere açıklayıcı değişkenler parametrik doğrusal kısım ($Z\gamma$) ve parametrik olmayan kısım ($f(t)$) olmak üzere iki kısım ile gösterilir.

Burada amaç parametre vektörü γ 'yı ve f fonksiyonunu etkin bir şekilde tahmin etmektir. Bu model son zamanlarda oldukça sık çalışılan bir modeldir [Green and Silverman, 1994; Eubank, 1999; Härdle ve ark., 2000; Härdle ve ark., 2004]. Bu modelin bu kadar sıklıkla çalışılmasının nedenlerinden biri kısmi doğrusal modelinin parametrik ve parametrik olmayan öğeleri birleştirdiğinden klasik regresyon modeline göre daha esnek olmasıdır, diğer bir sebebi de tamamen parametrik olmayan bir modele göre daha kolay yorumlanabilmesidir.

Engle ve ark. (1986)'nın çalışmasından esinlenerek Eş. 6.2 modelinin parametrelerini tahmin etmek amacıyla birçok çalışma yapılmıştır. Örnek olarak Rice (1986), Speckman (1988), Eubank ve ark. (1988), Schimek (2000) and Ruppert ve ark. (2003) kısmi doğrusal modellerle ilgili elde edilen sonuçların sistematik olarak bir özetini vermişlerdir.

Bölüm 1'de bahsedilen çoklu bağlantı problemi ile Eş. 6.2 modelinde de karşılaşılır. Bu bölümde semiparametrik regresyon modelinde $Z^T Z$ matrisinin kötü koşullu olması durumunda yanlış tahmin etme yöntemleri üzerinde durulacaktır. 1. bölümde olduğu gibi tahmin edicilerin özelliklerini daha kolay incelemek amacıyla Eş. 6.2 modeli kanonik formda yazılırsa

$$y = X\beta + f + \varepsilon \quad (6.3)$$

modeli elde edilir. Eş. 6.3 modelinde $X = ZT$ ve $\beta = T^T \gamma$, ve T kolonları $Z^T Z$ matrisinin özvektörlerinden oluşan ortogonal matristir. Bu durumda $X^T X = T^T Z^T Z T = \Lambda$ olur. Burada $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_p)$ dır ve $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p$ ler $Z^T Z$ matrisinin sıralı özdeğerleridir.

6.1. Semiparametrik Regresyon Modelinde Cezalı EKK ve Speckman Tahmini

Bu bölümde Eş. 6.3'deki kısmi doğrusal regresyon modelinin β ve f 'i tahmin etmek için cezalı en küçük kareler (CEKK) yöntemi kullanılacaktır. Eş. 6.3 modelinde β ve f

$$F_1(f, \beta) = (y - X\beta - f)^T (y - X\beta - f) + \alpha f^T K f, \quad (6.4)$$

fonksiyonunun minimize edilmesi ile elde edilir. Burada K simetrik, negatif olmayan tanımlı matris, $\alpha K = S^{-1} - I$ ve $S = (I + \alpha K)^{-1}$, S , α düzeltme parametresine

bağlı düzeltme matrisidir. Bu minimizasyon probleminin çözülmesi β ve f 'nin aynı anda tahmin edilmesini sağlar. Bu optimizasyon probleminin çözümü

$$\begin{pmatrix} X^T X & X^T \\ X & (I + \alpha K) \end{pmatrix} \begin{pmatrix} \beta \\ f \end{pmatrix} = \begin{pmatrix} X^T \\ I \end{pmatrix} y \quad (6.5)$$

eşitliğini sağlar [Green ve Silverman, 1994]. Eş. 6.3 modelinin cezalı en küçük kareler yöntemi ile tahmini

$$H_p = S + (I - S)X \left[X^T (I - S)X \right]^{-1} X^T (I - S)$$

olmak üzere

$$\begin{aligned} \hat{\beta}_p &= \left[X^T (I - S)X \right]^{-1} X^T (I - S)y \\ \hat{f}_p &= S(y - Z\hat{\beta}_p) \end{aligned} \quad (6.6)$$

ve

$$\mu_p = Z\hat{\beta}_p + \hat{f}_p = H_p y$$

bulunur. Burada $X^T (I - S)X$ 'nin tersi alınabilir olduğu varsayımı altında X , $n \times p$ tipinde gözlenmiş değerlerin matrisi olup Eş. 6.3 denkleminin parametrik kısmını oluşturur [Buja ve ark. 1989].

$y^o = (I - S)^{1/2} y$ ve $X^o = (I - S)^{1/2} X$ olarak alınırsa $\hat{\beta}_p$ tahmin edicisi EKK tahmin edicisi ile aynı biçimde yazılabilir.

$$\hat{\beta}_p = [X^{oT} X^o]^{-1} X^{oT} y^o \quad (6.7)$$

Speckman (1988)

$$\|(I - S)(y - X\beta)\|^2 \quad (6.8)$$

eşitliğinin minimize edilmesiyle elde edilen alternatif bir semiparametrik regresyon modeli önermiştir. Burada S herhangi bir düzeltme matrisidir. Bu durumda β 'nın Speckman tahmin edicisi EKK ile aynı biçimde yazılabilir.

$$H_s = S + (I - S)X \left[X^T (I - S)^T (I - S)X \right]^{-1} X^T (I - S)^T (I - S)$$

olmak üzere

$$\tilde{\beta}_s = \left[X^T (I - S)^T (I - S)X \right]^{-1} X^T (I - S)^T (I - S)y \quad (6.9)$$

$$\tilde{f}_s = S(y - X\tilde{\beta}_s)$$

ve

$$\mu_s = X\tilde{\beta}_s + \tilde{f}_s = H_s y$$

$\tilde{y} = (I - S)y$ ve $\tilde{X} = (I - S)X$ olarak alınırsa β 'nın Speckman tahmin edicisi EKK ile aynı biçimde yazılabilir.

$$\tilde{\beta}_s = [\tilde{X}^T \tilde{X}]^{-1} \tilde{X}^T \tilde{y}$$

Speckman yaklaşımı Buja ve ark. (1989), Eubank ve ark. (1998) ve Schimek (2000) tarafından incelenmiştir. Tahmin ediciler kısmi hataların regresyonu ile elde edilmiştir. Semiparametrik regresyon modelinin parametrik kısmının $\hat{\beta}_p$ ve $\tilde{\beta}_s$ tahminleri arasında küçük ama önemli bir fark olduğuna dikkat edilmelidir. S matrisi simetrik ve idempotent ise $\hat{\beta}_p$ ve $\tilde{\beta}_s$ tahminleri aynıdır.

6.2. Semiparametrik Regresyon Modelinde Liu-tipi Tahmin Edici

Birinci bölümde çoklu bağlantı ve sonuçları üzerinde durulmuştu. Bu bölümde kısmi doğrusal regresyon modelinde yeni bir Liu tahmin edicisi elde edilecek ve parametrik kısımdaki açıklayıcı değişkenler arasında çoklu bağlantı olması durumunda önerilen Liu tahmin edicisinin performansı incelenecektir.

Semiparametrik regresyon modelindeki β katsayılarının CEKK tahmin edicileri

$$(\tilde{y} - \tilde{X}\beta)^T (\tilde{y} - \tilde{X}\beta) = \|(I - S)(y - X\beta)\|^2. \quad (6.10)$$

fonksiyonunun minimize edilmesiyle bulunur. Eş. 6.10'daki fonksiyona regresyon katsayıları için $\|d\tilde{\beta} - \beta\|^2$ şeklinde bir ceza fonksiyonu eklenmesiyle

$$F(\beta) = \min_{\beta} \left\{ (\tilde{y} - \tilde{X}\beta)^T (\tilde{y} - \tilde{X}\beta) + (d\tilde{\beta}_s - \beta)^T (d\tilde{\beta}_s - \beta) \right\}. \quad (6.11)$$

koşullu amaç fonksiyonu elde edilir. $\tilde{y} = \tilde{X}\beta + \varepsilon^0$ eşitliğine $d\tilde{\beta}_s = \beta + \varepsilon^0$ eşitliğinin eklenmesiyle $\tilde{\beta}_s(d)$ elde edilir.

$$\begin{pmatrix} \tilde{y} \\ d\tilde{\beta}_s \end{pmatrix} = \begin{pmatrix} \tilde{X} \\ I \end{pmatrix} \beta + \begin{pmatrix} \tilde{\varepsilon} \\ \varepsilon^0 \end{pmatrix}.$$

$\varepsilon^* = \begin{pmatrix} \tilde{\varepsilon} \\ \varepsilon^0 \end{pmatrix}$ şeklinde tanımlansın, buradan

$$\varepsilon^{*T} \varepsilon^* = (\tilde{y} - \tilde{X}\beta)^T (\tilde{y} - \tilde{X}\beta) + (d\tilde{\beta}_s - \beta)^T (d\tilde{\beta}_s - \beta).$$

Eş. 6.11 fonksiyonunu minimize eden β vektörü

$$\frac{\partial F}{\partial \beta} = 0. \quad (6.12)$$

ile bulunur. Eş. 6.11'in β vektörüne göre türevi alınırsa

$$-2\tilde{X}^T \tilde{y} + 2\tilde{X}^T \tilde{X} \beta - 2d\tilde{\beta}_s + 2\beta = 0$$

veya

$$(\tilde{X}^T \tilde{X} + I)\beta = \tilde{X}^T \tilde{y} + d\tilde{\beta}_s. \quad (6.13)$$

elde edilir. Bu denklem sisteminin çözümü

$$\tilde{\beta}_s(d) = (\tilde{X}^T \tilde{X} + I)^{-1}(\tilde{X}^T \tilde{y} + d\tilde{\beta}_s), \quad 0 < d < 1 \quad (6.14)$$

olarak elde edilir; burada d yanlılık parametresidir [Liu, 1993; Akdeniz ve Kaçıranlar, 1995; Akdeniz ve ark., 1999]. Eş. 6.14 tahmin edicisi semiparametrik regresyon modeli için Liu-tipi tahmin edicidir. Eş. 6.14

$$f = S(y - X\beta) \quad (6.15)$$

eşitliği ile birleştirilirse f fonksiyonun Liu-tipi tahmin edicisi

$$\tilde{f}_s(d) = S(y - X\tilde{\beta}_s(d)) \quad (6.16)$$

olur.

Şimdi cezalı EKK ölçütüne göre Liu-tipi tahmin edicinin nasıl elde edildiğini açıklayalım. Semiparametrik durum için elde edilmesi doğrudan olur. EKK perspektifinden bakıldığında

$$F_2(f, \beta) = (y - X\beta - f)^T (y - X\beta - f) + (d\hat{\beta} - \beta)^T (d\hat{\beta} - \beta) + \alpha f^T K f. \quad (6.17)$$

eşitliğini minimum yapan β ve f katsayıları elde edilir. F_2 fonksiyonunun β ve f 'e göre türevleri alınır

$$\begin{aligned} \frac{1}{2} \frac{\partial F_2}{\partial \beta} &= -X^T y + (X^T X + I)\beta + X^T f - d\hat{\beta} = 0 \\ \frac{1}{2} \frac{\partial F_2}{\partial f} &= -(y - X\beta) + (I + \alpha K)f = 0 \end{aligned} \quad (6.18)$$

eşitlikleri elde edilir. Bu minimizasyon probleminin çözülmesi ile Eş. 6.3 modeli için $\hat{\beta}_p(d) = [X^T(I - S)X + I]^{-1} [X^T(I - S)y + d\hat{\beta}]$ Liu-tipi tahmin edicisi elde edilir. Bu optimizasyon probleminin çözümü

$$\begin{pmatrix} X^T X + I & X^T \\ X & (I + \alpha K) \end{pmatrix} \begin{pmatrix} \hat{\beta} \\ \hat{f} \end{pmatrix} = \begin{pmatrix} X^T y + d\hat{\beta} \\ y \end{pmatrix}. \quad (6.19)$$

eşitliğini sağlar [Akdeniz ve Akdeniz Duran, 2010].

4. bölümde klasik doğrusal regresyon modeli için $R\beta = r$ şeklindeki gibi eşitlik kısıtlaması şeklindeki parametrelerle ilgili önsel bilgileri kullanarak eşitlik kısıtı altında parametre tahminleri elde edilmişti. Devam eden bölümlerde de parametreler üzerinde

$$R\beta = r \quad (6.20)$$

şeklinde bir önsel bilgi olduğunda semiparametrik regresyon modelleri için tahminler elde edilecektir. Burada R bilinen ve rankı $m < p$ olan $m \times p$ tipinde bir matris r ise $m \times 1$ tipinde bilinen bir vektördür. $\beta^*(y)$ tahmin edicisi her $n \times 1$ tipinde y vektörü için $R\beta^*(y) = r$ kısıtını sağlarsa $\beta^*(y)$ 'ye β 'nın kısıtlı tahmin edicisi denir.

Przystalski and Krajewski (2007) bir tane düzeltme terimi olan semiparametrik regresyon modeli için bazı ek kısıtlar altında deney etkilerini tahmin etmek için bir metod önermiştir ve hipotez testlerinde uygulamasını anlatmıştır. Ahmet ve ark. (2007) Eş. 6.2 modelinin doğrusal kısmının β katsayı vektörünün $\beta = (\beta_1^T, \beta_2^T)$ şeklinde parçalandığı modelleri ele aldılar. Bu durumda araştırmacılar $R = (0, I)$ için $R\beta = 0$ kısıtı altında β_1 'in tahmini ile ilgilenirler.

6.3. Semiparametrik Regresyon Modelinde Eşitlik Kısıtlı Cezalı EKK ve Speckman Tahmin Edicileri

$R\beta = r$ kısıtının sağlandığı varsayımı altında takip eden iki teoremden Eş.6.3'teki semiparametrik regresyon modeli için β 'nin kısıtlı tahmin edicisi verilmiştir.

6.1. Teorem

Simetrik bir S düzeltme matrisi için Eş. 6.3 modelindeki β 'nin cezalı en küçük kareler yöntemi ile Eş. 6.20 kısıtı altında tahmini

$$D = X^T(I - S)X \text{ ve } \hat{\beta} = D^{-1}X^T(I - S)y \text{ için}$$

$$\hat{b}_p = \hat{\beta}_p + D^{-1}R^T(RD^{-1}R^T)^{-1}(r - R\hat{\beta}_p) \quad (6.21)$$

bulunur.

İspat

[Akdeniz ve Akdeniz Duran, 2010] $\hat{b}_p$ 'nin $R\hat{b}_p = r$ kısıtını sağladığı kolayca görülebilir. Eş. 6.18'de soldan R ile çarpım yapılırsa kısıtın sağlandığı görülmektedir.

Sonraki teoremden semiparametrik regresyon modelinde herhangi bir düzeltme matrisi S için $R\beta = r$ kısıtı altında β 'nin kısıtlı tahmin edicisi verilecektir.

6.2. Teorem

Herhangi bir S düzeltme matrisi için Eş. 6.20 kısıtı altında Eş. 6.3 modelindeki β 'nin Eş. 6.8 ölçütüne göre tahmini

$$\tilde{b}_s = \tilde{\beta}_s + M^{-1}R^T(RM^{-1}R^T)^{-1}(r - R\tilde{\beta}_s) \quad (6.22)$$

dir. Burada $M = X^T(I - S)^T(I - S)X$ ve $\tilde{\beta} = (\tilde{X}^T \tilde{X})^{-1} \tilde{X}^T \tilde{y}$ 'dir.

İspat

[Akdeniz ve Akdeniz Duran, 2010] Eş. 6.19'in soldan R ile çarpılmasıyla $\tilde{b}_s$ 'nin $R\tilde{b}_s = r$ kısıtını sağladığı açıktır.

6.4. Semiparametrik Regresyon Modelinde Eşitlik Kısıtlı Liu Tahmin Edicisi

Bu bölümde $R\beta = r$ kısıtı altında Eş. 6.3 modeli için kısıtlı Liu-tipi tahmin ediciyi bulmak istiyoruz.

6.3. Teorem

Herhangi bir S düzeltme matrisi için Eş. 6.3 modeli için Eş. 6.17 kısıtı altında Eş. 6.8 ölçütüne göre β 'nin tahmini

$$\tilde{b}_s(d) = \tilde{\beta}_s(d) + \tilde{Z}^{-1}R^T(R\tilde{Z}^{-1}R^T)^{-1}(r - R\tilde{\beta}_s(d)). \quad (6.23)$$

dır. Burada

$$\tilde{\beta}_s(d) = (\tilde{X}^T \tilde{X} + I)^{-1} (\tilde{X}^T \tilde{y} + d \tilde{\beta}_s); \quad \tilde{Z} = \tilde{X}^T \tilde{X} + I = X^T (I - S)^T (I - S) X + I$$

ve

$$\tilde{\beta}_s = (\tilde{X}^T \tilde{X})^{-1} \tilde{X}^T \tilde{y} = \left[X^T (I - S)^T (I - S) X \right]^{-1} X^T (I - S)^T (I - S) y$$

dir.

İspat

[Akdeniz ve Akdeniz Duran, 2010] $\tilde{b}_s(d)$ tahmin edicisinin $R\tilde{b}_s(d) = r$ kısıtını sağladığı görülebilir. Bu nedenle $\tilde{b}_s(d)$, $y = X\beta + f + \varepsilon$ modelindeki β 'nin kısıtlı Liu-tipi tahmin edicisidir.

Şimdi S düzeltme matrisinin simetrik olması durumunda $y = X\beta + f + \varepsilon$ semiparametrik regresyon modelinde $R\beta = r$ kısıtı altında cezalı en küçük kareler ölçütüne göre Liu-tipi tahmin ediciyi vereceğiz.

6.4. Teorem

Simetrik bir S düzeltme matrisi için Eş. 6.3'te β 'nin cezalı en küçük kareler ölçütüne göre tahmini

$$\|y - X\beta - f\|^2 + \alpha f^T K f + (d\hat{\beta} - \beta)^T (d\hat{\beta} - \beta) \quad (6.24)$$

$R\beta = r$ kısıtının sağlandığı varsayımı altında

$$\hat{b}_p(d) = \hat{\beta}_p(d) + V^{-1} R^T (R V^{-1} R^T)^{-1} (r - R \hat{\beta}_p(d)). \quad (6.25)$$

olarak elde edilir. Burada $\hat{\beta}_p = [X^T(I-S)X]^{-1} X^T(I-S)y$, $V = X^T(I-S)X + I$,
ve $\hat{\beta}_p(d) = V^{-1} [X^T(I-S)y + d\hat{\beta}_p]$ dir.

İspat

[Akdeniz ve Akdeniz Duran, 2010] Verilen $\hat{b}_p(d)$ tahmin edicisi $R\hat{b}_p(d) = r$ eşitliğini sağladığından $\hat{b}_p(d)$, Eş. 6.3'deki semiparametrik regresyon modeli için CEKK ölçütüne göre kısıtlı Liu-tipi tahmin edicidir.

7. KISMİ DOĞRUSAL MODELLERDE FARKA DAYALI TAHMİN EDİCİLER

Kısmi doğrusal modeller, istatistik ve ekonometride son zamanlarda çok çalışılan konular arasındadır. Bu modellerin biyomedikal ve ekonomi alanlarında birçok uygulaması vardır. Daha önce bahsedildiği gibi kısmi doğrusal modellerde bazı değişkenlerin yanıt değişkeni ile ilişkisi parametrik iken bazılarının parametrik değildir. Bu durumda esnek bir model olan kısmi doğrusal modellerdeki parametreleri tahmin etmek amacıyla ilk defa Yatchew (1997) tarafından önerilen parametrik olmayan terimleri yok eden fark yöntemi kullanılabilir. Fark yönteminin avantajı parametrik olmayan terimlerin etkisi gözönüne alınarak fakat onlar modelde yokmuş gibi doğrusal regresyon yöntemleriyle parametrelerin tahminlerinin yapılabilmesidir. Eş. 6.1’de verilen modelin kanonik biçimde yazılması ile elde edilen kısmi doğrusal model

$$y_i = x_i^T \beta + f(t_i) + \varepsilon_i, \quad i = 1, 2, \dots, n \quad (7.1)$$

şeklinde verilmiş olsun. Burada y_i ’ler yanıt değişkeniler, $x_i = (x_{i1}, x_{i2}, \dots, x_{ip})$ biçimindeki $x_1, x_2, \dots, x_n$ ’ler $p \leq n$ olacak şekilde p boyutlu vektörler, t_i ’ler ise parametrik olmayan değişkene ait gözlemler, $\beta = (\beta_1, \beta_2, \dots, \beta_p)^T$, p boyutlu bilinmeyen parametre vektörü, $f(\cdot)$ bilinmeyen sınırlı ve düzgün bir fonksiyon ve ε_i ’ler b.a.d. 0 ortalamalı, σ^2 varyanslı rasgele hata terimleridir. Burada $f(t)$ ile bağımlı gözlem arasında parametrik olmayan fonksiyonel bir ilişki vardır ve amaç β parametre vektörü ve parametrik olmayan $f(t)$ fonksiyonunu tahmin etmektir. Eş. 7.1’deki model vektör/matris biçiminde daha önce Eş. 6.3’te verildiği gibi

$$y = X\beta + f + \varepsilon, \quad (7.2)$$

şeklinde yazılabilir. Burada $y = (y_1, \dots, y_n)^T$, $X = [x_1, \dots, x_n]^T$, $f = (f(t_1), \dots, f(t_n))^T$ ve $\varepsilon = (\varepsilon_1, \dots, \varepsilon_n)^T$ 'dir. Bir sonraki bölümde parametrik olmayan değişkeni sıfıra yakınsatmak için kullanılacak fark matrisi anlatılacaktır.

7.1. Model ve Fark Matrisi

Bu bölümde semiparametrik regresyon modelindeki doğrusal değişkenlerin katsayılarını (β vektörü) tahmin etmek için farka dayalı yöntem anlatılacaktır. Bu yöntem semiparametrik modeldeki parametrik olmayan değişkeni yok etmek amacıyla birçok araştırmacı tarafından kullanılmıştır [Yatchew 1997, 2003; Klipple ve Eubank 2007; Brown ve Levine 2007]. Eş. 7.2'deki kısmi doğrusal regresyon modelini düşünelim. $d = (d_0, d_1, \dots, d_m)^T$, $(m+1)$ boyutlu bir vektör olsun ve m ise farkın derecesini göstere. Burada $d_0, d_1, \dots, d_m$ 'ler

$$\sum_{j=0}^m d_j = 0 \text{ ve } \sum_{j=0}^m d_j^2 = 1. \quad (7.3)$$

kısıtlarını sağlayan ve $\min_{d_0, \dots, d_m} \delta = \sum_{k=1}^m (\sum_{j=1}^{m-k} d_j d_{k+j})^2$ eşitliğini minimize eden fark ağırlıklarıdır. Eş. 7.2'ye fark matrisi uygulandıktan sonra

$$Dy = DX\beta + Df + D\varepsilon$$

elde edilir. Dy 'nin bir elemanı

$$d_0 y_i + \dots + d_m y_{i+m} = [d_0 x_i + \dots + d_m x_{i+m}] \beta + d_0 f(t_i) + \dots + d_m f(t_{i+m}) + d_0 \varepsilon_i + \dots + d_m \varepsilon_{i+m}$$

şeklinde olduğundan Eş. 7.3'teki kısıtların amacı bellidir. Birinci kısıt t değerleri birbirine yakın olduğunda parametrik olmayan değişkenin yok olmasını sağlar.

İkinci kısıt ise ağırlıklandırılmış hata terimlerinin varyansının σ^2 olarak kalmasını sağlar.

Optimum fark ağırlıkları ve bu ağırlıkların elde edilmesi sırasıyla Yatchew (2003) kitabında Bölüm 4, Tablo 4.1’de ve Ek C’de verilmiştir. Bu ağırlıklar 10. dereceye kadar aşağıdaki gibidir. Diğer dereceler için ağırlıklar Yatchew (2003)’te bulunabilir.

Çizelge 7.1. Optimum fark ağırlıkları

m	$(d_0, d_1, \dots, d_m)$
1	(0,7071; -0,7071)
2	(0,8090; -0,5000; -0,3090)
3	(0,8582; -0,3832; -0,2809; -0,1942)
4	(0,8873; -0,3099; -0,2464; -0,1901; -0,1409)
5	(0,9064; -0,2600; -0,2167; -0,1774; -0,1420; -0,1103)
6	(0,9200; -0,2238; -0,1925; -0,1635; -0,1369; -0,1126; -0,0906)
7	(0,9302; -0,1965; -0,1728; -0,1506; -0,1299; -0,1107; -0,0930; -0,0768)
8	(0,9380; -0,1751; -0,1565; -0,1389; -0,1224; -0,1069; -0,0925; -0,0791; -0,0666)
9	(0,9443; -0,1578; -0,1429; -0,1287; -0,1152; -0,1025; -0,0905; -0,0792; -0,0687; -0,0588)
10	(0,9494; -0,1437; -0,1314; -0,1197; -0,1085; -0,0978; -0,0877; -0,0782; -0,0691; -0,0606; -0,527)

D , $(n-m) \times n$ boyutlu fark matrisi olsun. 0_r tüm elemanları 0 olan r boyutlu bir vektör olmak üzere, D matrisinin birinci ve ikinci satırları sırasıyla $[d^T, 0_{n-m-1}^T]$, $[0_{n-m-1}^T, d^T]$ ve i . satırı $[0_i, d^T, 0_{n-m-i-1}^T]$, $i = 2, \dots, (n-m-1)$ ’dir. Buna göre elde edilen D matrisi aşağıda verilmiştir.

$$D = \begin{bmatrix} d_0 & d_1 & d_2 & . & . & . & d_m & 0 & 0 & . & . & . & 0 \\ 0 & d_0 & d_1 & d_2 & . & . & . & d_m & 0 & 0 & . & . & 0 \\ . & . & . & & & & & & & & & & \\ . & . & . & & & & & & & & & & \\ . & . & . & & & & & & & & & & \\ 0 & 0 & 0 & . & . & . & & d_0 & d_1 & . & . & d_m & 0 \\ 0 & 0 & 0 & . & . & . & & d_0 & d_1 & . & . & d_m \end{bmatrix}.$$

Eş. 7.2 modeline fark matrisinin uygulanması parametrik etkinin direk olarak tahmin edilmesini sağlar, Speckman (1988)'daki gelişmelerin bir sonucu olarak, Eş. 7.2 modelindeki β parametre vektörü etkin olarak tahmin edilebilir. Bu amaçla farka dayalı tahmin ediciler kullanıldığında parametrik olmayan değişken gözardı edilebilir duruma gelir. f birinci türevi sınırlı olan bilinmeyen bir fonksiyon olsun. Bu durumda Df , yaklaşık olarak 0'a eşit olur ve Eş. 7.2 kısmi doğrusal modeline fark matrisi uygulandığında Df gözardı edilebilir olduğundan model

$$Dy \approx DX\beta + D\varepsilon. \quad (7.4)$$

şeklinde yazılabilir. d_i 'ler üzerindeki kısıtların nedeni burada daha iyi anlaşılmaktadır [Yatchew, 2003: s.57; Klipple ve Eubank, 2007]. Parametrik olmayan değişkenlerin etkisi çıkarıldığından, model klasik doğrusal regresyon modeli gibi ele alınabilir. Bu durumda klasik doğrusal regresyon modelinde β 'nın tahmini için kullanılan EKK yöntemi ile parametre vektörünün tahmini

$$\hat{\beta}_{diff} = [(DX)^T (DX)]^{-1} (DX)^T Dy. \quad (7.5)$$

şeklinde yapılır. Bu durumda,

$$S_{diff}^2 = \frac{1}{n} (Dy - DX \hat{\beta}_{diff})^T (Dy - DX \hat{\beta}_{diff}), \quad (7.6)$$

$$= \frac{1}{n} (y - X \hat{\beta}_{diff})^T D^T D (y - X \hat{\beta}_{diff}).$$

dır. $\beta \neq 0$ durumunda serbestlik derecesine düzeltme yaparak hata teriminin varyansını

$$\hat{\sigma}^2 = \frac{y^T D^T [I - DX (X^T D^T DX)^{-1} X^T D^T] Dy}{\text{iz}[D^T (I - P) D]} = \frac{y^T D^T (I - P) Dy}{\text{iz}[D^T (I - P) D]} \quad (7.7)$$

ile tahmin edebiliriz. Burada $\text{iz}(\cdot)$ kare matrisin iz'ini ve P ise *izdüşüm matrisini* göstermektedir.

$$P = DX (X^T D^T DX)^{-1} X^T D^T,$$

[Eubank ve ark., 1988]. Farka dayalı bu tahmin edici ilk defa Yatchew (1997) tarafından önerilmiştir. Burada görüldüğü gibi, fark almak modelin yaklaşık olarak doğrusal regresyon modeline indirgenmesini ve modelde parametrik olmayan f değişkeni yokmuş gibi tahmin yapılmasına olanak sağlamıştır [Yatchew, 2003; Fan ve Wu, 2008]. Farka dayalı tahmin yapmak modelin parametrik kısmındaki parametrelerin tahminini modelde parametrik olmayan değişken yokmuş gibi yapmamızı sağlar. Ayrıca parametrik tahmin yapıldıktan sonra modelin parametrik kısmının bilindiği varsayılarak ve aşağıdaki eşitlikteki gibi y_i 'lerden çıkarılarak f fonksiyonun analizi herhangi bir parametrik olmayan teknikle yapılabilir.

$$y_i - x_i \hat{\beta}_{diff} = x_i (\beta - \hat{\beta}_{diff}) + f(t_i) + \varepsilon_i \equiv f(t_i) + \varepsilon_i$$

f fonksiyonuna kernel tahmini gibi klasik metotlar uygulandığında tutarlılık, optimum yakınsama hızı ve güven aralıkları özelliklerini korumasının nedeni $\hat{\beta}_{diff}$ 'in β 'ya yeteri kadar hızlı yakınsamasıdır [Yatchew, 2003].

7.2. Farka-dayalı Ridge Regresyon ve Liu-tipi Tahmin Ediciler

Bu bölümde kısmi doğrusal regresyon modellerindeki çoklu bağlantı problemini çözmek için ridge regresyon yöntemi ile farka dayalı tahmin ediciler verilecektir.

Tabakan and Akdeniz (2010), $\hat{\beta}_{diff} = [(DX)^T (DX)]^{-1} (DX)^T Dy$ tahmin edicisine alternatif olarak

$$\begin{aligned}\hat{\beta}_{diff}(k) &= [(DX)^T (DX) + kI]^{-1} (DX)^T Dy, \quad (k \geq 0) \\ &= (\tilde{X}^T \tilde{X} + kI)^{-1} \tilde{X}^T \tilde{y}.\end{aligned}\tag{7.8}$$

tahmin edicisini önermişlerdir. Burada $\tilde{y} = Dy$, $\tilde{X} = DX$; I , $p \times p$ tipinde birim matris ve k araştırmacı tarafından seçilen ridge yanlılık parametresidir. $\hat{\beta}_{diff}(k)$ semiparametrik regresyon modeli için farka dayalı ridge tahmin edicisidir. Geleneksel en küçük kareler tahmin edicisi ile aynı mantık doğrultusunda β katsayıları

$$(Dy - DX \beta)^T (Dy - DX \beta).\tag{7.9}$$

eşitliğini minimize edecek şekilde elde edilir. Eş. 7.9'daki amaç fonksiyonuna regresyon katsayıları için $\|l\hat{\beta}_{diff} - \beta\|^2$ karesel formu şeklinde bir ceza fonksiyonu eklenirse

$$L = \arg \min_{\beta} [(\tilde{y} - \tilde{X} \beta)^T (\tilde{y} - \tilde{X} \beta) + (l\hat{\beta}_{diff} - \beta)^T (l\hat{\beta}_{diff} - \beta)]\tag{7.10}$$

koşullu amaç fonksiyonu elde edilir. Lagrange çarpanlar yöntemi ile amaç fonksiyonunu β 'ya göre minimize ederek $\hat{\beta}_{diff}$ tahmin edicisine alternatif olarak

$$\hat{\beta}_{diff}(l) = (\tilde{X}^T \tilde{X} + I)^{-1}(\tilde{X}^T \tilde{y} + l\hat{\beta}_{diff}), \quad (7.11)$$

$\hat{\beta}_{diff}(l)$ farka dayalı Liu-tipi tahmin edici elde edilir. Burada l , $0 < l < 1$, yanlılık parametresidir ve $l=1$ için $\hat{\beta}_{diff}(l=1) = \hat{\beta}_{diff}$ eşitliği sağlanır. Eş. 7.11’de verilen tahmin edici ile Liu tahmin edicisinin [Liu, 1993; Akdeniz ve Kaçıranlar, 1995; Akdeniz ve ark., 2006; Hubert ve Wijekoon, 2006; Yang ve Xu, 2009] benzerliğinden dolayı bu tahmin ediciye farka dayalı Liu-tipi tahmin edici denilmiştir.

7.3. Farka-dayalı Liu-tipi Tahmin Edicisinin HKOM Ölçütüne Göre Farka Dayalı Tahmin Edici ile Karşılaştırılması

Bu bölümde $\hat{\beta}_{diff}(l)$ ve $\hat{\beta}_{diff}$ tahmin edicilerinin HKOM’leri arasındaki fark incelenecektir. Daha önceki bölümlerde bahsedildiği gibi β ’nın herhangi bir tahmin edicisi $\tilde{\beta}$ ’nın HKOM’si $HKOM(\tilde{\beta}) = V(\tilde{\beta}) + b(\tilde{\beta})b(\tilde{\beta})^T$ ’dir. Burada $V(\tilde{\beta})$ varyans-kovaryans matrisini ve $b(\tilde{\beta}) = E(\tilde{\beta}) - \beta$ ise yanlılık terimini göstermektedir. Farka dayalı Liu-tipi tahmin edici $\hat{\beta}_{diff}(l)$ ’nin beklenen değeri

$$E(\hat{\beta}_{diff}(l)) = \beta - (1-l)(\tilde{X}^T \tilde{X} + I)^{-1} \beta \quad (7.12)$$

olarak bulunur. Bu durumda $\hat{\beta}_{diff}(l)$ ’nin yanlılık terimi

$$b(\hat{\beta}_{diff}(l)) = -(1-l)(\tilde{X}^T \tilde{X} + I)^{-1} \beta. \quad (7.13)$$

olarak elde edilir. $F_l = (\tilde{X}^T \tilde{X} + I)^{-1}(\tilde{X}^T \tilde{X} + lI)$ alınırsa F_l ve $(\tilde{X}^T \tilde{X})^{-1}$ matrislerinin çarpımı çarpım sırasına göre değişmez olduğundan $\hat{\beta}_{diff}(l)$ tahmin edicisi

$$\hat{\beta}_{diff}(l) = F_l \hat{\beta}_{diff}(l=1) = F_l (\tilde{X}^T \tilde{X})^{-1} \tilde{X}^T \tilde{y} = (\tilde{X}^T \tilde{X})^{-1} F_l \tilde{X}^T \tilde{y}. \quad (7.14)$$

olarak yazılabilir. $S = (D^T \tilde{X})^T (D^T \tilde{X})$ ve $U = (\tilde{X}^T \tilde{X})^{-1}$ olarak alınsın. Bu durumda $\hat{\beta}_{diff}(l)$ tahmin edicisinin varyans kovaryans matrisi $V(\hat{\beta}_{diff}(l))$

$$V(\hat{\beta}_{diff}(l)) = \sigma^2 U F_l S F_l^T U. \quad (7.15)$$

olarak elde edilir. Aynı şekilde $l=1$ alırsa $\hat{\beta}_{diff}(l=1)$ 'in varyans kovaryans matrisi

$$V(\hat{\beta}_{diff}(l=1)) = \sigma^2 U S U \quad (7.16)$$

olarak elde edilir. Bu durumda tahmin edicilerin varyans kovaryans matrisleri farkı $V(\hat{\beta}_{diff}) - V(\hat{\beta}_{diff}(l))$

$$\Delta = V(\hat{\beta}_{diff}) - V(\hat{\beta}_{diff}(l)) = \sigma^2 U [S - F_l S F_l^T] U. \quad (7.17)$$

olarak elde edilir. Burada S simetrik matris, U ise tekil olmayan simetrik matristir. Böylece Δ farkı

$$\begin{aligned} \Delta &= V(\hat{\beta}_{diff}) - V(\hat{\beta}_{diff}(l)) = \sigma^2 F_l [F_l^{-1} U S U (F_l^T)^{-1} - U S U] F_l^T \\ &= \sigma^2 (1 - l^2) (U^{-1} + I)^{-1} \left[\frac{1}{1+l} (U S + S U) + U S U \right] (U^{-1} + I)^{-1} \end{aligned} \quad (7.18)$$

Burada $t = \frac{1}{1+l} > 0$, $M = U S U$ ve $N = U S + S U$ 'dir. $M = L^T L$ ve

$\text{rank}(L) = p < n - m$, olduğundan M , $p \times p$ tipinde pozitif tanımlı bir matristir.

Burada $N = U S + S U$ simetrik bir matristir. Bu durumda Eş. 7.17

$$\begin{aligned}
\Delta &= \sigma^2(1-l^2)H[M+tN]H \\
&= \sigma^2(1-l^2)H(Q^T)^{-1}[Q^T MQ+tQ^T NQ]Q^{-1}H \\
&= \sigma^2(1-l^2)H(Q^T)^{-1}[I+tE]Q^{-1}H, \tag{7.19}
\end{aligned}$$

şeklinde yazılabilir. Burada $I+tE = \text{diag}(1+te_{11}, \dots, 1+te_{pp})$ ve $H = (U^{-1} + I)^{-1}$ dir. $M > 0$ ve N simetrik bir matris olduğundan, $Q^T MQ = I$ ve $Q^T NQ = E$ koşullarını sağlayan tekil olmayan bir Q matrisi vardır. Burada E köşegen elemanları $|M^{-1}N - \lambda I| = 0$ polinomunun kökleri olan köşegenel bir matristir [Graybill, 1983: 408]. $N = US + SU \neq 0$ olduğundan E matrisinin en az bir tane sıfır olmayan köşegen elemanı vardır. $e_{ii} \neq 0$ olsun,

$$0 < t < \min_{e_{ii} \neq 0} \left| \frac{1}{e_{ii}} \right|$$

ve buradan her $i = 1, \dots, p$ için $1+te_{ii} > 0$ olur ve $I+tE$ pozitif tanımlı bir matristir. Buradan Δ da pozitif tanımlı bir matris olur. Buradan ancak ve ancak $\lambda_{\max}(M^{-1}N) \leq (l+1)$ veya $|\lambda_{\max}(USU)^{-1}(US+SU)| \leq (l+1)$ koşulunun sağlanması durumunda $\hat{\beta}_{diff}(l)$ tahmin edicisinin varyansı $\hat{\beta}_{diff}$ 'in varyansından daha küçük olur. Bir sonraki teoremden $\hat{\beta}_{diff}(l)$ tahmin edicisinin HKOM ölçütüne göre $\hat{\beta}_{diff}$ 'den üstün olması için gerekli ve yeterli koşullar verilecektir.

Teoremin ispatı için Farebrother (1976) tarafından önerilen Lemma 5.2'den yararlanılacaktır. $\hat{\beta}_{diff}(l)$ ve $\hat{\beta}_{diff}$ tahmin edicilerini HKOM ölçütüne göre karşılaştırmak için her iki tahmin edicinin HKOM'lerinin farkı $\tilde{\Delta} = \text{HKOM}(\hat{\beta}_{diff}) - \text{HKOM}(\hat{\beta}_{diff}(l))$ incelenmelidir. $\text{HKOM}(\tilde{\beta}) = V(\tilde{\beta}) + b(\tilde{\beta})b(\tilde{\beta})^T$ olduğundan, tahmin edicilerin HKOM'leri sırasıyla

$$\text{HKOM}(\hat{\beta}_{diff}(l)) = \sigma^2 F_l U S U F_l^T + (1-l)^2 (U^{-1} + I)^{-1} \beta \beta^T (U^{-1} + I)^{-1} \quad (7.20)$$

ve

$$\text{HKOM}(\hat{\beta}_{diff}) = V(\hat{\beta}_{diff}) = \sigma^2 U S U. \quad (7.21)$$

Eş. 7.20 ve Eş. 7.21'den HKOM'ler arasındaki fark

$$\begin{aligned} \tilde{\Delta} &= \text{HKOM}(\hat{\beta}_{diff}) - \text{HKOM}(\hat{\beta}_{diff}(l)), \\ \tilde{\Delta} &= \sigma^2 F_l (F_l^{-1} U S U F_l^{T-1} - U S U) F_l^T - (1-l)^2 (U^{-1} + I)^{-1} \beta \beta^T (U^{-1} + I)^{-1}, \\ &= W [\sigma^2 (1-l^2)(M + tN) - (1-l)^2 \beta \beta^T] W, \\ &= (1-l)^2 W [\sigma^2 (\frac{1+l}{1-l})(M + tN) - \beta \beta^T] W. \end{aligned} \quad (7.22)$$

biçiminde yazılabilir. Farebrother (1976) tarafından önerilen Lemma 5.2'den $\tilde{\Delta}$ ancak ve ancak

$$\beta^T (M + tN)^{-1} \beta \leq \sigma^2 \frac{1+l}{1-l}, \quad 0 < l < 1 \quad (7.23)$$

koşulunun sağlanması durumunda $\tilde{\Delta}$ pozitif tanımlı bir matris olabilir. Bu durumda aşağıdaki teorem verilebilir.

7.1. Teorem

β parametresi için $\hat{\beta}_{diff}(l)$ ve $\hat{\beta}_{diff}$ tahmin edicilerini düşünelim.

$W = (\frac{1+l}{1-l})(M + tN)$ pozitif tanımlı bir matris olsun. Bu durumda ancak ve ancak

$$\beta^T W^{-1} \beta \leq \sigma^2 \quad (7.24)$$

koşulunun sağlanması durumunda $\hat{\beta}_{diff}(l)$ yanlı tahmin edicisi HKOM ölçütüne göre $\hat{\beta}_{diff}$ 'den üstündür.

8. SİMÜLASYON ÇALIŞMALARI VE UYGULAMALAR

8.1. Jackknife Liu-tipi Tahmin Edici ile İlgili Simülasyon Çalışması

Bu kısımda, 5. bölümde önerilen jackknife tahmin edicileri Monte Carlo simülasyon çalışması ile karşılaştırılacaktır. Bu simülasyon çalışmasında çeşitli çoklu bağlantı dereceleri altında tahmin edicilerin SHKO ve ortalama mutlak sapma (OMS) ölçütüne göre performansları incelenecektir. Simülasyon çalışması için Eş. 1.1’de verilen model varsayılacaktır. Açıklayıcı değişkenler ilk defa McDonald and Galarneau (1975) tarafından önerilen

$$x_{ij} = (1 - \rho^2)^{1/2} w_{ij} + \rho w_{i(p+1)}, \quad i = 1, \dots, n; j = 1, \dots, p \quad (8.1)$$

denklemini ile üretilecektir. Burada w_{ij} ’ler bağımsız standart normal dağılıma sahip rasgele değişkenlerdir. ρ , her iki değişken arasındaki teorik korelasyon ρ^2 olacak biçimde seçilmiştir. Bu çalışmada, değişik derecelerdeki çoklu bağlantının tahmin ediciler üzerindeki etkisini görmek için 3 farklı korelasyon katsayısı ele alınmıştır, $\rho^2 = 0,64; 0,90; 0,98$. Bu 3 farklı korelasyon, koşul sayısı zayıf ilişkiden güçlü ilişkiye olacak şekilde seçilmiştir. Veriler $n = 15, 50$ ve 100 olmak üzere 3 farklı örnek çapı seçilerek üretilmiştir. Bu simülasyon çalışmasında ayrıca tahmin edicilerin hata varyansından etkilenip etkilenmediğini görmek için 3 farklı standart hata incelenmiştir $\sigma = 0,1; 1,0; 5,0$.

Yanıt değişkeni y_i , ($i = 1, \dots, n$)’ler aşağıdaki eşitlik kullanılarak üretilmiştir:

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_p x_{ip} + \varepsilon_i, \quad i = 1, \dots, n \quad (8.2)$$

Burada ε_i ’ler ortalaması 0 ve varyansı σ^2 olan, normal dağılıma sahip rasgele hatalardır. Veriler üretildikten sonra $X^T X$ ve $X^T y$ korelasyon matrisi biçiminde olacak şekilde standartlaştırılmıştır. Tahmin ediciler hesaplanırken orijinal model

önce kanonik forma dönüştürülmüş ve γ 'nın tahminleri bulunmuş ve daha sonra bu tahmin ediciler β 'nın tahminini verecek şekilde dönüştürülmüştür. Her bir ρ , σ^2 ve n için deney $B=1000$ kere tekrarlanmış ve SHKO ve ortalama mutlak sapmalar (OMS) aşağıdaki biçimde hesaplanmıştır:

$$\text{SHKO}(\hat{\beta}) = \frac{1}{B} \sum_{l=1}^B \left\| \hat{\beta}^{(l)} - \beta \right\|^2, \quad (8.3)$$

$$\text{OMS}(\hat{\beta}) = \frac{1}{B} \sum_{l=1}^B \sum_{i=1}^p | \hat{\beta}_{il} - \beta_i |, \quad (8.4)$$

Burada $\hat{\beta}^{(l)}$, l . simülasyondaki parametre tahminini ve $\hat{\beta}_{il}$, i . parametrenin l . simülasyondaki parametre tahminini göstermektedir. Her deneme için d değeri aşağıdaki denklem kullanılarak hesaplanmıştır

$$\hat{d}_{\min} = 1 - \hat{\sigma}^2 \left[\sum_{i=1}^p \frac{1}{\lambda_i (1 + \lambda_i)} / \sum_{i=1}^p \frac{\hat{\gamma}_i^2}{(1 + \lambda_i)^2} \right] \quad (8.5)$$

[Liu, 1993]. Simülasyon çalışmasının sonuçları Çizelge 8.1-8.6'de özetlenmiştir. Küçük örnek çapının ($n=15$) alındığı Çizelge 8.1'de, zayıf çoklu bağlantının ($\kappa \leq 10$) ve standart hatanın küçük ($\sigma = 0,1$) olduğu durumda $\hat{\beta}_{YJ}(d)$ 'nin SHKO değeri $\hat{\beta}(d)$ ve $\hat{\beta}_{JL}(d)$ tahmin edicilerinin SHKO değerlerine çok yakın olmasına rağmen standart hata ve koşul sayısı arttıkça aradaki fark $\hat{\beta}_{YJ}(d)$ lehine belirginleşmektedir. Orta büyüklükteki bir örnek çapı seçildiğinde ($n=50$) yine güçlü bağlantı durumunda $\hat{\beta}_{YJ}(d)$ tahmin edicisi diğer tahmin edicilerden üstündür. Büyük örnek çapı durumunda ($n=100$), $\hat{\beta}_{YJ}(d)$ ve $\hat{\beta}(d)$ hemen hemen eşit SHKO değerine sahiptir ve $\hat{\beta}_{JL}(d)$ 'nin SHKO'sundan küçüktür.

Çizelge 8.1. EKK, LT, JLT ve YJLT tahmin edicilerinin SHKO'ları, $n=15$

	$\sigma = 0,1$			$\sigma = 1,0$			$\sigma = 5,0$		
	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$
$\hat{\beta}$	0,022	0,022	0,024	0,782	0,750	0,765	1,132	1,092	1,114
$\hat{\beta}(d)$	0,022	0,016	0,011	0,759	0,684	0,653	1,090	1,012	0,987
$\tilde{\beta}_{JL}(d)$	0,022	0,018	0,011	0,774	0,715	0,672	1,121	1,053	1,017
$\hat{\beta}_{YJ}(d)$	0,022	0,015	0,010	0,756	0,674	0,644	1,085	1,000	0,979
κ	8,33	41,30	226,76	8,33	41,30	226,76	8,33	41,30	226,76

Çizelge 8.2. EKK, LT, JLT ve YJLT tahmin edicilerinin SHKO'ları, $n=50$

	$\sigma = 0,1$			$\sigma = 1,0$			$\sigma = 5,0$		
	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$
$\hat{\beta}$	0,047	0,048	0,047	0,798	0,792	0,787	1,005	0,997	0,996
$\hat{\beta}(d)$	0,048	0,048	0,042	0,798	0,784	0,766	1,000	0,987	0,971
$\tilde{\beta}_{JL}(d)$	0,047	0,047	0,044	0,798	0,790	0,775	1,004	0,994	0,981
$\hat{\beta}_{YJ}(d)$	0,048	0,048	0,041	0,798	0,784	0,763	1,000	0,986	0,967
κ	7,55	32,90	166,97	7,55	32,90	166,97	7,55	32,90	166,97

Çizelge 8.3. EKK, LT, JLT ve YJLT tahmin edicilerinin SHKO'ları, $n=100$

	$\sigma = 0,1$			$\sigma = 1,0$			$\sigma = 5,0$		
	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$
$\hat{\beta}$	0,097	0,097	0,097	0,845	0,844	0,846	0,995	0,995	0,997
$\hat{\beta}(d)$	0,098	0,098	0,095	0,846	0,843	0,838	0,994	0,992	0,988
$\tilde{\beta}_{JL}(d)$	0,097	0,097	0,095	0,845	0,844	0,843	0,995	0,994	0,994
$\hat{\beta}_{YJ}(d)$	0,097	0,098	0,094	0,846	0,842	0,837	0,994	0,991	0,987
κ	7,762	35,644	184,645	7,762	35,644	184,645	7,762	35,644	184,645

Çizelge 8.4. LT, JLT ve YJLT tahmin edicilerinin OMS'leri, $n=15$

	$\sigma = 0,1$			$\sigma = 1,0$			$\sigma = 5,0$		
	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$
$\hat{\beta}(d)$	0,208	0,179	0,157	1,326	1,302	1,317	1,631	1,618	1,643
$\tilde{\beta}_{JL}(d)$	0,206	0,176	0,141	1,316	1,299	1,311	1,629	1,617	1,643
$\hat{\beta}_{YJ}(d)$	0,208	0,175	0,154	1,328	1,301	1,317	1,632	1,617	1,643
κ	8,33	41,30	226,76	8,33	41,30	226,76	8,33	41,30	226,76

Çizelge 8.5. LT, JLT ve YJLT tahmin edicilerinin OMS'leri, $n=50$

	$\sigma = 0,1$			$\sigma = 1,0$			$\sigma = 5,0$		
	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$
$\hat{\beta}(d)$	0,329	0,338	0,327	1,494	1,488	1,482	1,679	1,678	1,675
$\tilde{\beta}_{JL}(d)$	0,324	0,326	0,319	1,486	1,484	1,481	1,677	1,676	1,673
$\hat{\beta}_{YJ}(d)$	0,329	0,338	0,327	1,490	1,487	1,482	1,679	1,677	1,675
κ	7,55	32,90	166,97	7,55	32,90	166,97	7,55	32,90	166,97

Çizelge 8.6. LT, JLT ve YJLT tahmin edicilerinin OMS'leri, $n=100$

	$\sigma = 0,1$			$\sigma = 1,0$			$\sigma = 5,0$		
	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$	$\rho = 0,8$	$\rho = 0,95$	$\rho = 0,99$
$\hat{\beta}(d)$	0,510	0,515	0,514	1,565	1,565	1,565	1,701	1,701	1,703
$\tilde{\beta}_{JL}(d)$	0,507	0,508	0,508	1,564	1,563	1,565	1,701	1,701	1,703
$\hat{\beta}_{YJ}(d)$	0,510	0,515	0,514	1,565	1,565	1,565	1,701	1,701	1,703
κ	7,762	35,644	184,645	7,762	35,644	184,645	7,762	35,644	184,645

Çizelge 8.4-Çizelge 8.6'deki sonuçlardan, her durumda jackknife Liu-tipi tahmin edicinin ($\hat{\beta}_{JL}(d)$) mutlak yanlışlık teriminin Liu-tipi tahmin edici ($\hat{\beta}(d)$) ve yeni Liu-tipi tahmin edicinkinden ($\hat{\beta}_{YJ}(d)$) daha küçük olduğu görülmektedir.

vektörü ve X_j ise j . değişken gözlemleri çıkarıldığında geriye kalan değişkenlerin oluşturduğu matristir.

Çizelge 8.7. Değişkenler arasındaki korelasyon katsayıları

	<i>şömine</i>	<i>garaj</i>	<i>lüks</i>	<i>gelir</i>	<i>uzaklık</i>	<i>brüt alan</i>	<i>oda</i>	<i>net alan</i>
<i>şömine</i>	1,0000	0,1456	0,0544	0,3815	0,1085	0,1595	0,3144	0,4633
<i>garaj</i>	0,1456	1,0000	0,0793	0,0271	0,0579	0,1699	0,1352	0,2231
<i>lüks</i>	0,0544	0,0793	1,0000	0,0008	0,0942	0,0020	0,1967	0,1479
<i>gelir</i>	0,3815	0,0271	0,0008	1,0000	-0,1096	0,1334	0,1165	0,2779
<i>uzaklık</i>	0,1085	0,0579	0,0942	-0,1096	1,0000	0,0815	0,0288	0,0230
<i>brüt alan</i>	0,1595	0,1699	0,0020	0,1334	0,0815	1,0000	0,3245	0,1545
<i>oda</i>	0,3144	0,1352	0,1967	0,1165	0,0288	0,3245	1,0000	0,5771
<i>net alan</i>	0,4633	0,2231	0,1479	0,2779	0,0230	0,1545	0,5771	1,0000

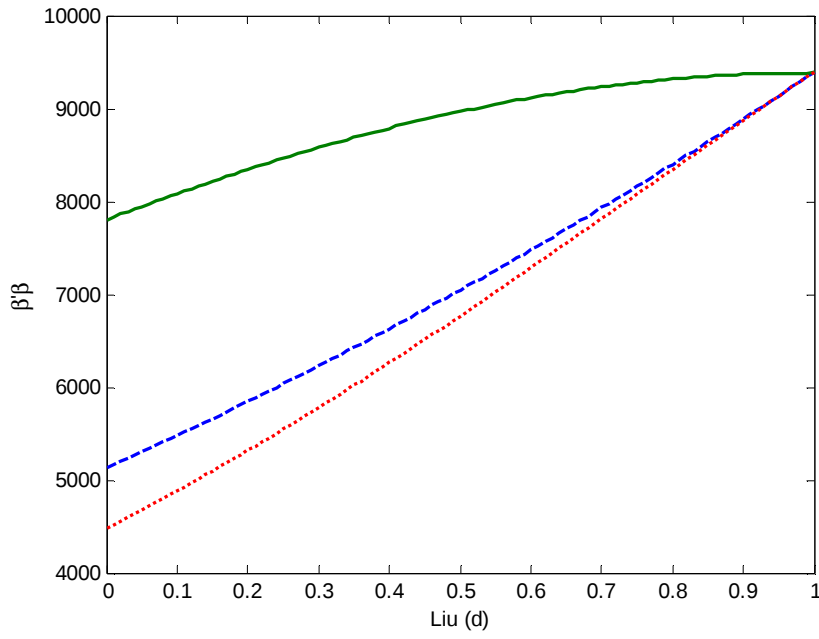
Değişkenlerin VIF değerleri sırasıyla 59,5610; 4,7576; 3,0264; 1,1213; 23,9085; 4,8217; 26,0818; 50,1542; 31,3564'dir. Bazı VIF değerleri 30'dan büyük olduğundan VIF değerleri de değişkenler arasında güçlü bir çoklu bağlantı olduğunu göstermektedir. Bu durumda EKK tahmin edicisi yerine çoklu bağlantı olması durumunda daha uygun bir sonuç veren Liu-tipi tahmin edicileri kullanmak daha uygun olacaktır.

Doğrusal regresyon modeli kanonik forma getirilmek istenirse T kolonları $Z^T Z$ matrisinin özvektörlerinden oluşan ortogonal matris olarak alınabilir. Bu durumda $X^T X = T^T Z^T Z T = \Lambda$ olur. Burada $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_p)$ dır ve $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p$ 'ler $Z^T Z$ matrisinin sıralı özdeğerleridir. $X = ZT$, $T^T \beta = \gamma$ olmak üzere önce γ 'nın EKK, Liu, jackknife Liu ve Yeni jackknife Liu tahmin edicileri bulunmuş daha sonra γ 'lar T matrisi ile çarpılarak β 'lar için EKK, Liu, jackknife Liu ve yeni jackknife Liu tahmin edicileri elde edilmiştir ve sonuçlar Çizelge 8.8'de verilmiştir.

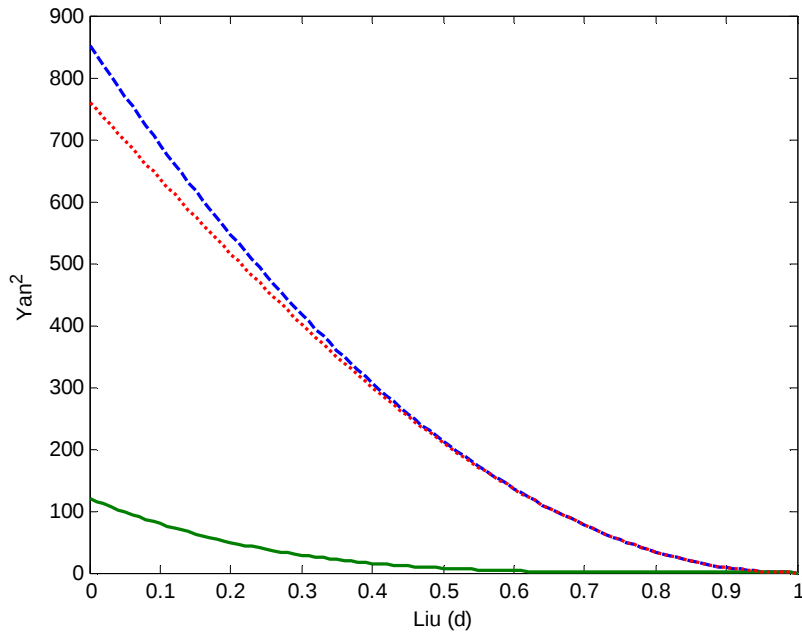
Çizelge 8.8. Farklı d değerleri için β vektörü tahminleri

	$d=0,55$			$d=0,75$			$d=1$
	$\hat{\beta}(d)$	$\tilde{\beta}_J(d)$	$\hat{\beta}_{YJ}(d)$	$\hat{\beta}(d)$	$\tilde{\beta}_J(d)$	$\hat{\beta}_{YJ}(d)$	$\hat{\beta}$
β_0	51,169	60,426	49,458	56,214	61,874	55,633	62,520
<i>şömine</i>	5,531	6,226	5,370	5,929	6,365	5,874	6,428
<i>garaj</i>	13,235	13,153	13,270	13,180	13,125	13,192	13,112
<i>lüks</i>	60,497	65,827	59,968	63,132	66,241	62,959	66,426
<i>gelir</i>	0,685	0,620	0,696	0,650	0,611	0,654	0,607
<i>uzaklık</i>	-9,731	-10,948	-9,546	-10,371	-11,102	-10,309	-11,171
<i>brüt alan</i>	1,048	0,792	1,097	0,907	0,749	0,924	0,730
<i>oda</i>	5,718	4,002	5,982	4,816	3,785	4,904	3,688
<i>netalan</i>	25,463	26,695	25,446	26,009	26,692	26,007	26,691
$\beta^T \beta$	7261,2	9045,7	7023,3	8163	9280	8075,4	9386
SHKO	759,999	740,380	737,941	715,806	756,060	708,205	764,718
Yan^2	172,623	4,937	170,873	53,279	0,470	53,186	-
Varyans	587,377	735,442	567,068	662,527	755,590	655,019	764,718

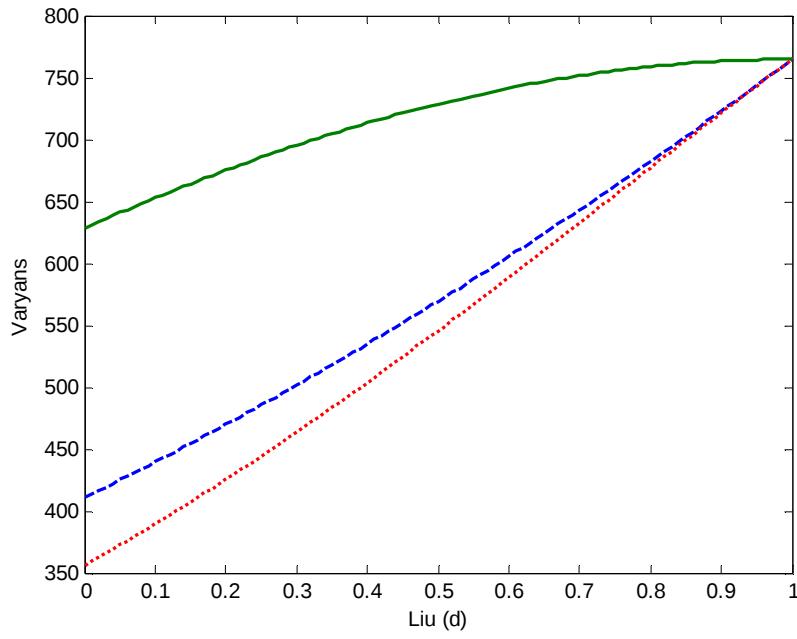
Çizelge 8.8’de farklı d değerleri için tahmin edicilerin normları, skaler hata kare ortalamaları, yanlışlık terimlerinin karesi ve varyansları verilmiştir. Buna göre EKK tahmin edicisi yansız bir tahmin edici olmasına rağmen en yüksek varyanslıdır. Özel olarak seçilen d değerleri için yeni jackknife Liu tahmin edicisinin diğer tahmin edicilere göre daha küçük norm, SHKO ve varyansa sahip olduğu görülmektedir. Yanlışlık terimi en küçük olan tahmin edici beklenildiği gibi tahmin jackknife Liu tahmin edicidir. Şekil 8.1-8.3’te sırasıyla tahmin edicilerin değişen d değerlerine göre norm, yanlışlık ve varyans özellikleri gösterilmiştir. Bu grafiklerde tüm d değerleri için bir tahmin edicinin diğerlerinden daha üstün olduğu görülmektedir. Ancak d ’nin seçimine göre tahmin edicilerin SHKO ölçütüne göre üstünlükleri değişmektedir. Bu durum Şekil 8.4’te açıkça gösterilmiştir. Şekil 8.4’e göre d ’nin $0 < d < 0,53$ değerleri için JLTE’nin SHKO’su YJLTE’nin SHKO’sundan daha küçüktür. $d > 0,53$ için ise yeni jackknife tahmin edicisi daha iyi sonuç vermiştir.



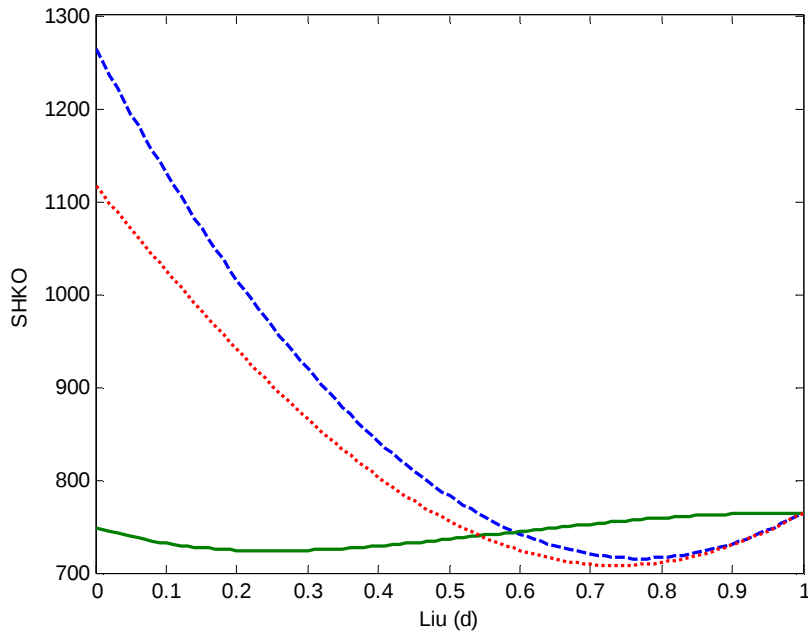
Şekil 8.1. Beta tahminlerinin normlarının karesi ve Liu d parametresi
LTE (mavi --), JLTE (yeşil -), YJLTE (kırmızı ...)



Şekil 8.2. Beta tahminlerinin normlarının karesi ve Liu d parametresi, LTE (mavi --),
JLTE (yeşil -), YJLTE (kırmızı ...)



Şekil 8.3. Beta tahminlerinin varyansları ve Liu d parametresi, LTE (mavi --), JLTE (yeşil -), YJLTE (kırmızı ...)



Şekil 8.4. Beta tahminlerinin SHKO'ları ve Liu d parametresi, LTE (mavi --), JLTE (yeşil -), YJLTE (kırmızı ...)

8.3. Speckman ve CEKK Tahmin Edicileri ile İlgili Simülasyon Çalışması

Bu bölümde 6. bölümde önerilen tahmin edicilerin SHKO ölçütüne göre performanslarını incelemek için Monte Carlo simülasyon çalışması yapılacaktır. Parametrelerle ilgili doğrusal kısıt şeklinde bir önsel bilgi olduğunda kısıtlı tahmin ediciler kullanmanın daha etkili tahminler sağladığından bahsedilmiştir. Bu simülasyon çalışmasında kısıtlı ve kısıtsız tahmin edicilerin değişik çoklu bağlantı derecelerinde SHKO ölçütüne göre performansları incelenecektir.

β 'nin herhangi bir $\hat{\beta}$ tahmin edicisinin SHKO performansı B tekrar için Eş. 8.4'teki gibi hesaplanmıştır. Aynı şekilde, f 'nin $\hat{f}$ tahmin edicisinin SHKO performansı Eş. 8.9 ile ölçülebilir [Bianco ve Boente, 2004].

$$\text{SHKO}(f^*) = \frac{1}{B} \sum_{i=1}^B [f^*(t_i) - f(t_i)]^2. \quad (8.7)$$

Eş. 6.1'de tanımlanan semiparametrik regresyon modelinin parametrik bileşeni bölüm 8.1'deki gibi McDonald and Galarneau (1975) çalışmasına dayanmaktadır. Açıklayıcı değişkenler Eş. 8.1'deki gibi üretilmiştir. Bu simülasyon çalışmasında özel olarak $n = 30, 60$, $p = 4$ ve $n = 80$, $p = 8$ alınmıştır.

Simülasyon çalışmasında $\rho = 0, 8; 0, 99; 0, 999; 0, 9999$ olmak üzere çeşitli derecelerdeki çoklu bağlantı ele alınmıştır. Buradaki ρ 'lar koşul sayısı zayıftan güçlüye olacak şekilde seçilmiştir. $X^T X$ matrisinin koşul sayıları Çizelge 8.9 ve Çizelge 8.10'da parantez içinde verilmiştir. β parametre vektörü sırasıyla $p = 4$ ve $p = 8$ tane açıklayıcı değişken için özel olarak $\beta = \{1; 2; 0, 5; 2\}$ ve $\beta = \{1; 2; 0, 5; 2; 0, 5; -0, 5; 1; 2\}$ seçilmiştir.

Yanıt değişkeni

$$y_i = \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_p x_{ip} + f(t_i) + \varepsilon_i, \quad i = 1, \dots, n \quad (8.8)$$

eşitliğinden türetilmiştir. Burada ε_i 'ler 0 ortalamalı ve σ^2 varyanslı bağımsız ve rasgele hatalardır. Eş. 8.8'deki f fonksiyonu $[0, 1]$ aralığında değerler alan $t_i = (i - 0.5)/n$ için $f(t_i) = m_n f_0(t_i)$ fonksiyonundan üretilmiştir. Burada $f_0 = 4.26(\exp(-3.25t_i) - 4(\exp(-6.5t_i)) + 3(\exp(-9.75t_i)))$ ve $\max |f(t)| = 9$ ($m_n = 9$) varsayılmıştır [Schimek, 2000; Liang, 2006].

Bu çalışmada σ için 0,01 ve 1 değerleri alınmıştır. Kısıtlı tahmin edicilerin performanslarını incelemek amacıyla $p = 4$ ve $p = 8$ için sırasıyla $R = [0, 1, 0, -1]$ ve $R = [0, 1, 0, -1, 0, 0, 0, 0]$ kısıtları varsayılmıştır. Bu kısıtların sağlandığı hipotezi 0.05 anlamlılık düzeyinde Kaçıranlar ve ark. (1999)'da verilen F-testine göre reddedilemediğinden seçilen kısıtlar uygundur. Bu simülasyon çalışmasında her bir durum için sırasıyla aşağıdaki tahmin edicilerin performansları karşılaştırılmıştır.

EKK: En küçük kareler ($\hat{\beta}_{OLS}$)

SPC: Speckman tahmin edicisi ($\tilde{\beta}$)

CEKK: Cezalı en küçük kareler ($\hat{\beta}$)

LSPC: Liu type Speckman's estimator ($\tilde{\beta}_L(d)$)

LCEKK: Liu-tipi backfitting estimator ($\hat{\beta}_L(d)$)

Kısıtlı tahmin ediciler,

KEKK: Kısıtlı en küçük kareler ($\hat{b}_{OLS}$)

KSPC: Kısıtlı Speckman tahmin edicisi ($\tilde{b}$)

KCEKK: Kısıtlı CEKK tahmin edicisi ($\hat{b}$)

KLSPC: Kısıtlı Liu-tipi Speckman tahmin edicisi ($\tilde{b}_L(d)$)

KLCEKK: Kısıtlı Liu-tipi CEKK tahmin edicisi ($\tilde{b}_L(d)$)

Her bir durum için X matrisi ve β vektörü sabit kalacak şekilde yeni hata terimleri üretilerek 1000 tane örnek elde edilmiştir. Örneklerden her bir tahmin edicinin SHKO'su hesaplanmıştır. Her bir simülasyonda düzeltme parametresi α ,

genelleştirilmiş çapraz geçerlilik yöntemi ile hesaplanmıştır. Simülasyon sonuçları Çizelge 8.9 ve 8.10’da gösterilmiştir.

Simülasyon sonuçlarına bakıldığında, EKK’nın en kötü performanslı tahmin edici olduğu görülmektedir. Parametre sayısının artması ve ayrıca ρ veya σ ’daki artış tüm tahmin edicilerin SHKO’larını arttırmıştır. Liu-tipi tahmin ediciler özellikle güçlü çoklu bağlantı olması durumunda SHKO bakımından diğer tahmin edicilere göre üstündür. Zayıf çoklu bağlantı olması durumunda, $\rho = 0.8$, (K)SPC, (K)CEKK, (K)LSPC ve (K)LCEKK tahmin edicilerinin SHKO’ları birbirine çok yakın olmasına rağmen, çoklu bağlantı arttıkça (K)SPC ve (K)CEKK tahmin edicilerinin SHKO’ları neredeyse bu tahmin edicileri kullanışsız yapacak kadar artmıştır.

Burada ayrıca parametre tahmininde kullanılan eşitlik kısıtlaması şeklindeki önsel bilginin de tüm tahmin edicilerin SHKO’larını azalttığı görülmektedir.

Simülasyon sonuçları semiparametrik regresyon modellerinde çoklu bağlantı olması durumunda Liu-tipi Speckman ve Liu-tipi cezalı en küçük kareler tahmin edicilerinin daha tercih edilebilir tahmin ediciler olduğunu göstermektedir. Ayrıca parametreler hakkında eşitlik kısıtı gibi bir önsel bilgi varsa bunlar da tahmin sürecine eklenerek daha etkin tahminler elde edilebilir.

Çizelge 8.9. Tahmin edicilerin SHKO’ları

(p, n)	$\rho, (\kappa)$	σ	EKK	SPC	CEKK	LSPC	LCEKK
(4, 30)	0,8 (3,55)	0,01	180	7,46	7,63	1,84	1,96
		1	191	17,85	17,89	4,20	4,31
	0,99 (17,73)	0,01	2660	133	132	2,83	2,93
		1	2864	319	317	41	40
	0,999 (56,59)	0,01	26209	1344	1339	9,77	9,88
		1	28256	3217	3187	451	451
	0,9999 (179,33)	0,01	261596	13512	13447	73	72
		1	282112	32272	31960	4781	4782
	(4, 60)	0,01	309	13	20	5,99	8,96
		1	320	23	30	10,84	13,90

Çizelge 8.9. (Devam) Tahmin edicilerin SHKO'ları

(8, 80)	0,99 (15,44)	0,01	4799	247	380	81	125
		1	4986	424	556	155	201
	0,999 (49,57)	0,01	45971	2456	3778	761	1170
		1	47826	4213	5518	1469	1903
	0,9999 (157,37)	0,01	453838	24533	37715	7456	11461
		1	472354	42073	55086	14475	18737
	0,8 (4,90)	0,01	537	26	41	8,37	12,91
		1	561	50	64	16,02	20,68
	0,99 (25,55)	0,01	9653	476	749	127	200
		1	10090	899	1169	245	322
	0,999 (77,71)	0,01	96399	4746	7457	1257	1979
		1	100734	8945	11629	2415	3169
	0,9999 (245,55)	0,01	964732	47459	74552	12592	19793
		1	1008022	89393	116207	24104	31633

Çizelge 8.10. Kısıtlı tahmin edicilerin SHKO'ları

(p,n)	$\rho, (\kappa)$	σ	KEKK	KSPC	KCEKK	KLSPC	KLCEKK
(4, 30)	0,8 (3,55)	0,01	179	3,19	2,97	1,24	1,26
		1	187	10,94	10,71	3,15	3,17
	0,99 (17,73)	0,01	2643	46	38	1,19	1,10
		1	2796	182	174	27	26
	0,999 (56,59)	0,01	26042	461	377	2,47	1,80
		1	27568	1814	1728	308	298
	0,9999 (179,33)	0,01	259929	4618	3769	21	16
		1	275191	18150	17275	3297	3202
	0,8 (3,04)	0,01	75	3,40	5,58	2,58	3,69
		1	81	9,38	11,43	5,34	6,49
	0,99 (15,44)	0,01	1222	56	93	25	42
		1	1308	153	188	61	79
(4, 60)	0,999 (49,57)	0,01	11837	548	916	228	385
		1	12673	1495	1842	562	726
	0,9999 (157,37)	0,01	117419	5457	9118	2192	3717
		1	125727	14879	18324	5491	7090
	0,8 (4,90)	0,01	532	38	52	6,44	11,50
		1	556	44	60	14,72	19,65
	0,99 (25,55)	0,01	9617	431	737	119	198
		1	9989	785	1090	225	305
	0,999 (77,71)	0,01	96020	4288	7336	1187	1960
		1	99709	7811	10841	2214	3012
	0,9999 (245,55)	0,01	960863	42869	73330	11895	19601
		1	997695	78049	108323	22114	30079

8.4. Speckman ve CEKK Tahmin Edicileri ile İlgili Uygulama

Bu kısımda müstakil evlerin barındırdığı özelliklerin satış fiyatına etkisini gösteren Bölüm 8.2'deki örnek incelenecektir. 8.2. bölümde yanıt değişkeni ile açıklayıcı değişkenler arasında doğrusal bir ilişki olduğu varsayılmıştı ancak daha gerçekçi bir yaklaşım bazı açıklayıcı değişkenlerin yanıt değişkeni ile doğrusal olmayan ilişkili olduğudur.

Klasik regresyon modeline alternatif olarak kısmi doğrusal model kullanılabilir. Evin brüt alanı (*brüt alan*) evin satış fiyatında oldukça etkili olmasına rağmen parametrik olarak modellenemediğinden parametrik olmayan değişken olarak ele alınacaktır. Aydın (2005) 'te de aynı veri seti *lüks* ve *oda* dışında kalan değişkenlerle ve evin brüt alanı parametrik olmayan değişken alınarak incelenmiştir. Burada da değişkenlerin evin satış fiyatı üzerine etkileri kısmi doğrusal regresyon modeli ile açıklanacaktır. Bu örnek için kısmi doğrusal model

$$y_i = \beta_0 + \beta_1(\text{\textit{\$}ömine})_i + \beta_2(\text{\textit{garaj}})_i + \beta_3(\text{\textit{lüks}})_i + \beta_4(\text{\textit{gelir}})_i + \beta_5(\text{\textit{uzaklık}})_i + \beta_6(\text{\textit{oda}})_i + \beta_7(\text{\textit{net alan}})_i + f(\text{\textit{brüt alan}})_i + \varepsilon_i \quad (8.7)$$

olarak yazılır. X , 92×7 tipinde *\\$ömine*, *garaj*, *lüks*, *gelir*, *uzaklık*, *oda* ve *net alan* değişkenlerinden oluşan veri matrisi olmak üzere $X^T X : 7 \times 7$ matrisinin özdeğerleri sırasıyla $\lambda_1 = 3,79$, $\lambda_2 = 4,58$, $\lambda_3 = 15,76$, $\lambda_4 = 18,59$, $\lambda_5 = 20,99$, $\lambda_6 = 90,73$ ve $\lambda_7 = 236956,33$. X matrisinin koşul sayısı $\kappa(X) = 250,069$ oldukça büyük olduğundan X kötü koşullu bir matristir. Bu durumda kısmi doğrusal regresyon modelinde kullanılan cezalı EKK ve Speckman tahmin edicileri ne alternatif olarak Liu-tipi cezalı EKK ve Liu-tipi Speckman tahmin edicileri hesaplanacaktır.

Çizelge 8.11 ve Çizelge 8.12'de parametrelerin sırasıyla kısıtlı ve kısıtsız EKK, cezalı EKK, Liu-tipi cezalı EKK, Speckman ve Liu-tipi Speckman tahmin edicileri

verilmiştir. Kısmi doğrusal modellerde gösterge değişkenler fonksiyonu kırımını etkilemediğinden modelin doğrusal kısmına eklenirler.

$S = (I + \alpha K)^{-1}$ düzeltme matrisinde yer alan düzeltme parametresi α genelleştirilmiş çapraz geçerlilik (GÇG) yöntemi ile $\text{iz}(S) < n$ varsayımının sağlanması koşulu altında yaklaşık olarak 18 bulunmuştur.

$$G\check{C}G(\alpha) = \frac{1/n \|I - H(\alpha)\|^2}{[1/n (1 - \text{iz}(H(\alpha)))]^2} \quad (8.8)$$

Burada α 'nın optimum çapraz geçerlilik tahmini Eş. 8.8'in minimum yapılmasıyla

$$\alpha_{opt} = \arg \min_{\alpha \in \mathbb{R}^+} G\check{C}G(\alpha)$$

şeklinde bulunur. Burada H her bir tahmin edici için şapka matrisidir.

Çizelge 8.11. Ev fiyatı örneği için parametre tahminleri ($d=0,628$)

Değişkenler	$\hat{\beta}$	$\hat{\beta}_p$	$\hat{\beta}_p(d)$	$\tilde{\beta}_s$	$\tilde{\beta}_s(d)$
β_0	62,520	-	-	-	-
<i>şömine</i>	6,428	6,026	6,075	5,736	5,786
<i>garaj</i>	13,112	12,832	12,845	12,688	12,698
<i>lüks</i>	66,426	65,928	61,493	65,706	61,279
<i>gelir</i>	0,607	0,626	0,635	0,638	0,647
<i>uzaklık</i>	-11,171	-11,562	-11,117	-11,761	-11,317
<i>oda sayısı</i>	3,688	3,830	4,498	3,904	4,564
<i>net alan</i>	26,691	25,543	24,012	25,183	23,661
<i>brüt alan</i>	0,730	-	-	-	-
R^2	0,541	0,555	0,554	0,555	0,554
$\text{iz}(\hat{V}(.))$	764,72	392,37	348,45	394,09	349,88
$\text{S\HKO}(.)$	764,72	392,72	365,70	394,19	369,86

Çizelge 8.11’de tahmin edicilerin SHKO performanslarına bakacak olursak kısmi doğrusal regresyon modelinin doğrusal regresyon modelinden daha küçük $iz(\hat{V}(\cdot))$ ve $\hat{SHKO}(\cdot)$ değerlerine sahip olduğu görülür. Ayrıca kısmi doğrusal regresyon modelinin belirtme katsayısı R^2 de parametrik regresyon modelinin belirtme katsayısından büyüktür.

$$\begin{aligned} iz(\hat{V}(\hat{\beta})) &= 764,72 > iz(\hat{V}(\hat{\beta}_p)) = 392,37 \cong iz(\hat{V}(\tilde{\beta}_s)) = 394,09 \\ &> iz(\hat{V}(\hat{\beta}_p(d))) = 348,45 \cong iz(\hat{V}(\tilde{\beta}_s(d))) = 349,88 \end{aligned}$$

ve

$$\begin{aligned} \hat{SHKO}(\hat{\beta}) &= 764,72 > \hat{SHKO}(\hat{\beta}_p) = 392,72 \cong \hat{SHKO}(\tilde{\beta}_s(d)) = 394,19 \\ &> \hat{SHKO}(\hat{\beta}_p(d)) = 365,70 \cong \hat{SHKO}(\tilde{\beta}) = 369,86 \end{aligned}$$

Kısmi doğrusal regresyon modelinin doğrusal regresyon modeline göre daha küçük bir SHKO ve $iz(\hat{V}(\cdot))$ ’ye sahip olduğu ve bu nedenle bu örnek için daha iyi bir model olduğu söylenebilir. Değişkenler arasında güçlü çoklu bağlantı olması nedeniyle beklenildiği gibi Liu-tipi cezalı EKK ve Liu-tipi Speckman tahmin edicileri belirgin biçimde $iz(\hat{V}(\cdot))$ ve SHKO ölçütlerine göre daha iyi sonuçlar vermişlerdir.

Kısıtlı tahmin edicilerin performanslarını karşılaştırmak için parametreler üzerinde $\beta_1 - \beta_2 - \beta_4 - \beta_5 - \beta_6 = 0$ veya $R = (1, -1, 0, -1, -1, -1, 0)$ olmak üzere $R\beta = 0$ kısıtı olduğu varsayılın.

Çizelge 8.10’un sonuçlarından tahmin ediciler örnek varyansları ve SHKO ölçütlerine göre karşılaştırıldığında aşağıdaki bulgular elde edilmiştir.

$$\begin{aligned} iz(\hat{V}(\hat{b})) &= 626,74 > iz(\hat{V}(\hat{b}_p)) = 343,35 \cong iz(\hat{V}(\tilde{b}_s)) = 345,02 \\ &> iz(\hat{V}(\hat{b}_p(d))) = 303,17 \cong iz(\hat{V}(\tilde{b}_s(d))) = 304,54 \end{aligned}$$

ve

$$\begin{aligned} \hat{S}HKO(\hat{\beta}_R) &= 626,74 > \hat{S}HKO(\hat{b}) = 343,35 \cong \hat{S}HKO(\tilde{b}_s) = 345,02 \\ &> \hat{S}HKO(\hat{b}_p(d)) = 319,31 \cong \hat{S}HKO(\tilde{b}_s(d)) = 323,11 \end{aligned}$$

Çizelge 8.12. Kısıt altında ev fiyatı örneği için parametre tahminleri, $(d=0,6846)$

Değişkenler	$\hat{\beta}_R$	$\hat{b}_p$	$\hat{b}_p(d)$	$\tilde{b}_s$	$\tilde{b}_s(d)$
β_0	62,378	-	-	-	-
<i>şömine</i>	6,363	5,925	6,339	5,646	6,057
<i>garaj</i>	13,150	12,892	12,685	12,742	12,534
<i>lüks</i>	66,401	65,889	61,578	65,672	61,366
<i>gelir</i>	0,608	0,628	0,630	0,640	0,642
<i>uzaklık</i>	-11,123	-11,487	-11,313	-11,695	-11,518
<i>oda sayısı</i>	3,728	3,892	4,337	3,959	4,398
<i>net alan</i>	26,667	25,506	24,094	25,151	23,745
<i>brüt alan</i>	0,7221	-	-	-	-
R^2	0,541	0,555	0,555	0,555	0,554
$iz(\hat{V}(\cdot))$	626,74	343,35	303,17	345,02	304,54
$\hat{S}HKO(\cdot)$	626,74	343,69	319,31	345,11	323,11

Çizelge 8.12'un sonuçlarına göre, parametreler üzerine kısıt konulması beklenildiği gibi tahmin edicilerin varyans $iz(\hat{V}(\cdot))$ ve SHKO tahminlerini azaltmıştır. Buna ek olarak kısıtlı Liu-tipi tahmin ediciler gözle görülür biçimde $iz(\hat{V}(\cdot))$ ve SHKO ölçütlerine göre diğer kısıtlı tahmin edicilerden daha iyi sonuçlar vermiştir.

8.5. Farka Dayalı Liu-tipi Tahmin Edici ile İlgili Simülasyon Çalışması

Bu bölümde Yatchew (2003)'da farka dayalı tahmin edici ile farka dayalı Liu tipi tahmin edici skaler hata kare ortalaması ölçütü baz alınarak karşılaştırılmıştır.

Simülasyon çalışması McDonald and Galarneau (1975) makalesine dayanmaktadır. Bu çalışmada beş açıklayıcı değişkene ait $n=100$ tane gözlem Eş. 8.15'deki denklemle üretilmiştir.

$$x_{ij} = (1 - \rho^2)^{1/2} z_{ij} + \rho z_{i(p+1)} \quad i = 1, \dots, n; j = 1, \dots, p \quad (8.9)$$

Burada z_{ij} 'ler birbirinden bağımsız standart normal dağılımlı rasgele değişkenler ve ρ^2 ise her bir açıklayıcı değişken arasındaki teorik korelasyondur. Bu çalışmada, $\rho = 0,7; 0,8; 0,9; 0,99$ değerlerine karşı gelen dört korelasyon katsayısı incelenmiştir. Yanıt değişkeni y aşağıdaki modelden elde edilmiştir.

$$y_i = x_{1i}\beta_1 + x_{2i}\beta_2 + x_{3i}\beta_3 + x_{4i}\beta_4 + x_{5i}\beta_5 + f(t_i) + \varepsilon_i, \quad i = 1, \dots, n \quad (8.10)$$

Burada özel olarak $\beta = (1.5, 2, 3, -5, 4)$ alınmıştır, $\varepsilon \sim N(0, \sigma^2 I)$ ve $f(t_i) = \sqrt{t_i(1-t_i)} \sin(2.1\pi / (t_i + 0.05))$, $t_i = (i - 0.5) / n$, $i = 1, \dots, n$ değerleri için Doppler fonksiyonudur. $1 \leq m \leq 10$ için optimum fark ağırlıkları Çizelge 7.1'de verilmiştir. Bu simülasyon çalışmasında $m = 3$. dereceden fark alınmıştır. Her bir ρ ve σ , kombinasyonu için X matrisi, β vektörü ve f fonksiyonunu sabit tutarak ve yeni hata terimleri üreterek deney 1000 defa tekrarlanmıştır. 1000 tane örnek üretildikten sonra her bir tahmin edici için skaler hata kareler ortalaması hesaplanmıştır. SHKO aşağıdaki şekilde tanımlanmıştır.

$$SHKO(\hat{\beta}) = \frac{1}{1000} \sum_{j=1}^{1000} \sum_{i=1}^5 (\hat{\beta}_{ij} - \beta_i)^2 \quad (8.11)$$

Burada $\hat{\beta}_{ij}$, i . parametrenin j . denemedeki tahminini ve $\beta_1, \beta_2, \beta_3, \beta_4, \beta_5$ 'ler ise doğru parametre değerlerini göstermektedir. Simülasyon sonuçları Çizelge 8.11'de verilmiştir, burada κ koşul sayısını göstermektedir.

Çizelge 8.13’de, $\sigma = 0,1$ için hiçbir l değeri için Teorem 7.1’deki koşul sağlanmadığından tüm l değerleri için $\hat{\beta}_{diff}(l)$ ’in SHKO’su $\hat{\beta}_{diff}$ ’in SHKO’sundan her zaman daha büyüktür. $\sigma = 5$ için ise Teorem 7.1’deki koşul her zaman sağlandığından $\hat{\beta}_{diff}(l)$ ’in SHKO’su $\hat{\beta}_{diff}$ ’in SHKO’sundan her zaman daha küçüktür. $\sigma = 1$ için durum biraz farklıdır çünkü bazı l değerleri için Teorem 7.1’deki koşul sağlanırken diğer l değerleri için sağlanmamaktadır.

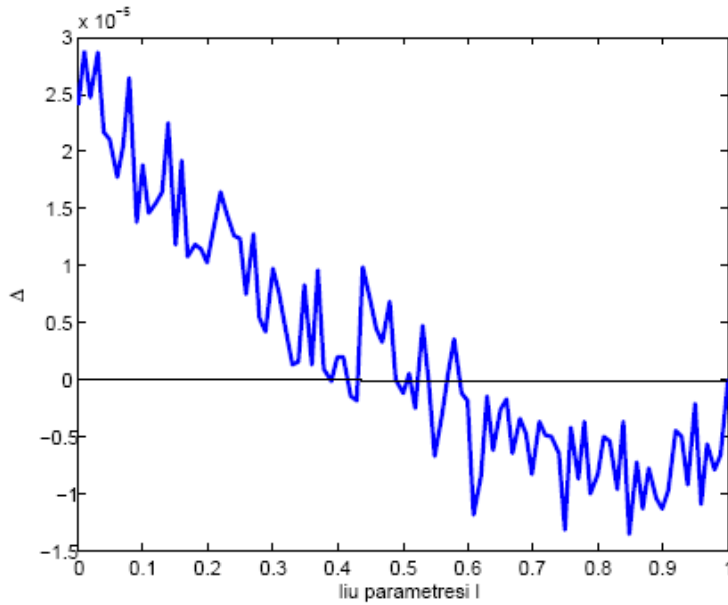
Çizelge 8.13. $\hat{\beta}_{diff}(l=1)$ ve $\hat{\beta}_{diff}(l)$ tahmin edicilerinin tahmini SHKO değerleri

ρ (κ)	σ	l değerleri					
		0,5	0,7	0,8	0,85	0,90	1
0,7(2,9)	0,1	0,0093	0,0060	0,0051	0,0048	0,0044	0,0044
	1	0,1120	0,1078	0,1080	0,1044	0,1061	0,1138
	5	2,6083	2,5838	2,5960	2,5059	2,5796	2,7374
0,8(3,9)	0,1	0,0161	0,0096	0,0077	0,0070	0,0064	0,0064
	1	0,1575	0,1504	0,1503	0,1452	0,1474	0,1588
	5	3,5982	3,5768	3,6016	3,4768	3,5801	3,8155
0,9(5,9)	0,1	0,0463	0,0237	0,0169	0,0145	0,0126	0,0121
	1	0,3026	0,2827	0,2807	0,2808	0,2746	0,2972
	5	6,5531	6,5803	6,6616	6,4450	6,6495	7,1395
0,99(20,1)	0,1	1,6088	0,6348	0,3350	0,2325	0,1585	0,1149
	1	3,2907	2,6467	2,5258	2,4376	2,4861	2,8140
	5	44,8710	51,4055	55,4927	55,4667	58,9922	67,5794

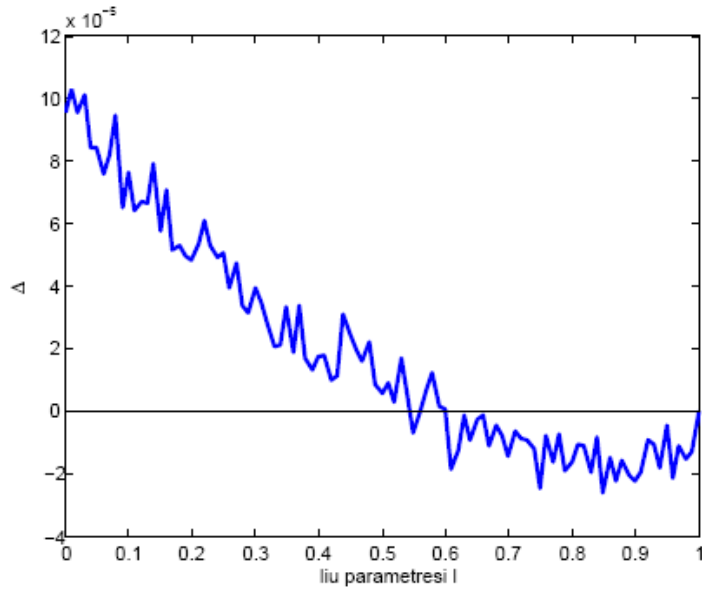
Çizelge 8.13’de Teorem 7.1’deki koşulun sağlandığı durumlarda $\hat{\beta}_{diff}(l)$ ’in SHKO’sunun $\hat{\beta}_{diff}$ ’in SHKO’sundan daha küçük olduğu görülmektedir. Tüm l değerleri için olan sonuçlar çizelgede gösterilemediğinden, Şekil 8.5-8.7 grafiklerinde tahmin edicilerin tüm l değerleri için SHKO grafikleri verilmiştir. Şekil 8.5-8.7’de sırasıyla $\rho = 0,8$, $\rho = 0,9$ ve $\rho = 0,99$ değerleri için

$\Delta = \hat{\text{SHKO}}(\hat{\beta}_{\text{diff}}(l)) - \hat{\text{SHKO}}(\hat{\beta}_{\text{diff}})$ 'nın l değerlerine karşı grafiği verilmiştir. Bu grafiklerde ayrıca koşul sayısı arttıkça $\hat{\beta}_{\text{diff}}(l)$ ve $\hat{\beta}_{\text{diff}}$ 'in SHKO'ları arasındaki mutlak farkın arttığı görülmektedir.

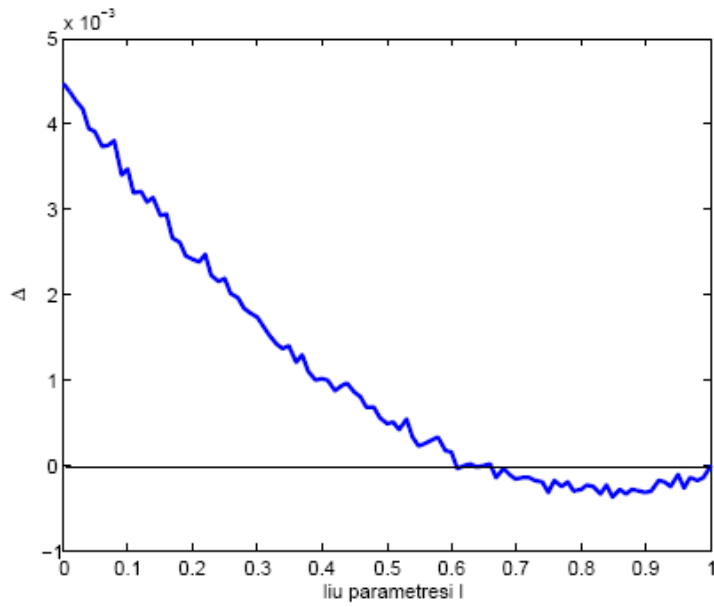
Şekil 8.5-8.7'de sıfır referans çizgisi altında kalan alanlar farka dayalı Liu tahmin edicisinin SHKO bakımından farka dayalı tahmin ediciden üstün olduğu yerleri göstermektedir. 3'ten farklı derece fark ağırlıkları için de simülasyon sonuçları benzer olduğundan diğer fark dereceleri için sonuçlar verilmemiştir.



Şekil 8.5. $\Delta = \hat{\text{SHKO}}(\hat{\beta}_{\text{diff}}(l)) - \hat{\text{SHKO}}(\hat{\beta}_{\text{diff}})$ 'nın $\rho = 0,8$ ve $\sigma = 1$ için l değerlerine karşı grafiği(mavi)



Şekil 8.6. $\Delta = \hat{S}\hat{H}KO(\hat{\beta}_{diff}(l)) - \hat{S}\hat{H}KO(\hat{\beta}_{diff})$ 'nın $\rho = 0,9$ ve $\sigma = 1$ için l değerlerine karşı grafiği(mavi)



Şekil 8.7 $\Delta = \hat{S}\hat{H}KO(\hat{\beta}_{diff}(l)) - \hat{S}\hat{H}KO(\hat{\beta}_{diff})$ 'nın $\rho = 0,99$ ve $\sigma = 1$ için l değerlerine karşı grafiği(mavi)

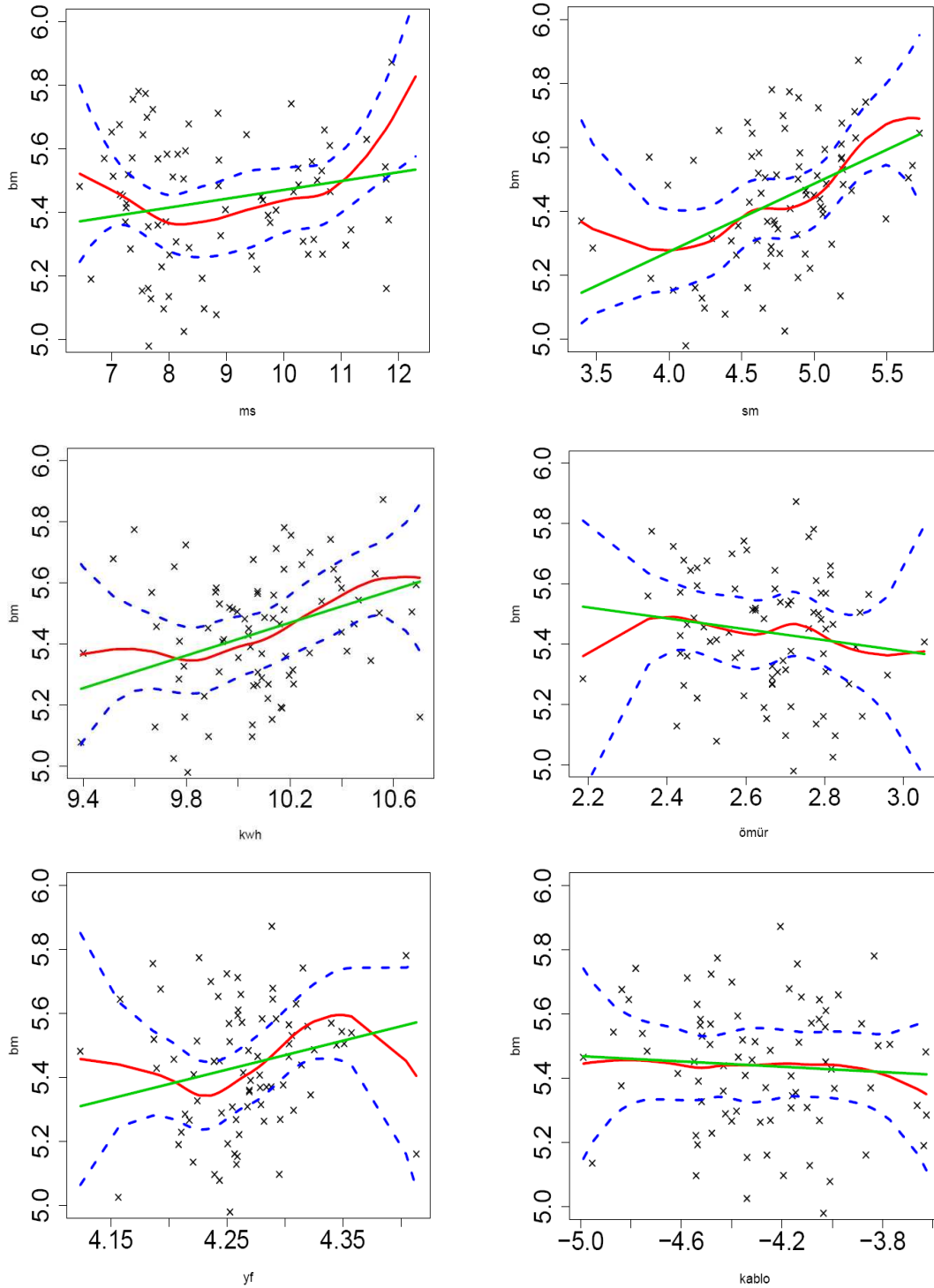
8.6. Farka Dayalı Liu-tipi Tahmin Edici ile İlgili Uygulama

Farka dayalı tahmin ediciler ile ilgili uygulama elektrik dağıtımındaki ölçek ekonomisi ile ilgilidir. Veri seti Kanada'nın Ontario şehrinde bulunan, eksik gözlemi olmayan ve 1993 yılında müşteri sayısı 600 ile 222000 arasında değişen belediyeye ait 81 dağıtımçıya aittir [Yatchew, 2000].

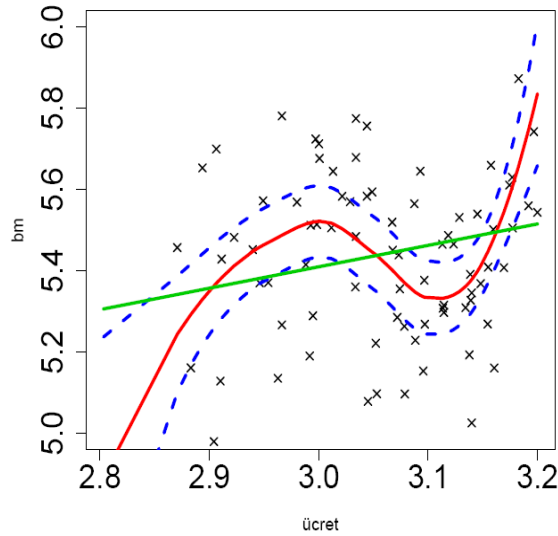
Bu uygulamanın temel amacı verinin parametrik olmayan ve çoklu bağlantılı yapısına uygun olarak elektrik dağıtımındaki ölçek ekonomisinin tahmin edilmesidir. Ölçek ekonomisi, bir firmada tesisler genişletilerek, üretim hacmini veya üretim fonksiyonunu değiştirerek, teknolojik yenilikler getirerek veya dış çevrede meydana gelen maliyet düşürücü faktörlerden yararlanılarak maliyet masraflarının düşürülmesi yoluyla elde edilen kazançtır. Müşteri sayısı (MS) ölçek verisidir. Müşteri sayısı ve birim maliyet (BM) arasındaki ilişki önceden bilinmemektedir.

Modelde kullanılacak diğer değişkenler işçilerin saat başı aldığı ücret (ÜCRET), sermaye maliyeti (SM), müşteri başına dağıtılan elektrik miktarı (KWH), sabit varlıkların kalan ömrü (ÖMÜR) ve talebin çok olduğu zamanlarda kapasite kullanımını ölçen yükleme faktörü (YF)'dür. Ayrıca müşteri yoğunluğu etkisini modele katmak için müşteri başına düşen dağıtım kablosu uzuluğu da eklenmiştir (KABLO). Dağıtıcı firmanın diğer kamu hizmetlerini verip vermemesi (KH). Bir dağıtım firmasının elektriğin yanında diğer kamu hizmetlerini de vermesinin, su dağıtımı gibi, firmanın dağıtım maliyetini düşürmesi beklenir. Modelin bir diğer amacı da bunu test etmektir. Logaritması alınan değişkenler küçük harfle ifade edilecektir. Buna göre model, *bm* yanıt değişkeni olmak üzere, *ms*, *ücret*, *sm*, *KH*, *kwh*, *ömür*, *yf* ve *kablo* değişkenlerinden oluşmaktadır. Şekil 8.5'de her bir açıklayıcı değişkenin yanıt değişkeni ile olan grafikleri verilmiştir. Şekil 8.5'e bakıldığında *bm* ile *ms* ve *bm* ile *ücret* arasında doğrusal olmayan bir ilişki olduğu görülmektedir, ayrıca tahmin edilmiş doğrusal model güven bantları dışında kalmıştır. Bu nedenle *ms* ve *ücret* değişkenleri potansiyel parametrik olmayan

değişkenlerdir. Diğer tüm değişkenlerle bm arasında yaklaşık doğrusal bir ilişki olduğu görülmektedir.



Şekil 8.8. Açıklayıcı değişkenlerin yanıt değişkene karşı grafikleri, doğrusal model (yeşil), kernel regresyonu (kırmızı), %95 güven bantları (mavi)



Şekil 8.8. (Devam) Açıklayıcı değişkenlerin yanıt değişkene karşı grafikleri, doğrusal model (yeşil), kernel regresyonu (kırmızı), %95 güven bantları (mavi)

Bu örnekte ms ve $ücret$ değişkenlerinin gerçekten parametrik olmayan değişkenler olup olmadıklarını test edilmesi gerekiyor. Burada Yatchew (2003)'de bölüm 4.3.1'de verilen berlileme (specification) testi semiparametrik regresyon modeline uyarlanarak kullanılacaktır.

$y = X\beta + f(t) + \varepsilon$ semiparametrik regresyon modelinin daha önceki $E(\varepsilon | t, X) = 0$ ve $V(\varepsilon | t, X) = \sigma^2$ koşullarını sağladığı ve f fonksiyonun birinci türevinin sınırlı olduğu varsayılacaktır. $h(t, \gamma)$, t ve bilinmeyen γ parametresinin bilinen bir fonksiyonu olsun. Burada semiparametrik regresyon modelinin parametrik olmayan kısmının $h(t, \gamma)$ olduğu sıfır hipotezi $f(t)$ alternatifine karşı test edilecektir. Kısıtlı modelin artık terimlerinin varyansı

$$s_{\text{kısıtlı}}^2 = \frac{1}{n} (y - X\hat{\beta} - h(t, \hat{\gamma}))^T (y - X\hat{\beta} - h(t, \hat{\gamma})). \quad (8.12)$$

olarak gösterilsin. Burada $\hat{\gamma}$, γ 'nın herhangi bir tahminidir. Örneğin EKK tahmini olabilir.

Yatchew (2003)'teki 4.3.1 önermesinde f fonksiyonunun bilinen parametrik fonsiyona sahip olduğu hipotezini $H_0 : f(t) = h(t, \gamma)$ test etmek için

$$(mn)^{1/2} \frac{(s_{kısıtlı}^2 - s_{kısıtsız}^2)}{s_{kısıtsız}^2} \xrightarrow{D} N(0,1) \quad (8.13)$$

test istatistiği verilmiştir. Tek yanlı bu test istatistiğine s^2 'deki *kısıtlı* indisi s^2 'nin parametrik fonksiyon varsayımı altında tahmin edildiğini, *kısıtsız* indisi ise f fonksiyonu için herhangi bir parametrik fonksiyon varsayımı yapılmadığı anlamına gelmektedir. Burada $s_{kısıtsız}^2$ daha önce verilen S_{diff}^2 'e karşılık gelmektedir ve

$$s_{kısıtsız}^2 = S_{diff}^2 = \frac{1}{n} (Dy - DX \hat{\beta}_{diff})^T (Dy - DX \hat{\beta}_{diff})$$

dir.

Müşteri sayısı değişkeninin bilinen $f(ms) = \gamma_0 + \gamma_1 ms + \gamma_2 ms^2$ yapısında olduğu hipotezini, parametrik olmayan bir fonksiyon yapısında olduğu hipotezine karşı test etmek için Eş. 8.13'de verilen test kullanılacaktır. Buna göre $s_{kısıtlı}^2 = 0,021$ ve $s_{kısıtsız}^2 = 0,0192$ olarak bulunduğundan test istatistiğinin hesaplanan değeri 1,568 olur. Test istatistiği asimptotik olarak normal dağılıma sahip ve $1,568 < 1,96$ olduğundan H_0 hipotezi 0,05 anlamlılık düzeyinde reddedilemez. Bu durumda ms değişkeninin kuadratik fonksiyonu da uygundur. Ancak, Yatchew (2001) makalesinde müşteri sayısı değişkeni parametrik olmayan değişken alındığından burada da sonuçları karşılaştırmak için ms değişkeni parametrik olmayan değişken olarak alınarak semiparametrik model kurulacaktır.

Veri seti parametrik olmayan ms değişkenine göre artan biçimde sıralandıktan sonra fark matrisi uygulanmıştır. Bu örnekte 3. derece optimum fark ağırlıkları kullanılmıştır [Yatchew, 2003]. Diğer derecedeki farklar için de benzer sonuçlar elde edildiğinden burada sadece 3. derece farklar için sonuçlar verilecektir. Farkı alınmış serinin koşul sayısı $\kappa = \sqrt{\lambda_{\max}((DX)^T DX) / \lambda_{\min}((DX)^T DX)} = 13$ bulunduğundan bu veride çoklu bağlantı problemi vardır. Bu durumda farka dayalı Liu-tipi tahmin edicileri kullanmanın daha iyi sonuçlar verebileceği düşünülebilir.

$\hat{\beta}_{diff}(l)$ and $\hat{\beta}_{diff}$ 'in HKOM'leri hata terimlerinin(ε) bilinmeyen varyansı σ^2 'ye bağlı olduğundan, σ^2 'yi tahmin etmek için S_{diff}^2 kullanılabilir. Eş. 7.3'deki kısıtları sağlayan keyfi fark katsayıları için $S_{diff}^2 \xrightarrow{P} \sigma^2$ sağlanır [Yatchew, 2003: s.71].

$\hat{\beta}_{diff}(l)$ tahmin edicisinin tahmin edilmiş SHKO

$$S\hat{H}KO(\hat{\beta}_{diff}(l)) = iz[S_{diff}^2 U F_{\eta} S F_{\eta}' U + (1-l)^2 (U^{-1} + I)^{-1} \hat{\beta}_{diff} \hat{\beta}_{diff}^T (U^{-1} + I)^{-1}] \quad (8.14)$$

Fark tahmin edicisi $\hat{\beta}_{diff}$ 'nin tahmin edilmiş HKOM'sinin iz'i

$$S\hat{H}KO(\hat{\beta}_{diff}) = iz[S_{diff}^2 U S U]. \quad (8.15)$$

şeklinde bulunur. Çizelge 8.13'te parametrelerin çeşitli l değerlerine karşı gelen farka dayalı Liu tahmin edicileri verilmiştir. Parantez içinde gösterilen değerlerin beta tahminlerinin standart hatalarıdır.

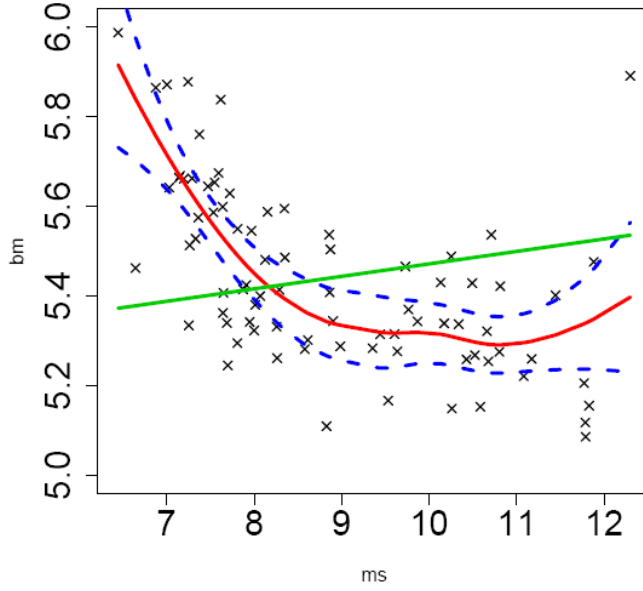
Çizelge 8.13'deki sonuçlar yorumlanacak olursa; KH değişkeninin beklenildiği gibi işareti negatiftir. Buna göre, diğer kamu hizmetlerini de veren dağıtım merkezlerinin elektrik dağıtım maliyeti düşmektedir. Ayrıca diğer değişkenlerin de işaretleri anlamlıdır. Bazı l değerleri için $S\hat{H}KO$ 'nun $\hat{\beta}_{diff}(l=1)$ 'in $S\hat{H}KO$ 'sundan daha düşük olduğu görülmektedir.

Çizelge 8.14. $\hat{\beta}_{diff}(l=1)$ ve $\hat{\beta}_{diff}(l)$ tahmin edicilerinin, standart hataları, tahmini varyans ve SHKO değerleri

değişkenler	$l = 0,6$	$l = 0,69$	$l = 0,75$	$l = 0,85$	$l = 0,95$	$l = 1$
<i>ücret</i>	0,404 (0,227)	0,453 (0,252)	0,486 (0,269)	0,541 (0,296)	0,596 (0,324)	0,623 (0,338)
<i>sm</i>	0,472 (0,065)	0,489 (0,067)	0,500 (0,068)	0,518 (0,070)	0,536 (0,072)	0,545 (0,073)
<i>KH</i>	-0,081 (0,039)	-0,080 (0,039)	-0,079 (0,039)	-0,077 (0,040)	-0,075 (0,040)	-0,075 (0,040)
<i>kwh</i>	0,040 (0,081)	0,033 (0,083)	0,028 (0,085)	0,020 (0,087)	0,012 (0,090)	0,008 (0,091)
<i>ömür</i>	-0,522 (0,104)	-0,546 (0,109)	-0,562 (0,111)	-0,588 (0,116)	-0,615 (0,120)	-0,628 (0,123)
<i>yf</i>	0,846 (0,298)	0,954 (0,335)	1,027 (0,360)	1,147 (0,402)	1,267 (0,443)	1,327 (0,464)
<i>kmwire</i>	0,344 (0,079)	0,360 (0,081)	0,370 (0,083)	0,387 (0,085)	0,404 (0,088)	0,413 (0,090)
s_{ε}^2	0,01992	0,01961	0,01945	0,01926	0,01916	0,01915
SHKO	0,4707	0,3880	0,3521	0,3264	0,3434	0,3680
R^2	0,639	0,645	0,647	0,651	0,653	0,653
L	1,5022	1,7902	1,9992	2,3776	2,7936	3,0157

Parametrik tahminler yapıldıktan sonra Yatchew(1997) ve Yatchew (2000)'deki gibi parametrik tahminler y 'den arındırılarak parametrik olmayan fonksiyonun analizi yapılabilir. Bunun için $(y_i - x_i \hat{\beta}_{diff}, t_i)$, $i = 1, \dots, n$ veri seti oluşturulur ve standart parametrik olmayan yöntemlerle f 'nin analizi yapılabilir. Bu yöntemlerin geçerli olmasının nedeni daha önce belirtildiği gibi $\hat{\beta}_{diff}$ 'in β 'ya yeteri kadar hızlı yakınsamasıdır. Bu örnek için parametrik etkiyi tahmin etmek için $l = 0,85$ için farka dayalı Liu-tipi tahmin edici kullanılacaktır. Şekil 8.9'da

sıralı $(y_i - x_i \hat{\beta}_{diff}(l = 0,85), t_i)$ veri çiftleri ve f 'nin kernel tahmini ve ilgili güven bantları verilmiştir.



Şekil 8.9. $y_i - x_i \hat{\beta}_{diff}(l = 0,85)$ 'in ms değişkenine karşı grafiği, doğrusal model (yeşil), kernel regresyonu (kırmızı), %95 güven bantları (mavi)

Şekil 8.6'ya göre *birim maliyet* ile *müşteri sayısı* arasında önce azalan daha sonra sabitleşen bir ilişki vardır. Müşteri sayısı arttıkça birim maliyetin düşmesi sonucu ölçek ekonomisi modelinin beklentisini de karşılamaktadır.

9. SONUÇ VE ÖNERİLER

Parametrik ve semiparametrik regresyon modellerinde çoklu bağlantı olması tahmin edicilerin etkinliğini ve geçerliliğini etkilemektedir. Bu çalışmada çoklu parametrik ve semiparametrik regresyon modellerinde çoklu bağlantı olması durumunda kullanılabilecek yanlı tahmin edicilerden önerilmiştir.

Çoklu bağlantının sonuçlarını ortadan kaldırmak için literatürde birçok tahmin edicinin olması bunların karşılaştırılma gerekliliğini doğurmuştur. Bu çalışmada önerilen tüm tahmin tahmin ediciler literatürde sıklıkla kullanılan diğer tahmin edicilerle hata kare ortalama matrisi ölçütüne göre karşılaştırılmıştır. Önerilen tahmin edicilerin literatürde var olan tahmin edicilere önemli bir alternatif olduğu gösterilmiştir.

Parametrik regresyon modellerinde tahminlerin etkinliğini arttırmak için parametreler hakkındaki önsel bilgileri de tahmin sürecine katmak literatürde sıklıkla kullanılan bir yöntemdir. Bu çalışmada çoklu bağlantı durumunda sıklıkla kullanılan Liu tahmin edicisine eşitlik şeklinde bir kısıt eklenerek yeni bir kısıtlı Liu tahmin edicisi önerilmiştir. Bu tahmin edici kısıtlı en küçük kareler yöntemi ile teorik olarak karşılaştırılarak kısıtlı EKK tahmin edicisine HKOM ölçütüne göre üstün olduğu teorik bölge verilmiştir.

Liu tahmin edicisi yanlı bir tahmin edicidir. Bu çalışmada ikinci olarak Liu tahmin edicisine jackknife yöntemi uygulanarak yanlılık terimi azaltılmıştır ve bu tahmin edicinin Akdeniz ve Kaçıranlar (1995) tarafından farklı bir yöntemle elde edilen hemen hemen yansız Liu-tipi tahmin edici ile aynı olduğu görülmüştür. Bu jackknife tahmin edicide tahmin edicideki EKK yerine Liu-tipi tahmin edici koyularak yeni bir jackknife Liu-tipi tahmin edici elde edilmiştir. Yeni jackknife Liu-tipi tahmin edici jackknife Liu ve genelleştirilmiş Liu tahmin edicileri ile HKOM ve mutlak yanlılık terimi ölçütüne göre karşılaştırılmış ve birbirlerine üstün olduğu yerler verilmiştir. Simülasyon çalışmasında önerilen yeni jackknife tahmin edicisinin örnek çapının 50

veya daha küçük olduğu durumlarda ve güçlü çoklu bağlantı olması durumunda diğer tahmin edicilerden üstünlüğü gösterilmiştir.

Doğrusal regresyon modellerinin literatürde sıklıkla kullanılmasının temel nedeni basit olmalarıdır. Ancak basit olması gerçekçi olduğu anlamına gelmez. Birçok durumda açıklayıcı değişkenler ve yanıt değişkeni arasındaki fonksiyonel ilişki tam olarak bilinmediğinden parametrik olmayan regresyon modelleri bazı durumlarda daha gerçekçi sonuçlar vermektedir. Parametrik olmayan regresyon modellerinde açıklayıcı değişken sayısı arttıkça tahminlerin gerçeğe yakınsama hızının düşmesi problemi olduğundan alternatif olarak semiparametrik regresyon modelleri önerilmiştir. Semiparametrik regresyon modelleri parametrik ve parametrik olmayan fonksiyonların bir kombinasyonudur. Parametrik regresyon modellerinde olduğu gibi semiparametrik regresyon modellerinde de çoklu bağlantı durumu ile sıklıkla karşılaşmaktadır. Çoklu bağlantı olması durumu semiparametrik modellerde de tahminlerin etkinliğini ve güvenilirliğini azaltır. Bu çalışmada bu problemi çözmek amacıyla semiparametrik regresyon modelindeki parametrik değişkenler arasında çoklu bağlantı olması durumunda kullanılabilecek tahmin ediciler önerilmiştir. İlk olarak, semiparametrik regresyon modellerinde sıklıkla kullanılan Speckman ve cezalı EKK tahmin edicilerine dayalı yeni Liu-tipi tahmin ediciler üretilmiş ve bu tahmin ediciler teorik olarak HKOM ölçütüne göre karşılaştırılmıştır. Simülasyon çalışmasında Liu-tipi tahmin edicilerin SHKO ölçütüne göre üstünlüğü gösterilmiştir. Güçlü çoklu bağlantı olması durumunda bu üstünlük daha belirgindir.

Son olarak semiparametrik regresyon modelinde parametrik değişkenler arasında çoklu bağlantı olması durumunda alternatif olarak kullanılabilecek farka dayalı Liu-tipi tahmin edici verilmiştir. Farka dayalı yöntemin avantajı fark alarak parametrik olmayan değişken sanki modelde hiç yokmuş gibi kolaylık sağlaması ve klasik parametrik olmayan regresyonda kullanılan yöntemlerin kullanılabilmesini sağlamaktır. Veri matrislerinin farkı alındıktan sonra da çoklu bağlantı problemi devam ettiğinden bu durumda da yanlış tahmin edicileri kullanmanın yerinde olacağı düşünülmüştür. Simülasyon çalışmasında farka dayalı Liu-tipi tahmin edicinin

Teorem 7.1'deki koşulun sağlandığı durumlarda diğer tahmin edicilerden SHKO bakımından üstün olduğu gösterilmiştir.

Son bölümde bu çalışmada elde edilen tahmin edicilerle ilgili birer uygulama ve simülasyon çalışması yapılmıştır. Tahmin edicilerin teorik olarak birbirlerine üstün olduğu yerler uygulamalarda gösterilmiştir. Simülasyon çalışmalarında çoklu bağlantı olması durumunda Liu-tipi tahmin edicilerin daha küçük SHKO'ya sahip oldukları görülmektedir.

Bu çalışmada parametrik ve semiparametrik regresyon modellerinde çoklu bağlantı olması durumunda tahmin edicilerin iyileştirilmesine yönelik bulgular elde edilmiştir. İlerideki çalışmalarda hata terimlerinin varyansının değişen varyanslı ve birbiriyle bağlantılı (autocorrelated) olduğu durumlar için tahmin ediciler geliştirilebilir.

KAYNAKLAR

Ahmet, S. E., Doksum, K. A., Hossain, S. ve You, J., “Shrinkage, pretest and absolute penalty estimators in partially linear models”, *Aust. N. Z. J. Stat.*, 49 (4): 435-454 (2007).

Akdeniz, F. ve Kaçıranlar, S., “On the almost unbiased ridge generalized Liu estimator and unbiased estimation of the bias and mse”, *Comm. Statist. Theory Methods*, 24 (7): 1789-1797 (1995).

Akdeniz, F., Erol, H. ve Öztürk, F., "MSE comparisons of some biased estimators in linear regression model", *Journal of Applied Statistical Science*, 9 (1): 73-85 (1999).

Akdeniz, F., “The examination and analysis of residuals for some biased estimators in linear regression”, *Comm. Statist. Theory Methods*, 30 (6): 1171-1183 (2001).

Akdeniz, F., “New biased estimators under the LINEX loss function”, *Statistical Papers*, 45 (2): 175-190 (2004).

Akdeniz, F. ve Erol, H., “Mean squared error comparisons of some biased estimators in linear regression”, *Commun. Statist. Theory Methods*, 32 (12), 2391-1415 (2003).

Akdeniz, F., Styan, G. P. ve Werner, H. J., “The general expressions for the moments of the stochastic shrinkage parameters of the Liu-type estimator”, *Comm. Statist. Theory Methods*, 35: 423-437 (2006).

Akdeniz, F., Wan, A.T. K. ve Akdeniz, E., “Generalized Liu-type estimators under Zellner’s balanced loss function”, *Comm. Statist. Theory Methods*, 34: 1725-1736 (2005).

Akdeniz, F. ve Öztürk, F., “The distribution of stochastic shrinkage biasing parameters of the Liu-type estimator”, *Applied Mathematics and Computation*, 163 (1): 29-38 (2005).

Akdeniz, F. , Akdeniz Duran, E., “Liu-type estimator in semiparametric regression models”, *Journal of Statistical Computation and Simulation*, 80(8): 853-871 (2010).

Alheety, M. I., Ramanathan, T. V. ,Gore, S. D.,“On the distribution of stochastic shrinkage parameters of Liu-type estimators”, *Brazilian Journal of Probability and Statistics*, 23(1): 57-67 (2009).

Aydın, D., "Semiparametrik regresyon modellemede splayn düzeltme yöntemi ile tahmin ve çıkarsamalar", Doktora tezi, **Anadolu Üniversitesi**, 23-35 (2005).

Baksalary, J. K., Kala, R., "Partial orderings between matrices one of which is of rank one", **Bull. Polish Acad. Sci. Math.**, 31: 5–7 (1983).

Batah, F. S. M., Ramanathan, T. V. ve Gore, S. D., "The efficiency of modified jackknife and ridge type regression estimators: a comparison", **Surveys in Mathematics and its Applications**, 3: 111-122 (2008).

Belsley, D. A., Kuh, E. and Welsch, R. E., "Regression Diagnostics", **Wiley, New York**, 30-50 (1980).

Bianco, A. ve Boente, G., "Robust estimators in semiparametric partly linear regression models", **Journal of Statistical Planning and Inference**, 122: 229-252, (2004).

Buja, A., Hastie, T. ve Tibshirani, R., "Linear smoothers and additive models", **The Annals of Statistics**, 17(2): 453-510 (1989).

Chatterjee, S. ve Hadi, A. S., "Sensitivity Analysis in Linear Regression", **Wiley, New York**, 65-73 (1988).

Chung, Y. ve Kim, C., "Simultaneous estimation of the multivariate normal mean under balanced loss function", **Commun. Statist-Theory Meth.**, 26: 1599–1611 (1997).

Dipak, D.K., Ghosh M. ve Strawderman, W.E., "On estimation with balanced loss functions", **Statistics and Probability Letters**, 45 (2): 97-102, (1999).

Engle, R. F., Granger, C. W. J., Rice, J. ve Weiss, A., "Semiparametric estimates of the relation between weather and electricity sales", **J. Amer. Statist. Assoc.** 81: 310-320 (1986).

Eubank, R. L., "Nonparametric regression and spline smoothing", **Marcel Dekker, New York**, 50-60 (1999).

Eubank, R. L., Kambour, E. L., Kim, J. T., K. Klipple, K., Reese, C. S., Schimek, M. G., "Estimation in partially linear models", **Computational Statistics and Data Analysis**, 29: 27-34 (1998).

Farebrother, R. W., "Further results on the mean square error of ridge regression", **Journal of the Royal Statistical Society B**, 38: 248-250 (1976).

Farrar, D.E. ve Glauber, R.R., "The problem of multicollinearity: The problem revisited", **Review of Economics and Statistics**, 49: 92-107 (1967).

Foucart, T., “Stability of the Inverse Correlation Matrix. Partial Ridge Regression”, *Journal of Statistical Planning and Inference*, 77: 141-154 (1999).

Giles, J.A. ve Giles, D.E.A., “Pre-Test Estimation and Testing in Econometrics: Recent Developments”, *Journal of Economic Surveys*, 7 (2): 145-197 (1993).

Golan, A., Judge, G. Ve Miller, D., “Maximum Entropy Econometrics: Robust Estimation with Limited Data”, *John Wiley & Sons, New York*, 64-73 (1996).

Graybill, F., “Matrices with Applications in Statistics”, *Wadsworth Publishing Company, California*, 85-100 (1983).

Green, P.J. ve Silverman, B.W., “Nonparametric Regression and Generalized Linear Models”, *Chapman and Hall*, New York, 25-35 (1994).

Gross, J., “Linear Regression”, *Springer-Verlag*, Berlin, 23 (2003).

Gruber, M. H. J., “The Efficiency of Shrinkage Estimators with respect to Zellner's Balanced Loss Function”, *Communications in Statistics-Theory and Methods*, 33(2): 235-250 (2004).

Härdle, W., Liang, H. ve Gao, J., “Partially Linear Models”, *Springer-Verlag Co.*, New York, 32-38 (2000).

Härdle, W., Müller, M., Sperlich, S. ve Werwatz, A., “Nonparametric and Semiparametric Models”, *Springer*, New York, 55-75 (2004).

Hinkley, D.V., “Jackknifing in Unbalanced Situations”, *Technometrics* 19: 285-292 (1977).

Hoerl, A.E. ve Kennard, R.W., “Ridge Regression: Biased Estimation for Non-orthogonal Problems”, *Technometrics*, 12: 55-67 (1970).

Hu, H., “Ridge Estimation of a Semiparametric Regression Model”, *Journal of Computational and Applied Mathematics*, 176: 215-222 (2005).

Hubert, M.H. ve Wijekoon, P., “Improvement of the Liu Estimator in Linear Regression Model”, *Statistical Papers*, 47: 471-479 (2006).

Kaçıranlar, S., Sakallıoğlu, S., Akdeniz, F., Styan, G.P.H. ve Werner, H.J., “A New Biased Estimator in Linear Regression and a Detailed Analysis of the Widely Analyzed Data Set on Portland Cement”, *Sankhya: Indian Journal of Stat. B*, 61(3): 443-459 (1999).

Kadiyala, K., “A Class of Almost Unbiased and Efficient Estimators of Regression Coefficients”, *Economic Letters*, 16: 293-296 (1984).

- Liang, H., "Estimation in Partially Linear Models and Numerical Comparisons", *Computational Statistics and Data Analysis*, 50: 675-687 (2006).
- Liu, K., "A New Class of Biased Estimate in Linear Regression", *Communications in Statistics- Theory Methods*, 22: 393-402 (1993).
- Liu, K., "Using Liu-type Estimator to Combat Collinearity", *Communications in Statistics- Theory and Methods*, 32 (5): 1009-1020 (2003).
- Liu, K., "More on Liu-type Estimator in Linear Regression", *Communications in Statistics- Theory and Methods*, 33: 2723-2733 (2004).
- McDonald, G.C. ve Galarnaeu, D. I., "A Monte Carlo Evaluation of Some Ridge-type Estimators", *Journal of the American Statistical Association*, 20: 407-416 (1975).
- Massy, W.F., "Principal Components Regression in Exploratory Statistical Research", *JASA*, 60: 234-260 (1965).
- Montgomery, D.C. ve Peck, E., "Introduction to Linear Analysis", *John Wiley and Sons*, New York, 62-77 (1992).
- Nyquist, H., "Applications of The Jackknife Procedure in Ridge Regression". *Computational Statistics & Data Analysis*, 6: 177-183 (1988).
- Ohtani, K., "On Small Sample Properties of the Almost Unbiased Generalized Ridge Estimator", *Commun. Statist.-Theor. Meth.*, 15(5): 1571-1578 (1986).
- Ohtani, K., "Generalized Ridge Regression Estimator Under The Linex Loss Function", *Statistical Papers*, 36: 99-110 (1995).
- Ohtani, K., "Risk Performance of a Pre- Test Estimator for Normal Variance with the Stein-Variance Estimator Under The Linex Loss Function", *Statistical Papers*, 40: 75-87 (1999).
- Przystalski, M. ve Krajewski, P., "Constrained Estimators of Treatment Parameters in Semiparametric Models", *Statistics and Probability Letters*, 77 (10): 914-919 (2007).
- Quenouille, M., "Approximate Tests of Correlation in Time-series", *J.R.Statist. Soc. B*, 11: 68-84 (1949).
- Quenouille, M., "Notes on Bias in Estimation", *Biometrika*, 43: 353-360 (1956).
- Rao, C. R., "Linear Statistical Inference and Its Application", *John Wiley and Sons, Inc.* (2002).

Rao, C. R., Toutenburg, H., Shalabh, Heumann, C.,. “Linear Models and Generalizations: Least Squares and Alternatives”, **Springer Verlag**, Berlin, 25 (2008).

Rice, J. A.,. “Convergence Rates for Partially Splined Models”, **Statistics and Probability Letters**, 4: 203-208 (1986).

Ruppert, D., Wand, M. P. ve Carroll, R. C.,. “Semiparametric Regression”, **Cambridge**, 43-64 (2003).

Sarkar, N.,. “A New Estimator Combining the Ridge Regression and the Restricted Least Squares Methods of Estimation”, **Comm. Statist. Theory Methods**, 21: 1987-2000 (1992).

Schaffrin, B., Toutenburg, H.,. “Weighted Mixed Regression”, **Zeitschrift für Angewandte Mathematik und Mechanik**, 70: 735-738 (1990).

Schimek, M. G.,. “Estimation and Inference In Partially Linear Models with Smoothing Splines”, **Journal of Statistical Planning and Inference**, 91: 525-540 (2000).

Singh, B., Chaubey, Y. P. ve Dwivedi, T. D.,. “An Almost Unbiased Ridge Estimator”, **Sankhya: The Indian Journal of Statistics, Series B**, 48: 342-346 (1986).

Speckman, P.,. “Kernel Smoothing In Partial Linear Models”, **J. Roy. Statist. Soc. Ser. B.**, 50 (3): 413-436 (1988).

Stein, C.,. “Inadmissibility of the Usual Estimator for Mean of Multivariate Normal Distribution”, **In: Neyman J (ed) Proceedings of the third Berkeley symposium on mathematical statistics and probability**, 1: 197-206 (1956).

Swindel, B.F.,. “Good Estimators Based On Prior Information”, **Comm. Statist. Theory and Methods**, 5: 1065-1075 (1976).

Tabakan, G. ve Akdeniz, F.,. “Difference based ridge estimator of parameters in partial linear model”, **Stat Papers**, 51: 357-368 (2010).

Thiel, H. ve Goldberger, A.S.,. “On Pure and Mixed Estimation in Economics”, **International Economic Review**, 2: 65-77 (1961).

Tibshirani, R.,. “Regression shrinkage and selection via the lasso”, **J. Royal Statist. Soc. B.**, 58 (1): 267-288 (1996).

Torigoe, N. ve Ujiie, K.,. “On the Restricted Liu Estimator in Gauss Markov Model”, **Comm. Statist. Theory and Methods**, 35 (9): 1713-1722 (2006).

Tukey, J.,. “Bias and Confidence in not Quite Large Samples”, *Annals of Statistics*, 29: 614 (1958).

Varian, H. R.,. “ A Bayesian approach to real estate assessment”, *In Studies in Bayesian Econometrics and Statistics in honor of Leonard J. Savage, eds. Stephan E. Fienberg and Arnold Zellner*, Amsterdam: North-Holland, 195-208 (1975).

Vinod, H. D.,. “A Survey for Ridge Regression and Related Techniques for Improvements over Ordinary Least Squares”, *The Review of Econometrics and Statistics*, 60: 121-131 (1978).

Vinod, H. D. ve Ullah, A.,. “Recent Advances in Regression Methods”, *Marcel Dekker Inc.*, New York, 30-45 (1981).

Wan, A.T. K.,. “Risk comparison of the inequality constrained least squares and the related estimators under balanced loss”, *Economics Letters*, 46: 203-210 (1994).

Wan, A. T. K.,. “A Note On Almost Unbiased Generalized Ridge Regression Estimator Under Asymmetric Loss”, *Journal of Statistical Computation and Simulation*, 62: 411-421 (1999).

Wan, A. T. K.,. “On Generalized Ridge Regression Estimators Under Collinearity and Balanced Loss”, *Applied Mathematics and Computation*, 129: 455-467 (2002).

Weisberg, S.,. “Applied Linear Regression”, 2nd edition, *Wiley*, New York, 43-51 (1985).

Wen, D. ve Levy,M.S.,. “Blinex: A Bounded Asymmetric Loss Function with Application To Bayesian Estimation”, *Comm. Statistics-Theory Methods*, 30 (1): 147-153 (2001).

Yang, H. ve Zhang, C.,. “The Research On Two Kinds of Restricted Biased Estimators Based On Mean Squared Error Matrix”, *Comm. Statist. Theory Methods*, 37: 70-80 (2008).

Yang, H. ve Xu, J.,. “An Alternative Stochastic Restricted Liu Estimator in Linear Regression”, *Statistical Papers*, 50: 639-647 (2009).

Yang, H., Chang, X. ve Liu, D.,. “Improvement of the Liu Estimator in Weighted Mixed Regression”, *Commun. Statistics-Theory and Methods*, 38: 285-292 (2009).

Yang H. ve Xu J.,. “Tuning Parameter Selection and Various Good Fitting Characteristics for the Liu-Type Estimator in Linear Regression”, *Commun. Statistics-Theory and Meth.*, 37: 3204-3215 (2008).

Yatchew, A.,. “Scale Economies in Electricity Distribution”, *Journal of Applied Economics*, 15: 187-210 (2000).

Yatchew, A.,. “Semiparametric Regression for the Applied Econometrician”, *Cambridge University Press* , 80-100 (2003).

Zellner, A.,. “Bayesian Estimation and Prediction Using Asymmetric Loss Functions”, *J. Amer. Statist. Assoc.*, 81: 446–451 (1986).

Zellner, A.,. “Bayesian and Non-Bayesian Estimation Using Balanced Loss Functions”, S. Gupta J. O. Berger (Eds.), *Statistical Decision Theory and Related Topics*. vol. V, *Springer*, New York, 377-390 (1994).

Zhong, Z. ve Yang, H.,. “Ridge Estimation to the Restricted Linear Model”, *Commun. Statistics-Theory and Meth.*, 36: 2099-2115 (2007).

EKLER

EK-1 Jackknife tahmin edicilerinin elde edilmesi

$y = X\beta + \varepsilon$ modelinde β 'nın EKK tahmin edicisi $\hat{\beta} = (X^T X)^{-1} X^T y$ olarak bulunmuştur. x_i^T , X matrisinin i . satırı olmak üzere (x_i^T, y_i) , i . gözleminin çıkarılmasıyla elde edilen X matrisi ve y vektörü X_{-i} ve y_{-i} ile gösterilsin. Bu durumda

$\hat{\beta}_{-i} = (X_{-i}^T X_{-i})^{-1} X_{-i}^T y_{-i}$ olur. Burada $X_{-i}^T X_{-i} = X^T X - x_i x_i^T$ ve $X_{-i}^T y_{-i} = X^T y - x_i y_i$ olduğundan

$\hat{\beta}_{-i} = (X^T X - x_i x_i^T)^{-1} (X^T y - x_i y_i)$ yazılabilir.

Teorem 5.1'den

$$(X^T X - x_i x_i^T)^{-1} = (X^T X)^{-1} + \frac{(X^T X)^{-1} x_i x_i^T (X^T X)^{-1}}{1 - x_i^T (X^T X)^{-1} x_i}$$

$w_i = x_i^T (X^T X)^{-1} x_i$ olmak üzere

$$\begin{aligned} \hat{\beta}_{-i} &= [(X^T X)^{-1} + \frac{(X^T X)^{-1} x_i x_i^T (X^T X)^{-1}}{1 - x_i^T (X^T X)^{-1} x_i}] [X^T y - x_i y_i] \\ &= (X^T X)^{-1} X^T y + \frac{(X^T X)^{-1} x_i x_i^T (X^T X)^{-1} X^T y}{1 - w_i} - (X^T X)^{-1} x_i y_i - \frac{(X^T X)^{-1} x_i x_i^T (X^T X)^{-1} x_i y_i}{1 - w_i} \end{aligned}$$

$w_i (= x_i^T (X^T X)^{-1} x_i)$ skaler bir değer olduğundan aşağıdaki gibi matris çarpımının başına alınabilir

$$\hat{\beta}_{-i} = \hat{\beta} + \frac{(X^T X)^{-1} x_i x_i^T \hat{\beta}}{1 - w_i} - (X^T X)^{-1} x_i y_i - \frac{w_i (X^T X)^{-1} x_i y_i}{1 - w_i}$$

EK-1 (Devam) $\hat{\beta}_{-i}$ 'nin elde edilmesi

$$\begin{aligned}
 &= \hat{\beta} + \frac{(X^T X)^{-1} x_i x_i^T \hat{\beta} - (1 - w_i)(X^T X)^{-1} x_i y_i - w_i (X^T X)^{-1} x_i y_i}{1 - w_i} \\
 &= \hat{\beta} - \frac{(X^T X)^{-1} x_i (y_i - x_i \hat{\beta})}{1 - w_i}
 \end{aligned}$$

$r_i = y_i - x_i \hat{\beta}$, i . gözlemin artık terimi olmak üzere

$$\hat{\beta}_{-i} = \hat{\beta} - \frac{(X^T X)^{-1} x_i r_i}{1 - w_i}$$

bulunur.

$\hat{\beta}_{-i}(D)$ 'nin elde edilmesi

$$\begin{aligned}
 \hat{\beta}(D)_{-i} &= (A - z_i z_i^T)^{-1} (Z^T y - z_i y_i) = (A^{-1} + \frac{A^{-1} z_i z_i^T A^{-1}}{1 - z_i^T A^{-1} z_i}) (Z^T y - z_i y_i), \\
 &= A^{-1} Z^T y - A^{-1} z_i y_i + \frac{A^{-1} z_i z_i^T A^{-1}}{1 - z_i^T A^{-1} z_i} Z^T y - \frac{A^{-1} z_i z_i^T A^{-1} z_i y_i}{1 - z_i^T A^{-1} z_i}, \\
 &= \hat{\beta}(D) - A^{-1} z_i y_i (1 + \frac{z_i^T A^{-1} z_i}{1 - z_i^T A^{-1} z_i}) + \frac{A^{-1} z_i z_i^T \hat{\beta}(D)}{1 - z_i^T A^{-1} z_i} \\
 &= \hat{\beta}(D) - \frac{A^{-1} z_i (y_i - z_i^T \hat{\beta}(D))}{1 - z_i^T A^{-1} z_i} \\
 &= \hat{\beta}(D) - \frac{A^{-1} z_i e_i}{1 - w_i}
 \end{aligned}$$

ÖZGEÇMİŞ

Kişisel Bilgiler

Soyadı, adı : AKDENİZ DURAN, Esra
 Uyruğu : T.C.
 Doğum tarihi ve yeri : 01.07.1980 Ankara
 Medeni hali : Evli
 Telefon : 0 (312) 202 14 82
 Faks : 0 (312) 222 14 79
 e-mail : esraakdeniz@gmail.com.

Eğitim Derece

Eğitim Birimi

Mezuniyet tarihi

Yüksek Lisans	Penn State Üniversitesi / İstatistik Bölümü	2005
Lisans	Orta Doğu Teknik Üniversitesi/ İstatistik Bölümü	2002
Lise	Kurttepe Anadolu Lisesi	1998

İş Deneyimi Yıl

Yer

Görev

2005-	Gazi Üniversitesi	Araştırma Görevlisi
2002-2003	Çukurova Üniversitesi	Araştırma Görevlisi

Yabancı Dil İngilizce

Yayınlar

1. Akdeniz, F., Wan, A.T.K., Akdeniz, E.,. “Generalized Liu-type estimators under Zellner’s balanced loss function”, *Comm. Statist. Theory Methods*, 34: 1725-1736 (2005).

2. Akdeniz, E., Akdeniz, F., “ Lawless Wang’s operational ridge regression estimator under the LINEX loss function”, ***Advances on Models, Characterizations and Applications***, Chapman and Hall, 201-213 (2005). (kitapta 13. bölüm)
3. Akdeniz, E., Akritas, M.G., “Wald-Type statistics in heteroscedastic linear models”, ***Selçuk Journal of Applied Mathematics***, 7 (2): 91-101 (2006).
4. Akdeniz, E., Akdeniz, F., Wan, A.T.K., Chen, T., “Further results on the generalized Liu-type estimators under the balanced loss function”, ***Model Assisted Statistics and Applications***, 2: 213-223 (2007).
5. Akdeniz, F. , Akdeniz Duran, E.,. “Liu-type estimator in semiparametric regression models”, ***Journal of Statistical Computation and Simulation***, 80(8): 853-871 (2010).
6. Akdeniz Duran, E., Akdeniz, F.,. “Efficiency of the modified jackknifed Liu-type estimator”, ***Statistical Papers***, DOI: 10.1007/s00362-010-0334-5.
7. Akdeniz Duran, E., Akdeniz, F., Hu, H.,. “Efficiency of a Liu-type estimator in semiparametric regression models”, ***Journal of Computational and Applied Mathematics***, 235 (5): 1418-1428 (2011).
8. Akdeniz Duran, E., Guo M., Härdle, K.W.,. “A Confidence Corridor for Expectile Functions”, ***Berlin Humboldt University SFB 649 Discussion Paper***, 004 (2011).