



**NORMAL DAĞILMAYAN ÇOK DEĞİŞKENLİ SÜREÇ YETERLİLİK İNDEKSİ
İÇİN BİR YÖNTEM ÖNERİSİ: KMH YÖNTEMİ**

Burcu DARENDE ŞİMŞEK

**DOKTORA TEZİ
İSTATİSTİK ANA BİLİM DALI**

**GAZİ ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ**

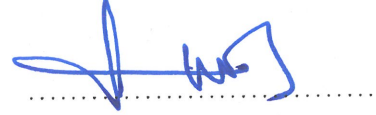
OCAK 2019

Burcu DARENDE ŞİMŞEK tarafından hazırlanan "NORMAL DAĞILMAYAN ÇOK DEĞİŞKENLİ SÜREÇ YETERLİLİK İNDEKSİ İÇİN BİR YÖNTEM ÖNERİSİ: KMH YÖNTEMİ" adlı tez çalışması aşağıdaki jüri tarafından OY BİRLİĞİ ile Gazi Üniversitesi Fen Bilimleri Enstitüsü İstatistik Ana Bilim Dalında DOKTORA TEZİ olarak kabul edilmiştir.

Danışman: Prof. Dr. M. Akif BAKIR

İstatistik Ana Bilim Dalı, Gazi Üniversitesi

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



Başkan: Prof. Dr. Hamza GAMGAM

İstatistik Ana Bilim Dalı, Gazi Üniversitesi

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



Üye : Doç. Dr. Özlem Müge AYDIN TESTİK

Endüstri Mühendisliği Ana Bilim Dalı, Hacettepe Üniversitesi

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



Üye : Doç. Dr. Rukiye DAĞALP

İstatistik Ana Bilim Dalı, Ankara Üniversitesi

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



Üye : Doç. Dr. Filiz KARDİYEN

İstatistik Ana Bilim Dalı, Gazi Üniversitesi

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



Tez Savunma Tarihi : 17.01.2019

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

.....

Prof. Dr. Sena YAŞYERLİ

Fen Bilimleri Enstitüsü Müdürü

ETİK BEYAN

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

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

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


Burcu DARENDE ŞİMŞEK

17.01.2019

NORMAL DAĞILMAYAN ÇOK DEĞİŞKENLİ SÜREÇ YETERLİLİK İNDEKSİ
İÇİN BİR YÖNTEM ÖNERİSİ: KMH YÖNTEMİ
(Doktora Tezi)

Burcu DARENDE ŞİMŞEK

GAZİ ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ
Ocak 2019

ÖZET

Süreç, bir ürün ya da hizmet üretimi sağlayan işletmelerde, istenen özelliklerde ürün üretmek için gerekli makine, işgücü, donanım ve malzeme gibi belirli girdileri içeren ve sonucunda bir çıktıya dönüştüren, tanımlanabilir, ölçülebilir ve birbirine bağlı değer yaratan faaliyetler dizisi olarak tanımlanabilir. Süreç Yeterliliği ise, en temel anlamda süreçten elde edilen ürünlerin istenen sınırlar içinde kalabilme yeteneğinin bir ölçüsüdür. Süreç yeterlilik araştırmasında kullanılan süreç yeterlilik indekslerinin tahmin edilmesi için bazı varsayımların sağlanması gerekir. Bu varsayımlardan en önemlisi sürecin normal dağılıma uygun olmasıdır. Klasik yeterlilik indekslerinin hesaplaması varsayım bozulmalarından olumsuz etkilendiğinden, hatalı tahminler elde edilir. Normallik varsayımının sağlanmadığı durumlar için süreç yeterlilik indekslerinin hesaplanmasında sınırlı sayıda çeşitli yöntemler mevcuttur. Bu çalışmada, çok değişkenli ve normal dağılım varsayımını sağlamayan süreçlerin yeterlilik analizinde kullanılmak üzere *KMH* yöntemi olarak adlandırılan yeni bir yaklaşım önerilmiştir. Bu amaçla her biri üretim sürecini temsil eden çarpık dağılım özelliklerine sahip 4 farklı senaryo verisi t-Copula yöntemi kullanılarak üretilmiştir. Çalışmada, kurgulanan senaryolar üzerinden normallik varsayımını sağlamayan süreçlerin yeterlilik araştırmasında kullanılan bazı yeterlilik indekslerinin uygulamalarına yer verilmiş ve bu indekslere alternatif olarak önerilen *KMH* yöntemine dayalı yeni bir yeterlilik indeksi tasarlanmıştır. Önerilen *KMH* süreç yeterlilik indeksi ile literatürde kullanılan bazı yaklaşımların kitle uygun ürün oranı p 'yi tahmin etmedeki hata yüzdeleri elde edilmiş, böylece önerilen yaklaşımın süreç yeterliliğine karar vermedeki başarısı irdelenmiştir. Bulgular, çalışmada önerilen indeksin, yine çalışmada yer verilen diğer yeterlilik indekslerine göre uygun ürün oranı p 'yi daha az hata ile tahmin ettiğini ortaya koymuştur. Çalışmada önerilen yaklaşımın özellikle örneklem hacminin küçük olduğu durumlarda daha iyi sonuçlar verdiği, örneklem hacmi büyüdükçe yaklaşımdan elde edilen hata yüzdelerinin diğer yaklaşımlara yakın olduğu görülmüştür.

Bilim Kodu : 20508
Anahtar Kelimeler : Çok Değişkenli Süreç Yeterlilik İndeksleri, Kernel Yoğunluk Tahmin Yöntemi, Metropolis Hastings Örnekleme
Sayfa Adedi : 176
Danışman : Prof. Dr. M. Akif BAKIR

A PROPOSED METHOD FOR NONNORMAL MULTIVARIATE PROCESS
CAPABILITY INDEX: KMH METHOD
(Ph. D. Thesis)

Burcu DARENDE ŞİMŞEK

GAZİ UNIVERSITY
GRADUATE SCHOOL OF NATURAL AND APPLIED SCIENCES
January 2019

ABSTRACT

A process can be defined as a series of activities which create definable, measurable and interconnected values that contain specific inputs such as machinery, labor, equipment and materials required to produce a product. Process capability, in the most basic sense, is a measure of the capability of the products obtained as an output of the process to remain within the desired limits. In order to estimate the process capability indices used in process capability research, some assumptions must be satisfied. The most important assumption is that the process must have a normal distribution. Since the calculation of classical indices is adversely affected by disturbance in assumptions, incorrect estimations are made. There are a variety of methods for calculating process capability indices for cases when the assumption of normality is not satisfied. In this study, a new approach called *KMH* method has been proposed to be used in the capability analysis of processes that do not satisfy the assumption of multivariate normal distribution. For this purpose, four different data set which are skewed distribution characteristics are simulated by using t-Copula method regarding to different scenarios representing the production process. In the study, implementation of some available capability indices used in the capability analysis of the nonnormal processes over the conceived scenarios, and then as an alternative to these indices, a new capability index based on the proposed *KMH* method was defined. The proposed process capability index *KMH* and some of the existing approaches have been used and obtained error rates in the estimation of population conforming product ratio p have been obtained to examine the performance of proposed method. Findings show that the suggested index estimates the conforming product ratio p with smaller error percentages than the considered available capability indices. It is observed that while the proposed approach gives better results when the sample size is small, the error percentages are close to those of considered approaches as the sample size increases.

Science Code : 20508
Keywords : Multivariate Process Capability Indices, Kernel Distribution
Estimation, Metropolis Hastings Sampling
Number of pages : 176
Supervisor : Prof. Dr. M. Akif BAKIR

TEŞEKKÜR

Tez çalışmam boyunca bana her türlü yardımını, katkısını, bilgi ve deneyimini esirgemeyen, danışmanım Sayın Prof. Dr. Mehmet Akif BAKIR'a, Tez İzleme Komitesinde yer alarak tez çalışmam esnasında bana sağladığı katkı ve yönlendirmelerinden dolayı Sayın Prof. Dr. Hamza GAMGAM'a ve Sayın Doç. Dr. Özlem Müge AYDIN TESTİK'e ve onların şahsında üzerimde emeği olan tüm hocalarıma teşekkürü bir borç bilirim.

Tez çalışmam süresince anlayış, yardım ve katkılarından dolayı Başkent Üniversitesi İstatistik ve Bilgisayar Bölüm Başkanı Sayın Prof. Dr. İsmail ERDEM'e ve Sosyal Güvenlik Kurumu yöneticilerim ve çalışma arkadaşlarımdan her birine ayrı ayrı teşekkür ederim.

Tez metninin düzeltilmesinde ve Latex formatında yazılmasında emeği geçen Atacan ERDİŞ'e değerli katkılarından dolayı teşekkür ederim.

Desteklerini her zaman arkamda hissettiğim eşim Burak ŞİMŞEK'e ve tez çalışmam sırasında doğan ve kimi zaman sevgisinden ve ilgisinden feragat ederek tezimle birlikte büyüttüğüm sevgili kızım Beril ŞİMŞEK'e teşekkür ederim.

Son olarak, çalışma azmimi destekleyen ve motivasyon kaynağım olan sevgili annem, Hilmiye DARENDE'ye, sevgili babam Mehmet DARENDE'ye, tez çalışmam esnasında kızım Beril ŞİMŞEK'ten ilgilerini esirgemeyen annem Ülkü ŞİMŞEK ve Özlem OKAN'a ve onların şahsında tüm aileme ve bugünlere gelmemde üzerimde emeği olan tüm sevenlerime teşekkür ederim.

İÇİNDEKİLER

	Sayfa
ÖZET	iv
ABSTRACT	v
TEŞEKKÜR	vi
İÇİNDEKİLER	vii
ÇİZELGE LİSTESİ	ix
ŞEKİL LİSTESİ	xii
SİMGELER VE KISALTMALAR	xv
1. GİRİŞ	1
2. SÜREÇ YETERLİLİK İNDEKSLERİ	11
2.1. Tek Değişkenli Süreç Yeterlilik İndeksleri	13
2.1.1. Normal dağılıma uyan süreçlerde tek değişkenli süreç yeterlilik indeksleri	13
2.1.2. Normal dağılmayan süreçlerde tek değişkenli süreç yeterlilik indeksleri	40
2.2. Çok Değişkenli Süreç Yeterlilik İndeksleri	46
2.2.1. Çok değişkenli normal dağılılan süreçlerde süreç yeterlilik indeksleri	47
2.2.2. Grup I’de yer alan başlıca çalışmaların incelenmesi	47
2.2.3. Grup II’de yer alan başlıca çalışmaların incelenmesi	52
2.2.4. Grup III’te yer alan başlıca çalışmaların incelenmesi	56
2.2.5. Grup IV’te yer alan başlıca çalışmaların incelenmesi	58
2.2.6. Çok değişkenli normal dağılmayan süreçlerde süreç yeterlilik indeksleri	59
3. NORMAL DAĞILMAYAN ÇOK DEĞİŞKENLİ SÜREÇ YETERLİLİK İNDEKSİ İÇİN BİR YÖNTEM ÖNERİSİ: KMH YÖNTEMİ	67
3.1. Çalışmada Önerilen Yöntem ve Kullanılan Yaklaşımlar	67
3.2. Çalışma Verisi	72

	Sayfa
3.3. Çalışma Adımları	74
3.4. Çalışma Verilerinin Dağılım Özellikleri	76
4. SİMÜLASYON SONUÇLARI VE DEĞERLENDİRİLMESİ	89
4.1. Senaryo I İçin Elde Edilen Sonuçlar	89
4.2. Senaryo II İçin Elde Edilen Sonuçlar	103
4.3. Senaryo III İçin Elde Edilen Sonuçlar	113
4.4. Senaryo IV İçin Elde Edilen Sonuçlar	123
5. ÖRNEKLEM KÜMELERİNİN KERNEL TAHMİN GRAFİKLERİ	135
5.1. Senaryo I İçin Tahmin Grafikleri	135
5.2. Senaryo II İçin Tahmin Grafikleri	139
5.3. Senaryo III İçin Tahmin Grafikleri	142
5.4. Senaryo IV İçin Tahmin Grafikleri	145
6. SONUÇ VE ÖNERİLER	149
KAYNAKLAR	159
EKLER	165
EK-1	166
ÖZGEÇMİŞ	176

ÇİZELGE LİSTESİ

Çizelge	Sayfa
Çizelge 2.1. Kernel fonksiyon türleri	61
Çizelge 3.1. Veri setlerinin dağılım özellikleri	73
Çizelge 3.2. Bant genişliği matrisleri	74
Çizelge 4.1. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 50$, $\ddot{o}ks = 10$) . .	90
Çizelge 4.2. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 50$, $\ddot{o}ks = 20$)	90
Çizelge 4.3. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 100$, $\ddot{o}ks = 10$) .	93
Çizelge 4.4. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 100$, $\ddot{o}ks = 20$)	93
Çizelge 4.5. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 250$, $\ddot{o}ks = 10$) .	96
Çizelge 4.6. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 250$, $\ddot{o}ks = 20$)	96
Çizelge 4.7. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 500$, $\ddot{o}ks = 10$) .	98
Çizelge 4.8. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 500$, $\ddot{o}ks = 20$)	98
Çizelge 4.9. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 1000$, $\ddot{o}ks = 10$)	101
Çizelge 4.10. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 1000$, $\ddot{o}ks = 20$)	101
Çizelge 4.11. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 50$, $\ddot{o}ks = 10$) .	104
Çizelge 4.12. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 50$, $\ddot{o}ks = 20$)	104
Çizelge 4.13. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 100$, $\ddot{o}ks = 10$) .	106
Çizelge 4.14. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 100$, $\ddot{o}ks = 20$)	106
Çizelge 4.15. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 250$, $\ddot{o}ks = 10$) .	108
Çizelge 4.16. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 250$, $\ddot{o}ks = 20$)	108
Çizelge 4.17. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 500$, $\ddot{o}ks = 10$) .	110

Çizelge	Sayfa
Çizelge 4.18. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 500$, $\ddot{o}ks = 20$)	110
Çizelge 4.19. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 1000$, $\ddot{o}ks = 10$)	112
Çizelge 4.20. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 1000$, $\ddot{o}ks = 20$)	112
Çizelge 4.21. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 50$, $\ddot{o}ks = 10$) .	114
Çizelge 4.22. Senaryo III için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 50$, $\ddot{o}ks = 20$)	114
Çizelge 4.23. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 100$, $\ddot{o}ks = 10$)	116
Çizelge 4.24. Senaryo III için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 100$, $\ddot{o}ks = 20$)	116
Çizelge 4.25. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 250$, $\ddot{o}ks = 10$)	118
Çizelge 4.26. Senaryo III için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 250$, $\ddot{o}ks = 20$)	118
Çizelge 4.27. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 500$, $\ddot{o}ks = 10$)	120
Çizelge 4.28. Senaryo III için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 500$, $\ddot{o}ks = 20$)	120
Çizelge 4.29. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 1000$, $\ddot{o}ks = 10$)	122
Çizelge 4.30. Senaryo III için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 1000$, $\ddot{o}ks = 20$)	122
Çizelge 4.31. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 50$, $\ddot{o}ks = 10$) .	124
Çizelge 4.32. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 50$, $\ddot{o}ks = 20$)	124
Çizelge 4.33. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 100$, $\ddot{o}ks = 10$)	126
Çizelge 4.34. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 100$, $\ddot{o}ks = 20$)	126
Çizelge 4.35. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 250$, $\ddot{o}ks = 10$)	128
Çizelge 4.36. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 250$, $\ddot{o}ks = 20$)	128
Çizelge 4.37. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 500$, $\ddot{o}ks = 10$)	130

Çizelge	Sayfa
Çizelge 4.38. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 500$, $\overset{\circ}{oks} = 20$)	130
Çizelge 4.39. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 1000$, $\overset{\circ}{oks} = 10$)	132
Çizelge 4.40. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 1000$, $\overset{\circ}{oks} = 20$)	132



ŞEKİL LİSTESİ

Şekil	Sayfa
Şekil 1.1. Süreç yeterliliğine karar vermek için kullanılan dört temel adım	2
Şekil 1.2. Süreç yeterlilik indekslerini hesaplamak için kullanılan akış şeması	6
Şekil 2.1. Normal dağılım gösteren bir süreç için uygun olmama alanları	14
Şekil 2.2. Normal dağılım gösteren bir süreç için tolerans aralığı	15
Şekil 2.3. Tolerans bölgelerinin gösterimi	49
Şekil 2.4. Yeni tolerans bölgesinin gösterimi	51
Şekil 3.1. Ürün kitle saçılım grafiği (Senaryo I)	77
Şekil 3.2. İlk faz ürün saçılım grafiği (Senaryo I)	78
Şekil 3.3. İlk faz ürün kernel yoğunluk tahmini kontur grafiği (Senaryo I)	79
Şekil 3.4. Kernel yoğunluk tahmin grafiği (Senaryo I)	79
Şekil 3.5. Ürün kitle saçılım grafiği (Senaryo II)	80
Şekil 3.6. İlk faz ürün saçılım grafiği (Senaryo II)	81
Şekil 3.7. İlk faz ürün kernel yoğunluk tahmini kontur grafiği (Senaryo II)	81
Şekil 3.8. Kernel yoğunluk tahmin grafiği (Senaryo II)	82
Şekil 3.9. Ürün kitle saçılım grafiği (Senaryo III)	83
Şekil 3.10. İlk faz ürün saçılım grafiği (Senaryo III)	84
Şekil 3.11. İlk faz ürün kernel yoğunluk tahmini kontur grafiği (Senaryo III)	84
Şekil 3.12. Kernel yoğunluk tahmin grafiği (Senaryo III)	85
Şekil 3.13. Ürün kitle saçılım grafiği (Senaryo IV)	85
Şekil 3.14. İlk faz ürün saçılım grafiği (Senaryo IV)	86
Şekil 3.15. İlk faz ürün kernel yoğunluk tahmini kontur grafiği (Senaryo IV)	87
Şekil 3.16. Kernel yoğunluk tahmin grafiği (Senaryo IV)	87
Şekil 4.1. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 50, öks = 10$)	91
Şekil 4.2. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 100, öks = 10$)	94

Şekil	Sayfa
Şekil 4.3. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 250, öks = 10$)	97
Şekil 4.4. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 500, öks = 10$)	99
Şekil 4.5. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 1000, öks = 10$)	102
Şekil 4.6. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 50, öks = 10$)	105
Şekil 4.7. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 100, öks = 10$)	107
Şekil 4.8. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 250, öks = 10$)	109
Şekil 4.9. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 500, öks = 10$)	111
Şekil 4.10. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 1000, öks = 10$)	113
Şekil 4.11. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 50, öks = 10$)	115
Şekil 4.12. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 100, öks = 10$)	117
Şekil 4.13. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 250, öks = 10$)	119
Şekil 4.14. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 500, öks = 10$)	121
Şekil 4.15. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 1000, öks = 10$)	123
Şekil 4.16. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 50, öks = 10$)	125
Şekil 4.17. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 100, öks = 10$)	127
Şekil 4.18. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 250, öks = 10$)	129
Şekil 4.19. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 500, öks = 10$)	131
Şekil 4.20. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 1000, öks = 10$)	133
Şekil 5.1. Senaryo I, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (RBG)	136
Şekil 5.2. Senaryo I, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (RBG)	137
Şekil 5.3. Senaryo I, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (MDY)	138
Şekil 5.4. Senaryo I, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (MDY)	139
Şekil 5.5. Senaryo II, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (RBG)	140

Şekil	Sayfa
Şekil 5.6. Senaryo II, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (RBG)	140
Şekil 5.7. Senaryo II, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (MDY)	141
Şekil 5.8. Senaryo II, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (MDY)	142
Şekil 5.9. Senaryo III, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (RBG)	143
Şekil 5.10. Senaryo III, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (RBG)	143
Şekil 5.11. Senaryo III, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (MDY)	144
Şekil 5.12. Senaryo III, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (MDY)	145
Şekil 5.13. Senaryo IV, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (RBG)	146
Şekil 5.14. Senaryo IV, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (RBG)	146
Şekil 5.15. Senaryo IV, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (MDY)	147
Şekil 5.16. Senaryo IV, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (MDY)	148

SİMGELER VE KISALTMALAR

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

Kısaltmalar	Açıklamalar
AGS	Alt güven sınırı
AMISE	Asimptotik bütünleşik hata kareler ortalaması
CPI	Yeterlilik potansiyel indeksi
ISE	Bütünleşik hata kareler
İSK	İstatistiksel süreç kontrolü
KMH	Kernel yoğunluk tahminine dayalı Metropolis Hastings örnekleme
LCL	Alt kontrol sınırı
LPL	Pearson ailesinin 99,865 yüzdeliği
LSL	Alt spesifikasyon sınırı
MDY	Maksimum düzeltirme yöntemi
MISE	Bütünleşik hata kareler ortalaması
öks	Örnekleme kümelerinin sayısı
PCR	Süreç yeterlilik oranı
RBG	Referans bant genişliği
UMP	Düzgün en güçlü test
UMVUE	Düzgün en küçük varyanslı yansız tahmin edici
UPL	Pearson ailesinin 0,135 yüzdeliği
USL	Üst spesifikasyon sınırı

1. GİRİŞ

Herhangi bir ürün veya hizmet çıktısının işletmelerin ilgili birimlerince, müşteri istekleri doğrultusunda ya da mühendis görüşlerine göre belirlenen bazı özellikleri sağlayabilmesi için kullanılan ve ürünlerin kalitesini belirleyen makine, alet, ekipman, yöntem, malzeme ve işgücü gibi faktörlerin oluşturduğu nedenler sisteminin tümüne süreç adı verilir. Süreç yeterliliği ise en temel anlamda süreçten elde edilen ürünlerin istenen sınırlar içinde kalabilme kabiliyetinin bir ölçüsü olarak nitelendirilebilir. Süreç yeterliliği, sürecin değişmezliği ya da istikrarlılığı ile ilgilendiği için süreçteki sözü geçen değişkenlik, çıktının istikrarlılığının bir ölçüsüdür (Işığışok, 2012: 257).

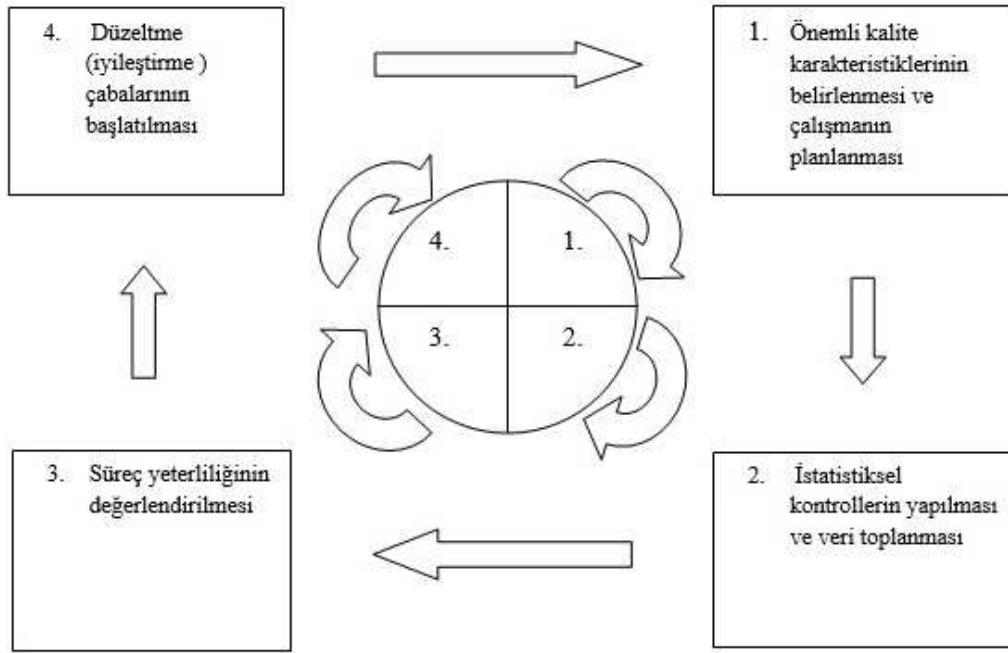
Süreç yeterliliği, Sullivan (1984, 1985) tarafından kalite güvence sisteminin değişen bir felsefesi olarak tanımlandıktan sonra, birçok araştırmacının dikkatini çeken bir konu haline gelmiştir. Bu konunun en büyük destekleyici sloganları, "ilk seferde doğru üretim" ve "*kaliteli ürün üretmek*" şeklinde ifade edilebilir. Sürecin hedeflenen özellikleri sağlamada yeterli olmadığı durumlarda kaynaklar, genellikle uygun olmayan (non-conforming) ürün üretmede kullanılmış olacak, uygun olmayan ürünlerin tanımlanması (denetim ya da müşteri memnuniyetsizliği gibi), değiştirilmesi ya da onarılması sonucunda bazı maliyetler oluşacaktır. Bu maliyetlerin bazıları somut maliyetler olduğu gibi (değiştirme maliyetleri ya da onarım maliyetleri gibi), bazıları da ölçülmesi daha zor olan (müşteri memnuniyetsizliği sonucu oluşan iş kaybı maliyeti gibi) soyut maliyetlerden oluşmaktadır (Spiring, 2010).

Süreç yeterliliği, ilk başlarda, süreç değişkenliği ya da süreç yayılımı ile eş anlamlı olarak kullanılmaktaydı. Daha sonra Dr. Genichi Taguchi tarafından geliştirilen bu düşünce, hedef değere olan yakınlık olarak algılanmaya başlamıştır. Taguchi'nin kayıp fonksiyonu, hedef değer etrafındaki küçük değişkenliği vurgulamaktadır. Taguchi'ye göre en iyi süreç, tüm ürünlerin hedeflenen değerde üretilmiş olmasıdır. Ürün kalitesini değerlendirirken, kayıp fonksiyon üzerinde hem süreç yayılımına hem de hedef değere olan yakınlığa dikkat etmek gerekir. Bir süreçte üretilen tüm ürünler, spesifikasyon sınırları içinde olup, istenen yayılım büyüklüğüne sahip ancak hedef değer üzerinde merkezlenmemiş de olabilir. Bu durumda "*en iyi süreç*" tanımı, tüm ürünlerin hedef değerlerde üretilmesi şeklinde yapılabilir. Ancak tüm ürünlerin hedef değerlerde üretilmesinin pratikte zor olması nedeniyle, hedef değer

etrafında küçük bir deęişkenlik gösteren süreçler için de "iyi bir süreç" tanımlamasını kullanmak mümkündür.

Süreç yeterlilięi, en yaygın biçimde süreç deęişkenlięinin çıktısı olan bir aralık olarak tanımlanır. Bu aralık, gerçek süreç yayılımı olarak da adlandırılır. Süreç yeterlilięi ile ilgili çalışmalar, sürecin yeterli olup olmadığına karar vermek amacı ile yapılır. Bu amaçla ilgili süreçten veri toplanır. Verinin elde edildięi sürecin istatistiksel olarak kontrol altında ya da duraęan olması gerekir.

Deleryd'e (1999) göre süreç yeterlilięi ile ilgili yapılan çalışmalarda dört temel adım bulunmaktadır. Bunlar Şekil 1.1'de görüldüğü gibidir:



Şekil 1.1. Süreç yeterlilięine karar vermek için kullanılan dört temel adım

Sekil 1.1'de görülen 4 adım kısaca aşağıdaki gibi açıklanabilir:

Yeterlilik çalışmalarına başlamadan önce dikkatli bir şekilde planlama yapılmalıdır.

İlk olarak, süreçte ölçülmesi gereken kalite karakteristikleri belirlenmeli, daha sonra bu ölçümlerin ne şekilde ve hangi ölçü aletleri kullanılarak gerçekleştirileceęine karar verilmelidir. İkinci olarak, bir sürecin yeterli olup olmadığına karar vermeden önce, sürecin istatistiksel açıdan kontrol altında olması ya da duraęan olması gerekir. Aksi halde, ileriki

dönemler için sürecin yeterliliğinin geçmişteki değeri ile aynı tutarlılığı gösterdiğini söylemek ve gelecekteki değerlerin de aynı olacağını beklemek mantıklı değildir. Bu nedenle, durağan olmayan bir sürecin yeterlilik çalışması o andaki süreç yeterliliğini ifade ettiği için, sürecin gelecekteki yeterliliği ile ilgili bir yargıya varılamaz. Bu durumda, sürecin kontrol altında olup olmadığına karar vermek amacıyla kalite kontrol grafiklerinin kullanılması iyi bir yöntemdir. Daha sonraki aşamada, sürecin yeterliliği değerlendirilir. Durağan bir sürecin yeterliliğini belirlemede kullanılan en basit yöntem, kontrol grafiklerinden elde edilen bireysel değerlerin histogramını çizmektir. Süreç yeterliliğinin görsel olarak elde etmenin bir başka yolu, kutu grafiğidir. Süreç yeterliliğinin sayısal olarak elde edilmesi için de yeterlilik indeksleri geliştirilmiştir. Son olarak, daha iyi bir süreç elde etmek amacı ile süreç iyileştirme imkanları tanımlanır (Deleryd, 1999).

Süreç yeterliliğine karar vermek amacıyla kullanılan yöntemlerden biri süreç yeterlilik indeksleridir ve süreç yeterliliğini ölçmede kullanılan çok sayıda yeterlilik indeksi mevcuttur.

Yeterlilik indeksleri, üretim sürecinin istenen özellikleri sağlamadaki başarısını ölçmede kullanılan bir ölçüm olarak karşımıza çıkar. Süreç yeterlilik indeksleri, izin verilen süreç yayılımı ve gerçek süreç yayılımı ile ilgili elde edilen bir oran değeridir ve en temel anlamda aşağıdaki şekilde ifade edilir:

$$\text{Yeterlilik oranı} = \frac{\text{izin verilen süreç yayılımı}}{\text{gerçek süreç yayılımı}}$$

bu oran ölçekten bağımsız olduğu için farklı kalite değişkenlerine sahip süreçleri karşılaştırmada kullanılabilen bir ölçümdür. Bu ölçüm değeri ile aynı zamanda süreçteki kalite değişkenlerine ve üretilen ürüne bakılmaksızın, yeterlilik değerlendirilmesinde çıkarımlar yapılabilir (Spiring, 2010).

Süreç yeterliliğinin tahmini için kullanılan ve süreç yeterlilik indeksleri (process capability indices) olarak bilinen standart teknikler genellikle parametrik varsayımlara dayanır. Bu varsayımlardan özellikle normallik varsayımı, süreç yeterlilik indekslerinin hesaplanması ve yorumlanmasında büyük bir öneme sahiptir. Polansky, Chou ve Mason (1998), Kotz ve Johnson (1993: 139) ve Sommerville ve Montgomery (1996), bu varsayımda meydana

gelen küçük bozulmalar sonucu süreç yeterlilik değerlendirmesinin nasıl yanıltıcı olabileceğini göstermişlerdir. Dolayısıyla, süreç kalitesini değerlendirmede kullanılan süreç yeterlilik indeksleri tartışmalı hale gelmiştir. Bu tartışmalar sonucunda bazı kalite uygulayıcıları bu gibi teknikleri saf dışı bırakmışlardır. Ancak, bir üretim sürecinin yeterliliğinin belirlenmesi endüstriler için yararlı sonuçlar sağlar. Bu nedenle, geçerli ve kullanışlı yöntemler geliştirilmesi büyük önem taşımaktadır.

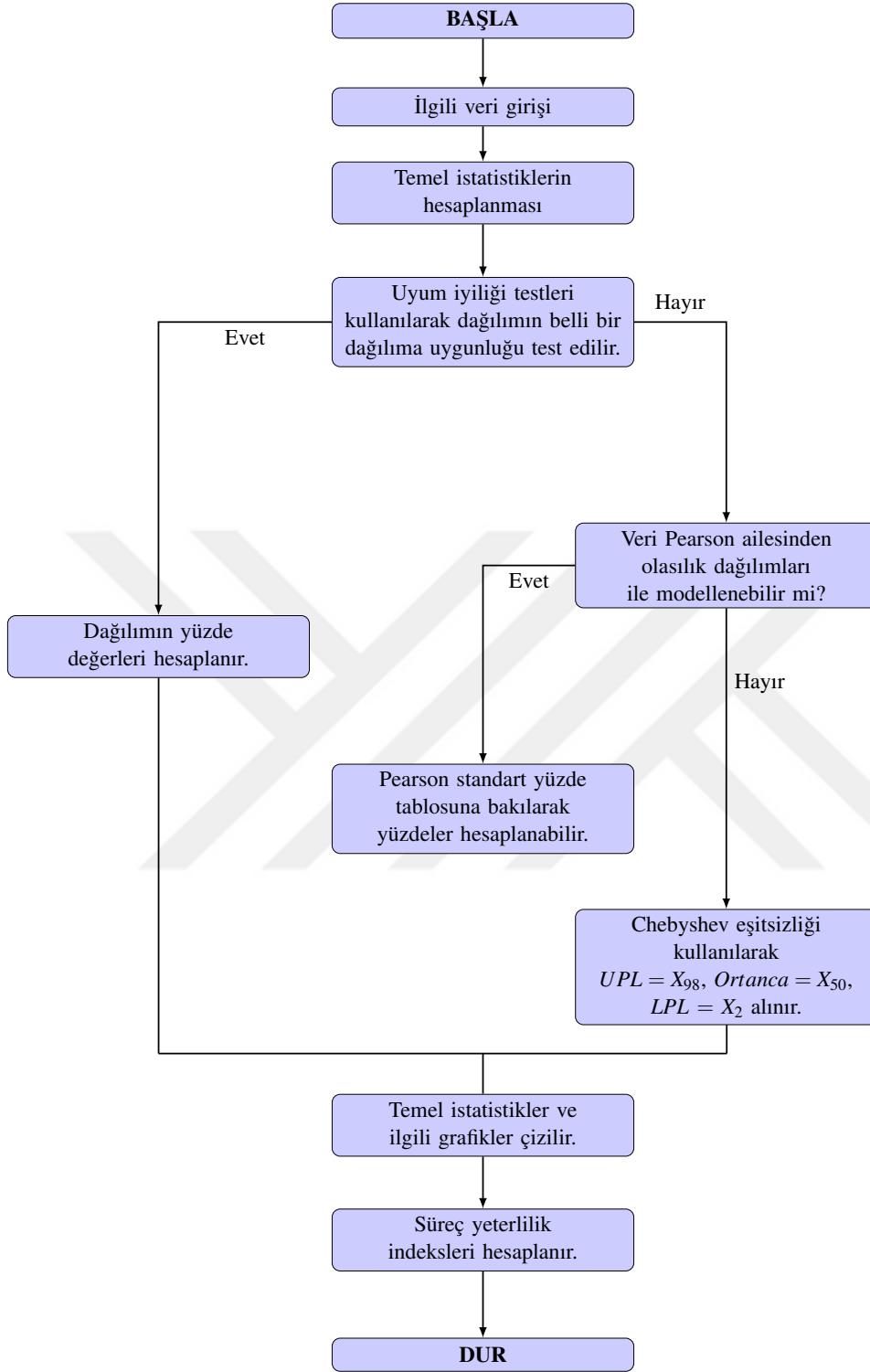
Birçok yazar, süreç yeterlilik analizinin uygulanmasındaki olasılıkla ilgili zorlukları ortadan kaldırmak için esnek parametrik model (flexible parametric model) kullanmayı önermiştir. Clement (1989), Pearson dağılımını uygulamıştır. Johnson dağılımı Farnum (1996) ve Chou, Polansky ve Mason (1998) tarafından önerilmiş ve Burr dağılımı Castagliola (1996) tarafından kullanılmıştır. Ancak, bu yöntemlerin her birinde yeterlilik tahminlerinin güvenilirliği sadece verinin varsayılan parametrik dağılımlara yaklaştığı durumda mümkündür. Dağılım probleminin üstesinden gelebilmek için henüz az sayıda parametrik olmayan yöntem çalışılmıştır. Ayrıca, yapılmış olan çalışmaların çoğu, tek değişkenli kalite karakteristikleri içindir. Çok değişkenli durumlar için standart indekslerin genişlemeleri, Chan, Cheng ve Spiring (1988), Chen (1994), Kocherlakota ve Kocherlakota (1991) ve Pearn, Kotz ve Johnson (1992) tarafından çalışılmıştır. Genellikle bu indeksler çok değişkenli normal dağılım varsayımına dayanır. Bu varsayım, bazı sorunlara neden olur. Bunlardan ilki, bu indekslerin çok değişkenli normallik varsayımına çok duyarlı olması, ikincisi, eldeki veri kümesinin çok değişkenli normal dağılıma uygunluğunun belirlenmesinin zor olması, son olarak, veri seti çok değişkenli normal dağılıma uymuyorsa süreç yeterliliğinin belirlenmesinde alternatif bir yöntemin bulunmamasıdır. Önerilen süreç yeterlilik tahmin edicisinin normallik varsayımına duyarsız olması ve bir alternatif sağlaması gerekir (Polansky, 2001).

Gunter (1989), ortalama ve standart sapmaları aynı olan bilinen üç farklı dağılıma sahip süreçlerden elde edilen kusurlu parça sayılarını karşılaştırmıştır. Bu dağılımlar 4,5 serbestlik derecesine sahip ki-kare, 8 serbestlik derecesine sahip ve uniform dağılımı göstermelerine rağmen, her birinin normal dağıldığı varsayılarak hesaplanan C_p ve C_{pk} indekslerinin aynı olduğu görülmüştür. Ancak $\pm 3\sigma$ sınırları dışına düşen kusurlu parça sayılarının önemli derecede farklı olduğu gözlenmiştir. Elde edilen kusurlu oranları, Ki-kare dağılımı için milyonda 14.000 (ppm); t dağılımı için milyonda 4.000 ve uniform

dağılımı için milyonda 0 iken; bu oran normal dağılım için 2.700'dür. Bu nedenle Gunter (1989), süreç dağılımı normal dağılımdan saptığı durumlarda, süreç yeterlilik indekslerinin güvenilir olmadığı sonucuna varmıştır (Pan ve Wu, 1997).

Pan ve Wu (1997), süreç dağılımı normal dağılıma uygun olmadığı durumlar için süreç yeterlilik indeksleri ile ilgili olarak çalışmışlardır. S Plus bilgisayar programını kullanarak normallik varsayımını sağlamayan süreçler için yeterlilik indeksleri elde ederek bu amaçla kullanılan bir işlemler dizisi geliştirmişlerdir. Bu işlemler dizisi Şekil 1.2'de görülmektedir.





Şekil 1.2. Süreç yeterlilik indekslerini hesaplamak için kullanılan akış şeması

Pan ve Wu (1997), Şekil 1.2’de akış şeması verilen işlemler dizisini üç farklı uygulama üzerinde çalıştırarak elde ettikleri sonuçları tartışmışlardır. Her birinin karakteristiklerine göre karar verilen dağılımları dikkate alarak elde edilen süreç yeterlilik indekslerinin normal dağılımda olduğu gibi, kusurlu oranlarının %0,27 olduğu görülmüştür.

Son zamanlarda, kalite değişkenlerinin normal olmadığı durumlarda kullanmak amacıyla klasik süreç yeterlilik indekslerinde bazı değişiklikler yapılması önerilmiştir. Bu öneriler iki temel yaklaşımda ele alınabilir. Bunlardan ilki, normal olmayan veriyi normal dağılıma dönüştürmek ve daha sonra klasik yöntemleri kullanmaktır. Johnson (1949), momentler yöntemine dayanan ve Johnson dönüşüm sistemi olarak adlandırılan bir normal olmayan veri dönüşüm sistemi geliştirmiştir. Box ve Cox (1964), normal olmayan veriyi normal veriye dönüştüren bir güç dönüşümler ailesi tanımlamıştır. Somerville ve Montgomery (1996), çarpık dağılımları, normal dağılıma dönüştüren bir karekök dönüşümü kullanmışlardır. Niaki ve Abbasi (2007), çarpık kesikli çok değişkenli veriyi çok değişkenli normal veriye dönüştüren bir kök yöntemi önermişlerdir. Hosseini, Abbasi, Ahmad ve Abdollahian (2009), tek değişkenli normal olmayan süreçler için süreç yeterlilik indekslerinin tahmininde bir kök dönüşüm yöntemi kullanmışlardır.

Bununla birlikte, dönüşüm yöntemlerinin kullanımının da bazı dezavantajları vardır. İlk olarak, dönüşüm yöntemleri işlem açısından zaman almakta ve uygulayıcılar hesaplanan sonuçların orijinal ölçeklerin çevrilmesiyle ilgili sorunlardan dolayı bu yöntemleri kullanmakta tereddüt etmektedirler.

Normal olmayan süreçler için kullanılan ikinci yaklaşım, veriyi bir dağılıma uydurmak ve daha sonra normal olmayan yüzdelikleri kullanarak süreç yeterlilik indekslerini tahmin etmektir. Bu yaklaşımda, normal dağılımda kullanılan 6σ terimi yerine üst yüzdelik değeri 99,865 ve alt yüzdelik değeri 0,135 arasındaki aralık uzunluğu kullanılır. Bu yaklaşımı kullanan araştırmacılardan birisi olan Clements (1989), yaygın olarak endüstri sektöründe kullanılan Pearson ailesine ait herhangi bir dağılımın süreç yeterlilik tahminleri için normal olmayan yüzdelikler yöntemini önermiştir. Pearn ve Kotz (1994), C_{pm} ve C_{pmk} indekslerini normal olmayan dağılımlara uygulamak için Clements yöntemi kullanmıştır. Clements yöntemi yaygın olarak kullanılmasına rağmen, Wu, Wang ve Liu (1998), bu yöntemin özellikle çarpık dağılımlar için doğru ölçümler vermediğini göstermişlerdir. Liu ve Chen (2006), normal olmayan veriler için süreç yeterlilik indeksi hesaplamada kullanılan Clements yöntemini değiştirerek, Pearson eğrisi yüzdeliklerini kullanmak yerine Burr XII dağılımı kullanmayı önermişlerdir. Burr XII olasılık yoğunluk dağılımının parametreleri normal, gama, beta, weibull ve log-normal dağılımlarına uyarlanabilmektedir (Abbasi ve Niaki, 2010).

Süreç yeterlilik analizinde, süreç yeterliliğini ölçmede etkili olan süreç değişkenlerinin birden fazla olması da mümkündür. Fabrika problemlerinde, tek bir değişkenin yanı sıra, birden fazla kalite değişkenine sahip problemler ile de karşılaşmaktadır. Örneğin, bir çelik malzemenin kalitesi, kalınlığına, dayanıklılığına ve hareket kabiliyetine bağlıdır. Bu durumda, çok değişkenli süreçlerin yeterlilik analizi, çok değişkenli süreç yeterlilik indekslerinin hesaplanması ile mümkün olacaktır. Çok değişkenli durumlar için süreç yeterlilik indeksleri, sürecin normallik varsayımına göre farklılıklar göstermektedir.

Çok değişkenli süreçler için literatürde yer alan çalışmalar dört grupta incelenmiştir. Bu gruplandırmaya göre, süreç yeterlilik indeksleri,

1. Tolerans bölgesi ile süreç bölgesini birbiri ile oranlamaya dayalı olarak;
2. Uygun olmayan ürün elde etme olasılığına dayalı olarak;
3. Temel bileşenler analizi kullanılarak değişkenler arası korelasyonu kaldırmaya yönelik olarak;
4. Diğer yaklaşımlar kullanılarak

elde edilmiştir (Shinde ve Khadse, 2009). Bahsi geçen çalışmalar normallik varsayımı sağlanması koşulu altında ele alınmıştır.

Çok değişkenli üretim süreçleri için normallik varsayımının sağlanmadığı durumlarda süreç yeterlilik indekslerinin hesaplanması, spesifikasyon bölgesi ve değişkenler arasındaki ilişki durumu da göz önünde bulundurulduğunda, daha güç ve karmaşık olabilmektedir. Dağılımın normal olmadığı çok değişkenli süreç yeterlilik analizinde, yeterlilik indekslerini az hata ile tahmin edebilmek ancak çalışılan verinin dağılım bilgisine uygun tahmin yöntemleri kullanmakla mümkündür. Ancak, süreç dağılımı bilinen çok değişkenli dağılımlara uygun olabileceği gibi, bilinen herhangi bir dağılıma uymayan bir süreç de olabilir.

Polansky (2001), Normal dağılım göstermeyen çok değişkenli süreçler için Kernel tahminine dayalı bir süreç yeterlilik tahmin yaklaşımı ($\hat{p}(H, S)$) önermiştir. Önerilen bu yaklaşımda $\hat{p}(H, S)$ Kernel tahmin edicisi, H bant genişliği matrisi ve S ise spesifikasyon kümesi olarak tanımlanmıştır. Kernel tahmin edicisinin performansını değerlendirmek

amacıyla iki değişkenli sekiz farklı dağılım simülasyonu üzerinden elde ettiği sonuçları normal varsayım altında parametrik normal tahmin edicisi ile elde edilen sonuçlarla karşılaştırmıştır. Çalışmasında yeterlilik değerlendirmesi yapmak için süreç dışına düşen ürün oranı p 'yi tahmin etmek amacıyla, simüle ettiği dağılımlar, iki değişkenli Normal Dağılım (aralarında ilişki bulunmayan), iki değişkenli Normal Dağılım (aralarında ilişki bulunan), basık bir dağılım, çarpık bir dağılım, Bimodal, Trimodal I, Trimodal II ve Quadrimodal olarak belirlenmiştir. Çalışmasının sonucunda, parametrik normal tahmin edicisinin Normal dağılımlar için daha iyi sonuçlar verdiğini, Kernel tahmin edicisinin büyük çaplı örnekler için ($n > 200$) kullanılmasının daha doğru olacağını, küçük çaplı örnekler açısından Kernel tahmin edicisinin iyi sonuçlar vermediğini ifade etmiştir.

Polansky'nin (2001) çalışmasında önerdiği Kernel tahmin edicisi $\hat{p}(H, S)$, $n < 200$ olan küçük çaplı örneklemelere uygulandığında p 'yi iyi bir şekilde tahmin edemediği ve sonuçların farklı dağılımlara sahip simülasyon verileri üzerinden elde edildiği göz önünde bulundurulduğunda, dağılımdan çekilen tek bir örneklem verisinin kitle parametresini tahmin etmede zayıf kalacağı düşünülmektedir. Nitekim bu çalışmada kullanılacak verilere Polansky'nin (2001) yöntemi uygulandığında elde edilen sonuçlar düşüncemizi doğrular niteliktedir. Bu nedenle bu tez çalışmasında, belirlenen senaryo verilerine Metropolis Hastings örneklemesi uygulanarak, örneklem kümeleri üzerinden Kernel tahminleri elde edilmiş ve böylece, dağılım bilgisinden uzaklaşmadan küçük çaplı örneklem kümeleri ile çalışma imkanı sağlayan yeni bir yöntem önerilmiştir. Bu yeni yöntem *KMH* olarak adlandırılmış olup, yöntemeye dayalı süreç yeterlilik indeksi ise $\tilde{p}_{KMH}$ olarak ifade edilmiştir.

Bu çalışmada, Normal dağılım varsayımı sağlanmayan çok değişkenli süreçlerin yeterlilik analizinde kullanılmak üzere t-Copula yöntemi ile üretilen çarpık dağılım özelliklerine sahip 4 farklı üretim süreci senaryosu kurularak, üretilen senaryo verilerinin dağılımı Kernel Yoğunluk Tahmin yöntemi ile tahmin edilerek, literatürde çalışılan yeterlilik indekleri $(Y_{pk}, \tilde{p})$ ve Polansky (2001)'nin çalışmasında önerdiği $\hat{p}(H, S)$ yaklaşımı ile $\tilde{p}_{KMH}$ indeksinin süreç yeterliliğine karar vermedeki başarısı değerlendirilmektedir. Bu bağlamda, senaryoların her biri için ($n = 50, 100, 250, 500, 1000$) çapında örneklemler çektirilmiş ve $\tilde{p}_{KMH}$ yaklaşımının örneklem büyüklüğüne olan duyarlılığının incelenmesi amaçlanmıştır.

Ayrıca, uygun ürün oranı p 'yi tahmin etmek amacıyla önerilen $\tilde{p}_{KMH}$ yaklaşımının $n < 200$ ($n = 50, 100$) küçük çaplı örneklemelerden elde edilen sonuçları ile $n > 200$ ($n = 250, 500, 1000$) büyük çaplı örneklemelerden elde edilen sonuçları karşılaştırılacak ve $\tilde{p}_{KMH}$ indeksinin performansı değerlendirilecektir.

Çalışmanın bu bölümünde, süreç yeterlilik indeksleri ve çalışmanın amacına yönelik bir çerçeve çizilmeye çalışılmış, çalışmada yapılacaklardan kısaca bahsedilmiştir.

Çalışmanın ikinci bölümü, süreç yeterlilik indekslerinin geniş bir literatür bilgisini içermektedir. Bu bölümde, hem tek değişkenli süreç yeterlilik indekslerinin, hem de çok değişkenli süreç yeterlilik indekslerinin normal ve normal olmayan durumlardaki literatür incelemesine yer verilmiştir.

Çalışmanın üçüncü bölümünde, tezde önerilen $\tilde{p}_{KMH}$ yaklaşımın dayandığı teorik bilgi ve kullanılan yöntemlerden bahsedilmiştir. $\tilde{p}_{KMH}$ yaklaşımının performansını değerlendirmek için kurgulanan senaryo verilerinin üretim aşamaları ve verilerinin dağılım özelliklerine yer verilmiştir.

Dördüncü bölümde, her bir senaryodan elde edilen sonuçlar yorumlanmış, beşinci bölümünde tahmin grafikleri ile sonuçların görsel açıdan verilmesi sağlanmıştır.

Çalışmanın son bölümünde ise çalışma sonuçları özetlenmiş ve yeterlilik indekslerine ilişkin yapılacak gelecek çalışmalar açısından önerilere vurgu yapılmıştır.

2. SÜREÇ YETERLİLİK İNDEKSLERİ

Süreç yeterlilik indekslerini inceleyebilmek için öncelikle süreç yeterlilik analizinde kullanılan bazı kavramları tanımlamak gerekir. Bir ürün ya da hizmetin olduğu her yerde kalitenin belirleyicisi kalite değişkenliklerinin ölçülmesi, özelliklerinin incelenmesi ve değişkenlik nedenleri ile hata kaynaklarının belirlenmesi için istatistiksel tekniklerden yararlanılır (Işığışık, 2012: 69). Bu tekniklerin sağlıklı bir şekilde kullanılabilmesi için kalite değişkenleri iyi tanımlanmalıdır. Kalite değişkenleri genellikle spesifikasyonlar ile ilgilidir. Spesifikasyonlar, bir ürünün kalite değişkeninin nihai üründe olduğu kadar ürün bileşenlerinde ve alt gruplarında da istenen ölçülerde olması anlamına gelir. Kalite değişkeninin istenen değerine "*hedef değer*" denir ve T ile gösterilir. Bu hedef değer genellikle hedefe yeteri kadar yakın olan bir aralık ile sınırlandırılır. Bir kalite değişkeni için izin verilen en büyük değer "*üst spesifikasyon sınırı*"; izin verilen en küçük değer "*alt spesifikasyon sınırı*"; bu sınırlar arası fark ise "*tolerans aralığı*" olarak adlandırılır. Bazı kalite değişkenleri hedef değer için sadece tek bir spesifikasyon sınırına sahip olabilir. Spesifikasyon sınırlarını sağlayan ürün ya da ürün bileşenleri için "*uygun*"; spesifikasyon sınırlarından en az birini sağlamayan ürün ya da ürün bileşenleri için "*uygun olmayan*" terimleri kullanılır (Montgomery, 2009: 8-9).

Spesifikasyon sınırlarından farklı olarak kontrol sınırları, sürecin dağılımının kontrol altında olup olmadığının belirlenmesi için çizilen ve süreç ortalaması ile süreç değişkenliğine bağlı olarak değişiklik gösteren sınırlar olarak tanımlanabilir. Kontrol sınırları "*doğal tolerans sınırları*" olarak da adlandırılır (Işığışık, 2012: 65).

Kontrol sınırlarına genellikle kontrol grafiklerinin çiziminde başvurulur. Kontrol grafikleri süreçlerden elde edilen ürünlerin gözlem sonuçlarına ilişkin olarak ortaya çıkan değişimleri gösteren ve sürecin kontrol altında olup olmadığına karar vermede kullanılan görsel araçlardır. Kontrol grafikleri, alt ve üst olmak üzere kontrol sınırlarından ve bir orta (merkez) çizgiden oluşmaktadır.

Spesifikasyon sınırlarında olduğu gibi, kontrol sınırları için de "*alt kontrol sınırı*" ve "*üst kontrol sınırı*" tanımlamasının yapılması mümkündür. Üst ve alt kontrol sınırları, değişkenin normal dağıldığı durumda süreçten alınan ürünlere ilişkin gözlem değerlerinden hesaplanan

ve kontrol grafiğindeki merkez çizgisine (genellikle $\pm 3\sigma$) eşit uzaklıkta bulunan sınırlardır. Merkez çizginin her iki yönündeki 3 standart sapma arasında kalan alan, eğrinin altında kalan toplam alanın %99,73'ünün göstermektedir. üst ve alt kontrol sınırları arasındaki farka, "*doğal tolerans aralığı*" denilmektedir. (Işığışok, 2012: 65).

Spesifikasyon sınırlarına bağlı olarak belirlenen "*tolerans aralığı*", müşteri isteklerine göre mühendislik çalışmaları veya üreticiler tarafından dışsal olarak belirlenirken; *doğal tolerans aralığı*, üretilen ürünlerin değişkenliğini ifade eder ve üretim sürecinden elde edilir.

Süreç yeterlilik indeksleri, bir ürün ya da hizmet sürecindeki kalite değişkenlerinin aldığı değerleri, spesifikasyon sınırları ile karşılaştırmayı sağlayan sayısal değerlerdir. Bu değerler, çoğunlukla "*yeterlilik ya da performans indeksleri*" ya da "*oranları*" olarak tanımlanır. Bir yeterlilik indeksi, müşterilerin sesi (spesifikasyon sınırları) ile sürecin sesini ilişkilendirir. Örneğin, yeterlilik indeksinin büyük bir değeri, mevcut süreçteki üretim birimlerinin, müşteri gereksinimlerini karşılamada yeterli olduğunu gösterir. Yeterlilik indeksleri, süreç hakkında elde edilen karmaşık bilgiyi tek bir sayıya indirgeyerek özetlediğinden, kullanımları yaygındır. Bu nedenle, sürecin ne kadar iyi çalıştığını belirlemede kullanılır. Durağan süreçler için, bu indekslerin aynı zamanda sürecin gelecekteki beklenen performansını ölçmede de kullanılacağı varsayılır. Tedarikçiler, farklı değişkenler için elde edilen yeterlilik indekslerini, iyileştirme faaliyetlerindeki önceliklerin belirlenmesinde de kullanmaktadırlar. Aynı zamanda, süreçteki değişikliklerin etkileri, süreç öncesinde ve sonrasındaki yeterlilik indekslerinin karşılaştırılması ile belirlenebilir.

Literatürde bulunan ve uygulamada kullanılan çok çeşitli süreç yeterlilik indeksi mevcuttur. Bu indeksler genellikle, kullanım amaçları, özellikleri ve hesaplama biçimleri bakımından farklılıklar gösterir. Buna rağmen, tasarımlarında benzerlikler görülebilir. Bir süreçten elde edilen yeterlilik indeksleri temel olarak iki kuşakta incelenebilir. İlk kuşak yeterlilik indeksleri olarak bilinen, C_p ve C_{pk} indeksleri, klasik istatistiksel süreç kontrol felsefesine dayanır. Bu felsefeye göre, istenilen tolerans aralığında tüm ölçüm sonuçlarının "*uygun*" olması amaçlanmaktadır. Tolerans aralığı dışında kalan ölçüm değerleri, "*uygun olmayan*" olarak kabul edilir.

İkinci kuşak yeterlilik indeksleri olarak bilinen, C_{pm} indeksi ise, kalite iyileştirme

çalışmalarına yeni bir yaklaşım getirmiştir (Taguchi yaklaşımı). Bu yaklaşımda, yalnızca ölçümlerin "*uygun*" olduğunu bilmek yeterli değil, aynı zamanda "*ne kadar uygun*" oldukları hakkında da bilgi sağlaması önemlidir (Kurekova, 2001).

Bu kısımda öncelikle bazı tek değişkenli yeterlilik indeksleri ele alınacak ve bunların dağılım özellikleri incelenecektir.

2.1. Tek Değişkenli Süreç Yeterlilik İndeksleri

Tek değişkenli süreç yeterlilik indeksleri Normal Dağılım gösteren ve Normal Dağılım göstermeyen süreçler olarak iki sınıfta incelenmiştir.

2.1.1. Normal dağılıma uyan süreçlerde tek değişkenli süreç yeterlilik indeksleri

Normal Dağılıma uyan süreçler için literatürde yer bulan süreç yeterlilik indeksleri "*Klasik Yeterlilik İndeksleri*" olarak adlandırılmaktadır. Bunun nedeni, yeterlilik indeksleri hesaplanırken süreç dağılımının Normallik koşulunu sağlandığı varsayımına dayalı olması ve ilk çalışılan yeterlilik indekslerinin temelini oluşturmalarıdır. Klasik yeterlilik indekslerinden C_p , C_{pk} ve C_{pmk} takip eden alt kısımlarda detaylı olarak tanımlanmaktadır.

C_p indeksi

Literatürde çalışılan ilk süreç yeterlilik indeksi Kane (1986) tarafından tanımlanan C_p indeksidir ve Eş. 2.1'de görüldüğü şekilde elde edilir.

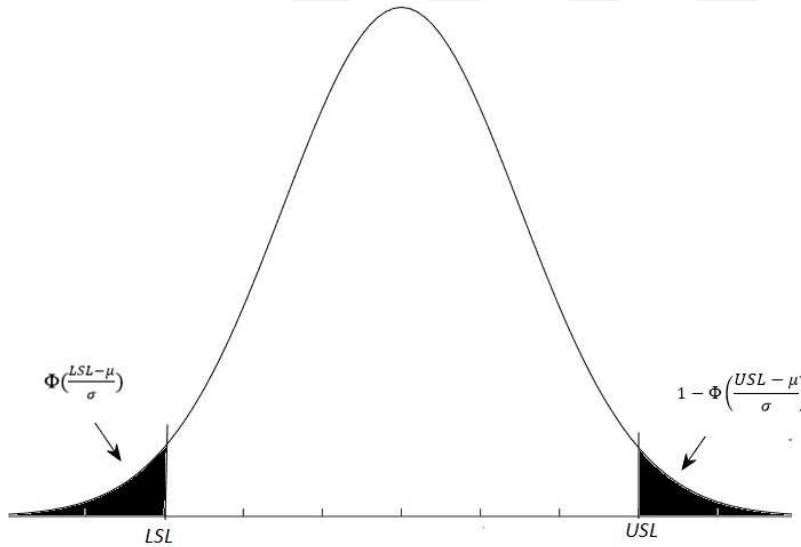
$$C_p = \frac{USL - LSL}{6\sigma} \quad (2.1)$$

Eş. 2.1'de verilen USL , (Upper Specification Limit) üst spesifikasyon sınırını; LSL , (Lower Specification Limit) alt spesifikasyon sınırını; $USL - LSL$, tolerans aralığını ve σ ise kalite değişkeninin standart sapmasını ifade eder. C_p indeksi, üretim süreci için dışsal olarak belirlenen tolerans aralığının, süreç değişkenliğine oranı olarak düşünülmüştür. C_p indeksinin ölçülmesi için süreç dağılımının Normal Dağılıma uyması, hesaplamada kullanılacak süreç verisinin rastgele bir şekilde elde edilmesi, diğer bir ifade ile gözlemlerin birbirinden bağımsız olması ve sürecin kontrol altında olması gerekmektedir. C_p indeksinin

hesaplanması sonucu elde edilen genellikle 1'den küçük değerler, sürecin doğal sınırları ile spesifikasyon sınırlarının örtüşmediği anlamına gelir. Bu nedenle C_p indeksinin büyük değerler alması daha çok tercih edilen bir durumdur. Finley (1992), C_p indeksi için CPI (Capability Potential ya da Capability Potential Index), diğer bir ifade ile "Yeterlilik Potansiyel İndeksi" tanımlamasını kullanmıştır. Montgomery (2009) ise C_p indeksine, PCR (Process Capability Ratio), diğer bir ifade ile, "Süreç Yeterlilik Oranı" adını vermiştir.

C_p indeksi ve uygun olmama yüzdesi (non-conforming, %NC)

İki yanlı spesifikasyon sınırlarına sahip süreçler için "uygun olmama" yüzdesi $1 - F(USL) + F(LSL)$ olarak hesaplanır. Burada $F(\cdot)$, süreç değişkeni olan X 'in dağılım fonksiyonudur. Süreç kalite değişkeninin normal dağılım göstermesi durumunda $\Phi(\cdot)$ birikimli standart normal dağılım fonksiyonu kullanılarak normal dağılım gösteren bir süreç için uygun olmama alanları şekil 2.1'de gösterilmiştir.



Şekil 2.1. Normal dağılım gösteren bir süreç için uygun olmama alanları

Sürecin normal dağılıma uygunluğu varsayımı altında % NC (% nonconforming), Şekil 2.1'den faydalanarak, Eş. 2.2'de görüldüğü gibi elde edilir:

$$\%NC = 1 - \Phi\left(\frac{USL - \mu}{\sigma}\right) + \Phi\left(\frac{LSL - \mu}{\sigma}\right) \quad (2.2)$$

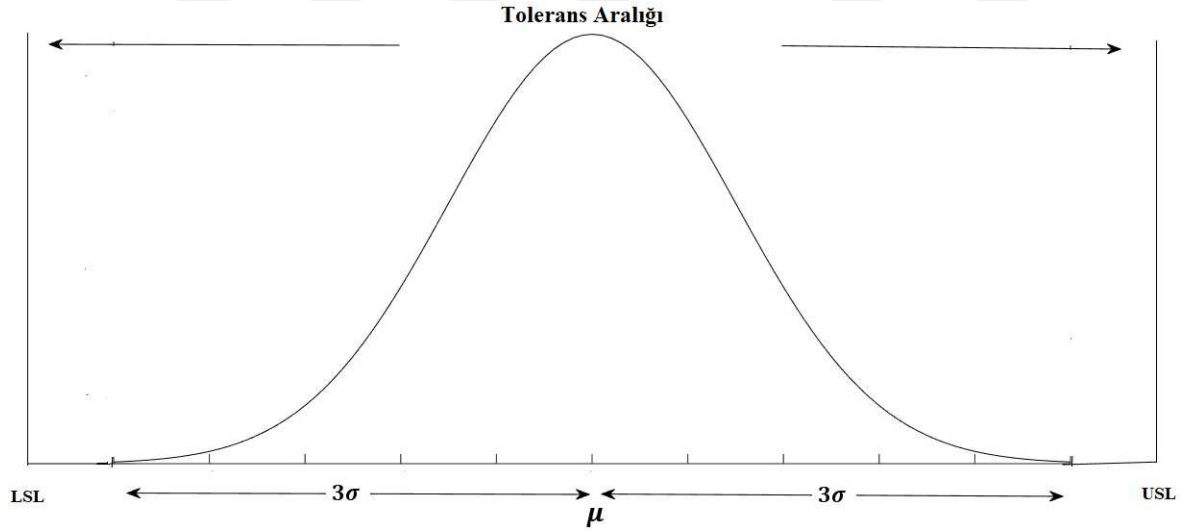
Burada, μ süreç ortalaması ve σ süreç standart sapmasıdır. $m = (USL + LSL) / 2$ üst ve alt spesifikasyon sınırlarının orta noktasıdır.

Süreç kalite değişkeni tam olarak tolerans aralığında merkezlenmişse ($m = \mu$), Eş. 2.2 yeniden yazılarak 3σ için $\%NC$ aşağıdaki şekilde elde edilir:

$$\begin{aligned}
 \%NC &= 1 - \Phi\left(\frac{USL - \left(\frac{USL + LSL}{2}\right)}{\sigma}\right) + \Phi\left(\frac{LSL - \left(\frac{USL + LSL}{2}\right)}{\sigma}\right) \\
 &= 1 - \Phi\left(\frac{2USL - USL - LSL}{2\sigma}\right) + \Phi\left(\frac{2LSL - USL - LSL}{2\sigma}\right) \\
 &= 1 - \Phi\left(\frac{USL - LSL}{2\sigma}\right) + \Phi\left(\frac{LSL - USL}{2\sigma}\right) \\
 &= 1 - \Phi(3C_p) + \Phi(-3C_p) \\
 &= 2\Phi(-3C_p)
 \end{aligned}$$

Böylece, Eş. 2.1'de verilen C_p yardımıyla $\%NC$ değerine ulaşmak mümkündür.

Şekil 2.2'de Normal dağılım gösteren bir süreç için tolerans aralığı görülmektedir.



Şekil 2.2. Normal dağılım gösteren bir süreç için tolerans aralığı

Örneğin Şekil 2.2'de Tolerans aralığının 6σ olarak benimsenmesi durumunda, $C_p = 1,00$ olarak elde edileceğinden birikimli standart normal dağılım tablosundan faydalanarak, $\%NC$ elde edilebilir. $\%NC = 2\Phi(-3C_p) = 2\Phi(-3) = 2 \times 0,00135 = 0,0027$ olarak hesaplanır. O halde, " $\%NC = \text{milyonda } 2700 \text{ parça (parts per million)}$ " ile $C_p = 1,33$ ise " $\%NC = \text{milyonda } 63 \text{ parça (ppm)}$ " olarak ifade edilir. Ancak $m \neq \mu$ ise C_p sürecin yeterliliğini doğru ifade

etmez ve %NC için sağlıklı bir ölçüm sağlamaz. Bu nedenle, C_p süreç ortalama değeri ile hedef değeri (T) örtüştüğünde sağlıklı bir yeterlilik ölçüsü olarak kullanılabilir. $\%NC = 2\Phi(-3C_p)$ de yeterlilik için sadece bir alt sınır sağlar (Pearn ve Kotz, 2006).

Tek bir örneklem ile C_p 'nin tahmin edilmesi

C_p indeksinin hesaplanmasında sadece σ parametresinin tahmin edilmesi gerekmektedir. $\{x_1, \dots, x_n\}$, tek değişkenli ve örneklem çapı n olan bir örneklem kümesi ise, C_p 'nin tahmin edicisi Eş. 2.3'te verildiği şekilde tahmin edilir.

$$\hat{C}_p = \frac{USL - LSL}{6s} \quad (2.3)$$

Burada, $s = \left(\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 \right)^{1/2}$ durağan, diğer bir ifade ile ortalama ve standart sapması zaman içinde sabit kalan bir süreçten elde edilen σ 'nın “bilinen” tahmin edicisidir. Ancak s , σ 'nın yansız bir tahmin edicisi olmadığından s düzeltilmiş kullanılmalıdır. $\frac{\sqrt{(n-1)}s}{\sigma} \sim \chi_{n-1}$ ile Ki-dağılımına sahiptir. Ki-dağılımının momentleri Eş. 2.4'ten elde edilebilir.

$$\mu_j = 2^{j/2} \frac{\Gamma((k+j)/2)}{\Gamma(k/2)} \quad (2.4)$$

Burada, $\Gamma(z)$ Gamma fonksiyonunu, k serbestlik derecesini göstermek üzere Ki-dağılımının birinci momenti Eş. 2.4 kullanılarak şu şekilde elde edilebilir.

$$\mu_1 = 2^{1/2} \frac{\Gamma((n-1+1)/2)}{\Gamma((n-1)/2)} \quad (2.5)$$

O halde, $E\left(\frac{\sqrt{n-1}s}{\sigma}\right) = 2^{1/2} \frac{\Gamma((n-1+1)/2)}{\Gamma((n-1)/2)}$ eşitliği yazılabilir. Buradan, $E(s) = \sqrt{\frac{2}{n-1}} \frac{\Gamma((n-1+1)/2)}{\Gamma((n-1)/2)} \sigma$ yazılabileceği açıktır. s 'nin σ 'nın yansız bir tahmin edicisi olabilmesi için, $c_4 = \sqrt{\frac{2}{n-1}} \frac{\Gamma(n/2)}{\Gamma((n-1)/2)}$ gibi bir standart sapma düzeltme faktörü kullanılır. Uygulamada süreçten elde edilen veri 60'tan az olduğu durumlarda kullanılan $s_{\text{düzeltilmiş}} = \frac{s}{c_4}$ olarak ifade edilir.

Eş. 2.3'te verilen $\hat{C}_p$ formülünü C_p cinsinden yazabilmek için, eşitlik $(n-1)^{1/2}$ ve σ ile çarpıp bölünerek yeniden düzenlenirse Eş. 2.6'ya ulaşılabilir.

$$\hat{C}_p = (n-1)^{1/2} \left(\frac{USL - LSL}{6\sigma} \right) \left[\frac{(n-1)s^2}{\sigma^2} \right]^{-1/2} = (n-1)^{1/2} C_p \left[\frac{(n-1)s^2}{\sigma^2} \right]^{-1/2} \quad (2.6)$$

Eş. 2.6 yeniden yazılırsa,

$$\hat{C}_p = (n-1)^{1/2} C_p \left[\frac{(n-1)s^2}{\sigma^2} \right]^{-1/2} = (n-1) \left(\frac{C_p}{\hat{C}_p} \right)^2 = \frac{(n-1)s^2}{\sigma^2} \quad (2.7)$$

elde edilir.

Eş. 2.7’de, normallik varsayımı altında, $(n-1)s^2/\sigma^2 \sim \chi_{n-1}^2$ ($n-1$) serbestlik derecesi ile Ki-kare dağılır (Chou ve Owen, 1989). Buradan hareketle Chou ve Owen (1989), $\hat{C}_p$ ’nin olasılık yoğunluk fonksiyonunu Eş. 2.8’deki gibi ifade etmişlerdir:

$$f(x) = \frac{2 \left[\sqrt{(n-1)/2} C_p \right]^{n-1}}{\Gamma[(n-1)/2]} x^{n-1} \exp \left[\frac{-(n-1)(C_p)^2}{2x^2} \right], \quad x > 0 \quad (2.8)$$

$\hat{C}_p$ ’nin r. momenti

$\hat{C}_p$ ’nin r. momenti Ki-kare dağılımının özellikleri kullanılarak,

$$E(\hat{C}_p^r) = \frac{\Gamma\left(\frac{n-r-1}{2}\right)}{\Gamma\left(\frac{n-1}{2}\right)} \left[\frac{(n-1)}{2} \right]^{r/2} C_p^r \quad (2.9)$$

şeklinde tanımlanır. Buna göre ilk iki moment aşağıdaki gibi elde edilir:

$$E(\hat{C}_p) = \frac{\Gamma\left(\frac{n-2}{2}\right)}{\Gamma\left(\frac{n-1}{2}\right)} \left[\frac{(n-1)}{2} \right]^{1/2} C_p \quad (2.10)$$

$$E(\hat{C}_p^2) = \frac{\Gamma\left(\frac{n-3}{2}\right)}{\Gamma\left(\frac{n-1}{2}\right)} \left[\frac{(n-1)}{2} \right] C_p^2 \quad (2.11)$$

Buradan, $E(\hat{C}_p^2) - E(\hat{C}_p)^2 = Var(\hat{C}_p)$ olduğundan,

$$Var(\hat{C}_p) = \left\{ \frac{n-1}{n-3} - \frac{n-1}{2} \left[\frac{\Gamma\left(\frac{n-2}{2}\right)}{\Gamma\left(\frac{n-1}{2}\right)} \right]^2 \right\} C_p^2 \quad (2.12)$$

eşitliğine ulaşılır.

$E(\hat{C}_p)$ 'de bilindiği gibi Gamma fonksiyonu özelliğinden $\Gamma(k) = \int_0^\infty t^{k-1} e^{-t} dt$ tüm n değerleri için 1'den büyüktür. $n \geq 15$ için, bu katsayı $(4n-7)/(4n-4)$ değerine yaklaşır. Bu nedenle $\hat{C}_p$, C_p 'nin yanlış ve aşırı tahmin edicisidir. Basit bir örnek verilecek olursa, Eş. 2.13'deki gibi yan oranının %1'den küçük olması için $n > 80$ olması gerekmektedir (Pearn ve Kotz, 2006:7-11).

$$|E(\hat{C}_p) - C_p| / \hat{C}_p \leq 0,01 \quad (2.13)$$

$\hat{C}_p$ tahmin edicisinin istatistiksel özellikleri

Pearn, Lin ve Chen (1998), C_p için Eş. 2.14'deki gibi yansız bir tahmin edici önermişlerdir.

$$\tilde{C}_p = b_{n-1} \hat{C}_p \quad (2.14)$$

Burada b_{n-1} düzeltme faktörüdür ve şu şekilde tanımlanır:

$$b_{n-1} = \sqrt{\frac{2}{n-1}} \frac{\Gamma\left(\frac{n-1}{2}\right)}{\Gamma\left(\frac{n-2}{2}\right)} \quad (2.15)$$

Pearn ve diğerleri (1998), süreç kalite değişkeninin normal dağılım göstermesi koşulu ile yansız $\tilde{C}_p$ tahmin edicisinin C_p 'nin UMVUE (Uniformly Minimum Variance Unbiased) tahmin edicisi olduğunu göstermişlerdir. $\tilde{C}_p$ asimptotik olarak etkin ve tutarlıdır. Ayrıca, $n^{\frac{1}{2}} (\tilde{C}_p - C_p) \xrightarrow{d} N(0, C_p^2/2)$ dır.

C_p için güven aralıkları

C_p için güven aralıkları, $(1 - \alpha)$ güven düzeyinde dağılım özelliklerini kullanarak $\hat{C}_p$ tahmin edicisine göre Eş. 2.16 görüldüğü gibi,

$$\frac{\sqrt{\chi_{n-1,1-\alpha/2}^2}}{\sqrt{n-1}}\hat{C}_p, \frac{\sqrt{\chi_{n-1,\alpha/2}^2}}{\sqrt{n-1}}\hat{C}_p \quad (2.16)$$

$\tilde{C}_p$ tahmin edicisine göre ise Eş. 2.17 verildiği biçiminde tanımlanır.

$$\frac{\sqrt{\chi_{n-1,1-\alpha/2}^2}}{\sqrt{n-1}b_{n-1}}\tilde{C}_p, \frac{\sqrt{\chi_{n-1,\alpha/2}^2}}{\sqrt{n-1}b_{n-1}}\tilde{C}_p \quad (2.17)$$

Burada, $\chi_{n-1,\alpha/2}^2$ ve $\chi_{n-1,1-\alpha/2}^2$ ($n-1$) serbestlik derecesine sahip $\alpha/2$ ve $1-\alpha/2$ için elde edilen Ki-kare değerleridir. Böylece, %100 $(1-\alpha)$ C_p 'nin alt güven sınırı (C_p^{ags}) aşağıdaki gibi gösterilebilir:

$$(C_p^{ags}) = \frac{\sqrt{\chi_{n-1,1-\alpha}^2}}{\sqrt{n-1}} \frac{\tilde{C}_p}{b_{n-1}} = \frac{\sqrt{\chi_{n-1,1-\alpha}^2}}{\sqrt{n-1}} \hat{C}_p \quad (2.18)$$

Ki-kare dağılımının yüzdelik noktaları genellikle formüller yardımı ile yaklaşık olarak elde edilir. Yaygın olarak kullanılan ve bilinen yaklaşımlardan olan Fisher'in yaklaşımı Eş. 2.19'da; Wilson-Hilferty'in yaklaşımı ise Eş. 2.20'de verilmiştir.

$$\chi_{v,1-\alpha} \cong \left(v - \frac{1}{2}\right)^{1/2} + \frac{z_{1-\alpha}}{\sqrt{2}} \quad (2.19)$$

$$\chi_{v,1-\alpha} \cong v^{1/2} \left[1 - \frac{2}{9v} + z_{1-\alpha} \left(\frac{2}{9v} \right)^{1/2} \right]^{3/2} \quad (2.20)$$

Burada z_α , standart normal dağılımın en üst kantil değeridir. Bu yaklaşımları kullanarak C_p 'nin %100 $(1-\alpha)$ güven aralıkları Fisher'in yaklaşımına göre Eş. 2.21 kullanılarak,

$$\left[\frac{1}{(n-1)^{1/2}} \left\{ \left(n - \frac{3}{2}\right)^{1/2} - \frac{z_{\alpha/2}}{\sqrt{2}} \right\} \hat{C}_p, \frac{1}{(n-1)^{1/2}} \left\{ \left(n - \frac{3}{2}\right)^{1/2} + \frac{z_{\alpha/2}}{\sqrt{2}} \right\} \hat{C}_p \right] \quad (2.21)$$

Wilson-Hilferty'in yaklaşımına göre de Eş. 2.22 yardımıyla,

$$\left[\left\{ 1 - \frac{2}{9(n-1)} - z_{\alpha/2} \left[\frac{2}{9(n-1)} \right]^{1/2} \right\}^{3/2} \hat{C}_p, \left\{ 1 - \frac{2}{9(n-1)} + z_{\alpha/2} \left[\frac{2}{9(n-1)} \right]^{1/2} \right\}^{3/2} \hat{C}_p \right] \quad (2.22)$$

biçiminde hesaplanabilir (Pearn ve Kotz, 2006: 13).

C_p tahmini için örneklem büyüklüğüne karar verilmesi

Franklin (1999), C_p yeterlilik indeksi için Wilson-Hilferty yaklaşımına dayanan ve yaklaşık örneklem büyüklüğünü hesaplamada kullanılan bir formül geliştirmiştir. Franklin (1999) tarafından önerilen formül Eş. 2.23'deki şekildedir.

$$n \cong 1 + \frac{2}{9} \frac{1}{\left[\frac{Z(\alpha)}{2} + \sqrt{1 + \left(\frac{Z(\alpha)}{2} \right)^2 - \left(\frac{C_p^L}{\hat{C}_p} \right)^{2/3}} \right]^2} \quad (2.23)$$

Eş. 2.23'de C_p 'ye ilişkin istenen alt güven sınırına (C_p^L) karar vermek için gerekli örneklem büyüklüğü hesaplanabilmektedir. Verilen bu formül aracılığı ile örneklem büyüklüğü 10 gibi küçük değerler için bile doğru sonuçlar elde edildiği gösterilmiştir.

Eş. 2.23'de görüldüğü gibi örneklem büyüklüğü ($C_p^L/\hat{C}_p$) oranına bağlı ve bu oranın artan bir fonksiyonudur.

C_p 'ye ilişkin hipotez testleri

Süreç yeterlilik çalışmalarında sürecin yeterliliğine karar vermede kullanılan hipotezler aşağıdaki gibi kurulabilir:

$$H_0 : C_p \leq C \text{ (süreç yeterli değil)} \quad (2.24)$$

$$H_1 : C_p > C \text{ (süreç yeterli)} \quad (2.25)$$

Burada C önceden belirlenmiş bir yeterlilik ölçütüdür. Örnek olması açısından, Pearn ve diğerleri (1998), yokluk hipotezine karar vermek amacı ile tek örneklem için bir $\phi^*(x)$ testi tanımlamışlardır. Hipotez testinde C_p 'nin tahmin edicisi olan ve Eş. 2.14'te verilen $\tilde{C}_p$ kullanılmıştır. $\phi^*(x)$ test fonksiyonu;

$$\phi^*(x) = \begin{cases} 1, & \tilde{C}_p > c_0 \\ 0, & d.d \end{cases}$$

şeklinde tanımlanmıştır. $\tilde{C}_p > c_0$ olması durumunda; Eş. 2.24 deki yokluk hipotezini I. tip hata oranı $\alpha(c_0) = \alpha$ ile reddetmekte; diğer bir ifade ile, yeterli olmayan bir süreci yeterli bir süreç olarak kabul edilme olasılığını göstermektedir. Hipotez testinin kritik değeri c_0 olarak gösterilmiştir. Pearn ve diğerleri (1998), ϕ^* testinin α düzeyinde tüm yansız testler arasında II. tip hatası minimum olan UMP (uniformly most powerful) test olduğunu göstermişlerdir. Kritik değeri c_0 olmak üzere I. tip hata, $P\{\tilde{C}_p > c_0 | C_p = C\} = \alpha$ olarak tanımlanabilir. Bu tanım, Eş. 2.7 ve Eş. 2.14 kullanılarak, yeniden düzenlenirse Eş. 2.26'daki gibi yazılabilir.

$$P\left\{b_{n-1}(n-1)^{1/2}C_p\left(\frac{(n-1)s^2}{\sigma^2}\right)^{-1/2} \geq c_0 | C_p = C\right\} = \alpha \quad (2.26)$$

$$P\left\{\left(\frac{(n-1)s^2}{\sigma^2}\right) \leq b_{n-1}^2(n-1)\left[\frac{C}{c_0}\right]^2\right\} = \alpha$$

Böylece,

$$b_{n-1}^2(n-1)\left[\frac{C}{c_0}\right]^2 = \chi_{n-1,\alpha}^2 \quad (2.27)$$

eşitliği elde edilir. Buradan uygun c_0 değeri,

$$c_0 = \frac{\sqrt{n-1} b_{n-1} C}{\sqrt{\chi_{n-1,1-\alpha}^2}} \quad (2.28)$$

biçiminde elde edilir.

Böylece, $\tilde{C}_p > c_0$ ise $\phi^*(x) = 1$ dir ve yokluk hipotezi reddedilerek sürecin yeterli olduğu yorumu yapılır (Pearn ve diğerleri, 1998).

C_p 'nin güçsüz yönü

C_p indeksinin başlıca güçsüz yönü, sürecin gerçek ortalamasını dikkate almadan, gerçek süreç yayılımına göre tanımlanan potansiyel yeterliliği ölçmesidir. Bu nedenle C_p , süreç

ortalama değeri ile hedef değeri örtüşmediği sürece sürecin gerçek performansını göstermez. Bunun sonucu olarak, C_{pk} indeksi geliştirilmiştir. C_p ve C_{pk} indeksleri birlikte kullanıldığında süreç yayılımı ve sürecin konumu bakımından uygun bir yeterlilik indeksi elde edilmiş olur.

C_{pk} indeksi

C_p indeksi, sürecin nerede merkezlendiğini dikkate almaz ve süreç ortalama değeri ile hedef değeri örtüşmediği durumlarda sürecin gerçek performansını göstermez. Süreç ortalamasının hedef değerden sapmasını hesaplamak için C_p 'ye benzer biçimde, süreç ortalamasını da dikkate alan ve Eş. 2.29'da verilen C_{pk} indeksi önerilmiştir.

$$C_{pk} = \min \left\{ \frac{USL - \mu}{3\sigma}, \frac{\mu - LSL}{3\sigma} \right\} \quad (2.29)$$

$\min \{a, b\} = \frac{(a+b) - |a-b|}{2}$ tanımından faydalanılarak, $C_{pk} = \frac{d - |\mu - m|}{3\sigma}$ ile elde edilebilir. Burada, $d = \frac{USL - LSL}{2}$, spesifikasyon aralık uzunluğunun yarısı, $m = \frac{USL + LSL}{2}$ ise üst ve alt spesifikasyon sınırları arasındaki orta noktadır. C_p ve C_{pk} indeksleri ile k arasındaki doğrudan ilişki aşağıdaki gibi ifade edilebilir:

$$C_{pk} = C_p * (1 - k) \quad (2.30)$$

Burada $k = \frac{|\mu - \frac{1}{2}(LSL + USL)|}{d}$ olarak tanımlanmıştır (Kotz ve Johnson, 1993).

Kane (1986), $T = \frac{1}{2}(LSL + USL)$ olduğunu varsayarak, yukarıda verilen k 'yi $k = \frac{|\mu - T|}{d}$ olarak tanımlamıştır.

μ , spesifikasyon aralığı dışında ise C_{pk} negatif bir değer alacak ve böylece kalite değişkeni X bakımından sürecin yeterliliğinin düşük olduğunu gösterecektir.

C_{pk} indeksi ve uygun olmama yüzdesi (%NC)

Süreç performansı, kalite standartlarını sağlayan çıktıların yüzdesi olarak tanımlanır. Kalite değişkeninden ölçülen değerlerin spesifikasyon sınırlarıyla karşılaştırılmasına bağlı olarak ürün, uygun (geçti) veya uygun olmayan (red) olmak üzere iki şekilde karakterize edilebilir.

O halde üretimdeki uygun çıktıların oranı olarak adlandırılan p , Eş. 2.31'deki gibi tanımlanır:

$$p = \int_{LSL}^{USL} f(x) dx = F(USL) - F(LSL) \quad (2.31)$$

Burada, $F(x)$, X kalite değişkeninin dağılım fonksiyonu, USL spesifikasyon sınırlarının üst sınırını, LSL ise alt sınırını ifade eder. Süreç rastgele değişkeni X normal dağılmışsa, $\%NC$ uygun olmama yüzdesi (eski terminolojide kusurlu yüzdesi) $\Phi(\cdot)$ kümülatif standart normal dağılım fonksiyonu kullanılarak aşağıdaki gibi elde edilebilir:

$$\%NC = 1 - \Phi\left(\frac{USL - \mu}{\sigma}\right) + \Phi\left(\frac{\mu - LSL}{\sigma}\right) \quad (2.32)$$

C_{pk} 'ya dayalı olarak uygun olma güvencesi

C_{pk} indeksi, belli bir C_{pk} değeri ile normal dağılan bir süreç için süreç çıktı sınırlarını (bounds on the process yield) sağladığından, çıktı indeksi (yield-based index) olarak da adlandırılır ve aşağıdaki gibi elde edilebilir:

$$2\Phi(3C_{pk}) - 1 \leq p \leq \Phi(3C_{pk}) \quad (2.33)$$

Pearn ve Kotz (2006: 43), süreç çıktı sınırları ile uygun olmama yüzdesi $\%NC$ 'yi ilişkilendirebilmek için C_{pk} indeksini yeniden Eş. 2.34'de görüldüğü gibi tanımlamışlardır:

$$C_{pk} = \frac{d - |\mu - m|}{3\sigma} = \frac{1 - |(\mu - m)/d|}{3(\sigma/d)} = \frac{1 - |\delta|}{3\gamma} = \begin{cases} \frac{1 + \delta}{3\gamma}, & LSL \leq \mu \leq m \\ \frac{1 - \delta}{3\gamma}, & m \leq \mu \leq USL \end{cases} \quad (2.34)$$

Burada, $m = (USL + LSL)/2$ spesifikasyon sınırlarının orta noktasını, $d = (USL - LSL)/2$ spesifikasyon aralığının uzunluğunun yarısını ifade eder. $\delta = (\mu - m)/d$ ve $\gamma = \sigma/d$ ile elde edilir.

$C_{pk} > 0$ için, C_{pk} ve C_p 'nin bir fonksiyonu olarak uygun olmama yüzdesi $\%NC$ beklenen değerinin aşağıdaki eşitlikler yardımıyla yazılması mümkündür.

$m \leq \mu \leq USL$ için, Eş. 2.34 kullanılarak $C_{pk} = \frac{1-\delta}{3\gamma} = \frac{USL-\mu}{3\sigma}$ olarak yazılabilir. Burada,

$$\frac{LSL - \mu}{3\sigma} = \frac{(USL - \mu) - (USL - LSL)}{3\sigma} = \frac{USL - \mu}{3\sigma} - \frac{USL - LSL}{3\sigma} = C_{pk} - 2C_p$$

olarak elde edilir.

$LSL \leq \mu \leq m$ için, Eş. 2.34 kullanılarak $C_{pk} = \frac{1+\delta}{3\gamma} = \frac{\mu - LSL}{3\sigma}$ olarak yazılabilir. Aynı şekilde,

$$\frac{USL - \mu}{3\sigma} = \frac{(USL - LSL) - (\mu - LSL)}{3\sigma} = \frac{USL - LSL}{3\sigma} - \frac{(\mu - LSL)}{3\sigma} = 2C_p - C_{pk}$$

olarak elde edilir.

Elde edilen bu eşitlikler, Eş. 2.32'de verilen uygun olmama oranı kullanılarak düzenlenirse, Eş. 2.35 elde edilir (Kotz ve Johnson, 1993: 53).

$$\%NC = \Phi(-3C_{pk}) + \Phi[-3(2C_p - C_{pk})] \quad (2.35)$$

$C_{pk} = 1$ için Eş. 2.35'de verilmiş olan uygun olmayan ürün oranının en çok milyonda 2.700 birim çıktı olması beklenir. $C_{pk} = 1,33$ ise bu oranın en çok milyonda 66 olması beklenir. Uygun olmayan ürün oranının milyonda yaklaşık 0,554'den daha az olması için C_{pk} değerinin 1,67 olması gerekir. C_{pk} değerinin 2 olması ise uygun olmayan ürün oranının milyarda 2 birim çıktı olduğunu ifade eder (Pearn ve Kotz, 2006: 45).

Tek bir örneklem ile C_{pk} 'nın tahmin edilmesi

Süreç istatistiksel olarak kontrol altında iken, $\hat{C}_{pk}$ tahmin edicisi Eş. 2.36'da olduğu gibi elde edilebilir.

$$\hat{C}_{pk} = \frac{d - |\bar{x} - m|}{3s} = \left\{ 1 - \frac{|\bar{x} - m|}{d} \right\} \hat{C}_p \quad (2.36)$$

Burada, $\bar{x} = \sum_{i=1}^n \frac{x_i}{n}$, $s = \left[\sum_{i=1}^n (x_i - \bar{x})^2 / (n - 1) \right]^{1/2}$ olarak tanımlıdır.

$\hat{C}_{pk}$ 'nın r . momenti

Normallik varsayımı altında, Kotz ve Johnson (1993: 56-57), $\hat{C}_{pk}$ 'nın r . momentini Eş. 2.37'de verildiği gibi elde etmişlerdir.

$$\begin{aligned} E(\hat{C}_{pk}^r) &= \frac{1}{3^r} E\left(\frac{1}{s^r}\right) \sum_{j=0}^r (-1)^j \binom{r}{j} d^{r-j} E(|\bar{x} - m|^j) \\ &= \left(\frac{d\sqrt{n-1}}{3\sigma}\right)^r E(X_{n-1}^{-r}) \sum_{j=0}^r (-1)^j \binom{r}{j} \left(\frac{\sigma}{d\sqrt{n}}\right)^j E\left[\frac{\sqrt{n}(\bar{x} - m)}{\sigma}\right]^j \end{aligned} \quad (2.37)$$

$r = 1$ için; $\hat{C}_{pk}$ 'nın beklenen değerine Eş. 2.38'deki gibi ulaşılabilir.

$$E(\hat{C}_{pk}) = \frac{1}{3b_{n-1}} \left\{ \frac{d}{\sigma} - \sqrt{\frac{2}{\pi n}} e^{-\lambda/2} - \frac{|\mu - m|}{\sigma} [1 - 2\phi(-\sqrt{\lambda})] \right\} \quad (2.38)$$

$r = 2$ için ikinci moment Eş. 2.39'da olduğu gibi elde edilerek, Eş. 2.40'da verilen varyans değeri kolaylıkla hesaplanabilir.

$$E(\hat{C}_{pk}^2) = \left(\frac{d\sqrt{n-1}}{3\sigma}\right)^2 E(X_{n-1}^{-2}) \sum_{j=0}^2 (-1)^j \binom{2}{j} \left(\frac{\sigma}{d\sqrt{n}}\right)^j E\left[\frac{\sqrt{n}(\bar{x} - m)}{\sigma}\right]^j \quad (2.39)$$

$$\begin{aligned} Var(\hat{C}_{pk}) &= \frac{n-1}{9(n-3)} \left\{ \left(\frac{d}{\sigma}\right)^2 - 2\frac{d}{\sigma} \left[\sqrt{\frac{2}{\pi n}} e^{-\lambda/2} + \sqrt{\frac{\lambda}{n}} (1 - 2\phi(\sqrt{\lambda})) \right] \right. \\ &\quad \left. + \frac{\lambda + 1}{n} \right\} - [E(\hat{C}_{pk})]^2 \end{aligned} \quad (2.40)$$

Burada, merkezi olmama faktörü (non-centrality factor) $\lambda = n(\mu - m)^2/\sigma^2$ olarak ve düzeltme faktörü b_{n-1} ise Eş. 2.41'deki gibi tanımlanır.

$$b_{n-1} = (2/n - 1)^{1/2} \Gamma[(n-1)/2] / \Gamma[(n-1)/2] \quad (2.41)$$

Kotz ve Johnson (1993: 58), $\hat{C}_{pk}$ 'nin C_{pk} için yanlış bir tahmin edici olduğunu göstermişlerdir. $\mu = m$ ise, bu yan değeri $n \leq 10$ iken pozitif ancak n 'nin büyük değerleri için negatif değerler alır. n değeri sonsuza gittikçe yan değeri de sıfıra yaklaşır. $d/\sigma = 3$ veya 4 değeri olduğunda $n \geq 20$ için, $d/\sigma = 5$ olduğunda $n \geq 30$ için ve $d/\sigma = 6$ olduğunda $n \geq 40$ için yan değerleri negatiftir. n 'nin büyük değerleri için $\mu = m$ olduğunda, gözlemlere dayanarak yan değerinin negatif olduğuna sezgisel bir yorum getirilemez.

Ayrıca, yan değerini sıfır yapan tam bir n değerine ulaşılması henüz mümkün gözükmemektedir (Pearn ve Kotz, 2006).

$\hat{C}_{pk}$ 'nın dağılım özellikleri

Chou ve Owen (1989), Eş. 2.42'de verilen $\hat{C}_{pk}$ 'nın dağılımını elde etmeden önce $\hat{C}_{pu}$ ve $\hat{C}_{pl}$ 'nin dağılımını elde etmişlerdir.

$$\hat{C}_{pk} = \min(\hat{C}_{pu}, \hat{C}_{pl}) \quad (2.42)$$

Burada, $\hat{C}_{pu} = (USL - \bar{x})/3s$ ve $\hat{C}_{pl} = (\bar{x} - LSL)/3s$ olarak ifade edilir. $3\sqrt{n}\hat{C}_{pu}$ ve $3\sqrt{n}\hat{C}_{pl}$, normallik varsayımı altında $(n-1)$ serbestlik derecesi ile merkezi olmayan t dağılımına sahiptir ve merkezi olmama parametreleri $\sqrt{n}(USL - \mu)/\sigma$ ve $\sqrt{n}(\mu - LSL)/\sigma$ 'dır (Pearn ve Kotz, 2006: 48).

$\hat{C}_{pu} \leq y$ ya da $\frac{USL - \bar{x}}{3s} \leq y$ olarak yazılır ve düzenlenirse, $\frac{\sqrt{n}(\bar{x} - \mu)/\sigma - \sqrt{n}(USL - \mu)/\sigma}{s/\sigma} \geq -3\sqrt{ny}$ elde edilir. Bu ifadeyi daha sade biçimde, $\frac{Z + \delta_1}{s/\sigma}$ ile yazmak mümkündür. Burada, Z , standart normal dağılımını ifade etmekte, δ_1 ise, $\delta_1 = -\frac{\sqrt{n}(USL - \mu)}{\sigma} = -3\sqrt{n}C_{pu}$ eşitliğini göstermektedir.

$\frac{Z + \delta_1}{s/\sigma}$ değişkeni, $T_f(\delta_1)$ ile gösterilmiştir. Burada, $T_f(\delta_1)$, f serbestlik derecesine sahip merkezi olmayan t dağılımı göstermektedir ve δ_1 merkezi olmama parametresidir.

Böylece, $\hat{C}_{pu} \leq y$ eşitsizliği ancak ve ancak $T_f(\delta_1) \geq -3\sqrt{ny}$ olduğunda yazılabilir.

Benzer olarak, $\hat{C}_{pl} \leq y$ eşitsizliğini de, $T_f(\delta_2) \leq 3\sqrt{ny}$ olduğunda yazmak mümkündür. Burada, $\delta_2 = \frac{\sqrt{n}(\mu - LSL)}{\sigma} = 3\sqrt{n}C_{pl}$ eşitliğini göstermektedir.

$\hat{C}_{pu}$ 'nun dağılım fonksiyonu,

$$F_{\hat{C}_{pu}}(y) = P[T_f(\delta_1) \geq -3\sqrt{ny}] \quad (2.43)$$

$\hat{C}_{pu}$ 'nun olasılık yoğunluk fonksiyonu,

$$g_{\hat{C}_{pu}}(y) = 3\sqrt{n}g_{T_f(\delta_1)}(-3\sqrt{ny}), \quad -\infty \leq y \leq \infty \quad (2.44)$$

$\hat{C}_{pl}$ 'nin dağılım fonksiyonu,

$$F_{\hat{C}_{pl}}(y) = P[T_f(\delta_2) \leq 3\sqrt{ny}] \quad (2.45)$$

$\hat{C}_{pl}$ 'nin olasılık yoğunluk fonksiyonu,

$$g_{\hat{C}_{pl}}(y) = 3\sqrt{n}g_{T_f(\delta_2)}(3\sqrt{ny}), \quad -\infty \leq y \leq \infty \quad (2.46)$$

şeklinde tanımlıdır.

Eş. 2.42'ye denk olarak, $T_f(\delta_1) \leq -3\sqrt{ny}$ ve $T_f(\delta_2) \geq 3\sqrt{ny}$ yazılabilir.

Böylece $\hat{C}_{pk}$ 'nin dağılım fonksiyonu,

$$F_{\hat{C}_{pk}}(y) = 1 - P[T_f(\delta_1) \leq -3\sqrt{ny}, T_f(\delta_2) \geq 3\sqrt{ny}] \quad (2.47)$$

olarak elde edilmiştir. $(T_f(\delta_1), T_f(\delta_2))$ bir korelasyon katsayısı ile iki değişkenli merkezi olmayan t dağılımına sahiptir. Chou ve Owen (1989), $\hat{C}_{pk}$ 'nin dağılımını elde etmek için iki bağımlı merkezi olmayan t değişkenin bileşik fonksiyonu için Owen (1965)'in formülünü kullanmışlardır. $\hat{C}_{pk}$ 'nin olasılık yoğunluk fonksiyonu $g_{\hat{C}_{pk}}(y)$ daha karmaşık olarak Eş. 2.48'de verilmektedir.

$$g_{\hat{C}_{pk}}(y) = \begin{cases} 3\sqrt{n} \sum_{i=1}^2 g_{T_{n-1}, \delta_i}(t_i) & y \leq 0 \text{ için,} \\ 3\sqrt{n} \sum_{i=1}^2 \frac{n-1}{t_i} \left[Q_{n+1} \left(\sqrt{\frac{n+1}{n-1}} t_i, \delta_i; O, R \right) - Q_{n-1}(t_i, \delta_i; O, R) \right] & y > 0 \text{ için,} \end{cases} \quad (2.48)$$

Burada, $t_1 = -t_2 = 3\sqrt{ny}$ ve $\delta_1 = -3\sqrt{n}C_{pu}$, $\delta_2 = 3\sqrt{n}C_{pl}$,

$$R = \sqrt{n-1} (\delta_2 - \delta_1) / (t_2 - t_1)$$

ve $g_{T_{n-1}, \delta_i}(\cdot)$ ise $n-1$ serbestlik derecesine ve δ_i parametresine sahip merkezi olmayan t

dağılımını gösterir. $Q_f(t, \delta; O, R)$ fonksiyonu Eş. 2.49'da verilmiştir (Chou ve Owen, 1989).

$$Q_f(t, \delta; O, R) = C(f) \int_0^R \left[\phi \left(\frac{t_x}{f} - \delta \right) x^{f-1} \phi(x) \right] dx \quad (2.49)$$

C_{pk} için güven aralıkları

C_{pk} 'nın güven aralığını tam olarak elde edebilmek, $\hat{C}_{pk}$ 'nın dağılımının merkezi olmayan t dağılımına sahip iki değişkenin ortak dağılımını içermesi nedeniyle zordur.

Heavlin (1988), %100(1 - α) güven aralığını yaklaşık olarak Eş. 2.50'de verildiği gibi önermiştir.

$$\left[\hat{C}_{pk} - z_{\alpha/2} \sqrt{\frac{n-1}{9n(n-3)} + \frac{\hat{C}_{pk}^2}{2(n-3)} \left(1 + \frac{6}{n-1}\right)}, \hat{C}_{pk} + z_{\alpha/2} \sqrt{\frac{n-1}{9n(n-3)} + \frac{\hat{C}_{pk}^2}{2(n-3)} \left(1 + \frac{6}{n-1}\right)} \right] \quad (2.50)$$

Zhang, Stenback ve Wardrop, (1990), normal dağılım varsayımına dayalı olarak $\hat{C}_{pk}$ 'nın örnekleme dağılımının varyansı ve ortalamasına ilişkin tanımlamalar elde etmiş ve %100(1 - α) güven aralığını Eş. 2.51'deki biçimde formüle etmişlerdir.

$$\left[\left(1 - z_{\alpha/2} \sqrt{\frac{n-1}{n-3} - \frac{1}{b_{n-1}^2}}\right) \hat{C}_{pk}, \left(1 + z_{\alpha/2} \sqrt{\frac{n-1}{n-3} - \frac{1}{b_{n-1}^2}}\right) \hat{C}_{pk} \right] \quad (2.51)$$

Franklin ve Wasserman (1992) ise, $n \geq 30$ olduğunda, Bissell'in (1990) yönteminin %95 alt kontrol limitini doğru bir şekilde elde etmedeki performansını $\hat{C}_{pk}$ 'nın birkaç değeri için simülasyon yaparak göstermiştir. Bissell'in (1990) yöntemi Eş. 2.52'de verilmektedir.

$$C_{pk}^{lcl} = \hat{C}_{pk} - z_{\alpha} \sqrt{\frac{1}{9n} + \frac{(\hat{C}_{pk})^2}{2(n-1)}} \quad (2.52)$$

Burada, z_{α} standart normal dağılımın üst α yüzdesini göstermektedir (Pearn ve Kotz, 2006: 51). Nagata ve Nagata (1994) Eş. 2.52'de verilen formüle $1/30\sqrt{n}$ ifadesini ekleyerek, iki yanlı güven aralıklarını hesaplamada simülasyon tekniği ile doğru sonuçlar elde etmişlerdir. Nagata ve Nagata (1992) tarafından C_{pk} için iki yanlı güven aralığı hesaplamada kullanılan başka bir formül Eş. 2.53'de görüldüğü gibi verilmiştir.

$$\left[\hat{C}_{pk} - z_{\alpha/2} \sqrt{\frac{1}{9n} + \frac{C_{pk}^2}{2(n-1)}}, \hat{C}_{pk} + z_{\alpha/2} \sqrt{\frac{1}{9n} + \frac{C_{pk}^2}{2(n-1)}} \right] \quad (2.53)$$

Eş. 2.53'de verilen formülden $1/(9n)$ ifadesi silindiğinde, C_{pk} 'nın yaklaşık güven aralığı Eş. 2.54'e indirgenebilir.

$$\left(1 - \frac{z_{\alpha}}{\sqrt{2(n-1)}} \right) \hat{C}_{pk}, \left(1 + \frac{z_{\alpha}}{\sqrt{2(n-1)}} \right) \hat{C}_{pk}, \quad (2.54)$$

Pearn ve Chen (1998), asimetrik tolerans sınırlarına sahip süreçler için $T = [3(USL) - LSL]/4$ olan A ve B gibi iki süreci ele almışlardır. A sürecinin ortalaması $\mu_A = T$, B sürecinin ortalaması $\mu_B = LSL + d/2$; standart sapmaları $\sigma_A = \sigma_B = d/6$ olduğu bilinen bu iki süreç için Eş. 2.29'da $C_{pk} = 1$ olduğu görülebilir. Her iki süreç için uygun olmama beklenen yüzdeleri %0,135 ise; A sürecinin üretimi hedef değerde gerçekleşirken; B sürecinin üretiminde hedef değerden uzaklaştığı dikkate alındığında, C_{pk} indeksi bu iki süreci ayırmada başarılı değildir. Bu problemi çözmek amacı ile C_{pk}'' olarak adlandırdıkları bir indeks önermişler ve bu indeks diğer genelleştirilmiş C_{pk} indeksleri ile karşılaştırmıştır. Eş. 2.55'te önerilen C_{pk}'' indeksinin diğerlerine göre daha üstün olduğu görülmüş ve

$$C_{pk}'' = \frac{d^* - A^*}{3\sigma} \quad (2.55)$$

olarak tanımlanmıştır.

Burada, $A^* = \max \left\{ \frac{d^*(\mu - T)}{D_u}, \frac{d^*(\mu - T)}{D_l} \right\}$; $D_u = USL - T$; $D_l = T - LSL$ ve $d^* = \min \{D_l, D_u\}$ ile ifade edilir (Pearn ve Chen, 1998).

C_{pk} tahmini için örneklem büyüklüğüne karar verilmesi

C_{pk} 'nın tahmin edicisinin örneklem dağılımının C_p 'ye göre daha karmaşık olması, ilgili tahminlerin yapılmasını güçleştirmektedir. C_{pk} 'nın tahmini için örneklem büyüklüğüne karar verilmesi ile ilgili olarak Kushler ve Hurley (1992), Bissell'in yaklaşımının (1990) diğer yaklaşımlara göre hesaplama bakımından daha kolay olduğunu ve daha doğru sonuçlar verdiğini belirtmektedir.

Franklin ve Wasserman (1992) $n \geq 30$ için, Bissell'in yaklaşımının (1990) simülasyonla C_{pk} 'nin alt %95 güven limiti için doğru sonuçlar verdiğini göstermiştir.

Eş. 2.52'de verilen Bissell'in yaklaşımının alt güven sınırı $n \geq 30$ için, $2n - 2 \cong 2n$ olarak alınarak, C_{pk}^{ags} değeri Eş. 2.56'daki gibi elde edilmiştir.

$$C_{pk}^{ags} \cong \hat{C}_{pk} - \frac{z_{\alpha}}{\sqrt{n}} \sqrt{\frac{1}{9} + \frac{(\hat{C}_{pk})^2}{2}} \quad (2.56)$$

Eş. 2.56 kullanılarak örneklem büyüklüğü,

$$n = \frac{(z_{\alpha})^2 \left[\frac{1}{9} + \frac{(\hat{C}_{pk})^2}{2} \right]}{(\hat{C}_{pk} - C_{pk}^{lcl})^2} = (z_{\alpha})^2 \frac{\left[\frac{1}{9(\hat{C}_{pk})^2} + \frac{1}{2} \right]}{(1 - C_{pk}^{ags} / \hat{C}_{pk})^2} \quad (2.57)$$

olur. Burada n değeri, $\hat{C}_{pk} - C_{pk}^{ags}$, α , $\hat{C}_{pk}$ ve $C_{pk}^{ags} / \hat{C}_{pk}$ değerlerine bağlı olarak elde edilmektedir (Franklin, 1999).

C_{pk} 'ya ilişkin hipotez testleri

C_{pk} ile ilgili hipotezler aşağıdaki şekilde kurulur.

$$H_0 : C_{pk} \leq C \text{ (süreç yeterli değil)} \quad (2.58)$$

$$H_1 : C_{pk} > C \text{ (süreç yeterli)} \quad (2.59)$$

H_0 hipotezi hakkında karar vermek için kullanılan p – değeri yeterli olmayan bir sürecin ($H_0 : C_{pk} \leq C$), yeterli bir süreç olarak ($H_1 : C_{pk} > C$) yorumlanma riskini ifade eder. Bu durumda p – değeri $< \alpha$ ise yokluk hipotezi reddedilir ve süreç yeterli olarak değerlendirilir.

Pearn ve Lin (2002), $\hat{C}_{pk}$ 'nin dağılım fonksiyonuna dayalı olarak p – değeri hesaplamada uygulama kolaylığı sağlayan pratik bir yöntem geliştirmiştir.

Pearn ve Kotz (2006: 57), Eş. 2.58'de verilen yokluk hipotezi hakkında karar vermek için Eş. 2.60'daki $\phi^*(x)$ testini tanımlamıştır.

$$\phi^*(x) = \begin{cases} 1, & \hat{C}_{pk} > c_0 \\ 0, & d.d \end{cases} \quad (2.60)$$

Eş. 2.60'da $\hat{C}_{pk} > c_0$ ise, yokluk hipotezi ($H_0 : C_{pk} \leq C$), I. tip hata olasılığı $\alpha(c_0) = \alpha$ ile reddedileceği ifade edilmektedir. α ve C değerinin verilmesi durumunda c_0 değeri, $P(\hat{C}_{pk} > c_0 | C_{pk} = C) = \alpha$ eşitliği çözülerek elde edilebilir.

Spesifikasyon sınırlarının orta noktası olarak ifade edilen m değeri ile hedef değeri T 'nin aynı olduğu ($T = m$) süreçler için Eş. 2.34'te yeniden tanımlanan C_{pk} indeksi, $C_{pk} = (\frac{d}{\sigma} - |\xi|) / 3$ olarak yazılabilir. Burada, $\xi = (\mu - m) / \sigma$ 'dır. $C_{pk} = C$ olarak verilirse, $b = d / \sigma$ eşitliği $b = 3C + |\xi|$ ile ifade edilebilir. Yeterlilik koşulu C bilindiğinde c^* a ilişkin p – değeri,

$$\begin{aligned} p - \text{değeri} &= P(\hat{C}_{pk} > c^* | C_{pk} = C) \\ &= \int_0^{b\sqrt{n}} G\left(\frac{(n-1)(b\sqrt{n}-t)^2}{9n(c^*)^2}\right) [\phi(t + \xi\sqrt{n}) + \phi(t - \xi\sqrt{n})] dt \end{aligned}$$

eşitliği çözülerek hesaplanabilir.

Benzer olarak, yeterlilik koşulu C , parametre ξ , örneklem büyüklüğü n ve I. tip hata olasılığı α değerleri için c_0 kritik değeri,

$$\int_0^{b\sqrt{n}} G\left(\frac{(n-1)(b\sqrt{n}-t)^2}{9n(c_0)^2}\right) [\phi(t + \xi\sqrt{n}) + \phi(t - \xi\sqrt{n})] dt = \alpha \text{ eşitliği çözülerek elde edilebilir.}$$

C_{pm} indeksi

Hsiang ve Taguchi (1985) tarafından tanımlanan C_{pm} indeksi, Taguchi'nin kayıp fonksiyon fikrine dayanır ve Taguchi indeksi olarak da adlandırılır. C_{pm} indeksi, hedef değer (T) etrafındaki kümelenmeyi dikkate aldığından, sürecin merkezlenme derecesini daha iyi yansıtır. Bu nedenle, süreç ortalamasının konumu hakkında C_p ve C_{pk} indekslerine göre daha sağlıklı bir bilgi sağlar. C_{pm} indeksi;

$$C_{pm} = \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} = \frac{USL - LSL}{6\tau} = \frac{d}{3\tau} \quad (2.61)$$

olarak ifade edilir. Burada, $USL - LSL$, sürecin kabul edilebilir tolerans aralığını; $d = USL - LSL/2$, spesifikasyon aralık uzunluğunun (spesifikasyon bandının) yarısını; τ ise hedef değerden (T) ortalama sapmayı ifade etmektedir. $\tau^2 = \sigma^2 + (\mu - T)^2 = E[(X - T)^2]$ terimi, iki varyans bileşeninden oluşur. Bunlardan ilki, süreç ortalamasına göre, diğeri de süreç ortalamasının hedeften farkına göre belirlenir. $E[(X - T)^2]$ kayıp değerinin beklenen değeri olduğu için, süreç değişkeni X 'in kayıp fonksiyonunun, ($Kayıp(X) = k(X - T)^2$) bir k pozitif sabit değerleri ile simetrik bir karesel hata fonksiyonuna yaklaştığı varsayılır. Bu nedenle C_{pm} yeterlilik indeksi kayıp değerine dayalı bir indeks olarak tanımlanır.

C_{pm} indeksinin Eş. 2.61'de görüldüğü gibi, süreç varyansı arttığında payda değeri de büyüyecek, böylece C_{pm} indeksi küçülecektir. Ayrıca, süreç ortalaması hedef değerden uzaklaştığında C_{pm} indeksinin payda değerinin artacağı, böylece C_{pm} indeks değerinin azalacağı; aynı şekilde, süreç ortalaması hedef değere yaklaştığında C_{pm} indeksinin payda değerinin azalacağı, böylece C_{pm} indeks değerinin artacağı söylenebilir.

C_{pm} indeksi süreç değişkenliğini ölçmesi bakımından C_p indeksinden farklıdır. C_{pm} indeksi C_p ve C_{pk} indeksleri cinsinden Eş. 62'de görüldüğü gibi yazılabilir (Pearn ve Kotz, 2006: 67-68).

$$C_{pm} = \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} = \frac{C_p}{\sqrt{1 + \left(\frac{\mu - T}{\sigma}\right)^2}} = \frac{C_{pk}}{\left(1 + \frac{|\mu - m|}{\sigma}\right) \sqrt{1 + \left(\frac{\mu - T}{\sigma}\right)^2}} \quad (2.62)$$

Kotz ve Johnson (1999) $T = m$ varsayımı altında tanımlanan $k = |\mu - m|/d$ ifadesini kullanarak,

$$C_{pk} = C_p(1 - k) \quad (2.63)$$

$$C_{pm} = C_p(1 + 9C_p^2 k^2)^{-1/2} \quad (2.64)$$

eşitliklerini yazmıştır.

Eş. 2.63 ve Eş. 2.64'de verildiği gibi C_{pk} ve C_{pm} indekslerini C_p indeksi cinsinden

karşılaştırmışlardır. Böylece, $0 \leq k \leq 1$ için, C_p 'nin küçük değerleri için C_{pm} indeksi C_{pk} indeksinden büyük olur. C_p 'nin büyük değerleri için C_{pm} indeksi C_{pk} indeksinden küçük olur. Aynı zamanda, C_{pm} indeksini, C_p ve C_{pk} indekslerini içerecek şekilde Eş. 2.65'de verildiği gibi de formüle etmişlerdir.

$$C_{pm} = C_p \left\{ 1 + 9(C_p - C_{pk})^2 \right\}^{-1/2} \quad (2.65)$$

C_{pm} indeksi ve uygun olmama yüzdesi (%NC)

$T = m$ varsayımı altında, Eş. 2.66'da verilen C_{pm} indeksi $\gamma = \sigma/d$ ve $\xi = (\mu - m)/d$ 'nin bir fonksiyonu olarak Eş. 2.66'daki gibi yeniden yazılabilir.

$$C_{pm} = \frac{d}{3\sqrt{\sigma^2 + (\mu - m)^2}} = \frac{1}{3\sqrt{\gamma^2 + \xi^2}} \quad (2.66)$$

Eş. 2.66'dan ξ değeri,

$$\gamma^2 = \frac{1}{(3C_{pm})^2} - \xi^2 = \left(\frac{1}{3C_{pm}} + \xi \right) \left(\frac{1}{3C_{pm}} - \xi \right) \quad (2.67)$$

ya da

$$\gamma = \sqrt{\left(\frac{1}{3C_{pm}} + \xi \right) \left(\frac{1}{3C_{pm}} - \xi \right)} = \sqrt{\left(\frac{1}{3C_{pm}} + |\xi| \right) \left(\frac{1}{3C_{pm}} - |\xi| \right)} \quad (2.68)$$

olarak elde edilir. Burada,

$0 \leq |\xi| \leq 1/(3C_{pm})$ 'dir. Örneğin, $1 - \frac{1}{3C_{pm}} \leq C_a \leq 1$ için uygun olmama yüzdesi, normal dağılım varsayımı altında,

$$\begin{aligned} \%NC &= 1 - \phi \left(\frac{USL - \mu}{\sigma} \right) + \phi \left(\frac{\mu - LSL}{\sigma} \right) \\ &= \phi \left[-\frac{2 - C_a}{\sqrt{\frac{1}{3C_{pm}^2} - (1 - C_a)^2}} \right] + \phi \left[-\frac{C_a}{\sqrt{\frac{1}{3C_{pm}^2} - (1 - C_a)^2}} \right] \end{aligned} \quad (2.69)$$

olarak elde edilir. Burada, C_a sürecin merkezlenme derecesini ölçmede kullanılan ve Eş. 2.70'te ifade edilen doğruluk (accuracy) indeksidir.

$$C_a = 1 - \frac{|\mu - m|}{d} \quad (2.70)$$

C_a indeksi, merkez etrafında kümelenme yeteneğini gösterirken, kullanıcıyı süreç ortalamasının genellikle hedef değere eşit olan m değerinden sapmasına karşı uyarır (Pearn ve Kotz, 2006: 70-72).

Tek bir örneklem ile C_{pm} 'nin tahmin edilmesi

Chan ve diğerleri (1988), C_{pm} indeksi için Eş. 2.71'deki,

$$\tilde{C}_{pm} = \tilde{C}_{pm(CCS)} = \frac{d}{3\tilde{\tau}_{CCS}} = \frac{d}{3\sqrt{\sum_{i=1}^n (x_i - T)^2 / (n-1)}} = \frac{d}{3\sqrt{s^2 + \frac{n}{n-1}(\bar{x} - T)^2}} \quad (2.71)$$

tahmin ediciyi önermiştir. Bir başka tahmin edici de Boyles (1991) tarafından

$$\hat{C}_{pm} = \hat{C}_{pm(B)} = \frac{d}{3\tilde{\tau}_B} = \frac{d}{3\sqrt{\sum_{i=1}^n (x_i - T)^2 / (n)}} = \frac{d}{3\sqrt{s_n^2 + \frac{n}{n-1}(\bar{x} - T)^2}} \quad (2.72)$$

biçiminde önerilmiştir.

Süreç değişkeninin normal dağıldığı ve $T = m$ varsayımı altında Chan ve diğerleri (1988), $\tilde{C}_{pm}$ tahmin edicisinin olasılık yoğunluk fonksiyonunu, Eş. 2.73'de verildiği gibi tanımlamışlardır.

$$f_Y(y) = \frac{a}{2^{\frac{n}{2}-1}y^3} \exp \left[-\frac{1}{2} \left(\frac{a}{y^2} + \lambda \right) \right] \sum_{j=0}^{\infty} \frac{\lambda^j \left(\frac{a}{y^2} \right)^{\frac{n}{2}+j-1}}{j! \Gamma(\frac{n}{2} + j) 2^{2j}}, \quad y > 0 \quad (2.73)$$

Burada, $Y = \tilde{C}_{pm}$, $a = C_{pm}^2 \left(1 + \frac{\lambda}{n} \right) (n-1)$ ve $\lambda = n(\mu - T)^2 / \sigma^2$ olarak ifade edilmiştir. $Y = \tilde{C}_{pm}$ tahmin edicisinin dağılım fonksiyonu ise Eş. 2.74'te

$$F_Y(y) = 1 - \exp\left(-\frac{\lambda}{2}\right) \sum_{j=0}^{\infty} \frac{\left(\frac{\lambda}{2}\right)^j}{j! \Gamma\left(\frac{n}{2} + j\right)} \int_0^{\frac{a}{2y^2}} t^{\frac{n}{2}+j-1} e^{-t} dt, \quad y > 0 \quad (2.74)$$

olarak tanımlanmıştır. Eş. 2.71 ve Eş. 2.72’de tanımlanan tahmin ediciler incelendiğinde asimptotik olarak eşit olduğu görülebilir (Pearn ve Kotz, 2006: 74-75).

$Y = \tilde{C}_{pm}$ ’nin beklenen değeri Eş. 2.75’de verildiği gibidir.

$$E(\tilde{C}_{pm}) = \frac{\left(\frac{n-1}{2}\right)^{1/2} \Gamma\left(\frac{n-1}{2}\right)}{\Gamma\left(\frac{n}{2}\right)} C_{pm} \quad (2.75)$$

$\tilde{C}_{pm}$ ’nin Ortalama Hata Karesi (MSE) ise,

$$MSE(\tilde{C}_{pm}) = C_{pm}^2 \left(\frac{n-1}{2}\right) \left\{ \frac{\Gamma\left(\frac{n}{2}-1\right)}{\Gamma\left(\frac{n}{2}\right)} - \frac{\Gamma^2\left(\frac{n-1}{2}\right)}{\Gamma^2\left(\frac{n}{2}\right)} \right\} \quad (2.76)$$

olarak tanımlanır. $\tilde{C}_{pm}$ tahmin edicisi C_{pm} indeksi için yanlış ancak, asimptotik olarak yansız bir tahmin edicidir.

$\tilde{C}_{pm}$ ’nin "Yan" değeri;

$$\begin{aligned} Yan(\tilde{C}_{pm}) &= E(\tilde{C}_{pm}) - C_{pm} \\ &= \frac{\left(\frac{n-1}{2}\right)^{1/2} \Gamma\left(\frac{n-1}{2}\right)}{\Gamma\left(\frac{n}{2}\right)} C_{pm} - C_{pm} \\ &= C_{pm} \left(\frac{\left(\frac{n-1}{2}\right)^{1/2} \Gamma\left(\frac{n-1}{2}\right)}{\Gamma\left(\frac{n}{2}\right)} - 1 \right) \end{aligned} \quad (2.77)$$

olarak elde edilir. Buradan,

$$\lim_{n \rightarrow \infty} Yan(\tilde{C}_{pm}) = \lim_{n \rightarrow \infty} C_{pm} \left(\frac{\left(\frac{n-1}{2}\right)^{1/2} \Gamma\left(\frac{n-1}{2}\right)}{\Gamma\left(\frac{n}{2}\right)} - 1 \right) = C_{pm} (1 - 1) = 0$$

olur ki bu da $\tilde{C}_{pm}$ tahmin edicisinin asimptotik olarak yansız olduğunu göstermektedir (Stoumbos, 2002).

$\hat{C}_{pm}$ 'nin r. momenti

Kotz ve Johnson (1993: 99-100), $\hat{C}_{pm}$ 'nin r. momenti için

$$\begin{aligned}\mu'_r(\hat{C}_{pm}) &= \left(\frac{d\sqrt{n}}{3\sigma}\right)^r \exp\left(-\frac{\lambda}{2}\right) \sum_{j=0}^{\infty} \frac{\left(\frac{\lambda}{2}\right)^j}{j!} E\left[\left(\chi_{n+2j}^2\right)^{-r/2}\right] \\ &= \left(\frac{d\sqrt{n}}{3\sigma}\right)^r \exp\left(-\frac{\lambda}{2}\right) \sum_{j=0}^{\infty} \frac{\left(\frac{\lambda}{2}\right)^j}{j!} \frac{\Gamma\left(\frac{n-1}{2} + j\right)}{\Gamma\left(\frac{n}{2} + j\right)}\end{aligned}\quad (2.78)$$

formülasyonunu türetmiştir. Burada, ilk iki moment,

$$E(\hat{C}_{pm}) = \frac{d\sqrt{n}}{3\sigma} \exp\left(-\frac{\lambda}{2}\right) \sum_{j=0}^{\infty} \frac{\left(\frac{\lambda}{2}\right)^j}{j!} \frac{\Gamma\left(\frac{n-1}{2} + j\right)}{\Gamma\left(\frac{n}{2} + j\right)} \quad (2.79)$$

$$Var(\hat{C}_{pm}) = \left(\frac{d\sqrt{n}}{3\sigma}\right)^2 \exp\left(-\frac{\lambda}{2}\right) \sum_{j=0}^{\infty} \left\{ \frac{\left(\frac{\lambda}{2}\right)^j}{j!} \frac{1}{n-2+2j} \right\} - E[\hat{C}_{pm}]^2 \quad (2.80)$$

ile elde edilir.

Eş. 2.79'da, $d/3\sigma$ yerine $d/3\sigma = C_p$ eşitliğini ve λ yerine $\lambda = n(\mu - T)^2/\sigma^2$ eşitliğini yazılarak yeniden düzenleyen Boyles (1991), $E(\hat{C}_{pm})$ 'nin doğru bir yaklaşımını önermiş ve $n > 100$ için normal bir yaklaşım geliştirmiştir.

$$E(\hat{C}_{pm}) = \left(\frac{n(n+\lambda)}{2(n+2\lambda)} \frac{\Gamma\left[\frac{(n-1+\lambda)^2 + (n-1)}{2(n+2\lambda)}\right]}{\Gamma\left(\frac{(n+\lambda)^2}{2(n+2\lambda)}\right)} \right) C_p \quad (2.81)$$

Eş. 2.81'de görüldüğü gibi, $\hat{C}_{pm}$, C_{pm} 'nin yanlı bir tahmin edicisidir. $\hat{C}_{pm}$ 'nin "Yan" değeri;

$Yan(\hat{C}_{pm}) = E(\hat{C}_{pm}) - C_{pm} = C_p \left[\sqrt{\frac{n}{2}} \frac{\Gamma\left(\frac{n-1}{2}\right)}{\Gamma\left(\frac{n}{2}\right)} - 1 \right]$ olarak elde edilebilir (Pearn ve Kotz, 2006: 76).

C_{pm} için güven aralıkları

$$C_{pm} = \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} = \frac{USL - LSL}{6\sqrt{\sigma^2 + \left(\frac{\lambda}{n}\right)\sigma^2}} \quad (2.82)$$

ifadesi,

$$C_{pm} = \frac{USL - LSL}{6\sigma\sqrt{1 + \left(\frac{\lambda}{n}\right)}} \text{ ya da } C_{pm}\sqrt{1 + \left(\frac{\lambda}{n}\right)} = \frac{USL - LSL}{6\sigma} \text{ olarak yazılabilir. Böylece,}$$

$$\hat{C}_{pm} \sim C_{pm}\sqrt{1 + \left(\frac{\lambda}{n}\right)}\sqrt{\frac{n}{\chi^2_{n,\lambda}}} \text{ şekilde ifade edilebilir (Zimmer, Hubele ve Zimmer, 2001).}$$

Boyles (1991) ve Pearn ve diğerleri (1992) $\hat{C}_{pm}$ 'nin,

$$\hat{C}_{pm} \sim \frac{USL - LSL}{6\sigma} \sqrt{\frac{n}{\chi^2_{n,\lambda}}} \quad (2.83)$$

olarak dağıldığını göstermiştir. Burada, $\chi^2_{n,\lambda}$, n serbestlik derecesine ve $\lambda = \frac{n(\mu - T)^2}{\sigma^2}$ merkezi olmama parametresi ile merkezi olmayan Ki-kare dağılımına sahiptir (Kotz ve Johnson, 1993: 99).

C_{pm} ile merkezi olmayan Ki-kare dağılımı arasındaki ilişki Zimmer ve diğerleri (2001) tarafından gösterilmiştir.

Zimber ve Hubele (1997), $\hat{C}_{pm}$ tahmin edicisinin örnekleme dağılımı için yüzdeler tabloları elde etmişlerdir. C_{pm} 'nin $\% a(1 - \alpha)$ 100 güven aralığı,

$$\left(\frac{\hat{C}_{pm}}{\sqrt{1 + \frac{\lambda}{n}} \sqrt{\frac{n}{\chi^2_{(1-\alpha/2),n,\lambda}}}}, \frac{\hat{C}_{pm}}{\sqrt{1 + \frac{\lambda}{n}} \sqrt{\frac{n}{\chi^2_{(\alpha/2),n,\lambda}}}} \right) \quad (2.84)$$

ile ifade edilmiştir. (Zimmer ve diğerleri, 2001).

C_{pm} tahmini için örneklem büyüklüğüne karar verilmesi

Örneklem büyüklüğüne karar verilmesi, maliyetle ilişkili olduğundan veri toplama işinin önemli bir adımıdır. Franklin (1999), örneklem büyüklüğüne karara vermek için Boyles (1991) tarafından önerilen alt kontrol sınırını kullanarak bir formül önermiştir.

$$n = \frac{2}{9} \frac{(1 + 2\xi^2)}{(1 + \xi^2)^2} \frac{1}{\frac{z_\alpha}{2} \sqrt{1 + \left(\frac{z_\alpha}{2}\right)^2 - \left(\frac{C_{pm}^{lcl}}{\hat{C}_{pm}}\right)^{2/3}}} \quad (2.85)$$

Bu eşitliği elde edebilmek için C_{pm} 'nin istenen alt kontrol sınırına (C_{pm}^{lcl}) bağlı olarak ($C_{pm}^{lcl}/\hat{C}_{pm}$) oranına ve $\xi = (\mu - T)/\sigma$ değerine ihtiyaç duyulmaktadır.

C_{pm} 'ye ilişkin hipotez testleri

C_{pm} indeksi kullanılarak, bir sürecin yeterli olup olmadığına karar vermek için Eş. 2.86-Eş. 2.87'deki hipotezler test edilebilir.

$$H_0 : C_{pm} \leq C \text{ (süreç yeterli değil)} \quad (2.86)$$

$$H_1 : C_{pm} > C \text{ (süreç yeterli)} \quad (2.87)$$

Karar kuralına göre, $\hat{C}_{pm} > c_0$ ise yokluk hipotezi reddedilir. Burada, c_0 kritik değerdir.

Hedef değeri, spesifikasyon sınırlarının ortasında yer alan ($T = m$) süreçler için C_{pm} indeksi yeniden Eş. 2.88'deki gibi yazılabilir.

$$C_{pm} = \frac{b}{\left(3(1 + \xi^2)^{1/2}\right)} \quad (2.88)$$

$C_{pm} = C$ ve $b = d/\sigma$ olarak verildiğinde, $b = 3C(1 + \xi^2)^{1/2}$ olarak ifade edilebilir. p – değeri ise,

$$p - \text{değeri} = \int_0^{b\sqrt{n}/3c^*} G\left(\frac{b^2n}{9(c^*)^2} - t^2\right) \phi(t + \xi\sqrt{n}) + \phi(t - \xi\sqrt{n}) dt \quad (2.89)$$

ile hesaplanır. Burada, $G(\cdot)$ χ_{n-1}^2 in dağılım fonksiyonu ve $\phi(\cdot)$ ise standart normal olasılık yoğunluk fonksiyonunu göstermektedir. c^* , örneklem verisinden hesaplanan $\hat{C}_{pm}$ nin belli bir değeridir.

Verilen α ve C değerleri için, c_0 kritik değeri,

$$P(\hat{C}_{pm} \geq c_0 | C_{pm} = C) = \alpha \quad (2.90)$$

eşitliği çözülerek elde edilebilir.

Yeterlilik gereksinimi olarak adlandırılan C , ξ parametresi, örneklem büyüklüğü n ve I. tip hata olasılığı α verildiğinde, Eş. 2.91 çözülerek kritik değer c_0 elde edilebilir.

$$\int_0^{b\sqrt{n}/3c_0} G\left(\frac{b^2n}{9(c_0)^2} - t^2\right) \phi(t + \xi\sqrt{n}) + \phi(t - \xi\sqrt{n}) dt = \alpha \quad (2.91)$$

$p < \alpha$ için yokluk hipotezi reddedilir ve sürecin yeterli olduğu söylenebilir (Pearn ve Kotz, 2006).

Vannman (1995), tek değişkenli ve C_p , C_{pk} , C_{pm} ve C_{pmk} indekslerini içeren, $C_p(u, v)$ olarak adlandırdığı başka bir yeterlilik indeksi önermiştir. Bu indeks Eş. 2.92'de verilmiştir.

$$C_p(u, v) = \frac{d - u|\mu - m|}{3\sqrt{\sigma^2 + v(\mu - T)^2}} \quad (2.92)$$

Burada, μ bilindiği gibi süreç ortalamasını, σ süreç standart sapmasını, $d = (USL - LSL)/2$ spesifikasyon aralık uzunluğunun yarısını, $m = (USL + LSL)/2$ üst ve alt spesifikasyon sınırları arasındaki orta noktayı ve T , hedef değerini ifade etmektedir.

Eş. 2.92'de görülen $u, v \geq 0$ ise, $C_p(u, v)$ yeterlilik indeksinin u ve v 'nin belli değerleri için C_p , C_{pk} , C_{pm} ve C_{pmk} indeksleri ile ilişkisi, aşağıdaki şekilde yazılmaktadır.

$$C_p(0,0) = C_p, C_p(1,0) = C_{pk}, C_p(0,1) = C_{pm}, C_p(1,1) = C_{pmk}$$

Bu dört indeks arasındaki ilişki de Eş. 2.93 ve Eş. 2.94’de verildiği gibi gösterilmiştir (Pearn ve Chen, 1997).

$$C_{pm} = C_p \left\{ 1 + [(\mu - T) / \sigma]^2 \right\}^{-1/2} \quad (2.93)$$

$$C_{pmk} = C_{pk} \left\{ 1 + [(\mu - T) / \sigma]^2 \right\}^{-1/2} \quad (2.94)$$

Pearn ve Chen (1997), çalışmasında süreç ortalamasının hedef değerden ayrılışına göre yeterlilik indekslerinin duyarlılığını sıralamış ve en hassastan daha az hassasa doğru sırası ile C_{pmk} , C_{pm} , C_{pk} ve C_p olacağını belirtmiştir (Pearn ve Chen, 1997).

2.1.2. Normal dağılmayan süreçlerde tek değişkenli süreç yeterlilik indeksleri

Süreç yeterlilik tahmini için kullanılan standart teknikler genellikle parametrik varsayımlara dayanır ve süreç yeterlilik indeksleri olarak adlandırılır. Sözü geçen parametrik varsayımların en önemlisi genellikle sürecin normal dağılmasıdır. Bu varsayımda meydana gelen bozulmalar süreç yeterliliğinde yanıltıcı değerlendirmelere neden olabilmektedir. Bu nedenle, bir üretim süreç kalitesini değerlendirmede kullanılan süreç yeterlilik indeksleri tartışmalı hale gelmiştir. Bu tartışma bazı kalite uygulamacılarının bu indeksleri saf dışı bırakmalarına neden olmuştur. Ancak, bir üretim sürecinin yeterliliğinin belirlenmesi endüstriler için yararlı sonuçlar sağlar. Bu nedenle sürecin normal dağılmaması durumu için geçerli yöntemler geliştirilmesi büyük önem taşımaktadır (Polansky, 2001).

Son yıllarda kalite değişkenlerinin normallik varsayımını sağlamadığı durumlarda süreç yeterlilik analizi için önerilen iki yaklaşım göze çarpmaktadır. Bu yaklaşımlardan ilki, normal dağılıma uygun olmayan veriyi normal dağılıma dönüştürmek ve daha sonra yeterlilik indeksleri için önerilen klasik yaklaşımları kullanmaktır. Bu amaçla en çok kullanılan dönüşüm teknikleri, Johnson (1949) ve Box-Cox (1964) Güç dönüşüm teknikleridir. İkinci yaklaşım, dağılımın bilinmediği durumlarda veriyi bir dağılıma uydurmak ve daha sonra normal olmayan yüzdeler kullanılarak süreç yeterlilik

indekslerini tahmin etmektir. Bu amaçla en çok kullanılan yöntemlere örnek olarak, Burr (1942) ve Clements (1989) yüzde metodu verilebilir.

Kalite değişkeninin dağılımı Normal, Lognormal, t, F, Beta ve Gamma dağılımlarını içeren Pearson ailesi olasılık eğrilerinden birine sahipse;

$$P(LPL \leq \mu \leq UPL) = 1 - 0,0027 = 0,9973 \quad (2.95)$$

şeklinde tanımlanır. Clements (1989), normal olmayan süreç yeterlilik indekslerini tahmin etmek için Pearson dağılım eğrisini kullanarak Eş. 2.104'de süreç ortalaması (μ) yerine süreç ortanca değerini (M) kullanmıştır. Bunun nedeni, süreç ortancasının, süreç ortalamasına göre, özellikle uzun kuyruklu çarpık dağılımlara daha az duyarlı (robust) olmasıdır (Pearn ve Chen, 1997).

Clements (1989) yüzde metodunda, UPL , Pearson ailesinin 99,865 yüzdeliği; LPL , Pearson ailesinin 0,135 yüzdeliği olarak ifade edilmiştir.

Klasik süreç yeterlilik indekslerinde kullanılan 6σ yerine $UPL - LPL$ aralık uzunluğu kullanılır. Böylece, normal olmayan süreç yeterlilik indeksi C_p ,

$$C_p = \frac{USL - LSL}{x_{99,865} - x_{0,135}} \quad (2.96)$$

olarak, C_{pk} ise

$$C_{pu} = \frac{USL - x_{50}}{x_{99,865} - x_{50}} \quad (2.97)$$

$$C_{pl} = \frac{x_{50} - LSL}{x_{50} - x_{0,135}} \quad (2.98)$$

$$C_{pk} = \min \{C_{pu}, C_{pl}\} \quad (2.99)$$

olarak tanımlanmaktadır.

Burada, $x_p = p \times 100$. yüzdelerik değeri olarak tanımlanmıştır.

Pearn ve Kotz (1994), Clements (1989) yöntemini esas alarak C_{pm} ve C_{pmk} indekslerini Eş. 2.100 ve Eş. 2.101’de verildiği gibi önermiştir.

$$C_{pm} = \left\{ \frac{USL - LSL}{6 \sqrt{\frac{[x_{99,865} - x_{0,135}]^2}{6} + (M - T)^2}} \right\} \quad (2.100)$$

$$C_{pmk} = \min \left\{ \frac{USL - M}{3 \sqrt{\frac{[x_{99,865} - M]^2}{3} + (M - T)^2}}, \frac{M - LSL}{3 \sqrt{\frac{[M - x_{0,135}]^2}{3} + (M - T)^2}} \right\} \quad (2.101)$$

Pearn ve Chen (1997), Vannman (1995) tarafından tanımlanan $C_p(u, v)$ indeksini normal dağılım varsayımı olmadığı durumlar için $C_{Np}(u, v)$ ve $C'_{Np}(u, v)$ olarak adlandırdığı iki indekse genellemiştir. $C_{Np}(u, v)$, Eş. 2.102’de görüldüğü gibi tanımlanmıştır:

$$C_{Np}(u, v) = \frac{d - u |M - m|}{3 \sqrt{\left[\frac{F_{99,865} - F_{0,135}}{6} \right]^2 + v(M - T)^2}} \quad (2.102)$$

Burada, F_α , α ’ncı yüzdelik (percentile); M , dağılımın ortancası $m = (USL - LSL)/2$; $u, v \geq 0$ şeklindedir. $C_p(u, v)$ tanımının bulunduğu Eş. 2.92’de süreç standart sapması σ ’nın yerine $\frac{F_{99,865} - F_{0,135}}{6}$ konularak Eş. 2.102’ nin elde edildiği açıktır. Böyle bir yerine koymadaki esas amaç, ortalamadan (μ) $\mp 3\sigma$ sınırları dışında kalan uygun olmama oranı %0,27 olan normal dağılım özelliğini taklit edebilmektir. Böylece, $C_{Np}(u, v) = 1$ olarak hesaplandığında (süreç iyi bir şekilde merkezlenmişse); spesifikasyon sınırları (LSL, USL) dışına çıkma olasılığı önemsiz derecede küçük olacaktır. Parametre değerleri (u, v) ’nin $(0, 0)$; $(0, 1)$; $(1, 0)$ ve $(1, 1)$ değerleri için dört farklı forma genellenmiştir. Normal dağılım varsayımı olmaksızın kullanılabilen bu indeksler Eş. 2.103-Eş. 2.106’da görülmektedir:

$$C_{Np} = \frac{USL - LSL}{F_{99,865} - F_{0,135}} \quad (2.103)$$

$$C_{Npk} = \min \left\{ \frac{USL - M}{\frac{F_{99,865} - F_{0,135}}{2}}, \frac{M - LSL}{\frac{F_{99,865} - F_{0,135}}{2}} \right\} \quad (2.104)$$

$$C_{Npm} = \left\{ \frac{USL - LSL}{6\sqrt{\left[\frac{F_{99,865} - F_{0,135}}{6}\right]^2 + (M - T)^2}} \right\} \quad (2.105)$$

$$C_{Npmk} = \min \left\{ \frac{USL - M}{3\sqrt{\left[\frac{F_{99,865} - F_{0,135}}{6}\right]^2 + (M - T)^2}}, \frac{M - LSL}{3\sqrt{\left[\frac{F_{99,865} - F_{0,135}}{6}\right]^2 + (M - T)^2}} \right\} \quad (2.106)$$

Yukarıda tanımlanan dört indeks arasındaki ilişki $C_{Npm} = C_{Np} \left\{ 1 + [(M - T) / \sigma^*]^2 \right\}^{-1/2}$ ve $C_{Npmk} = C_{Npk} \left\{ 1 + [(M - T) / \sigma^*]^2 \right\}^{-1/2}$ şeklindedir.

Burada, $\sigma^* = \frac{F_{99,865} - F_{0,135}}{6}$ olarak verilmiştir.

Pearn ve Chen (1997), süreç ortancasının hedef değerden ayrılışına göre yeterlilik indekslerinin en duyarlıdan az duyarlıya doğru sıralanışının C_{Npmk} , C_{Npm} , C_{Npk} ve C_{Np} şeklinde olacağını belirtmiştir. Ayrıca, simetrik tolerans sınırları için $C_{Npk} = (1 - k)C_{Np}$ ve $C_{Npmk} = (1 - k)C_{Npm}$ olduğunu göstermişlerdir. Burada, $k = |M - T| / d$ "ayrılma oranı" olarak ifade edilmiştir. Süreç ortalaması hedef değere eşit olduğunda, $k=0$ ve $C_{Np} = C_{Npk} = C_{Npm} = C_{Npmk} = d/3\sigma^*$ olduğu açıkça görülmektedir. Diğer taraftan, süreç dağılımı normal dağılım olduğu durumda $\mu = M$ ve $\sigma^* = \sigma$ dır. Bu durumda, genellenen $C_{Np}(u, v)$ indekslerinin parametrik $C_p(u, v)$ indekslerine indirgendiği görülmektedir.

Pearn ve Kotz (1994) ve Pearn ve Chen (1995), Clements yüzde metodu, normal olmayan Pearson dağılımlarına uygulamış ve $C_p(u, v)$ için tahmin ediciler önermişlerdir. Önerilen tahmin ediciler ortalama, standart sapma, basıklık ve çarpıklık tahminlerinden yararlanılarak iki yüzdelik olan $F_{99,865}$ ve $F_{0,135}$ için U_p ve L_p tahminlerine dayalıdır. Bu tahmin edici, Eş. 2.107'de verilmiştir.

$$C'_{Np}(u, v) = (1-u) * \frac{USL - LSL}{6\sqrt{\left[\frac{F_{99,865} - F_{0,135}}{6}\right]^2 + v(M-T)^2}} + u * \min \left\{ \frac{USL - M}{3\sqrt{\left[\frac{F_{99,865} - F_{0,135}}{3}\right]^2 + v(M-T)^2}}, \frac{M - LSL}{3\sqrt{\left[\frac{F_{99,865} - F_{0,135}}{3}\right]^2 + v(M-T)^2}} \right\} \quad (2.107)$$

önerilen indeksler parametre değerleri (u, v) 'nin $(0, 0)$; $(0, 1)$; $(1, 0)$ ve $(1, 1)$ değerleri için dört farklı forma genellenmiştir. Bu genelleştirilmiş formlar Eş. 2.108-Eş. 2.111'de verilmektedir (Pearn ve Chen, 1997).

$$C'_{Np} = \frac{USL - LSL}{F_{99,865} - F_{0,135}} \quad (2.108)$$

$$C'_{Npk} = \min \left\{ \frac{USL - M}{F_{99,865} - M}, \frac{M - LSL}{M - F_{0,135}} \right\} \quad (2.109)$$

$$C'_{Npm} = \left\{ \frac{USL - LSL}{6\sqrt{\left[\frac{F_{99,865} - F_{0,135}}{6}\right]^2 + (M - T)^2}} \right\} \quad (2.110)$$

$$C'_{Npmk} = \min \left\{ \frac{USL - M}{3\sqrt{\left[\frac{F_{99,865} - M}{3}\right]^2 + (M - T)^2}}, \frac{M - LSL}{3\sqrt{\left[\frac{M - F_{0,135}}{3}\right]^2 + (M - T)^2}} \right\} \quad (2.111)$$

Pearn ve Chen (1997) A, B ve C gibi Ki-kare dağılımına sahip üç süreci Eş. 2.103-Eş. 2.106 ve Eş. 2.108-Eş. 2.111'de verilen indeksleri karşılaştırmak amacı ile incelemiştir. Önerilen C_{Np} indeksinin süreç yeterliliğini ölçmede C'_{Np} ve $C_p(u, v)$ indekslerine göre daha tutarlı ve daha doğru sonuç verdiğini göstermiştir.

Liu ve Chen (2006), Clements (1989) yüzde metodunda Pearson dağılımının yüzdelerliğini kullanmak yerine Burr dağılımının yüzdelerliğini kullanmışlardır. Burr (1942), X rastgele değişkeninin yüzdelerliğini elde etmek için Burr XII dağılımı olarak adlandırılan bir dağılım önermiştir. X rastgele değişkenin dağılım fonksiyonu,

$$F(x) = \begin{cases} 1 - (1 + x^c)^{-k} & , \quad x \geq 0, c > 0, k > 0 \\ 0 & , \quad x < 0 \end{cases}$$

olarak tanımlanmıştır. Karşılıklı dönüşümler $\left(1 - F\left(\frac{1}{y}\right)\right)$ ile elde edilen yeni Burr XII değişkeni Y 'nin dağılım fonksiyonu ise,

$$G(y) = \begin{cases} (1 + y^{-c})^{-k} & , \quad y \geq 0, c > 0, k > 0 \\ 0 & , \quad y < 0 \end{cases} \quad (2.112)$$

olarak tanımlıdır. Burada, c ve k sırasıyla çarpıklık ve basıklık katsayılarıdır.

Burr yüzde metodunun adımları (Ahmad, Abdollahian ve Zeephongsesekul, 2008) tarafından aşağıdaki gibi açıklanmıştır.

Adım 1'de örneklem verisinden örneklem ortalaması $\bar{x}$, standart sapması s , çarpıklık katsayısı s_3 ve basıklık katsayısı s_4 tahmin edilir.

$$\alpha_3 = \frac{(n-2)}{\sqrt{n(n-1)}} s_3$$

$$\alpha_4 = \frac{(n-2)(n-3)}{\sqrt{(n^2-1)}} s_4 + 3 \frac{(n-1)}{(n+1)}$$

Adım 2'de yukarıdaki eşitlikler kullanılarak standartlaştırılmış α_3 çarpıklık ve α_4 basıklık momentleri hesaplanır.

Adım 3'te, hesaplanan α_3 ve α_4 değerleri kullanılarak Burr parametrelerinden en uygun c ve k seçilir. $Z = (Y - \mu)/\sigma$ ile standartlaştırılan değişkenlerin dağılımını elde etmek için BurrXII dağılımı kullanılır. Burada Y , Burr rastgele değişkenidir.

Burr tablosu kullanılarak standartlaştırılmış değerler $Z_{0,00135}$, $Z_{0,5}$ ve $Z_{0,99865}$ elde edilir. X 'in ilgili yüzdelikleri, iki standartlaştırılmış değerın aşağıdaki gibi eşleştirilmesi sonucu elde edilir.

$$\frac{(X - \bar{x})}{s} = \frac{(Y - \mu)}{\sigma}$$

Üst, medyan ve alt yüzdelik değerleri aşağıdaki gibi hesaplanır.

$$x_{0,135} = \bar{x} + sZ_{0,00135}$$

$$x_{50} = \bar{x} + sZ_{0,5}$$

$$x_{99,865} = \bar{x} + sZ_{0,99865}$$

Son adımda ise, yukarıda elde edilen eşitlikler kullanılarak Eş. 2.96-Eş. 2.99'da verilen C_p , C_{pu} , C_{pl} ve C_{pk} süreç yeterlilik indeksleri hesaplanabilir.

Ahmad ve diğerleri (2008), Box-Cox dönüşümü, Clements ve Burr yüzde yöntemlerinin normal olmayan süreçlerde yeterlilik indeksi hesaplamadaki doğruluğunu incelemek amacıyla Weibull, Gamma ve Lognormal dağılımlarına sahip her birinin örneklem büyüklüğü 100 olan 30 örneklem üreterek elde ettikleri verilere Box-Cox dönüşümü ve Clements ve Burr yüzde metotlarını uygulamışlardır. Karşılaştırma kriteri olarak $\hat{C}_{pu}$ ele alınmıştır. Her üç yöntem için elde edilen $\hat{C}_{pu}$ değerlerinin ortalama ve standart sapmaları hesaplanmıştır. Sonuç olarak, Burr yüzde yönteminden elde edilen $\hat{C}_{pu}$ değerleri belirlenen $\hat{C}_{pu}$ değerlerinden, en az sapmayı göstermiştir. Burr yüzde metodu kullanılarak elde edilen $\hat{C}_{pu}$ değerlerinin standart sapması, Clements yüzde metodu kullanılarak elde edilen $\hat{C}_{pu}$ değerlerinin standart sapmasından daha küçüktür. Box-Cox dönüşüm yöntemi ile elde edilen $\hat{C}_{pu}$ değerleri belirlenen C_{pu} değerlerine yakın değildir. Büyük örneklemelerin her üç metot için de daha iyi sonuçlar sağladığı görülmüştür. Böylece, Burr yüzde metodunun diğer iki metoda göre daha doğru sonuçlar verdiğini göstermişlerdir.

2.2. Çok Değişkenli Süreç Yeterlilik İndeksleri

Bu bölümde, çok değişkenli süreç yeterlilik indeksleri normal dağılım gösteren ve normal dağılım göstermeyen süreçler olarak iki sınıfta incelenmiştir.

2.2.1. Çok değişkenli normal dağılan süreçlerde süreç yeterlilik indeksleri

Çoğu süreçte uygulamada birden fazla kalite değişkeni bulunmaktadır. Bu değişkenler genellikle birbirleri ile ilişkili olduğu için çok değişkenli yeterlilik indekslerine gereksinim duyulmuştur. Çok değişkenli süreçler için literatürde yer alan çalışmalar dört grupta incelenmiş buna göre çok değişkenli süreç yeterlilik indeksleri,

1. Tolerans bölgesi ile süreç bölgesini birbiri ile oranlamaya dayalı olarak,
2. Uygun olmayan ürün elde etme olasılığına dayalı olarak,
3. Temel bileşenler analizi kullanılarak değişkenler arası korelasyonu kaldırmaya yönelik olarak,
4. Diğer yaklaşımlar kullanılarak

elde edilmiştir (Shinde ve Khadse, 2009).

2.2.2. Grup I'de yer alan başlıca çalışmaların incelenmesi

Hubele, Shahriari ve Cheng (1991), üç bileşenden oluşan bir $[C_{pm}, PV, LI]$ yeterlilik vektörü tanımlamıştır. Yeterlilik vektörünün ilk bileşeni C_{pm} , Eş. 2.113'de verildiği gibi tanımlanmıştır.

$$C_{pm} = \left(\frac{\sum_{i=1}^v (USL_i - LSL_i)}{\sum_{i=1}^v (UPL_i - LPL_i)} \right)^{1/v} \quad (2.113)$$

Burada,

USL_i : i . kalite değişkeni için üst spesifikasyon sınırını, LSL_i : i . kalite değişkeni için alt spesifikasyon sınırını, UPL_i : i . kalite değişkeni için modifiye edilmiş süreç bölgesinin üst spesifikasyon sınırını, LPL_i : i . kalite değişkeni için modifiye edilmiş süreç bölgesinin alt spesifikasyon sınırını ve v : kalite değişkeni sayısını göstermektedir.

Yeterlilik vektörünün ikinci bileşeni PV ise,

$$PV = P \left[T^2 > \frac{v(n-1)}{n-v} F_{v,n-v} \right] \quad (2.114)$$

olarak tanımlanmıştır.

Burada, $T^2 = n(\bar{\mathbf{X}} - \mathbf{T})' \mathbf{S}^{-1} (\bar{\mathbf{X}} - \mathbf{T})$ olarak ifade edilmiştir. $\mathbf{T}$ hedef vektörünü, $\bar{\mathbf{X}}$ örneklem ortalama vektörünü, $\mathbf{S}$, örneklem varyans-kovaryans matrisini, n örneklem büyüklüğünü ve $F_{v,n-v}$ ise $(v, n-v)$ serbestlik derecelerine sahip F dağılımını ifade etmektedir.

Yeterlilik vektörünün üçüncü bileşeni LI , modifiye edilmiş süreç bölgesi ile tanımlanan tolerans bölgesini karşılaştırır. Bu bileşen modifiye edilmiş süreç bölgesinin tolerans bölgesinin dışına düşüp düşmediğini gösterir. $LI = 1$ olursa, modifiye edilmiş süreç bölgesi tolerans bölgesinin içine; $LI = 0$ ise dışına düştüğü söylenir (Pan ve Lee, 2010).

Taam, Subbaiah ve Liddy (1993), MC_p ve MC_{pm} olarak gösterdikleri çok değişkenli iki yeterlilik indeksi önermiştir.

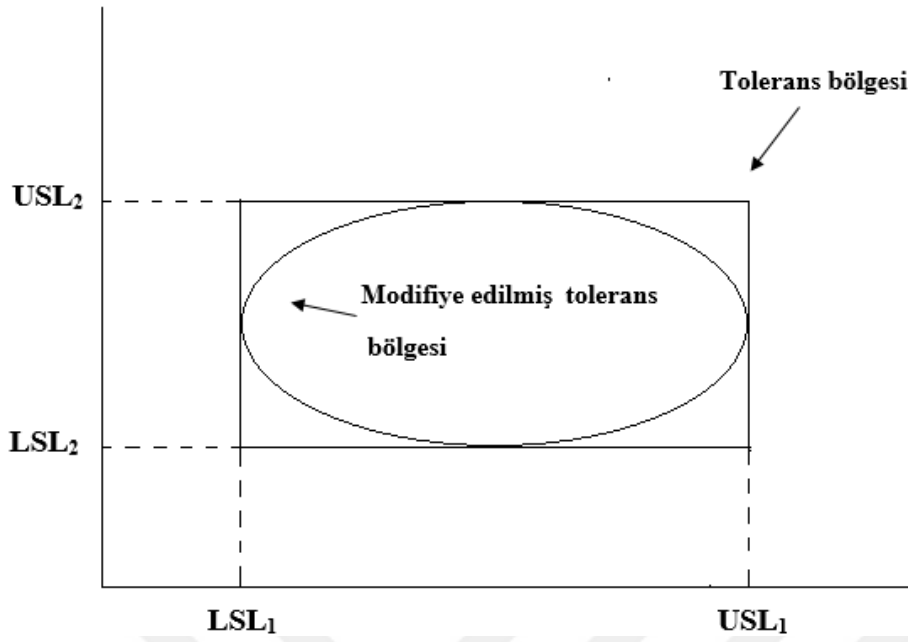
$$MC_{pm} = \frac{hacim(R_1)}{hacim(R_2)}$$

Burada,

R_1 : modifiye edilmiş tolerans bölgesini

R_2 : çok değişkenli normal dağılıma sahip olduğu varsayılan süreç bölgesinin hesaplanan %99,73'ünü ifade etmektedir.

Modifiye edilmiş tolerans bölgesi, Pan ve Lee (2010) tarafından şekil 2.3'te çizildiği şekilde gösterilmiştir.



Şekil 2.3. Tolerans bölgelerinin gösterimi

Modifiye edilmiş tolerans bölgesi hedef değerde merkezlenmiş ve orijinal tolerans bölgesinin içine düşen en geniş elipsoid olduğundan MC_{pm} Eş. 2.115’de verildiği gibi tanımlanmıştır.

$$MC_{pm} = MC_p \frac{1}{D} \quad (2.115)$$

Burada, $D = \left(1 + (\mu - \mathbf{T})' \Sigma^{-1} (\mu - \mathbf{T})\right)^{1/2}$ olarak ifade edilen düzeltme faktörüdür ve süreç ortalaması hedef değerden sapmış ise kullanılır. MC_p indeksi Eş. 2.116’da olduğu gibi tanımlanmıştır.

$$MC_p = \frac{\left(\prod_{i=q}^v r_i\right) \pi^{v/2} [\Gamma(v/2) + 1]^{-1}}{|\Sigma|^{1/2} (\pi K(v))^{v/2} [\Gamma(v/2) + 1]^{-1}} \quad (2.116)$$

Burada, $i = 1, 2, \dots, v$ ’nci kalite değişkeni için,

$r_i = (USL_i - LSL_i)/2$ ve $\Gamma(.)$ ise Gamma fonksiyonunu göstermektedir.

MC_p indeksi, Pearn, Wang ve Yen (2007) tarafından Eş. 2.117’ de verildiği gibi önerilmiştir.

$$\begin{aligned}
MC_p &= \frac{\text{hacim}(\text{modifiye edilmiş tolerans bölgesi})}{\text{hacim} \left[(\mathbf{X} - \boldsymbol{\mu})' \boldsymbol{\Sigma}^{-1} (\mathbf{X} - \boldsymbol{\mu}) \leq Q^* \right]} \\
&= \frac{\text{hacim}(\text{modifiye edilmiş tolerans bölgesi})}{(\pi \chi_{v,0.9973}^2)^{v/2} |\boldsymbol{\Sigma}|^{1/2} [\Gamma(v/2) + 1]^{-1}}
\end{aligned} \tag{2.117}$$

Burada, Q^* v serbestlik derecesine sahip χ^2 dağılımının 99,73'ncü yüzdeliğini, $|\boldsymbol{\Sigma}|$ ise, $\boldsymbol{\Sigma}$ 'nin determinantını göstermektedir.

MC_p tahmin edicisi $\widehat{MC}_p$ Eş. 2.117'de görülen $|\boldsymbol{\Sigma}|$ yerine $|\mathbf{S}|$ yazılarak,

$$\begin{aligned}
\widehat{MC}_p &= \frac{\text{hacim}(\text{modifiye edilmiş tolerans bölgesi})}{\text{hacim}(\%99.73 \text{ süreç bölgesi})} \\
&= \frac{\text{hacim}(\text{modifiye edilmiş tolerans bölgesi})}{(\pi \chi_{v,0.9973}^2)^{v/2} |\mathbf{S}|^{1/2} [\Gamma(v/2) + 1]^{-1}}
\end{aligned} \tag{2.118}$$

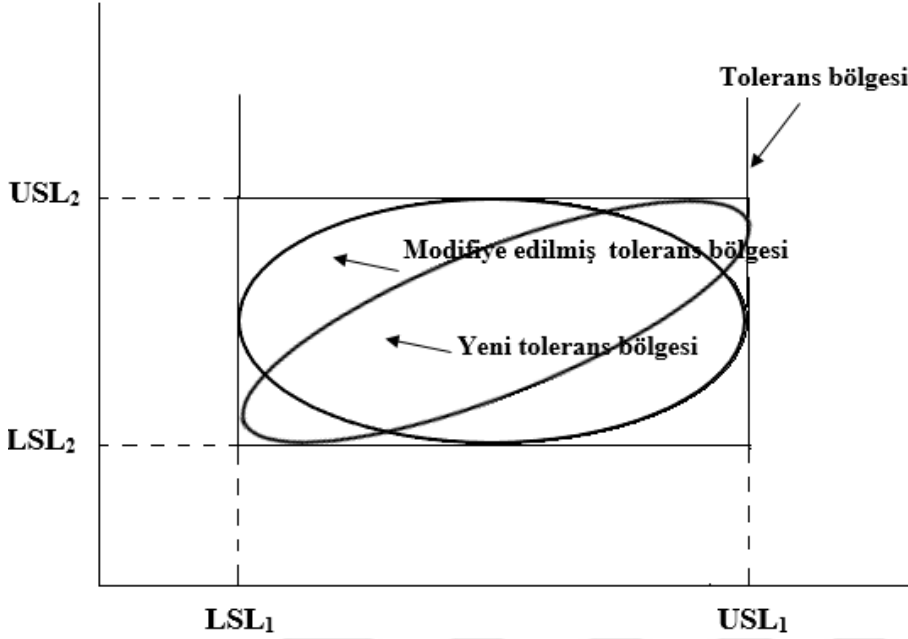
biçiminde elde edilir. Burada, $\mathbf{S}$ örnek varyans-kovaryans matrisi ve $|\mathbf{S}|$, $\mathbf{S}$ 'nin determinantıdır.

Pan ve Lee (2010), MC_p ve MC_{pm} indekslerinin tahmin edilmesinde Taam ve diğerleri (1993)'nin çoklu kalite değişkenleri arasındaki korelasyonu dikkate almamaları ve Hubele ve diğerleri'nin (1991) önerdiği, üç-bileşenli yeterlilik vektöründeki uygulama zorluğu nedeniyle NMC_p ve NMC_{pm} olarak adlandırdıkları iki yeterlilik indeksi tanımlamıştır. Eş. 2.117'de verilen MC_p indeksinin daha basit olarak Eş. 2.119'da verilen biçimde hesaplanabileceğini göstermişlerdir.

$$MC_p = |\rho|^{-1/2} \tag{2.119}$$

Burada, ρ korelasyon matrisini göstermektedir. Korelasyon matrisinin determinant değeri 0 ile 1 arasında bir değer alacağından Eş. 2.119'da görülen MC_p indeksi 1'den büyük değerler alabilir. Diğer bir ifade ile, kalite değişkenleri bağımlı ise, MC_p indeksi 1'den büyük olacaktır. Bu durumda, süreç performans tahminleri gerçek değerlerini aşabilir. Benzer biçimde, çoklu kalite değişkenleri bağımlı ise $MC_{pm} = MC_p/D$ indeksi için de aynı sorun söz konusudur. Bu durumu göz önünde bulunduran Pan ve Lee (2010), Şekil 2.3'te

tanımlanan tolerans bölgesini revize ederek Şekil 2.4’de görüldüğü gibi yeni bir bölge önermiştir.



Şekil 2.4. Yeni tolerans bölgesinin gösterimi

Yeni tolerans bölgesinin, kalite değişkenleri arasındaki kolerasyonu dikkate alınarak çizilmiş olması nedeniyle ilişki yönünde eğimli bir elipsoid şeklinde olduğu Şekil 2.4’te görülmektedir. Pan ve Lee (2010) tarafından tanımlanan yeni tolerans bölgesi, Eş. 2.120’de verildiği gibi önerilmiştir.

$$E_{d,A^*,T} = \left\{ \mathbf{X} \in \mathbb{R}^V \mid (\mathbf{X} - \mathbf{T})' (\mathbf{A}^*)^{-1} (\mathbf{X} - \mathbf{T}) = \mathbf{d}^2 \right\} \quad (2.120)$$

$\mathbf{A}^*$ Matrisinin elemanları aşağıdaki gibidir:

$$\rho_{ij} \left(\frac{USL_i - LSL_i}{2d} \right) \left(\frac{USL_j - LSL_j}{2d} \right), \quad i, j = 1, \dots, v$$

$\mathbf{T}$, hedef vektörünü, ρ_{ij} , i . kalite değişkeni ile j . kalite değişkeni arasındaki korelasyon katsayısını ve $USL_i - LSL_i$ ise $E_{d,A^*,T}$ elipsoidini çevreleyen dikdörtgenin her bir tarafındaki spesifikasyon genişliğini göstermektedir. Böylece Pan ve Lee (2010) tarafından tanımlanan yeni indeks, Eş. 2.121’deki biçimde yazılabilir.

$$NMC_{pm} = \frac{Hacim(E_{d,A^*,T})}{Hacim(E_{d,\Sigma,\mu})} \left(1 + (\mathbf{X} - \mathbf{T})' \Sigma^{-1} (\mathbf{X} - \mathbf{T}) \right)^{-1/2} \quad (2.121)$$

$d^2 = \chi^2_{1-\alpha, v}$ ise, NMC_{pm} Eş. 2.122'deki gibi yeniden yazılabilir.

$$\begin{aligned} NMC_{pm} &= \frac{|\mathbf{A}^*|^{1/2} \left(\pi \chi^2_{1-\alpha, v} \right)^{v/2} \left[\Gamma \left(\frac{v}{2} + 1 \right) \right]^{-1}}{|\Sigma|^{1/2} \left(\pi \chi^2_{1-\alpha, v} \right)^{v/2} \left[\Gamma \left(\frac{v}{2} + 1 \right) \right]^{-1}} \left(1 + (\mathbf{X} - \mathbf{T})' \Sigma^{-1} (\mathbf{X} - \mathbf{T}) \right)^{-1/2} \\ &= NMC_p / D \end{aligned} \quad (2.122)$$

Burada, $NMC_p = (|\mathbf{A}^*| / |\Sigma|)^{1/2}$ ve $D = (1 + (\mathbf{X} - \mathbf{T})' (\Sigma)^{-1} (\mathbf{X} - \mathbf{T}))^{1/2}$

olarak ifade edilmiştir. D , süreç ortalaması ile $\mathbf{T}$ hedef vektörü arasındaki uzaklığın Mahalanobis fonksiyonudur.

$\mathbf{X}_1, \dots, \mathbf{X}_n$, v kalite değişkenine sahip çok değişkenli bir süreçten rastgele alınan n büyüklüğündeki örneklem için NMC_p indeksinin tahmin edicisi Eş. 2.123'deki gibi hesaplanmıştır.

$$\widehat{NMC}_p = (|\mathbf{A}^*| / |\mathbf{S}|)^{1/2} \quad (2.123)$$

Burada, $\mathbf{S}$ bilindiği gibi örneklem varyans-kovaryans matrisi,

$$\mathbf{S} = (\mathbf{n} - \mathbf{1})^{-1} \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})' (\mathbf{X}_i - \bar{\mathbf{X}}); \bar{\mathbf{X}}, \text{ örneklem ortalama vektörüdür.}$$

$$\bar{\mathbf{X}} = n^{-1} \sum_{i=1}^n \mathbf{X}_i \text{ biçiminde tanımlanır.}$$

2.2.3. Grup II'de yer alan başlıca çalışmaların incelenmesi

Gonzales ve Sanchez'e (2009) göre, "Yeterlilik indeksleri uygun olmayan ürün oranı ile yakından ilişkili olabilir."

Bu yorumlama tek değişkenli ve merkezlenmiş süreçler için klasik C_p indeksini hesaplarken geçerli olabilir. Örneğin $C_p = 0,5$ ise sürecin yeterli olması için ($C_p = 1$) standart sapmanın %50 gibi bir değere kadar indirilmesi gerekir. Ancak, merkezlenmemiş

ve çok değişkenli süreçler için böyle bir yorumlamayı mümkün kılan bir yeterlilik indeksi mevcut değildir. Buradan yola çıkarak Gonzales ve Sanchez (2009), sürecin hedeflenen yeterliliğini koruyabilmek için süreçteki değişkenliğin ne kadar azaltılması ya da artırılması gerektiğini doğrudan yorumlayabilen tek değişkenli ve çok değişkenli süreç yeterlilik indeksleri önermiştir. Tek değişkenli süreçler için önerdikleri yeni indeks uygun olmayan ürün oranı p ile doğrudan ilişkilidir ve C_n ile gösterilmiştir. Burada n , uygun olmama oranını (nonconforming propotion) ifade etmektedir. $C_n = 1$ ise, $P(X \in TB) = 1 - \alpha$ olarak tanımlanmıştır. Burada, $X \sim N(\mu, \Sigma^2)$ ve TB ise tolerans bölgesini göstermektedir. C_n indeksi, gerçekte süreçten elde edilen değişkenlik ile tolere edilen değişkenliğin oranlanmasına dayanır ve Eş. 2.124'de görüldüğü gibi ifade edilmiştir.

$$C_n = \frac{\Sigma_{maks}}{\Sigma} \quad (2.124)$$

Σ_{maks}^2 , X ile aynı yoğunluk fonksiyonuna sahip X_{maks} sürecinin $P(LSL \leq X_{maks} \leq USL) = 1 - \alpha$ ifadesini doğrulayan varyans değeridir. Σ_{maks} ise, süreç ortalaması aynı kalmak koşulu ile $p = \alpha$ olmasını mümkün kılan en büyük standart sapmadır.

C_n 'nin çok değişkenli süreçler için tanımlanması

Eş. 2.124'de tanımlanan tek değişkenli C_n indeksinin çok değişkenli süreçler için uyarlaması Eş. 2.125'de görüldüğü gibi ifade edilmiştir.

$$C_n = f\left(\frac{|\Sigma|^{maks}}{|\Sigma|}\right)^{1/2} \quad (2.125)$$

Burada, $f(\cdot)$ indeksin özel kullanımlarına bağlı olarak alternatif kullanımı Eş. 2.126'da verilecektir.

Burada, Σ , sürecin varyans-kovaryans matrisi; Σ^{maks} , $P(\mathbf{X}^{maks} \notin TB) = \alpha$ ile $\mathbf{X}^{maks}$ çok değişkenli sürecinin izin verilen maksimum varyans-kovaryans matrisidir.

$\mathbf{X}$ çok değişkenli sürecinin değişkenleri birbirinden bağımsız değildir. Bu değişkenler, Y bağımsız faktörlerin kombinasyonu olarak düşünülebilir. Çok değişkenli analizde bu bağımsız faktörler genellikle latent (gizli) faktörler olarak bilinir ve Y bağımsız

faktörlerinin varyanslarında meydana gelen bir değişim, Σ değerinde bir değişime neden olacağı için süreç değişkenliğinin birincil nedeni olarak sayılabilir. Y bağımsız faktörlerinin varyanslarındaki değişimler incelenerek çok değişkenli problemin çözümlenme sorunu, tek değişkenli problemin çözümlenmesi şeklinde ele alınabilir (Gonzales ve Sanchez, 2009).

Gonzales ve Sanchez (2009), normallik varsayımı altında temel bileşenler analizi kullanarak Y bağımsız faktörlerini elde edip, yeterlilik analizini temel bileşenlerin varyanslarındaki değişim olarak ele almışlardır. Tanımlamalarında, bilinen temel bileşenler analizi terminolojisini kullanmışlardır. Buna göre,

$$\Sigma = CDC'$$

Burada, C : Σ 'nın özvektörler matrisini, D : köşegenleri özdeğerlerinden oluşan (λ_i $i = 1, 2, \dots, m$) köşegen matrisini göstermektedir.

Y_i , temel bileşenler analizi sonucu elde edilen i . bağımsız faktörünü gösterebilir. Bu faktörün varyansı, λ_i özdeğeri ile b_i^2 faktörü çarpılarak değiştirilirse, elde edilen yeni özdeğerler matrisi $D_{(i)}^*$ ile gösterilir.

$$D_{(i)}^* = \text{diag}(\lambda_1, \lambda_2, \dots, b_i^2 \lambda_i, \dots, \lambda_m)$$

i . bileşenin varyansında meydana gelen bu değişimin $\mathbf{X}$ 'de neden olduğu değişimin sonucu $X_{(i)}^*$ ile ve dağılımı da $X_{(i)}^* \sim N(\mu, \Sigma_{(i)}^*)$ ile gösterilsin. Burada, $\Sigma_{(i)}^* = CD_i^*C'$ ile ifade edilir. Bu yeni dağılım $p_{(i)}^* = P(\mathbf{X}_{(i)}^* \notin TB)$ uygun olmayan ürün oranını içerir. Eş. 2.125'de verilen C_n indeksi genel olarak Eş. 2.126'da görüldüğü gibi ifade edilmiştir.

$$C_{n,i} = \left(\frac{|\Sigma_{(i)}^{maks}|}{|\Sigma|} \right)^{1/2} = \left(\frac{(b_i^{maks})^2 \prod_{i=1}^m \lambda_i}{\prod_{i=1}^m \lambda_i} \right)^{1/2} = b_i^{maks} \quad (2.126)$$

b_i^{maks} , $p_{(i)}^* = \alpha$ ile i . faktörde $\Sigma_{(i)}^{maks}$ değişime neden olan değerdir ve $C_{n,i} \geq 1$, $i = 1, 2, \dots, m$ ise süreç yeterlidir.

Süreç yeterli değilse, $C_{n,i}$ indeksi hangi faktörün düzeltilmesi gerektiğini gösterebilir.

Diğer bir çalışmada ise, Castagliola ve Castellanos (2005), iki değişkenli normal süreçler için BC_p ve BC_{pk} olarak tanımladıkları yeni yeterlilik indeksleri önermiştir. Süreç kalite değişkenlerinin (X_1, X_2) ve tolerans sınırlarının X_1 için $[L_1, U_1]$; X_2 için $[L_2, U_2]$ olarak gösterildiği varsayıldığında, bu sınırların oluşturduğu tolerans bölgesi A ile gösterilmiştir. Σ 'nın özvektörleri kullanarak süreç bölgesini A_i , $i=1,2,3,4$ olan dört bölgeye ayırmışlardır. Normal dağılımın simetrik özelliği nedeniyle, her bir A_i alanında gözlemlenenlerin oranının $1/4$ olduğu ifade edilebilir. A_1, A_2, A_3, A_4 bölgelerinin tolerans bölgesi A ile kesişen poligon konveksleri Q_1, Q_2, Q_3, Q_4 olarak tanımlanmıştır (örneğin $Q_i = A \cap A_i$).

Her bir A_i ve tolerans bölgesi içinde gözlemlenenlerin oranı q_i , $i=1,2,3,4$ olarak ifade edilirse, A_i bölgesinde oluşan uygun olmayan birimlerin oranı $p_i = \frac{1}{4} - q_i$ olmak üzere,

$$BC_{pk} = \frac{1}{3} \min \{ -\phi^{-1}(2p_1), -\phi^{-1}(2p_2), -\phi^{-1}(2p_3), -\phi^{-1}(2p_4) \} \quad (2.127)$$

biçiminde ifade edilebilir.

Castagliola ve Castellanos (2005) Eş. 2.127'de tanımladıkları bu indeksin $\hat{BC}_{pk}$ tahmin edicisini elde edebilmek için bir prosedür geliştirmişlerdir. Bu prosedürün algoritması aşağıdaki gibi özetlenebilir.

$X_1, \dots, X_n$, kontrol altında bir süreçten çekilen n büyüklüğünde bir örneklem ve $X_k = (X_{k,1}, X_{k,2})'$ ise,

1. örneklemden beklenen değer vektörünün ve varyans-kovaryans matrisinin tahminleri elde edilir. $\hat{\mu} = \frac{1}{n} \sum_{k=1}^n X_k$
 $\hat{\Sigma} = \frac{1}{n-1} \sum_{k=1}^n (X_k - \hat{\mu})(X_k - \hat{\mu})^T$
2. $\hat{\Sigma}$ 'nin özdeğerleri ve özvektörleri hesaplanır.
3. Q_1, Q_2, Q_3, Q_4 konveks polinomlarının köşeleri hesaplanır.
4. $\hat{q}_1, \hat{q}_2, \hat{q}_3, \hat{q}_4$ oranları hesaplanır ve $\hat{p}_i = \frac{1}{4} - \hat{q}_i$ elde edilir.
5. Buradan,
 $\hat{BC}_{pk} = \frac{1}{3} \min \{ -\phi^{-1}(2\hat{p}_1), -\phi^{-1}(2\hat{p}_2), -\phi^{-1}(2\hat{p}_3), -\phi^{-1}(2\hat{p}_4) \}$
 tahmini elde edilebilir.

Castagliola ve Castellanos (2005), benzer biçimde yeni BC_p indeksini Eş. 2.128’de görüldüğü gibi tanımlamışlardır.

$$BC_p = BC_{pk}^* = maks BC_{pk} \quad (2.128)$$

BC_{pk} , $-\Phi^{-1}(2\hat{p}_1)$, $-\Phi^{-1}(2\hat{p}_2)$, $-\Phi^{-1}(2\hat{p}_3)$ ve $-\Phi^{-1}(2\hat{p}_4)$ değerleri arasında minimum değere sahip olan olarak tanımlandığı için, BC_{pk} maksimum değerine ancak şu durumlarda ulaşır:

- $-\Phi^{-1}(2\hat{p}_1) = -\Phi^{-1}(2\hat{p}_2) = -\Phi^{-1}(2\hat{p}_3) = -\Phi^{-1}(2\hat{p}_4)$ maksimum olduğunda,
- $p_1 = p_2 = p_3 = p_4 = \frac{p}{4}$ minimum olduğunda ya da
- $q_1 = q_2 = q_3 = q_4 = \frac{q}{4}$ maksimum olduğunda.

Böylece,

$$BC_p = BC_{pk}^* = -\frac{1}{3}\Phi^{-1}\{2 \times (p/4)\} = -\frac{1}{3}\Phi^{-1}(p/2) \quad (2.129)$$

olarak elde edilebilir.

Burada, uygun olmayan ürün oranı p Eş. 2.130’da,

$$p = 2\Phi(-3BC_p) \quad (2.130)$$

olarak tanımlanmıştır.

2.2.4. Grup III’te yer alan başlıca çalışmaların incelenmesi

Wang ve Chen (1998), normal dağılıma sahip çoklu kalite değişkenleri arasındaki korelasyonu ortadan kaldırmak için temel bileşenler analizi kullanarak kısa dönemli üretimler için çok değişkenli bir süreç yeterlilik indeksi önermiştir. Önerdikleri yeni indeks Eş. 2.131’de görüldüğü gibi elde edilir.

$$MC_p = \left[\prod_{i=1}^v C_{p:PC_i} \right]^{1/v} \quad (2.131)$$

$$\text{Burada, } C_{p:PC_i} = \frac{USL_{PC_i} - LSL_{PC_i}}{6\sqrt{\lambda_{PC_i}}}$$

$C_{p:PC_i}$, i . bileşenin tek değişkenli C_p yeterlilik indeks değerini; LSL_{PC_i} ve USL_{PC_i} i . bileşenin alt ve üst spesifikasyon sınırlarını ifade eder. λ_{PC_i} , i . bileşenin özdeğeri; v , seçilen bileşen sayısıdır. Benzer biçimde $C_{pk:PC_i}$ ve $C_{pm:PC_i}$ indekslerini de tanımlamışlardır.

Wang (2005), çok değişkenli süreç yeterlilik indekslerini, temel bileşenler analizine dayalı yeni bir prosedür kullanarak elde etmeye çalışmıştır. Bu prosedüre göre, çok değişkenli süreç yeterlilik indekslerine ağırlıklandırılmış geometrik ortalamalara dayalı olarak karar verilir. İlk bileşen çoklu kalite değişkenlerinin varyansının en büyük kısmını; ikinci bileşen bu varyansın ikinci en büyük kısmını açıklar ve bu şekilde diğer bileşenlerin varyansı açıklama oranları belirlenebilir. Wang'ın (2005) önerdiği yeterlilik indeksleri Eş. 2.132 ve Eş. 2.133'de görülmektedir.

$$MC_p = \left[\prod_{i=1}^t C_{pv}^{e_v} \right]^{\frac{1}{\sum_{i=1}^t e_v}} \quad (2.132)$$

$$MC_{pk} = \left[\prod_{i=1}^t C_{pkv}^{e_v} \right]^{\frac{1}{\sum_{i=1}^t e_v}} \quad (2.133)$$

Burada, e_v , v . bileşenin özdeğerini; t , bileşenlerin sayısını gösterir. MC_p ve MC_{pk} , C_p ve C_{pk} 'nin kısa dönemli çok değişkenli süreç yeterlilik indekslerini ifade eder.

Shinde ve Khadse (2009), Wang ve Chen (1998) tarafından tanımlanan LSL_{PC_i} ve USL_{PC_i} sınırlarının uygulamada yanlış sonuç verdiğini bir örnek yardımı ile ispatlamışlardır. Wang ve Chen (1998) tarafından hesaplanan yeterlilik indekslerinin yanlış olma nedeni olarak, farklı temel bileşenlerin spesifikasyon sınırlarının birbirinden bağımsız olduğunu varsaymaları olarak açıklamışlardır. Temel bileşenlerin sadece dağılımlarının bağımsız olduğunu, ancak spesifikasyon sınırlarının birbirleri ile bağlantılı olduğunu ispatlamışlardır. Bu nedenle, elde edilen temel bileşenlere ait sınırları yeniden tanımlayarak MP_1 ve MP_2 olarak adlandırdıkları çok değişkenli süreç yeterlilik indekslerini önermişlerdir. Bu indeksler için temel bileşenleri deneysel olasılık dağılımlarına dayalı olarak tanımlamışlardır.

2.2.5. Grup IV'te yer alan başlıca çalışmaların incelenmesi

Literatürdeki birçok yeterlilik indeksi, tolerans bölgesi içine düşen tüm ürünleri eşit derecede "*kabul edilebilir*" olarak değerlendirmektedir (Goethals ve Cho, 2011). Tolerans bölgesinde hedef değerden ölçümlerin uzaklığını gösterebilmesi amacıyla Goethals ve Cho (2011), parabolik bir kayıp fonksiyonu tanımlamışlardır.

$X = (X_1, \dots, X_v)$ ile v değişkenden oluşan ve çok değişkenli normal dağılıma sahip kalite değişkenler vektörü ise, $L(X)$, X kalite değişkeninin kalitesindeki kaybın bir ölçüsü olsun. $L(X)$ 'in hedef vektörü T komşuluğunda Taylor Serisine açılımı Eş. 2.134'de verildiği gibi tanımlanmıştır.

$$L(\mathbf{X}) = L(\mathbf{T}) + \frac{[DL(\mathbf{T})']}{1!}(\mathbf{X}-\mathbf{T}) + (\mathbf{X}-\mathbf{T})' \frac{[D^2L(\mathbf{T})']}{2!} + \dots + \frac{[D^vL(\mathbf{T})']}{v!}(\mathbf{X}-\mathbf{T})^v + R_v(\mathbf{X}) \quad (2.134)$$

$L(\mathbf{X})$ 'in yaklaşığı ise,

$$L(\mathbf{X}) \approx \frac{1}{2}(\mathbf{X}-\mathbf{T})' [D^2L(\mathbf{T})'] (\mathbf{X}-\mathbf{T}) \quad (2.135)$$

biçimindedir.

Kayıp fonksiyonu yaklaşımı genel terimleri ile Eş. 2.136'da yeniden yazılırsa

$$L(\mathbf{X}) \approx \sum_{i=1}^v \sum_{j=1}^i \delta_{ij}(x_i - \tau_i)(x_j - \tau_j) \quad (2.136)$$

olur.

Burada, $\delta_{ij} = \left(\frac{\partial^2 L}{\partial x_i \partial x_j} \Big|_{\mathbf{X}=\mathbf{T}} \right) \forall i \neq j$, i ve j kalite değişkenleri arasındaki kayıp katsayısını göstermektedir.

Goethals ve Cho (2011), tanımladıkları MC_{pmc} çok değişkenli yeterlilik indeksini elde edebilmek için kayıp fonksiyonunun beklenen değerini Eş. 2.137'de verildiği gibi elde etmişlerdir.

$$E[L(\mathbf{X})] \approx \sum_{i=1}^v \delta_{ii} [\Lambda_{ii} + (\mu_i - \tau_i)^2] + \sum_{i=2}^v \sum_{j=1}^{i-1} \delta_{ij} [\Lambda_{ij} + (\mu_i - \tau_i)(\mu_j - \tau_j)] \quad (2.137)$$

Çok değişkenli yeterlilik indeksi Eş. 2.138 'de verildiği gibi tanımlanmıştır.

$$MC_{pmc} = \left[\frac{\prod_{i=1}^v (USL_i - LSL_i)}{Vol_v(PR_{\%99,73})} \right]^{1/v} \frac{1}{D} = \frac{V_R}{D} \quad (2.138)$$

Burada, $Vol_v(PR_{\%99,73}) = (\pi \chi_{v,0.0027}^2)^{v/2} [\Gamma(\frac{v}{2} + 1)]^{-1} \left| \sum_{i=1}^v \sum_{j=1}^i \Lambda_{ij} \right|^{1/2}$ ve D , Eş. 2.139'daki şekilde ifade edilmiştir.

$$D = \left[\delta_{ii} \left[1 + \sum_{i=1}^v (\mu_i - \tau_i) \Lambda_{ii}^{-1} (\mu_i - \tau_i) \right] + \delta_{ij} \left[1 + \sum_{i=1}^v (\mu_i - \tau_i) \Lambda_{ii}^{-1} (\mu_i - \tau_i) \right] \right] \quad (2.139)$$

Önerilen MC_{pmc} indeksi diğer indekslere göre birçok düzeltme sunmaktadır. İlk olarak, dikdörtgensel tolerans bölgesini göz ardı etmeden, endüstriyel uygulamalarda kullanımı kolay ve daha gerçekçi bir yaklaşım sergilemektedir. Ayrıca, süreç hedef değerden sapan süreç ortalamasını göz önünde bulunduran kalite kayıp fonksiyonunu dahil etmesi, hem üretici ve hem de müşterinin ürünü değerlendirmesi bakımından önemlidir (Goethals ve Cho, 2011).

2.2.6. Çok değişkenli normal dağılmayan süreçlerde süreç yeterlilik indeksleri

Klasik süreç yeterlilik indekslerinin tahmin edilmesi için bazı varsayımların sağlanması gerekir. Bu varsayımlardan en önemlisi sürecin normal dağılıma uygun olmasıdır. Klasik C_p ve C_{pk} gibi indekslerin hesaplanması varsayım bozulmalarından olumsuz etkilenir ve bu nedenle de hatalı tahminler elde edilir. Normallik varsayımının sağlanmadığı durumlar için süreç yeterlilik indekslerinin hesaplanmasında çeşitli yöntemler vardır ve bunlardan biri de Kernel yoğunluk tahmin yöntemidir. Kernel yoğunluk tahmin yöntemi, X rastgele değişkenin olasılık yoğunluk fonksiyonu $f(X)$ 'i tahmin etmek için kullanılan parametrik olmayan bir tahmin yöntemidir. Genellikle süreçten elde edilen gözlem değerlerinin dağılımı bilinmediği ya da bilinen bir dağılıma uymadığı durumlarda gözleme dayalı olarak Kernel yoğunluk tahminleri elde edilebilir.

Kernel tahmin edicisi, Silverman (1986: 15) tarafından Eş. 2.140'da verildiği gibi tanımlanır.

$$\hat{f}(x) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{x - X_i}{h}\right) \quad (2.140)$$

Burada, n gözlem sayısını, h düzleştirme ya da bant genişliği parametresini ifade eder. Kernel yoğunluk tahmin yöntemi, parametrik olmayan regresyon, yoğunluk tahmini ve risk fonksiyonunu içeren bir parametrik olmayan eğri tahminidir. Bu tahminler, eğrinin düzgünlüğünü kontrol eden bant genişliği parametresine ve Kernel fonksiyon türüne bağlıdır. K , Kernel fonksiyonu olarak adlandırılır ve Eş. 2.141'de verilen koşulu sağlayan genellikle simetrik olan bir olasılık yoğunluk fonksiyonu, örneğin normal dağılım fonksiyonu olarak belirlenebilir.

$$\int K(x) dx = 1 \quad (2.141)$$

Kernel fonksiyonları arasında en çok kullanılan Kernel fonksiyonu, Normal (Gaussian) Kernel fonksiyonudur ve Eş. 2.142'de verilmiştir.

$$K(x) = \frac{1}{\sqrt{2\pi}} \exp(-x^2/2) \quad (2.142)$$

Normal Kernel fonksiyonu, tanım aralığı gereği sınırsız bir fonksiyon türü olduğundan bazı durumlarda Epanechnikov Kernel fonksiyonu da tercih edilir. Epanechnikov Kernel fonksiyonu Eş. 2.143'deki gibi tanımlanır.

$$K(x) = \frac{3}{4} (1 - x^2) I_{[-1,1]}(x) \quad (2.143)$$

Burada, $I_{[-1,1]}$ fonksiyonun tanımlı aralığı olan $[-1,1]$ aralığını temsil etmektedir. Diğer bilinen Kernel fonksiyonları Çizelge 2.1.'de verilmiştir.

Çizelge 2.1. Kernel fonksiyon türleri

Kernel Fonksiyon Türü	$K(x)$
Uniform (Düzgün)	$\frac{1}{2}I_{[-1,1]}(x)$
Triangle (üçgen)	$(1 - x)I_{[-1,1]}(x)$
Quartic (dördüncü dereceden)	$\frac{15}{16}(1 - x^2)^2I_{[-1,1]}(x)$

Çizelge 2.1’de verilen Kernel fonksiyon türleri için verilebilecek örnekleri arttırmak mümkündür.

Çok değişkenli durumlar için ise, K fonksiyonu $K : \mathbb{R}^d \rightarrow \mathbb{R}$, d -değişkenli Kernel fonksiyonu olarak adlandırılan sürekli bir fonksiyondur ve Eş. 2.144-Eş. 2.146’de verilen koşulları sağlamalıdır.

$$\int_{\mathbb{R}^d} K(\mathbf{x}) d\mathbf{x} = 1 \quad (2.144)$$

$$\int_{\mathbb{R}^d} \mathbf{x}K(\mathbf{x}) d\mathbf{x} = 0 \quad (2.145)$$

$$\int_{\mathbb{R}^d} \mathbf{x}\mathbf{x}' K(\mathbf{x}) d\mathbf{x} = \mu_2(K)\mathbf{I} \quad (2.146)$$

Burada,

$$\mu_2(K) = \int_{\mathbb{R}^d} x_i^2 K(\mathbf{x}) dx \quad (2.147)$$

olarak tanımlanır.

$\mathbf{I}$, tüm i ’ler için $d \times d$ boyutlu birim matristir. Bu varsayımlar ile K , "0" ortalama vektörü ve $\mu_2(K)$ kovaryans matrisine sahip çok değişkenli bir yoğunluk fonksiyonudur.

H bant genişliği matrisi, simetrik pozitif tanımlı $d \times d$ boyutlu bir matristir ve çok değişkenli Kernel fonksiyonu Eş. 2.148’de verildiği gibi tanımlanır.

$$K_H(x) = |H|^{-1/2} K(H^{-1/2}x) \quad (2.148)$$

$f(x; H)$ fonksiyonunun tahmini Eş. 2.149'da verildiği gibi ifade edilir.

Tüm $\mathbf{x} \in \mathbb{R}^d$ için

$$\hat{f}(x; H) = n^{-1} \sum_{i=1}^n K_H(x - X_i) \quad (2.149)$$

şeklinde tanımlanır.

Çok değişkenli Kernel tahmin yönteminde en çok tercih edilen Kernel fonksiyon türü olan çok değişkenli Normal Kernel fonksiyonu Eş. 2.150'de tanımlandığı şekilde verilmiştir (Silverman, 1986: 76).

$$K(x) = (2\pi)^{-d/2} \exp\left(-\frac{1}{2} X^T X\right) \quad (2.150)$$

Kernel Tahmin Yöntemi normal dağılım göstermeyen verilerin olasılık yoğunluk fonksiyonu tahmininde kullanılan önemli bir parametrik olmayan tahmin yöntemidir. Kernel Tahmin Yönteminde kullanılan ve yoğunluk tahmininde rol oynayan iki önemli unsur; Kernel fonksiyon türü ve bant genişliği matrisi olarak sınıflandırılabilir. Özellikle, bant genişliği matrisi, kernel tahmininin performans değerlendirilmesinde önemli bir role sahiptir. Performans ölçütü, Kernel yoğunluk tahmininin hedef yoğunluk fonksiyonuna ne kadar yaklaştığı ile ilgilidir. Bu amaçla, Kernel tahmininin performansını ölçmede Bütünleşik Hata Kareler (ISE), Bütünleşik Hata Kareler Ortalaması (MISE) ve Hata Kareler Ortalamasına bir asimptotik yaklaşım olan (AMISE) gibi hata kriter ölçütleri kullanılmaktadır.

Bant genişliği seçimi

Bant genişliği matrisi seçiminde çeşitli yöntemler olmakla birlikte en sık kullanılanları cross validation (çapraz sağlama) ve Plug-in seçim yöntemidir. Daha az kullanımı olan ve Terrell (1990) tarafından tanımlanan Maksimum Düzleştirme Yönteminde (MDY), veri ölçeği ile tutarlı maksimum düzleştirme sağlayan yoğunluk tahmini elde edilmektedir. Terrell matris

parametrizasyonu olarak $h^2 H'$ kullanmıştır. Burada, $|H'| = 1$ olarak tanımlanmıştır. Hata kriteri AMISE değerini minimize eden h değeri Eş. 2.151'den hesaplanabilir.

$$h = \left[\frac{dV(K)}{n \int_{\mathbb{R}^d} tr^2(\mathbf{H}' Df(\mathbf{x})) d\mathbf{x}} \right]^{1/(d+4)} \quad (2.151)$$

Duong (2005), Terrell (1990) tarafından tanımlanan Maksimum Düzleştirme Yöntemine minimaks bir yaklaşım önererek Eş. 2.152'de verilen eşitliği elde etmiştir.

$$\hat{H}_{MS} = \left[\frac{(d+8)^{(d+6)/2} \pi^{d/2} V(K)}{16(d+2)n\Gamma\left(\frac{d}{2}+4\right)} \right]^{2/(d+2)} \mathbf{S} \quad (2.152)$$

Burada $\mathbf{S}$, örneklem varyans kovaryans matrisini, Γ ise Gamma fonksiyonunu ifade etmektedir.

Normal dağılıma uygun olmayan süreçlerde karşılaşılan en büyük zorluk, dağılımın bilinen bir dağılıma uygun olmaması nedeniyle uygun ürün olasılık tahminlerinde karmaşık integral alma işlemi ile karşı karşıya kalma durumudur.

Huang, Pahwa ve Kong (2012) dağılımı bilinmeyen tek değişkenli parametrik olmayan süreçler için yeterlilik indekslerini, Kernel yoğunluk tahmini ve Metropolis Hastings (*MH*) örneklem algoritması kullanarak elde etmişlerdir. Normallik varsayımı gerektiren tek değişkenli C_p ve C_{pk} indeksleri yerine uygunluk (ürün) temeline dayalı olarak elde edilen Y_p , Y_{pk} indekslerini kullanmışlardır. Y_p (potansiyel uygunluk), Y_{pk} (ürün) indekslerinin tanımı Eş. 2.153 - Eş. 2.156'da verilmiştir.

Ürün, spesifikasyon sınırları ya da bölgesi içine düşme olasılığı olarak tanımlanır.

Tek değişkenli süreçler için, ürün $= P[USL \leq X \leq LSL]$ olarak ve çok değişkenli süreçler için ürün $= P[X \in \Omega_S]$ olarak ifade edilir.

Tek değişkenli süreçler için bu indeksler,

$$Y_{pk} = \int_{LSL}^{USL} f(X) dx \quad (2.153)$$

$$Y_p = \text{Max}_{\mu} \{Y_{pk}\} = \text{Max}_{\mu} \left\{ \int_{LSL}^{USL} f(X) dx \right\} \quad (2.154)$$

Çok değişkenli süreçler için bu indeksler,

$$Y_{pk} = \int_{\Omega_s} f(X) dX \quad (2.155)$$

$$Y_p = \text{Max}_{\mu} \{Y_{pk}\} = \text{Max}_{\mu} \left\{ \int_{\Omega_s} f(X) dX \right\} \quad (2.156)$$

olarak tanımlanır. Burada μ , X 'in ortalama vektörünü, Ω_s spesifikasyon bölgesini göstermektedir.

Polansky (2001), normal dağılım göstermeyen süreçler için Kernel tahminine dayalı bir süreç yeterlilik tahmin yaklaşımı önermiştir. Süreç yeterlilik fonksiyonu olan, uygun olmayan ürün oranı p 'nin Kernel'e dayalı tahmini Eş. 2.157'de verildiği şekilde tanımlamıştır.

$$\hat{p}(H, S) = 1 - \int_{\Omega_s} \hat{f}(x, H) dx = 1 - n^{-1} \sum_{i=1}^n K_H(x - X_i) \quad (2.157)$$

Polansky'in (2001) çalışmasında önerdiği Kernel tahmin edicisinin $n < 200$ olan küçük çaplı örneklemelere uygulanması sonucu p 'yi iyi bir şekilde tahmin edemediği ve sonuçların farklı dağılımlara sahip tek bir simülasyon verisi üzerinden elde edildiği göz önünde bulundurulduğunda, dağılımdan çekilen tek bir örneklem verisinin kitle parametresi tahmin etmede zayıf kalacağı düşünülmektedir. Bu düşünceden yola çıkarak, tez çalışmasında Polansky (2001)'in çalışmasından farklı olarak, tek bir örneklem kümesi ile çalışmak yerine aynı dağılımlı birden fazla örneklem kümesi ile çalışılmıştır. Çalışmada, Polansky'nin küçük çaplı örneklemelerde zayıf kalan yönünü güçlendiren bir indeks önerilmiştir. Önerilen indeks, örneklem yöntemlerinin kitleyi tahmin etmedeki gücünü kullanmakta ve Kernel tahmin yöntemi ile örneklem yöntemlerinden biri olan Metropolis Hastings örneklemesini bir araya getirerek, Kernel tahmin yöntemi ile dağılımı tahmin edilen süreç ile aynı dağılımlı örneklem kümeleri üzerinden tahmin yapılmasına olanak sağlamaktadır.

Bu tez çalışmasında önerilen indeks, $\tilde{p}_{KMH}$ ile gösterilmiş ve Eş. 2.157’de verilen $\hat{p}(H, S)$ indeksine farklı bir yaklaşım olarak Eş. 2.158’te verildiği şekilde tanımlanmıştır.

$$\tilde{p}_{KMH} = \int_{\Omega_s} \hat{f}(x, H) dx = n^{-1} \int_{\Omega_s} \sum_{i=1}^n K_{HMH}(x - X_i) \quad (2.158)$$

Burada Ω_s spesifikasyon bölgesini ifade etmektedir. $\hat{f}(x, H)$, Kernel Tahmin Yöntemi kullanılarak elde edilen dağılım fonksiyonunun tahminini göstermektedir. Önerilen $\tilde{p}_{KMH}$ indeksi, Kernel tahmin yöntemine dayalı (K) Metropolis Hastings (MH) örnekleme uygun ürün olasılık tahminini ifade etmektedir.





3. NORMAL DAĞILMAYAN ÇOK DEĞİŞKENLİ SÜREÇ YETERLİLİK İNDEKSİ İÇİN BİR YÖNTEM ÖNERİSİ: KMH YÖNTEMİ

Süreç yeterlilik indeksleri ile ilgili olarak verilen geniş bir literatür anlatımından sonra çalışmamızda normal dağılmayan çok değişkenli süreçler için uygun ürün oranı p 'yi tahmin etmek amacıyla Kernel Yoğunluk Tahmin yöntemi ve (K) Metropolis Hastings (MH) örnekleme kullanılarak yeni bir yöntem önerilmiştir. Bu yöntem KMH Yöntemi olarak adlandırılmıştır. KMH yönteminin süreç yeterliliğine karar vermedeki başarısını ölçmek amacı ile bu yönteme dayalı yeni bir süreç yeterlilik indeksi tasarlanmıştır. Bu indeks ise $(\tilde{p}_{KMH})$ olarak ifade edilmiştir. Önerilen $\tilde{p}_{KMH}$ yaklaşımının uygun ürün oranı p 'yi tahmin etmedeki başarısını ölçmek için t-Copula yöntemi kullanılarak bir kıyaslama verisi üretilmiştir. Kıyaslama verisinden uygun ürün oranı p hesaplanarak Eş. 3.1 yardımıyla hata yüzdeleri elde edilmiştir.

$$\% \varepsilon = \left| \frac{\tilde{p}_{KMH} - p}{p} \right| \times \%100 \quad (3.1)$$

Çalışmada ayrıca $\tilde{p}_{KMH}$ değerleri ile Huang ve diğerleri'nin (2012) Y_{pk} (ürün) indeksi ve Polansky (2001)'in önerdiği Kernel tahminine dayalı süreç yeterlilik tahmin yaklaşımı olan $\hat{p}(H, S)$ değerleri hesaplanarak her birinin p tahminlerine ilişkin hata yüzdeleri elde edilmiş ve sonuçlar karşılaştırılmıştır.

3.1. Çalışmada Önerilen Yöntem ve Kullanılan Yaklaşımlar

Kernel tahmin yöntemi daha önce de belirtildiği gibi normal dağılmayan çok değişkenli süreçlerin dağılımını tahmin etmek amacıyla kullanılan parametrik olmayan bir yöntemdir. Kernel tahmin yönteminde, Kernel yoğunluk tahmin fonksiyonu $\hat{f}(x; H)$, Kernel fonksiyonu $K_H(\mathbf{x})$ ve Kernel fonksiyon türü $K(\mathbf{x})$ ile gösterilmiştir. Kernel yoğunluk tahmin yaklaşımını daha anlaşılır bir şekilde ifade etmek için aşağıda verilen eşitlikleri yazmakta fayda vardır:

$K_H(x) = |H|^{-1/2} K(H^{-1/2}x)$ ise burada, H bant genişliği matrisi olmakla birlikte $K(H^{-1/2}x)$, Eş. 3.2' de verildiği şekilde yazılır.

$$K\left(H^{-1/2}x\right) = |H|^{-1/2}K(x) \quad (3.2)$$

Bir Kernel fonksiyon türü olan çok değişkenli Normal Kernel fonksiyonu $K(\mathbf{x}) = (2\pi)^{-d/2} \exp(-1/2\mathbf{x}'\mathbf{X}\mathbf{x})$ ile ifade edilir. Kernel yoğunluk tahmin yönteminde çok değişkenli Normal Kernel fonksiyon türünün kullanılması durumunda Eş. 3.2 yeniden yazılırsa $K\left(H^{-1/2}x\right)$, Eş. 3.3'deki gibi elde edilir.

$$K\left(H^{-1/2}x\right) = |H|^{-1/2} \left[(2\pi)^{-d/2} \exp\left(-\frac{1}{2}\left(H^{-1/2}x\right)' \left(H^{-1/2}x\right)\right) \right] \quad (3.3)$$

Böylece, Kernel yoğunluk tahmin fonksiyonu $\hat{f}(x;H) = n^{-1} \sum_{i=1}^n K_H(x - X_i)$, Eş. 3.3 yardımıyla daha açık olarak,

$$\hat{f}(x;H) = n^{-1} \sum_{i=1}^n |H|^{-1/2} \left[(2\pi)^{-d/2} \exp\left(-\frac{1}{2}(x - X_i)' H^{-1}(x - X_i)\right) \right] \quad (3.4)$$

Eş. 3.4'te verildiği şekilde yazılır.

Çalışma verisinin dağılımı bilinmediği ve herhangi bilinen bir dağılıma uymadığı varsayılmıştır. Bu nedenle süreç verisinin dağılımı Eş. 3.4'te verilen Kernel yoğunluk tahmin fonksiyonu ile tahmin edilmiştir. Normal dağılmayan çok değişkenli süreçler için uygun ürün oranı p' 'yi tahmin etmek amacıyla çalışmamızda önerdiğimiz *KMH* yöntemine dayalı olarak tasarlanan yeterlilik indeksi $\tilde{p}_{KMH}$, Eş. 3.5'de verildiği şekilde tanımlanmıştır.

$$\tilde{p}_{KMH} = \int_{\Omega_s} \hat{f}(x, H) dx = n^{-1} \int_{\Omega_s} \sum_{i=1}^n K_{HMH}(x - X_i) dx \quad (3.5)$$

Burada H bant genişliği matrisini, Ω_s , spesifikasyon bölgesini göstermektedir. Eş. 3.5'te görülen $K_{HMH}(x - X_i)$ Eş. 3.6'da verildiği gibi ifade edilmiştir.

$$K_{HMH}(x - X_i) = |H|^{-\frac{1}{2}} K\left(H^{-\frac{1}{2}}(x - X_i)\right) \quad (3.6)$$

$K_{HMH}(x - X_i)$ gösteriminin kullanılma nedeni, süreç dağılımının Kernel yoğunluk yöntemi ile tahmin edilmesinden sonra aynı dağılıma sahip örneklem kümeleri elde etmek için süreç verisine Metropolis Hastings örnekleme yönteminin uygulanmasıdır.

Metropolis Hastings örnekleme

Metropolis Hastings (*MH*) örnekleme yaygın olarak kullanılan bir Markov Zincirleri Monte Carlo tekniğidir. *MH* örneklemesinin en temel kullanım amacı, doğrudan örneklem çekmesi zor dağılımlardan benzetim yapılarak büyük çaplı rastgele örneklem elde etmektir (Hitchcock, 2003). Bu çalışmada, normal dağılıma uygun olmayan çok değişkenli gözlem değerlerinin Kernel tahmininden yola çıkarak aynı Kernel dağılımlı ($\hat{f}(\mathbf{x})$) örneklem elde etmek amacı ile *MH* örnekleme kullanılacak, böylece, sürece ait normal dağılıma uygun olmayan çok değişkenli yeterlilik indeksleri tahmin edilmeye çalışılacaktır. *MH* örneklemesinin tercih edilmesinin nedeni, söz konusu tekniğin, Kernel tahmin yönteminde kullanılan fonksiyon türünden ve sürecin spesifikasyon bölgesinin şeklinden etkilenmemesi ve özellikle, düzensiz ve doğrudan integral işlemlerini zorlaştıran spesifikasyon bölgesine sahip çok değişkenli uygulamalar için kolay uygulanabilir olmasıdır. Bilinen parametrik dağılımlara (Normal, Uniform, Gamma) uygulanan Monte Carlo benzetimi ile $\hat{f}(X)$ dağılımından doğrudan örneklem çekmek dağılım yapısının karmaşık bir özellik taşıması nedeniyle kolay değildir. Bu nedenle, *MH* örnekleme uygulanarak $\hat{f}(X)$ dağılımından örneklem çekilmektedir.

MH örneklemesinde bir "hedef dağılımı (target distribution)" ve bir "öneri dağılımı (proposal distribution)" kullanılır. Süreç yeterlilik indekslerinin hesaplanması amacı ile süreçten elde edilen Kernel tahmini $\hat{f}(X)$ "hedef dağılımını"; benzetim yapılması kolaylığı açısından $q(\cdot | X)$ koşullu yoğunluk dağılımı ise "öneri dağılımını" ifade etmektedir.

Öneri dağılımı olarak en yaygın kullanılan dağılımlar, tek değişkenli durumlar için Uniform ($U(a,b)$) ve Normal dağılımdır. Uniform dağılımı x 'den bağımsız olduğundan Uniform öneri dağılımı $q(x^* | x) = q(x^*)$ şeklinde tanımlanır. Yine, Normal öneri dağılımı ise, $q(x^* | x) = \frac{1}{2\pi\sqrt{\sigma}} \exp\left(-\frac{(x^*-x)^2}{2\sigma^2}\right)$ olarak ifade edilir (Huang ve diğerleri, 2012).

Çok değişkenli *MH* örneklemesinde, $X = (X_1, X_2, \dots, X_d)$, d değişkenli rastgele bir örneklem vektörü ve $X^{(t)}$ örneklemin t . durumunu göstermek üzere;

Çok değişkenli *MH* örnekleme algoritması aşağıdaki adımları içerir:

Adım 1: $t = 1$ olarak alınır.

Adım 2: Bir $U = (u_1, u_2, \dots, u_d)$ başlangıç değeri vektörü üretilir ve $X^{(t)} = U$ olarak alınır.

Adım 3: $t = t + 1$ olarak alınır.

$q(X/X^{t-1})$ dağılımından bir öneri X^* üretilir.

Kabul olasılığı $\alpha = \min(1, \frac{p(X^*)}{p(X^{(t-1)})} \frac{q(X^{(t-1)}/X^*)}{q(X^*/X^{(t-1)})})$ hesaplanır.

$U(0, 1)$ dağılımından bir u noktası üretilir.

$u \leq \alpha$ ise önerilen X^* kabul edilir ve $X^{(t)} = X^*$ olarak alınır.

$u > \alpha$ ise önerilen X^* kabul edilmez ve $X^{(t)} = X^{(t-1)}$ olarak kalır.

Adım 4: 3. Adım $t = T$ olana kadar aynı işlemler devam ettirilir.

MH örnekleme yönteminin özelliklerinden ve yordam adımlarından kısaca bahsedildikten sonra, *MH* örneklemesinin kullanımında izlenecek adımlar şu şekildedir:

Öncelikle örneklem çekilmesi hedeflenen iki değişkenli Kernel tahmin dağılımı $\hat{f}(\mathbf{x})$ "hedef dağılımı" olarak; iki değişkenli normal dağılım ise "öneri dağılımı" olarak alınmıştır.

İki değişkenli normal dağılımın simetrik bir yapıda olması nedeniyle $\frac{q(\mathbf{x}^{(t-1)}/\mathbf{x}^*)}{q(\mathbf{x}^*/\mathbf{x}^{(t-1)})}$ oranı 1 değerini alır ve çok değişkenli *MH* örneklemesinde hesaplanacak kabul olasılığı

$\alpha = \min \left\{ 1, \frac{p(\mathbf{x}^*)}{p(\mathbf{x}^{(t-1)})} \right\}$ ile basitçe elde edilebilir (Martinez ve Martinez, 2002).

"Hedef dağılımı olasılık fonksiyonu"

$$\hat{f}(x; H) = n^{-1} \sum_{i=1}^n K_H(x - X_i) \quad (3.7)$$

"Öneri dağılımı olasılık fonksiyonu (iki değişkenli)"

$$f(x_1, x_2) = \frac{1}{2\pi\sqrt{\det\Sigma}} \exp\left(-\frac{1}{2}(X - \mu)' \Sigma^{-1} (X - \mu)\right) \quad (3.8)$$

olarak gösterilir ve burada bilindiği gibi,

$$\mu = \begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix}$$

$$\Sigma = \begin{bmatrix} \sigma_1^2 & \sigma_{12} \\ \sigma_{21} & \sigma_2^2 \end{bmatrix}$$

ile ifade edilir.

Önerilen $\tilde{p}_{KMH}$ yaklaşımı ile karşılaştırılmak üzere literatürde çalışılmış iki farklı indeks üzerinde durulmuştur.

$\tilde{p}_{KMH}$ yaklaşımı ile karşılaştırılan ilk indeks, dağılımı bilinmeyen tek değişkenli parametrik olmayan süreçler için Huang ve diğerleri'nin (2012) önerdiği Y_{pk} (ürün) indeksidir.

Eş. 2.155'de verilen Y_{pk} (ürün) indeksi, çok değişkenli süreçler için Eş. 3.9 yardımı ile tahmin edilmiştir.

$$Y_{pk} = 1 - \frac{\text{uygun olmayan ürün sayısı}}{\text{toplam örnek sayısı}} \quad (3.9)$$

$\tilde{p}_{KMH}$ yaklaşımı ile karşılaştırılan ikinci indeks, Polansky (2001)'in Kernel tahminine dayalı süreç yeterlilik tahmin yaklaşımıdır ve Eş. 2.157'de verilen $\hat{p}(H, S)$ değerleri çalışma verisi kullanılarak tahmin edilmiştir. Elde edilen sonuçlar için hata yüzdeleri hesaplanmıştır.

Ayrıca, çalışma verisinin çok değişkenli normal dağılıma uyduğu varsayımı altında klasik parametrik uygun ürün oranı p tahmin edilmiş ve eldeki sonuçlar karşılaştırılmıştır. Ortalaması μ , kovaryans matrisi Σ olan çok değişkenli normal dağılım için tanımlanan p parametresi, Eş. 3.10'daki gibidir.

$$p = \int_{\Omega_s} \Phi_{\mu, \Sigma}(X) dx \quad (3.10)$$

Eş. 3.10'da verilen uygun ürün oranı p 'nin tahmin edilmesi için Eş. 3.11 kullanılmıştır.

$$\tilde{p} = \int_{\Omega_s} \Phi_{\hat{\mu}, \hat{\Sigma}}(X) dx \quad (3.11)$$

Çalışma verisi çok değişkenli ve normal dağılmadığı aynı zamanda herhangi bilinen bir dağılıma uymadığı için verinin olasılık fonksiyonu üzerinden tahmin yapmaya uygun olmaması ve Kernel tahmin yöntemi ile tahmin edilen yoğunluk fonksiyonu üzerinden integral alma zorluğu nedeniyle, $\tilde{p}_{KMH}$ tahmininde, numerik integral hesaplama yöntemlerinden biri olan Trapezoidal Numerik İntegrasyon yöntemine başvurulmuştur.

3.2. Çalışma Verisi

Üretim işletmelerinin seri üretim öncesinde süreç kalite kontrolünü sağlayarak hatalı ürün üretimine mahal vermeden minimum maliyet, maksimum kapasite ile çalışmayı amaçlamaları gerekmektedir. Bu da üretim öncesi iyi bir planlama ile mümkündür. Aksi durumda, seri üretim esnasında alınan ürün özelliklerinin hedeflenen ürün özelliklerinden sapma göstermesine, uygun olmayan ürün üretimine, müşteri şikâyetlerine ve boşa giden harcamalara neden olacağı açıktır. Bu nedenle, çalışmada tasarlanan senaryolar seri üretime geçmeden önceki planlama aşaması olan çevrimdışı (offline) üretim olarak düşünülmüş ve üretimi yapılacak ürün kalite değişkenleri hedef vektörü ve spesifikasyon sınırları belirlenerek 4 farklı senaryo tasarlanmıştır. Senaryolar için üretilecek ürün kalite değişkenlerinin $X_1 = x_{11}, x_{12}, \dots, x_{1n}$ ve $X_2 = x_{21}, x_{22}, \dots, x_{2n}$ olduğu ve X_1 ile X_2 arasında $\rho = 0,7$ gibi yüksek bir ilişki bulunduğu varsayılmıştır. Her bir senaryo için üretimi yapılacak ürün hedef vektörü $[6, 5; 0, 65]$ ve spesifikasyon sınırları $0,2505 \leq X_1 \leq 10,2495$; $0,02505 \leq X_2 \leq 0,9999$ olarak belirlenmiştir.

Dağılımı bilinmeyen ve çok değişkenli normal dağılmayan süreçlere örnek olması açısından üretim senaryolarında kullanılacak süreç verisi MATLAB Programı yardımı ile üretilmiş olup gerekli kod dizisi EK-1'de verilmektedir. Süreç verileri üretilirken t-Copula yöntemi kullanılmıştır. Öncelikle, üretim senaryolarının her biri için işletmede üretilmesi mümkün

olan tüm ürün kapasitesinin (ürün kitlesi) bir milyon üründen oluştuğu varsayılmıştır. Ürün kalite değişkenleri olan X_1 ile X_2 'nin marjinal dağılımları $X_1 \sim \text{Gamma Dağılımı}(k, \theta)$ ve $X_2 \sim \text{Beta Dağılımı}(\alpha, \beta)$ olarak belirlenmiş ve belirlenen parametreleri Çizelge 3.1'de verilmiştir. X_1 ile X_2 'nin dağılım özellikleri belirlenirken, çok değişkenli ve normal dağılıma uygun olmayan süreçleri temsil etmesi amacıyla marjinal dağılımların çarpık özellik göstermesine dikkat edilmiştir. Bilindiği üzere, Gamma dağılımı parametrelerinden olan k , şekil parametresi; θ ise ölçü parametresi olarak adlandırılmaktadır. Gamma dağılımının çarpıklık değerinin $(\frac{2}{\sqrt{k}})$, k şekil parametresine bağlı olması nedeniyle, k büyüdükçe dağılım özelliklerinin normal dağılıma yaklaştığı söylenebilir. Ayrıca, Gamma dağılımının $[0, \infty]$ aralığında tanımlı olması nedeniyle senaryo verilerinin hedef vektörü ve spesifikasyon bölgesi açısından birbirine yakın ve mantıksal sonuçlar vermesi amacıyla Gamma dağılım parametreleri k ve θ her senaryo için aynı alınmıştır. İkinci marjinal dağılım olan Beta dağılımı için şekil parametre değerlerine göre bazı sınıflandırmalar mevcuttur. Örneğin;

- $\alpha < 1$; $\beta < 1$ için U şeklinde olduğu;
- $\alpha = 1$; $\beta > 1$ için kesinlikle düşüş gösterdiği;
- $\alpha = \beta$ için 1/2 etrafında simetrik olduğu;
- $\alpha > 1$; $\beta > 1$ tek modlu olduğu;
- $\alpha > 2$; $\beta = 1$ kesinlikle konveks olduğu bilinmektedir.

Beta dağılım parametreleri α ve β belirlenirken yukarıda belirtilen dağılım şekillerinden yola çıkarak tek modlu ve çarpık özellikte olanlar tercih edilmiştir. Çalışmada senaryolar için belirlenen çalışma verilerinin marjinal dağılım özellikleri Çizelge 3.1'de görülmektedir.

Çizelge 3.1. Veri setlerinin dağılım özellikleri

Veri Setleri	Marjinal Dağılımlar	
	$X_1 \sim \text{Gamma Dağılımı}(k, \theta)$	$X_2 \sim \text{Beta Dağılımı}(\alpha, \beta)$
Senaryo I	$k = 3,0; \theta = 2,0$	$\alpha = 2,0; \beta = 5,0$
Senaryo II	$k = 3,0; \theta = 2,0$	$\alpha = 0,5; \beta = 0,5$
Senaryo III	$k = 3,0; \theta = 2,0$	$\alpha = 2,0; \beta = 7,0$
Senaryo IV	$k = 3,0; \theta = 2,0$	$\alpha = 2,0; \beta = 6,0$

Çizelge 3.1’te verilen dağılım özelliklerine göre t-Copula yöntemi ile üretilen ürün kitle verileri kıyaslama verisi olarak kullanılmıştır. Üretim senaryoları için fabrikanın çevrimdışı aşamasında olduğu varsayılarak ilk fazda üretilen ürün miktarı 1000 olarak alınmıştır. Her bir senaryo için ilk faz ürün üretimi Çizelge 3.1’de verilen marjinal dağılım özelliklerine göre t-Copula yöntemi kullanılarak üretilmiştir.

3.3. Çalışma Adımları

Adım 1

Çalışmamızın ilk adımında, çalışma verilerinin çok değişkenli Kernel yoğunluk tahminleri $\hat{f}(X;H)$ ’yı elde etmek için Horová, Koláček ve Zelinka (2012) tarafından geliştirilen düzleştirme aracı kullanılmıştır. Tahmin edilen Kernel tahmin grafikleri Senaryo I için Şekil 3.3 ve Şekil 3.4’te, Senaryo II için Şekil 3.7 ve Şekil 3.8’de, Senaryo III için Şekil 3.11 ve Şekil 3.12’de, Senaryo IV için Şekil 3.15 ve Şekil 3.16’da verilmiştir. Çalışmada kullanılan Kernel fonksiyon türü $K(x)$, çok değişkenli Normal Kernel fonksiyonu olarak belirlenmiş ve bant genişliği matrisleri H Maksimum Düzleştirme Yöntemi (MDY) ile tahmin edilmiştir. Ayrıca, Kernel yoğunluk tahminlerinin bant genişliği parametresine olan duyarlılığını görebilmek amacıyla bant genişliği matrisleri H , çalışma verisinin Standart Normal Dağılıma uygunluğu varsayımı altında Referans Bant Genişliği matris tahmin yöntemi kullanılarak da tahmin edilmiş, sonuçlar Çizelge 3.2’de verilmiştir.

Çizelge 3.2. Bant genişliği matrisleri

Veri setleri	Maksimum Düzleştirme Yöntemi ile Bant Genişliği Matris Tahmini	Referans Bant Genişliği Matris Tahmini
Senaryo I - Faz I	$\begin{bmatrix} 1,4335 & 0,0457 \\ 0,0457 & 0,0031 \end{bmatrix}$	$\begin{bmatrix} 1,2185 & 0,0388 \\ 0,0388 & 0,0026 \end{bmatrix}$
Senaryo II - Faz I	$\begin{bmatrix} 1,4535 & 0,0898 \\ 0,0898 & 0,0148 \end{bmatrix}$	$\begin{bmatrix} 1,2355 & 0,0763 \\ 0,0763 & 0,0126 \end{bmatrix}$
Senaryo III - Faz I	$\begin{bmatrix} 1,3080 & 0,0354 \\ 0,0354 & 0,0021 \end{bmatrix}$	$\begin{bmatrix} 1,1119 & 0,0301 \\ 0,0301 & 0,0018 \end{bmatrix}$
Senaryo IV - Faz I	$\begin{bmatrix} 1,5108 & 0,0420 \\ 0,0420 & 0,0025 \end{bmatrix}$	$\begin{bmatrix} 1,2843 & 0,0357 \\ 0,0357 & 0,0021 \end{bmatrix}$

Çizelge 3.2’den de görüldüğü gibi Maksimum Düzleştirme Yöntemi (MDY) ile tahmin edilen bant genişliği matris elemanlarının Referans Bant Genişliği (MBG) ile tahmin edilen matris elemanlarından daha büyük değerlere sahip olduğu, bu nedenle Maksimum Düzleştirme Yöntemi ile yapılan tahminlerin süreç verisine daha fazla bir düzgünleştirme işlemi yaparak Kernel yoğunluk tahminlerinin elde edildiği söylenebilir.

Adım 2

Her bir senaryo için çalışma verisinin Kernel yoğunluk tahminleri $\hat{f}(X;H)$ elde edildikten sonra 2. Adımda, üretim senaryoları için ilk fazda üretilen ürün dağılım özelliklerinin her faz için aynı olacağı varsayılarak üretime geçmeden önce ilk fazdan alınan ürün verisi ile aynı dağılımlı örneklem çekilmesi ve sürecin yeterliliğine karar vermek için kullanımı önerilen $\tilde{p}_{KMH}$ değerlerinin örneklem kümeleri üzerinden elde edilmesi hedeflenmiştir. Bu amaçla MH örneklem algoritması kullanılmıştır.

Seri üretim öncesinde sadece ilk faz ürün verisinden tahmin edilecek $\tilde{p}_{KMH}$ değerinin kitle parametresi p ile karşılaştırılarak sürecin yeterli olup olmadığına karar vermede yanıltıcı sonuçlara neden olacağı düşünülmektedir. Bu nedenle ilk faz ürün dağılımı Kernel dağılım tahmini $\hat{f}(X;H)$ ile aynı dağılımlı örneklem kümelerinin her biri $n = 50$, $n = 100$, $n = 250$, $n = 500$ ve $n = 1000$ büyüklüğünde olacak şekilde belirlenmiş ve 10 tane örneklem kümesi elde edilmiştir. Ayrıca örneklem küme sayısı artırıldığında sonuçlar arasında ortaya çıkan varyans değişimlerini gözlemlemek amacıyla 20 tane örneklem kümesi çektilerilerek aynı işlem her bir örneklem kümesine de uygulanmıştır. MH örnekleme için MATLAB programında gerekli kodlar yazılarak program çıktısı elde edilmiş ve kod dizisi EK-1’de verilmiştir.

Adım 3

Çalışmamızın bu aşamasında, kıyaslama verisi olarak kullanılan ürün kitlesinden hesaplanan parametrik uygun ürün oranı p ile süreçten çekildiği varsayılan ve ilk faz ürün verisi ile aynı dağılımlı örneklem kümelerinden tahmin edilecek $\tilde{p}_{KMH}$, Y_{pk} (ürün), $\tilde{p}$ ve $\hat{p}(H,S)$ değerleri karşılaştırılacaktır. Hata yüzdeleri Eş. 3.1 ile verilen $\% \varepsilon = \left| \frac{\tilde{p}_{KMH} - p}{p} \right| \times \%100$ yardımıyla

hesaplanmıştır. Hesaplanan hata yüzdeleri "Simülasyon Sonuçları ve Değerlendirilmesi" bölümünde Çizelge 4.1 - Çizelge 4.40'da verilmiş ve yorumlanmıştır.

3.4. Çalışma Verilerinin Dağılım Özellikleri

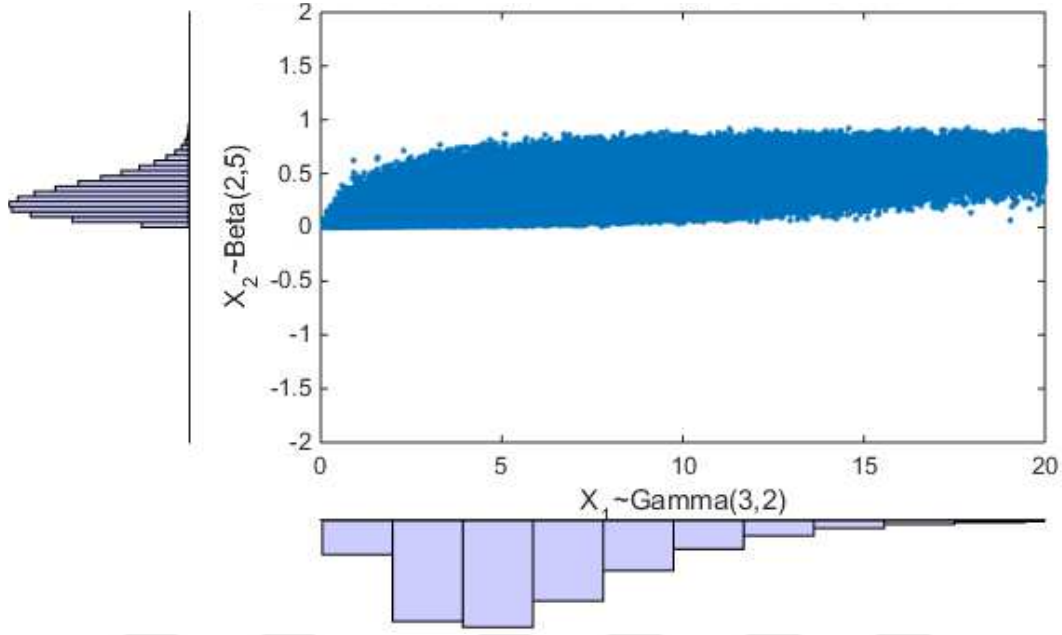
Çalışmada kıyaslama verisi olarak kullanılacak ürün kitlesi $N = 1\ 000\ 000$ olarak belirlendikten sonra, ürün uygunluğuna karar vermek amacı ile spesifikasyon sınırları belirlenmiştir. Daha önce de ifade edildiği gibi spesifikasyon sınırları dışsal olarak belirlenen, kalite mühendislerince teknik bilgi ve deneyim sonucunda üretilen ürün niteliğine göre karar verilen sınırlardır. Bu nedenle çalışmada spesifikasyon sınırları belirlenirken üretilen dağılım ortalamaları ve maksimum, minimum değerleri dikkate alınarak mantıksal çerçevede belirlenmiş ve aşağıda verilmiştir:

Spesifikasyon Alt Sınırı: $[X_1 = 0,2505; X_2 = 0,02505]$,

Spesifikasyon Üst Sınırı: $[X_1 = 10,2495; X_2 = 0,9999]$.

Senaryo I

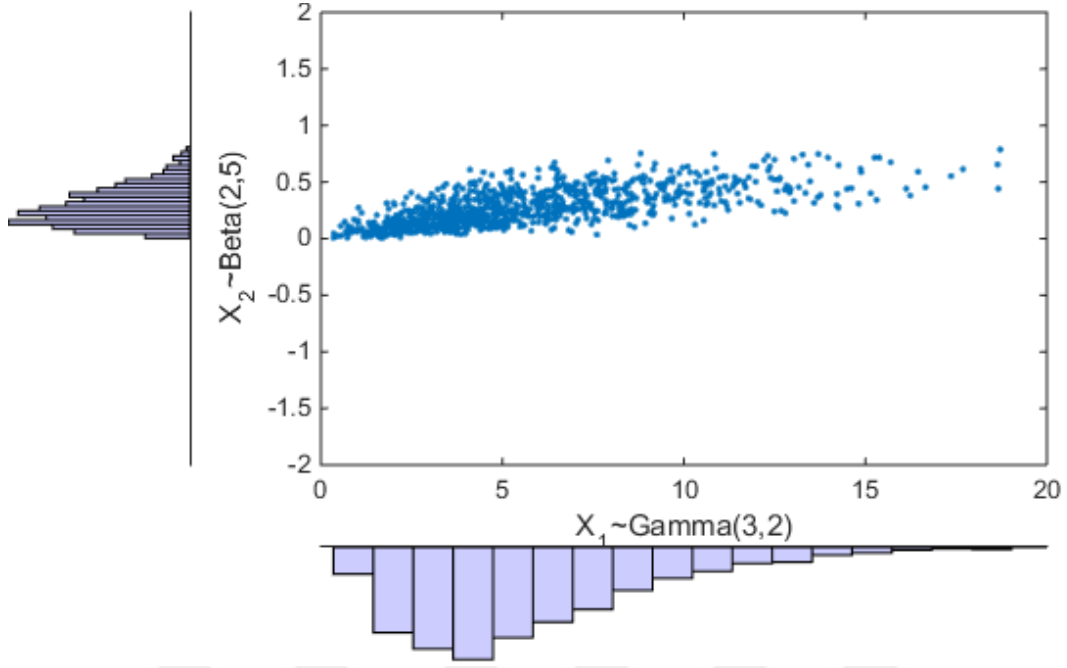
Senaryo I için, Marjinal dağılımları daha önce belirlenen ve Çizelge 3.1.'de verilen $X_1 \sim Gamma(3,2)$, $X_2 \sim Beta(2,5)$ dağılım özelliklerine sahip değişkenler arası $\rho = 0,7$ ilişki bulunan çok değişkenli ürün kitlesi üretilmiştir. Kitle Saçılım Grafiği Şekil 3.1'de verilmiştir.



Şekil 3.1. Ürün kitle saçılım grafiği (Senaryo I)

Şekil 3.1’de görülen ürün kitle saçılım grafiği üretilen verinin dağılım özelliklerini ve verinin düzlem üzerindeki saçılımını göstermektedir.

Senaryo I için ürün kitle verisinden hesaplanan uygun ürün sayısı 876 694 olarak elde edilmiştir. Uygun ürün oranı $p = 0,8767$ şeklinde hesaplanmıştır. Daha sonra yine, t-Copula yöntemi ile $X_1 \sim Gamma(3,2)$, $X_2 \sim Beta(2,5)$ dağılım özelliklerine sahip ve değişkenler arası $\rho = 0,7$ ilişki bulunan çok değişkenli ilk faz ürün verisi üretilmiştir. Üretilen verinin saçılım grafiği Şekil 3.2’de verilmiştir.

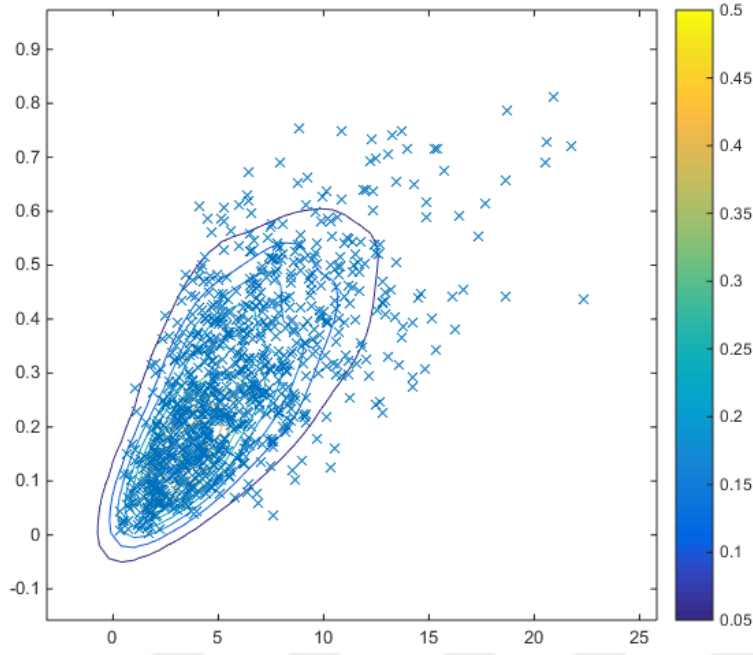


Şekil 3.2. İlk faz ürün saçılım grafiği (Senaryo I)

Şekil 3.2’de verilen ilk faz ürün saçılım grafiğinden, iki değişken arasında pozitif yönlü doğrusal bir ilişki olduğu görülmektedir.

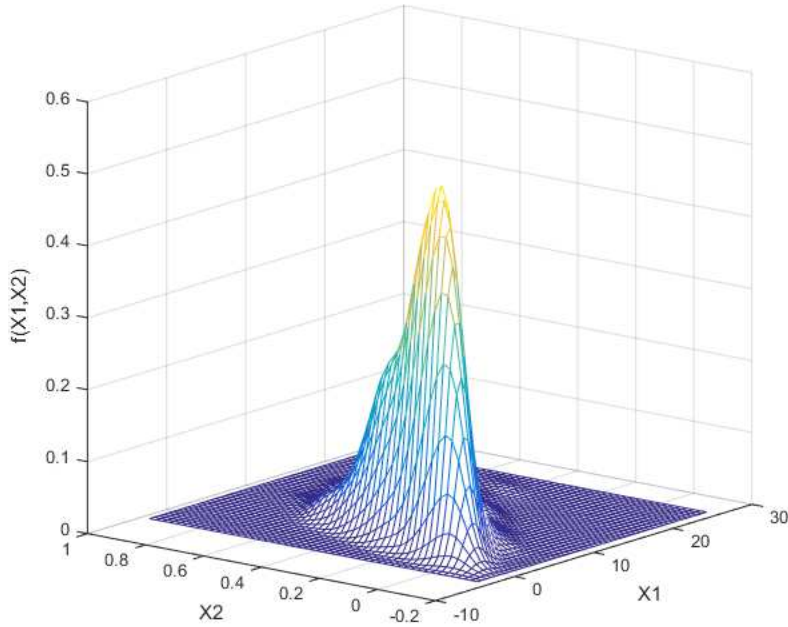
Ürün özellikleri belirlenerek elde edilen ilk faz ürün verisinin dağılımı normal dağılıma uygun olmadığı için $f(X)$ dağılımının $\hat{f}(X;H)$ tahmini Kernel tahmin yöntemi ile elde edilmiştir.

Kernel yoğunluk tahmin grafikleri ilk faz ürün verisi için Şekil 3.3-Şekil 3.4’de görülmektedir.



Şekil 3.3. İlk faz ürün kernel yoğunluk tahmini kontur grafiği (Senaryo I)

Şekil 3.3'te ilk faz ürün verisinden elde edilen Kernel yoğunluk tahmin kontur grafiği görülmektedir. Kontur grafiği, verinin Kernel yoğunluk tahmini sonrasındaki üstten yoğunluk kesitini göstermektedir. Her bir kontür çizgisi, eşit yoğunluk yüksekliğine sahip noktaları ifade etmekte ve şekilden verinin tek modlu bir yapıya sahip olduğu anlaşılmaktadır.

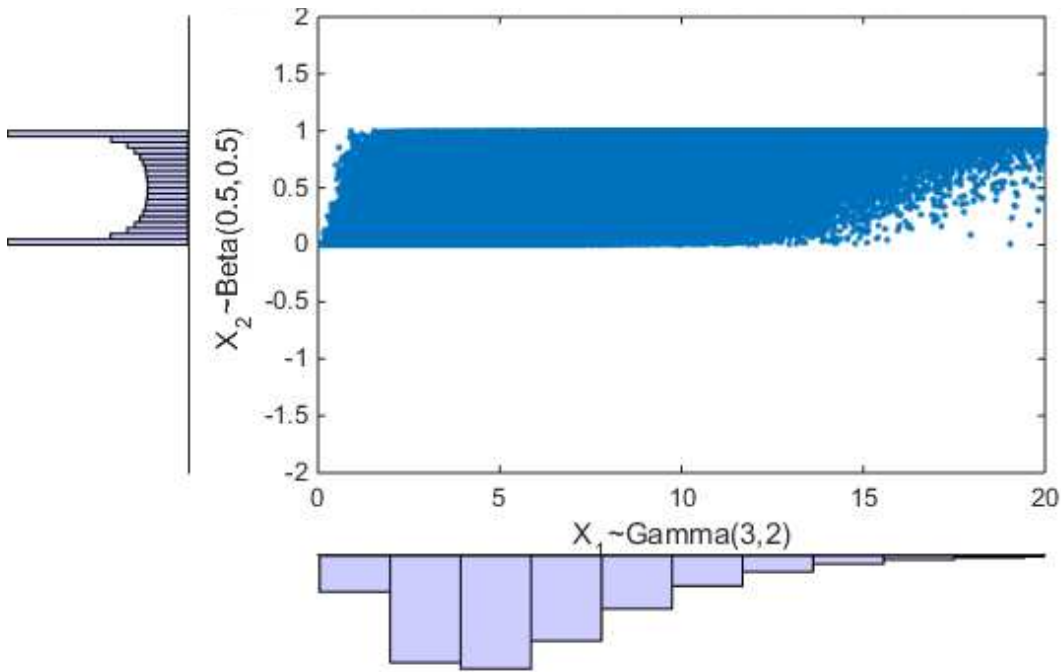


Şekil 3.4. Kernel yoğunluk tahmin grafiği (Senaryo I)

Şekil 3.4'te ilk faz ürün verisinin Kernel yoğunluk tahmin grafiği görülmektedir. Şekil 3.3'teki kontur grafiğinin 3-boyutlu görüntüsü olan grafikten verinin tek modlu olduğu daha açık görülmektedir.

Senaryo II

Senaryo II için, Çizelge 3.1'de Marjinal dağılımları $X_1 \sim \text{Gamma}(3,2)$ $X_2 \sim \text{Beta}(0,5, 0,5)$ olarak verilen çok değişkenli ve aralarında $\rho = 0,7$ ilişki bulunan ürün kitlesi üretilmiştir. Üretilen Kitle Saçılım Grafiği Şekil 3.5'de verilmiştir.

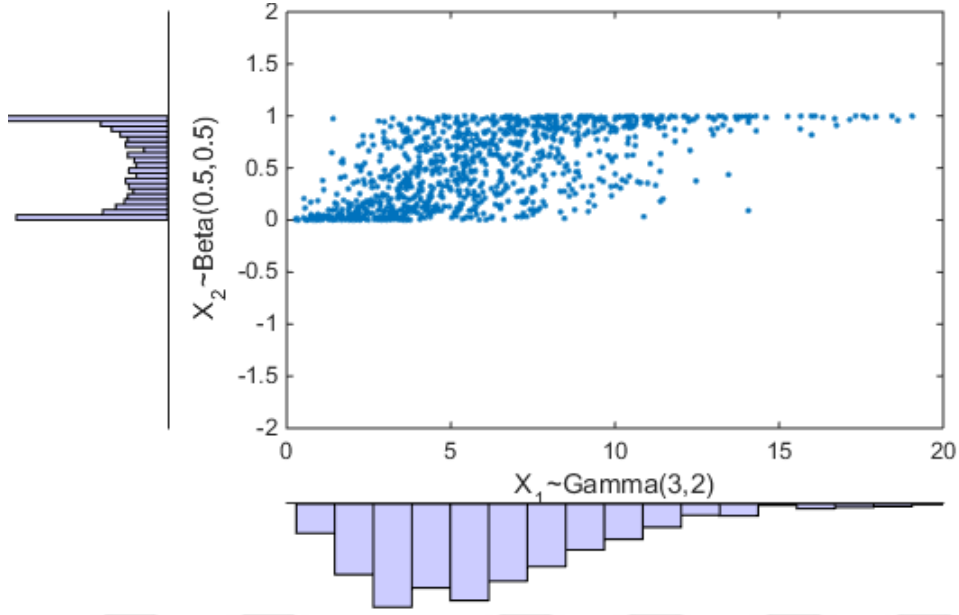


Şekil 3.5. Ürün kitle saçılım grafiği (Senaryo II)

Şekil 3.5'de görülen ürün kitle saçılım grafiği üretilen verinin dağılım özelliklerini ve verinin saçılımını göstermektedir.

Senaryo II için uygun ürün sayısı 783 703 ürün olarak elde edilmiştir. Uygun ürün oranı $p = 0,7837$ şeklinde hesaplanmıştır.

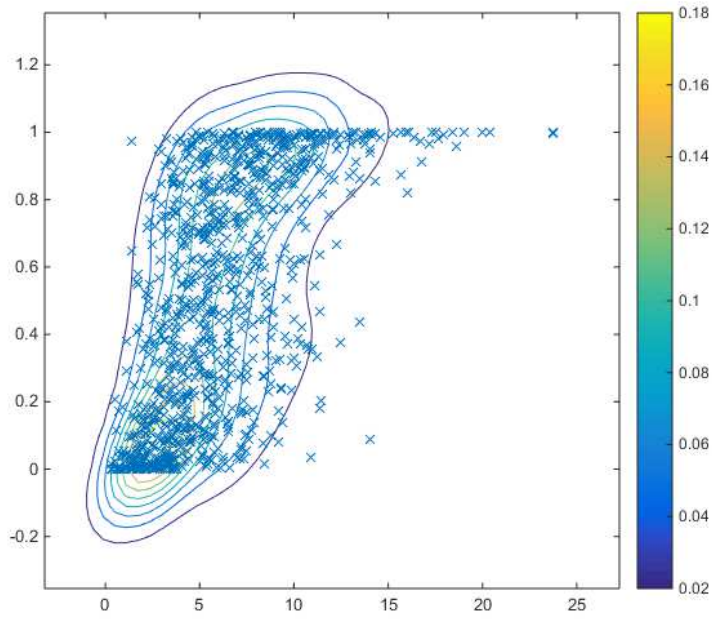
Daha sonra, ilk faz ürün verisi üretilmiş bunun için yine $X_1 \sim \text{Gamma}(3,2)$ $X_2 \sim \text{Beta}(0,5,0,5)$ dağılım özellikleri kullanılarak çok değişkenli ve aralarında $\rho = 0,7$ ilişki bulunan veri aynı şekilde elde edilmiştir. Üretilen verinin saçılım grafiği Şekil 3.6'da verilmiştir.



Şekil 3.6. İlk faz ürün saçılım grafiği (Senaryo II)

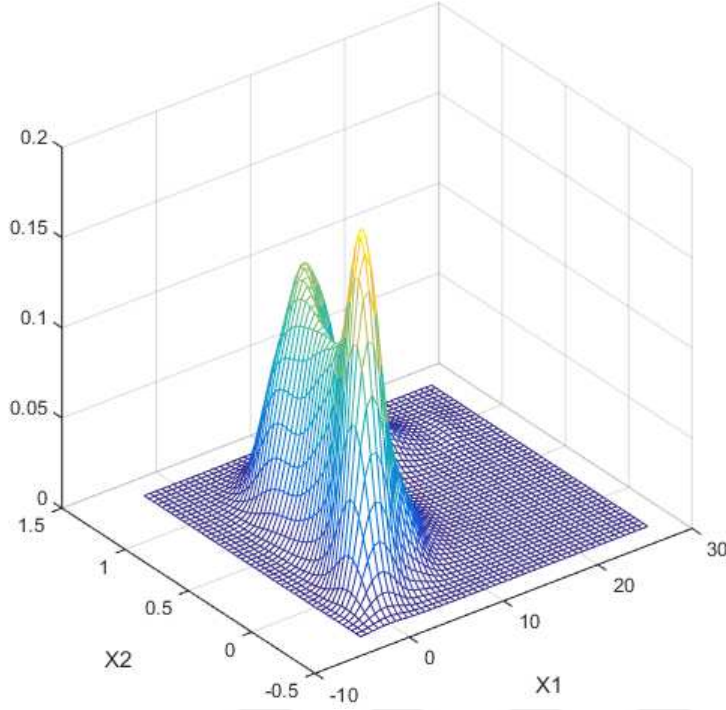
Şekil 3.6’da verilen ilk faz ürün saçılım grafiğinden, iki değişken arasında pozitif yönlü doğrusal bir ilişki olduğu görülmektedir.

Kernel yoğunluk tahmin grafikleri İlk faz ürün verisi için Şekil 3.7-Şekil 3.8’de görülmektedir.



Şekil 3.7. İlk faz ürün kernel yoğunluk tahmini kontur grafiği (Senaryo II)

Şekil 3.7, ilk faz ürün verisinin Kernel Yoğunluk tahmin kontur grafiğini göstermektedir. Kontur grafiğinden verinin iki modlu olduğu söylenebilir.

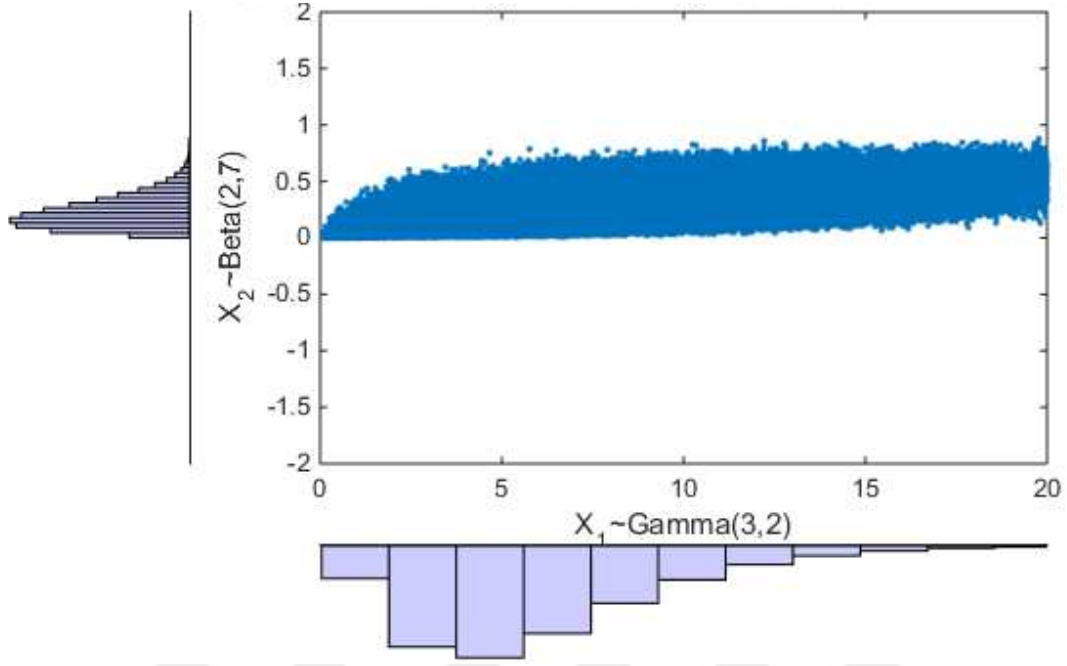


Şekil 3.8. Kernel yoğunluk tahmin grafiği (Senaryo II)

Şekil 3.8’de ilk faz ürün verisinin MATLAB programı yardımı ile çizilen Kernel Yoğunluk tahmin grafiği görülmektedir. Şekil 3.7’deki kontur grafiğinin 3-boyutlu grafiği olan yoğunluk grafiğinden verinin iki modlu olduğu görülmektedir.

Senaryo III

Senaryo III için t-Copula yöntemi kullanılarak MATLAB yazılımı yardımıyla marjinal dağılımları Çizelge 3.1’de verilen $X_1 \sim \text{Gamma}(3,2)$ $X_2 \sim \text{Beta}(2,7)$ olan ürün kitlesi üretilmiştir. Üretilen Kitle Saçılım Grafiği Şekil 3.9’da verilmiştir.



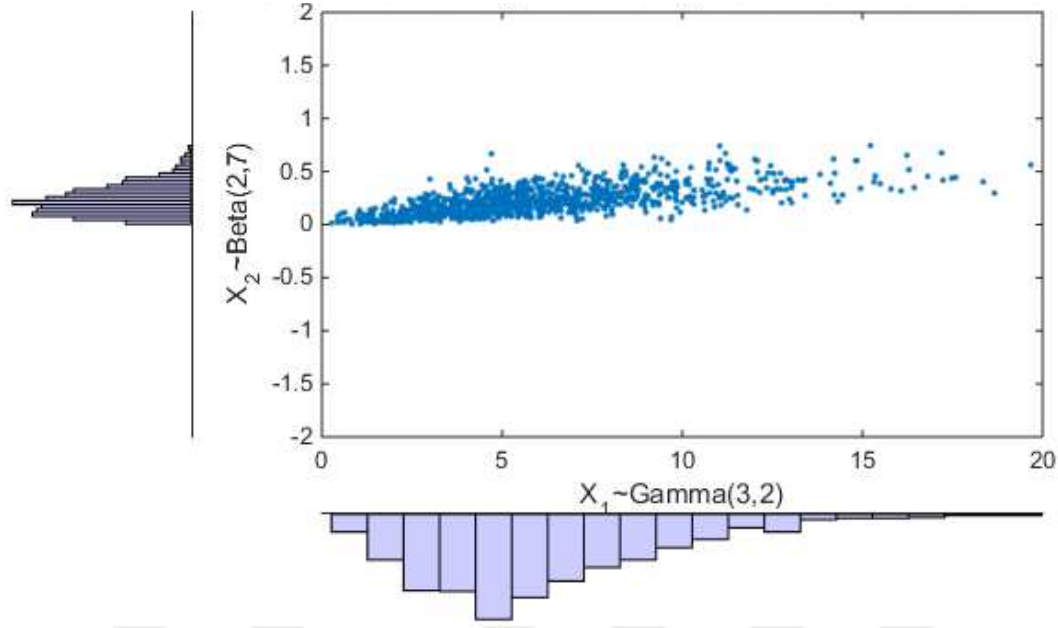
Şekil 3.9. Ürün kitle saçılım grafiği (Senaryo III)

Şekil 3.9’da görülen ürün kitle saçılım grafiği, üretilen verinin dağılım özelliklerini ve verinin saçılımını göstermektedir.

Uygun ürün sayısı Senaryo III için 869 378 ürün olarak elde edilmiştir. Uygun ürün oranı $p = 0,8694$ şeklinde hesaplanmıştır.

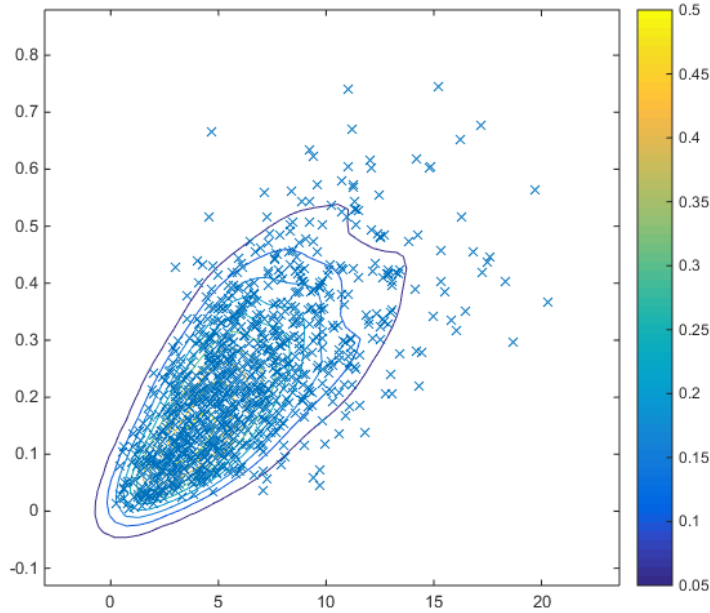
$X_1 \sim Gamma(3,2)$ $X_2 \sim Beta(2,7)$ marjinal dağılım özelliklerine sahip iki değişkenli aralarında ilişki $\rho = 0,7$ bulunan ilk faz ürün verisi üretilmiştir.

Üretilen verinin saçılım grafiği Şekil 3.10’da verilmiştir.



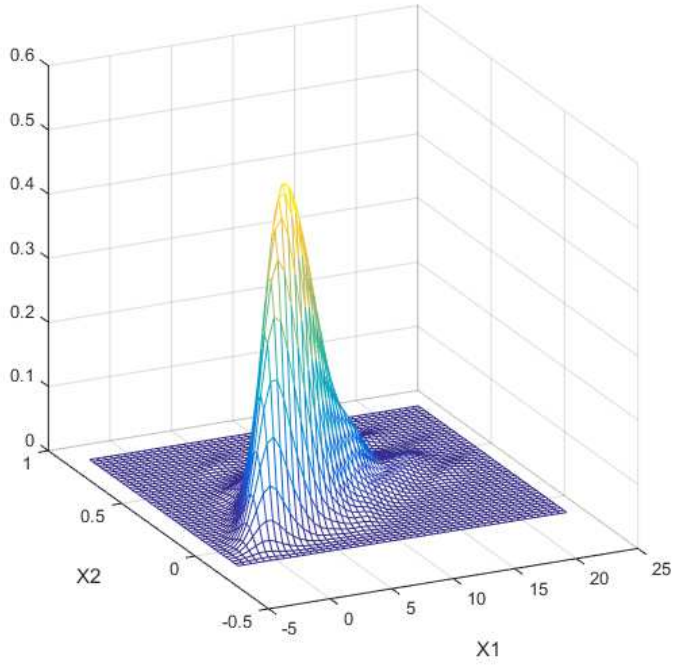
Şekil 3.10. İlk faz ürün saçılım grafiği (Senaryo III)

Şekil 3.10’da verilen ilk faz saçılım grafiğinden, iki değişken arasında pozitif yönlü doğrusal bir ilişki olduğu görülmektedir.



Şekil 3.11. İlk faz ürün kernel yoğunluk tahmini kontur grafiği (Senaryo III)

Şekil 3.11, ilk faz ürün verisinin MATLAB programı yardımı ile elde edilen Kernel Yoğunluk tahmin kontur grafiğini göstermektedir.

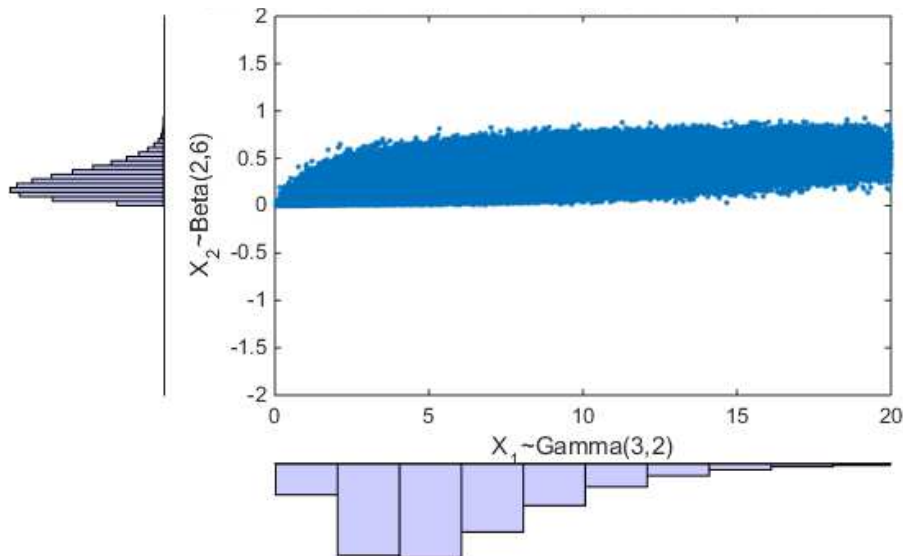


Şekil 3.12. Kernel yoğunluk tahmin grafiği (Senaryo III)

Şekil 3.12’de ilk ürün verisinin MATLAB programı yardımı ile elde edilen Kernel Yoğunluk tahmin grafiği görülmektedir.

Senaryo IV

Senaryo IV için, Marjinal dağılımları $X_1 \sim \text{Gamma}(3,2)$ $X_2 \sim \text{Beta}(2,6)$ olan ürün kitlesi üretilmiştir. Üretilen Kitle Saçılım Grafiği Şekil 3.13’de verilmiştir.

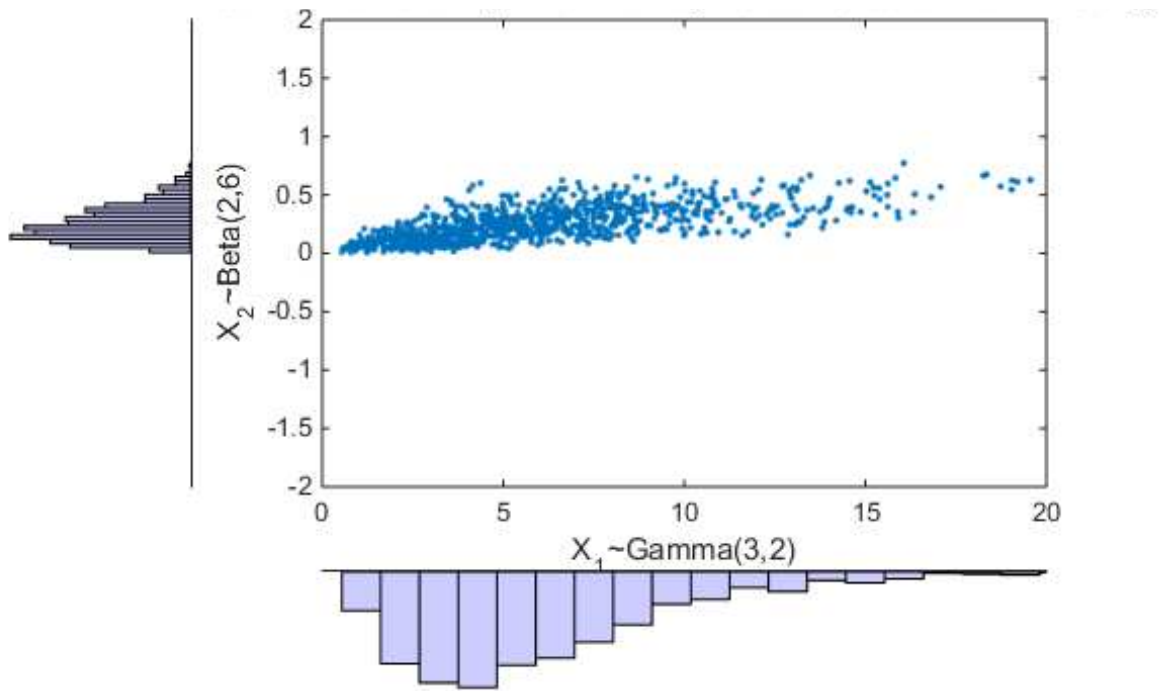


Şekil 3.13. Ürün kitle saçılım grafiği (Senaryo IV)

Şekil 3.13’de görülen ürün kitle saçılım grafiği üretilen verinin dağılım özelliklerini ve verinin saçılımını göstermektedir.

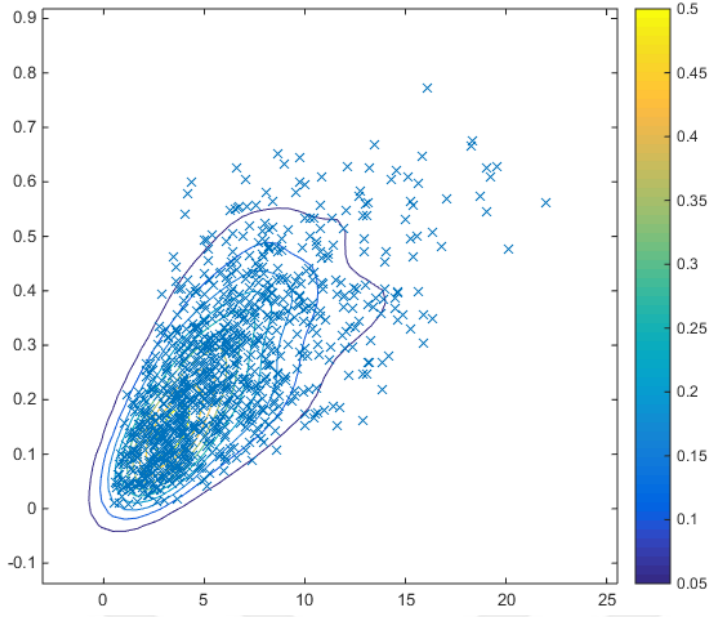
Uygun ürün sayısı Senaryo IV için 873 499 ürün olarak elde edilmiştir. Uygun ürün oranı $p = 0,8735$ şeklinde hesaplanmıştır.

Uygun ürün oranı p hesaplandıktan sonra, $X_1 \sim \text{Gamma}(3,2)$ $X_2 \sim \text{Beta}(2,6)$ dağılım özelliklerine sahip ilk faz ürün verisi üretilmiştir. üretilen verinin saçılım grafiği Şekil 3.14’de verilmiştir.



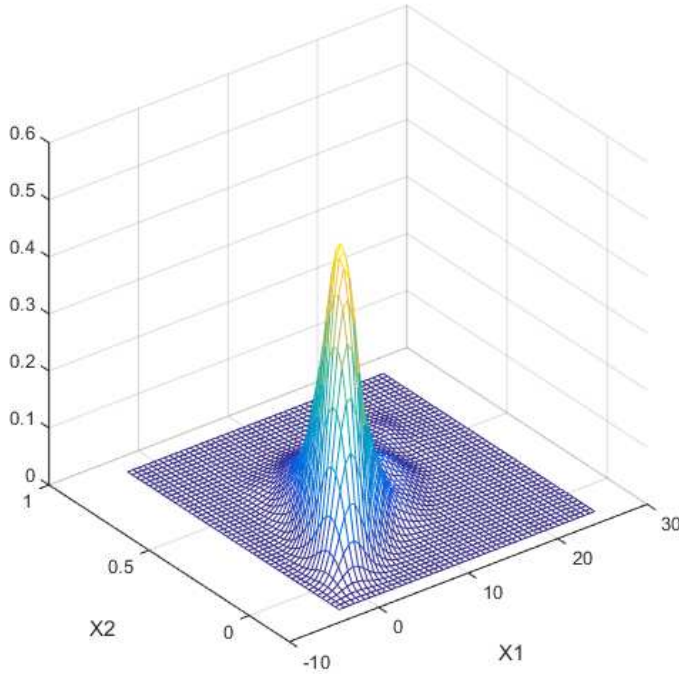
Şekil 3.14. İlk faz ürün saçılım grafiği (Senaryo IV)

Şekil 3.14’de verilen ilk faz ürün saçılım grafiğinden, iki değişken arasında pozitif yönlü doğrusal bir ilişki olduğu görülmektedir.



Şekil 3.15. İlk faz ürün kernel yoğunluk tahmini kontur grafiği (Senaryo IV)

Şekil 3.15, ilk faz ürün verisinin Kernel Yoğunluk tahmin kontur grafiğini göstermektedir. İlk faz ürün verisinin tek modlu olduğu görülmektedir. Kontur grafiğinin 3- boyutlu görüntüsü Şekil 3.16’da verilmiştir.



Şekil 3.16. Kernel yoğunluk tahmin grafiği (Senaryo IV)

Şekil 3.16’dan Kontur grafiğine benzer olarak ilk faz ürün verisinin tek modlu olduğu görülmektedir.

Metropolis Hastings yöntemi ile çekilen örneklem kümeleri ve özellikleri

Daha önce de belirtildiği gibi, çevrimdışı üretim aşamasında ilk faz veriden hesaplanan parametre tahmin sonuçlarının uygun ürün oranı p 'yi iyi bir şekilde tahmin ettiğini söylemek ve tek bir örneklem kümesi üzerinden yakalanan tahmin başarısını kıstas alarak çevrimiçi üretimde güvenilir sonuçlar elde edildiğini düşünmek süreç yeterliliği açısından yanıltıcı sonuçlar verecektir. Bu nedenle Polansky (2001)' in aksine ilk faz ürün verisi ile çalışmak yerine, ilk faz verisi ile aynı dağılımlı örneklem kümeleri üzerinden yaklaşım sonuçlarını yorumlamanın daha doğru olacağı düşünülmüştür. Böylece Kernel dağılımı tahmin edilen ilk faz verisinden örneklem çekmek ve daha duyarlı tahminler elde etmek amacı ile Metropolis Hastings örnekleme yöntemi kullanılmıştır.

MH örnekleme ile her birinin gözlem sayısı $n = 50, n = 100, n = 250, n = 500$ ve $n = 1000$ olan 10'ar ve 20'şer tane örneklem küme seti oluşturulmuş ve örneklem kümelerinden Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H, S)$ tahminleri ve hata yüzdeleri hesaplanmıştır. Sonuçlar 4. "Simülasyon Sonuçları ve Değerlendirilmesi" kısmında çizelgeler halinde verilmiş ve örneklem küme setlerinden elde edilen hata yüzdelerini gösteren çizgi grafiklerine de aynı bölümde yer verilmiştir. *MH* örnekleme, tahmin hesaplamaları ve çizgi grafiklerinin elde edilmesi için Matlab Programında yazılan kod dizileri Ek-1 de verilmektedir.

4. SİMÜLASYON SONUÇLARI VE DEĞERLENDİRİLMESİ

4.1. Senaryo I İçin Elde Edilen Sonuçlar

Senaryo I için Kernel yoğunluk tahmin yöntemine dayalı MH örnekleme ile $n = 50$, $n = 100$, $n = 250$, $n = 500$ ve $n = 1000$ gözlemlili örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H, S)$ değerleri, kitleden hesaplanan uygun ürün oranı p ile karşılaştırılarak hata yüzdeleri elde edilmiştir. Kernel yoğunluk tahminleri yapılırken Maksimum Düzleştirme Yöntemi (MDY) bant genişliği matrisi ve örneklem kümelerinin Normal dağılım gösterdiği varsayımına dayalı Referans Bant Genişliği (RBG) matrisi kullanılmıştır. Her bir gözlem sayısı için her iki bant genişliği matrisi kullanılarak hesaplanan tahmin sonuçları $n = 50$ gözlem sayısına sahip örneklem kümeleri için Çizelge 4.1’de verilmiştir.

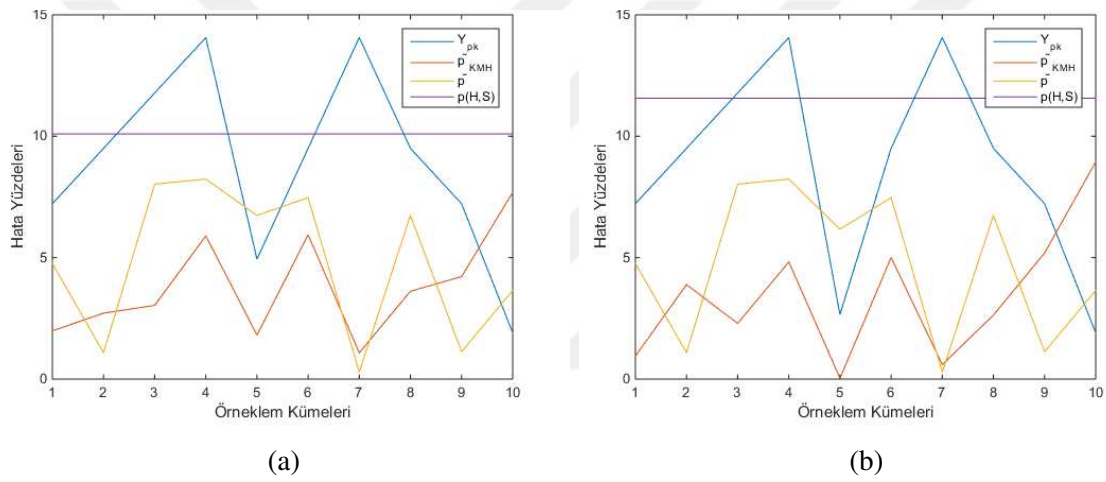
Çizelge 4.1. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 50$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
1	0,9400	0,8941	0,9185	0,7882	%7,220	%1,985	%4,768	%10,095	0,9400	0,8850	0,9185	0,7753	%7,220	%0,947	%4,768	%11,566
2	0,9600	0,8529	0,8862	0,7882	%9,502	%2,715	%1,084	%10,095	0,9600	0,8426	0,8862	0,7753	%9,502	%3,890	%1,084	%11,566
3	0,9800	0,9033	0,9471	0,7882	%11,783	%3,034	%8,030	%10,095	0,9800	0,8968	0,9471	0,7753	%11,783	%2,293	%8,030	%11,566
4	1,0000	0,9284	0,9489	0,7882	%14,064	%5,897	%8,235	%10,095	1,0000	0,9190	0,9489	0,7753	%14,064	%4,825	%8,235	%11,566
5	0,9200	0,8926	0,9358	0,7882	%4,939	%1,814	%6,741	%10,095	0,9000	0,8764	0,9308	0,7753	%2,658	%0,034	%6,171	%11,566
6	0,9600	0,9288	0,9422	0,7882	%9,502	%5,943	%7,471	%10,095	0,9600	0,9206	0,9422	0,7753	%9,502	%5,007	%7,471	%11,566
7	1,0000	0,8862	0,8792	0,7882	%14,064	%1,084	%0,285	%10,095	1,0000	0,8715	0,8792	0,7753	%14,064	%0,593	%0,285	%11,566
8	0,9600	0,9084	0,9358	0,7882	%9,502	%3,616	%6,741	%10,095	0,9600	0,8998	0,9358	0,7753	%9,502	%2,635	%6,741	%11,566
9	0,9400	0,8397	0,8866	0,7882	%7,220	%4,220	%1,129	%10,095	0,9400	0,8312	0,8866	0,7753	%7,220	%5,190	%1,129	%11,566
10	0,8600	0,8093	0,8447	0,7882	%1,905	%7,688	%3,650	%10,095	0,8600	0,7983	0,8447	0,7753	%1,905	%8,943	%3,650	%11,566
ORTALAMA	0,9520	0,8844	0,9125	0,7882	%8,970	%3,799	%4,814	%10,095	0,9500	0,8741	0,9120	0,7753	%8,742	%3,436	%4,756	%11,566
VARYANS	0,0017	0,0015	0,0013	0,0000	%0,148	%0,045	%0,095	%0,000	0,0019	0,0016	0,0013	0,0000	%0,173	%0,073	%0,093	%0,000

Çizelge 4.2. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 50$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
ORTALAMA	0,9550	0,8844	0,9108	0,7882	%9,122	%4,857	%5,600	%10,095	0,9515	0,8730	0,9108	0,7753	%8,913	%4,911	%5,704	%11,566
VARYANS	0,0017	0,0029	0,0020	0,0000	%0,178	%0,131	%0,086	%0,000	0,0020	0,0037	0,0021	0,0000	%0,184	%0,224	%0,091	%0,000

Çizelge 4.1'den $n = 50$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. Koyu olarak görülen hata yüzde değerleri, önerilen yaklaşımın diğer yaklaşımlara göre daha küçük bir hata ile ürün kitle parametresi olan uygun ürün oranı p 'yi tahmin ettiğini ifade etmektedir. Çizelgeden ayrıca örnekleme kümeleri üzerinden elde edilen yaklaşım değerlerinin ortalamaları karşılaştırıldığında, en iyi sonucun önerilen $\tilde{p}_{KMH}$ yaklaşımı ile elde edildiği, örnekleme kümeleri arasında oluşan varyans değerlerinin 0,000'a çok yakın olması nedeniyle tahmin ortalamaları üzerinden yaklaşımları karşılaştırmanın anlamlı olacağı söylenebilir. Çizelge 4.1'de görülen hata yüzde değerlerinin karşılaştırma kolaylığı sağlaması açısından çizgi grafiği çizdirilerek Şekil 4.1'de verilmiştir.



Şekil 4.1. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 50$, $\text{öks} = 10$)

Şekil 4.1 de verilen grafikler Senaryo I'e ait faz I verisi ile aynı dağılımlı örneklem kümelerinden elde edilen ϵY_{pk} , $\epsilon \tilde{p}_{KMH}$, $\epsilon \tilde{p}$ ve $\epsilon \hat{p}(H,S)$ hata yüzdeleri göstermektedir. Grafiklerden ilki daha önce de belirtildiği gibi dağılımı bilinmeyen çok değişkenli süreç verisinin dağılımının Referans Bant Genişliği (RBG) tahmin yöntemi yardımıyla Kernel yoğunluk tahmini yapılarak aynı dağılımdan çektilen 10 tane örneklem kümesinden, ikincisi ise Maksimum Düzleştirme Yöntemi (MDY) yardımıyla Kernel yoğunluk tahmini yapılarak aynı yoğunluk dağılımdan çektilen 10 tane örneklem kümesinden hesaplanmıştır. Grafikler incelendiğinde, her iki grafikte de $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdeleri diğer üç yöntemle göre daha iyi sonuç verdiği görülmektedir. Çizelge 4.1'de 10 örneklem kümesi üzerinden elde edilen ortalama ve varyans değerleri ile karşılaştırarak yaklaşımların örneklem kümesine karşı duyarlılığına karar vermek amacıyla ayrıca 20 örneklem kümesi ile de çalışılarak sonuçlar Çizelge 4.2'de verilmiştir.

Çizelge 4.2’de $n = 50$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri görülmektedir. Bu değerler, Çizelge 4.1’de verilen 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri ile karşılaştırıldığında, örneklem kümesi sayısındaki artış sonucu kümeler arasında oluşan ortalama ve varyans değerlerinde çok fazla bir fark olmadığı görülmektedir. Sonuçlar karşılaştırıldığında $n = 50$ gözlem sayısına sahip 10 örneklem kümesi ile 20 örneklem kümesinden elde edilen hata yüzdeleri için $\epsilon \tilde{p}_{KMH}$ değerlerinin diğer yaklaşımlara göre daha iyi sonuç verdiği görülmektedir.

$n = 100$ gözlem sayısına sahip örneklem kümeleri için elde edilen tahmin sonuçları Çizelge 4.3’de verilmiştir.

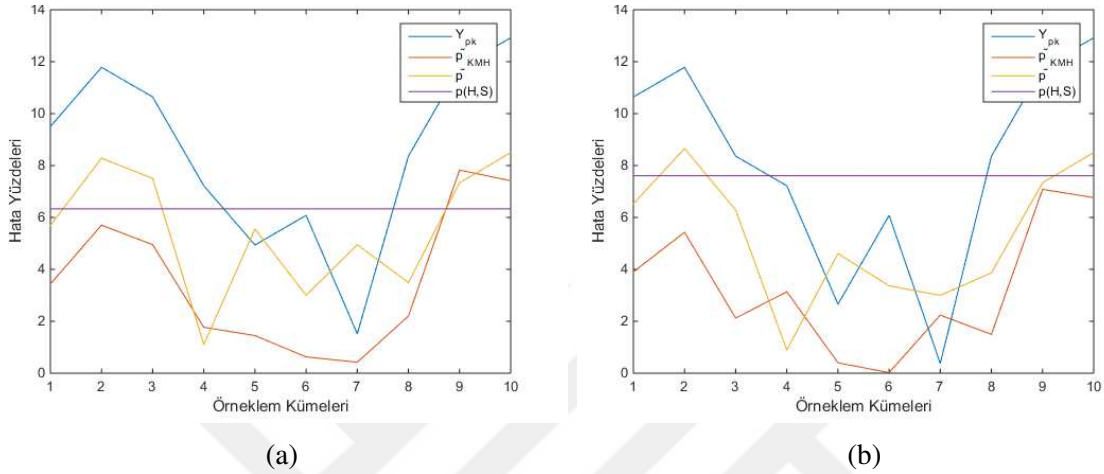
Çizelge 4.3. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 100$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
1	0,9600	0,9068	0,9266	0,8212	%9,502	%3,433	%5,692	%6,331	0,9700	0,9109	0,9339	0,8100	%10,642	%3,901	%6,524	%7,608
2	0,9800	0,9267	0,9493	0,8212	%11,783	%5,703	%8,281	%6,331	0,9800	0,9243	0,9526	0,8100	%11,783	%5,429	%8,657	%7,608
3	0,9700	0,9201	0,9425	0,8212	%10,642	%4,950	%7,505	%6,331	0,9500	0,8953	0,9318	0,8100	%8,361	%2,122	%6,285	%7,608
4	0,9400	0,8612	0,8864	0,8212	%7,220	%1,768	%1,106	%6,331	0,9400	0,8492	0,8845	0,8100	%7,220	%3,137	%0,890	%7,608
5	0,9200	0,8894	0,9254	0,8212	%4,939	%1,449	%5,555	%6,331	0,9000	0,8732	0,9171	0,8100	%2,658	%0,399	%4,608	%7,608
6	0,9300	0,8822	0,9030	0,8212	%6,080	%0,627	%3,000	%6,331	0,9300	0,8765	0,9062	0,8100	%6,080	%0,023	%3,365	%7,608
7	0,8900	0,8804	0,9201	0,8212	%1,517	%0,422	%4,950	%6,331	0,8800	0,8571	0,9030	0,8100	%0,376	%2,236	%3,000	%7,608
8	0,9500	0,8960	0,9073	0,8212	%8,361	%2,201	%3,490	%6,331	0,9500	0,8898	0,9106	0,8100	%8,361	%1,494	%3,867	%7,608
9	0,9800	0,9453	0,9410	0,8212	%11,783	%7,825	%7,334	%6,331	0,9800	0,9388	0,9410	0,8100	%11,783	%7,083	%7,334	%7,608
10	0,9900	0,9417	0,9512	0,8212	%12,923	%7,414	%8,498	%6,331	0,9900	0,9360	0,9512	0,8100	%12,923	%6,764	%8,498	%7,608
ORTALAMA	0,9510	0,9050	0,9253	0,8212	%8,475	%3,579	%5,541	%6,331	0,9470	0,8951	0,9232	0,8100	%8,019	%3,259	%5,303	%7,608
VARYANS	0,0010	0,0008	0,0005	0,0000	%0,129	%0,075	%0,060	%0,000	0,0013	0,0010	0,0005	0,0000	%0,168	%0,062	%0,066	%0,000

Çizelge 4.4. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 100$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
ORTALAMA	0,9525	0,9070	0,9256	0,8212	%8,646	%3,644	%5,581	%6,331	0,9505	0,8997	0,9256	0,8100	%8,418	%3,203	%5,582	%7,608
VARYANS	0,0006	0,0006	0,0003	0,0000	%0,082	%0,060	%0,036	%0,000	0,0008	0,0006	0,0003	0,0000	%0,102	%0,046	%0,037	%0,000

Çizelge 4.3’de $n = 100$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. Koyu renkle ifade edilen değerler, ürün kitle parametresi olan uygun ürün oranı p ’yi diğer yaklaşımlara göre daha iyi temsil eden yaklaşımı göstermektedir. $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelerinin diğer üç yöntemle göre daha iyi sonuç verdiği söylenebilir.



Şekil 4.2. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 100$, $\text{öks} = 10$)

Şekil 4.2’de verilen grafikler Senaryo I’e ait faz I verilerinden elde edilen ϵY_{pk} , $\epsilon \tilde{p}_{KMH}$, $\epsilon \tilde{p}$ ve $\epsilon \hat{p}(H,S)$ hata yüzdelerini göstermektedir. Grafiklerden ilkinde Referans Bant Genişliği tahmin yöntemi yardımıyla Kernel yoğunluk tahmini yapılarak aynı dağılımdan çektilen 10 tane örneklem kümesinden, ikincisinde ise Maksimum Düzleştirme Yöntemi yardımıyla Kernel yoğunluk tahmini yapılarak aynı dağılımdan çektilen 10 tane örneklem kümesinden hesaplanmıştır. Grafikler incelendiğinde, her iki grafikte de $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelerinin diğer üç yöntemle göre daha iyi sonuç verdiği görülmektedir. Her iki bant genişliği tahmin yöntemi kullanılarak elde edilen sonuçların benzer bir eğilim gösterdiği söylenebilir.

Çizelge 4.4’de $n = 100$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri görülmektedir. Bu değerler, Çizelge 4.3’de verilen 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri ile karşılaştırıldığında, örneklem küme sayısındaki artış sonucu kümeler arasında oluşan ortalama ve varyans değerlerinde çok fazla bir fark olmadığı söylenebilir. Çizelge 4.4’te verilen hata yüzdelerinin ortalama değerleri karşılaştırıldığında koyu olan değerlere

sahip yaklaşım olan $\tilde{p}_{KMH}$ yaklaşımının daha küçük bir hata ile p 'yi tahmin ettiği, yaklaşımlara göre daha iyi sonuçlar verdiği görülmektedir.

Aynı şekilde elde edilen $n = 250$ gözlem sayısına sahip örneklem kümeleri için sonuçlar Çizelge 4.5'de verilmiştir.



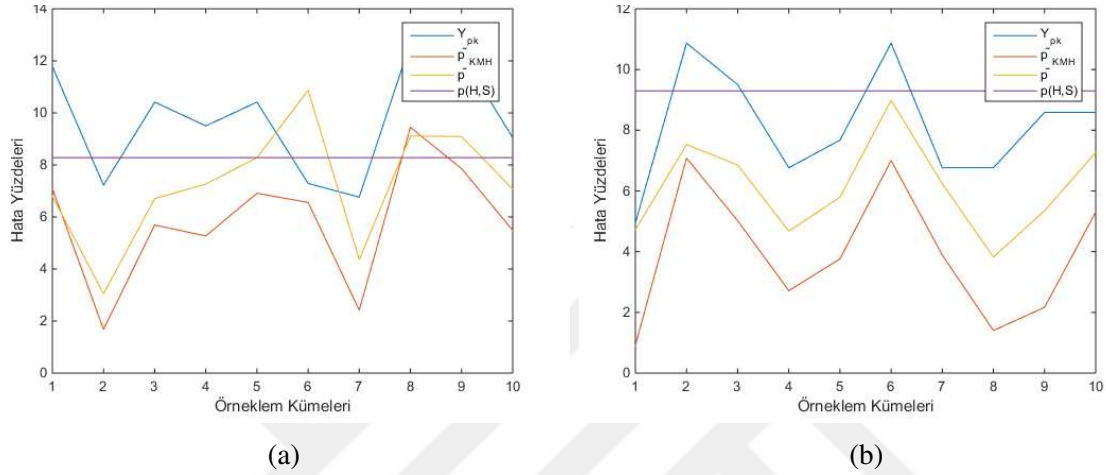
Çizelge 4.5. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 250$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
1	0,9800	0,9389	0,9363	0,8041	%11,783	%7,095	%6,798	%8,2810	0,9200	0,8847	0,9181	0,7952	%4,939	%0,913	%4,722	%9,2962
2	0,9400	0,8914	0,9034	0,8041	%7,220	%1,677	%3,046	%8,2810	0,9720	0,9388	0,9427	0,7952	%10,870	%7,083	%7,528	%9,2962
3	0,9680	0,9266	0,9355	0,8041	%10,414	%5,692	%6,707	%8,2810	0,9600	0,9207	0,9368	0,7952	%9,502	%5,019	%6,855	%9,2962
4	0,9600	0,9229	0,9404	0,8041	%9,502	%5,270	%7,266	%8,2810	0,9360	0,9005	0,9177	0,7952	%6,764	%2,715	%4,677	%9,2962
5	0,9680	0,9373	0,9493	0,8041	%10,414	%6,912	%8,281	%8,2810	0,9440	0,9097	0,9275	0,7952	%7,677	%3,764	%5,794	%9,2962
6	0,9406	0,9342	0,9720	0,8041	%7,289	%6,559	%10,870	%8,2810	0,9720	0,9381	0,9555	0,7952	%10,870	%7,004	%8,988	%9,2962
7	0,9360	0,8979	0,9150	0,8041	%6,764	%2,418	%4,369	%8,2810	0,9360	0,9108	0,9314	0,7952	%6,764	%3,890	%6,239	%9,2962
8	0,9880	0,9595	0,9566	0,8041	%12,695	%9,445	%9,114	%8,2810	0,9360	0,8890	0,9102	0,7952	%6,764	%1,403	%3,821	%9,2962
9	0,9840	0,9456	0,9564	0,8041	%12,239	%7,859	%9,091	%8,2810	0,9520	0,8957	0,9236	0,7952	%8,589	%2,167	%5,350	%9,2962
10	0,9560	0,9248	0,9387	0,8041	%9,045	%5,486	%7,072	%8,2810	0,9520	0,9232	0,9404	0,7952	%8,589	%5,304	%7,266	%9,2962
ORTALAMA	0,9621	0,9279	0,9404	0,8041	%9,737	%5,841	%7,261	%8,281	0,9480	0,9111	0,9304	0,7952	%8,133	%3,926	%6,124	%9,296
VARYANS	0,0004	0,0004	0,0004	0,0000	%0,046	%0,055	%0,053	%0,000	0,0003	0,0004	0,0002	0,0000	%0,037	%0,047	%0,025	%0,000

Çizelge 4.6. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 250$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
ORTALAMA	0,9524	0,9159	0,9286	0,8041	%8,635	%4,527	%5,922	%8,281	0,9483	0,9068	0,9294	0,7952	%8,167	%4,142	%6,011	%9,296
VARYANS	0,0004	0,0004	0,0003	0,0000	%0,047	%0,047	%0,033	%0,000	0,0005	0,0008	0,0003	0,0000	%0,063	%0,042	%0,034	%0,000

Çizelge 4.5'te $n=250$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelerinin diğer üç yöntemle göre daha iyi sonuç verdiği söylenebilir. Maksimum Bant Genişliği tahmin yöntemi ile elde edilen sonuçların Referans Bant Genişliği ile yapılan tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ürün kitle parametresi olan uygun ürün oranı p 'ye daha yakın sonuçlar verdiği görülmektedir.



Şekil 4.3. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 250$, $\text{öks} = 10$)

Şekil 4.3'de görülen grafikler Senaryo I'e ait faz I verilerinden elde edilen ϵY_{pk} , $\epsilon \tilde{p}_{KMH}$, $\epsilon \tilde{p}$ ve $\epsilon \hat{p}(H,S)$ hata yüzdelerini göstermektedir. Grafiklerden ilkinde Referans Bant Genişliği tahmin yöntemi yardımıyla, ikincisinde ise Maksimum Düzleştirme Yöntemi yardımıyla Kernel yoğunluk tahmini yapılarak aynı dağılımdan çektirilen 10 tane örneklem kümesinden hesaplanmıştır. Grafikler incelendiğinde, her iki grafikte de $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelerinin diğer üç yöntemle göre daha iyi sonuç verdiği görülmektedir. Her iki bant genişliği tahmin yöntemi kullanılarak elde edilen sonuçların benzer bir eğilim gösterdiği söylenebilir.

Çizelge 4.6'da $n = 250$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri görülmektedir. Bu değerler, Çizelge 4.3'de verilen 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri ile karşılaştırıldığında, örneklem küme sayısındaki artış sonucu kümeler arasında oluşan ortalama ve varyans değerlerinde çok fazla bir fark olmadığı söylenebilir. $n = 500$ gözlem sayısına sahip örneklem kümeleri için sonuçlar Çizelge 4.7'de verilmiştir.

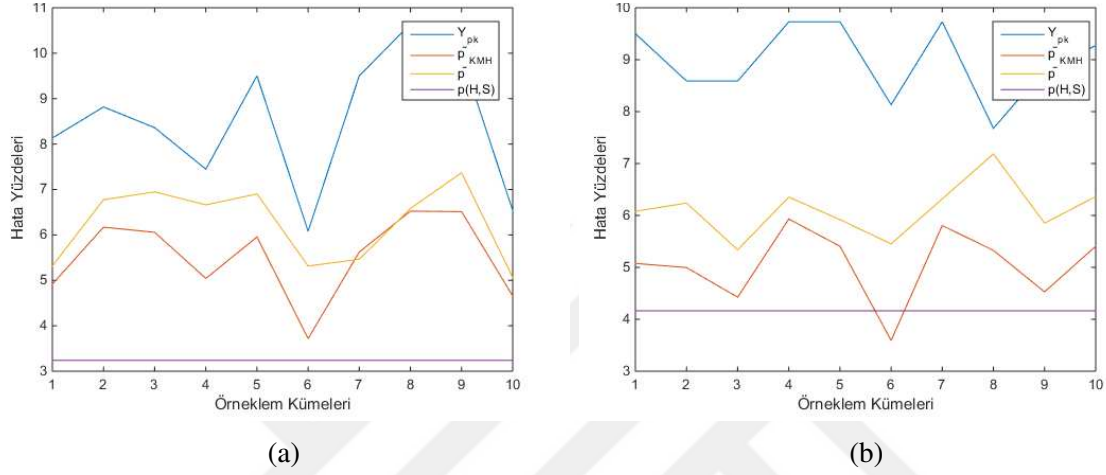
Çizelge 4.7. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 500$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
1	0,9480	0,9198	0,9233	0,8483	%8,133	%4,916	%5,315	%3,239	0,9600	0,9212	0,9300	0,8402	%9,502	%5,076	%6,080	%4,163
2	0,9540	0,9308	0,9361	0,8483	%8,817	%6,171	%6,775	%3,239	0,9520	0,9205	0,9314	0,8402	%8,589	%4,996	%6,239	%4,163
3	0,9500	0,9298	0,9376	0,8483	%8,361	%6,057	%6,947	%3,239	0,9520	0,9155	0,9235	0,8402	%8,589	%4,426	%5,338	%4,163
4	0,9420	0,9209	0,9351	0,8483	%7,448	%5,042	%6,661	%3,239	0,9620	0,9287	0,9324	0,8402	%9,730	%5,931	%6,353	%4,163
5	0,9600	0,9289	0,9372	0,8483	%9,502	%5,954	%6,901	%3,239	0,9620	0,9241	0,9286	0,8402	%9,730	%5,407	%5,920	%4,163
6	0,9300	0,9093	0,9233	0,8483	%6,080	%3,718	%5,315	%3,239	0,9480	0,9082	0,9245	0,8402	%8,133	%3,593	%5,452	%4,163
7	0,9600	0,9260	0,9246	0,8483	%9,502	%5,623	%5,464	%3,239	0,9620	0,9276	0,9321	0,8402	%9,730	%5,806	%6,319	%4,163
8	0,9700	0,9339	0,9344	0,8483	%10,642	%6,524	%6,581	%3,239	0,9440	0,9234	0,9397	0,8402	%7,677	%5,327	%7,186	%4,163
9	0,9640	0,9338	0,9413	0,8483	%9,958	%6,513	%7,369	%3,239	0,9540	0,9164	0,9280	0,8402	%8,817	%4,528	%5,851	%4,163
10	0,9340	0,9175	0,9211	0,8483	%6,536	%4,654	%5,064	%3,239	0,9580	0,9241	0,9325	0,8402	%9,273	%5,407	%6,365	%4,163
ORTALAMA	0,9512	0,9251	0,9314	0,8483	%8,498	%5,517	%6,239	%3,239	0,9554	0,9210	0,9303	0,8402	%8,977	%5,050	%6,110	%4,163
VARYANS	0,0002	0,0001	0,0001	0,0000	%0,022	%0,008	%0,007	%0,000	0,0000	0,0000	0,0000	0,0000	%0,005	%0,005	%0,003	%0,000

Çizelge 4.8. Senaryo I için tahmin değerleri ve hata yüzdelерinin ortalama ve varyansları ($n = 500$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
ORTALAMA	0,9537	0,9249	0,9309	0,8483	%8,783	%5,502	%6,181	%3,239	0,9535	0,9201	0,9308	0,8402	%8,760	%4,948	%6,176	%4,163
VARYANS	0,0001	0,0001	0,0001	0,0000	%0,010	%0,014	%0,012	%0,000	0,0001	0,0001	0,0001	0,0000	%0,009	%0,016	0,0001	%0,000

Çizelge 4.7’de $n = 500$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. $\hat{p}(H,S)$ yaklaşımı ile elde edilen hata yüzdelerinin diğer üç yöntemle göre daha iyi sonuç verdiği söylenebilir. Referans Bant Genişliği ile yapılan sonuçların Maksimum Düzleştirme Yöntemi ile elde edilen tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ürün kitle parametresi olan uygun ürün oranı p ’ye daha yakın sonuçlar verdiği görülmektedir.



Şekil 4.4. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 500$, $\tilde{\sigma}_{ks} = 10$)

Şekil 4.4’te görülen grafikler Senaryo I’e ait faz I verilerinden elde edilen ϵY_{pk} , $\epsilon \tilde{p}_{KMH}$, $\epsilon \tilde{p}$ ve $\epsilon \hat{p}(H,S)$ hata yüzdelerini göstermektedir. İlk Grafik, Referans Bant Genişliği tahmin yöntemi yardımıyla Kernel yoğunluk tahmini yapılarak aynı dağılımdan çektilen 10 tane örneklem kümesi üzerinden elde edilen hata yüzdelerini göstermektedir. İlk grafikte $\hat{p}(H,S)$ yaklaşımı ile p ’nin daha iyi tahmin edildiği görülmektedir. İkinci grafikte ise Maksimum Düzleştirme Yöntemi yardımıyla Kernel yoğunluk tahmini yapılarak aynı dağılımdan çektilen 10 tane örneklem kümesinden hesaplanmıştır. İkinci grafik incelendiğinde de, $\hat{p}(H,S)$ yaklaşımının başarılı olduğu görülmektedir. Ayrıca her iki grafikten de önerilen $\tilde{p}_{KMH}$ yaklaşımı sonuçları ile $\hat{p}(H,S)$ yaklaşımı sonuçlarının genellikle diğer iki yaklaşıma göre birbirine yakın sonuçlar verdiği açıkça görülmektedir.

Çizelge 4.8’de $n = 500$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri görülmektedir. Bu değerler, Çizelge 4.7’de verilen 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri ile karşılaştırıldığında, 10 örneklem kümesi için hesaplanan hata yüzdelerinde olduğu gibi

20 örneklem kümesinden hesaplananlar için de geçerli olduğu, diğer bir ifade ile $\hat{p}(H, S)$ yaklaşımının daha iyi sonuç verdiği söylenebilir.

$n = 1000$ gözlem sayısına sahip örneklem kümeleri için sonuçlar Çizelge 4.9'da verilmiştir.



Çizelge 4.9. Senaryo I için tahmin değerleri ve hata yüzdeleri ($n = 1000$, $\ddot{o}ks = 10$)

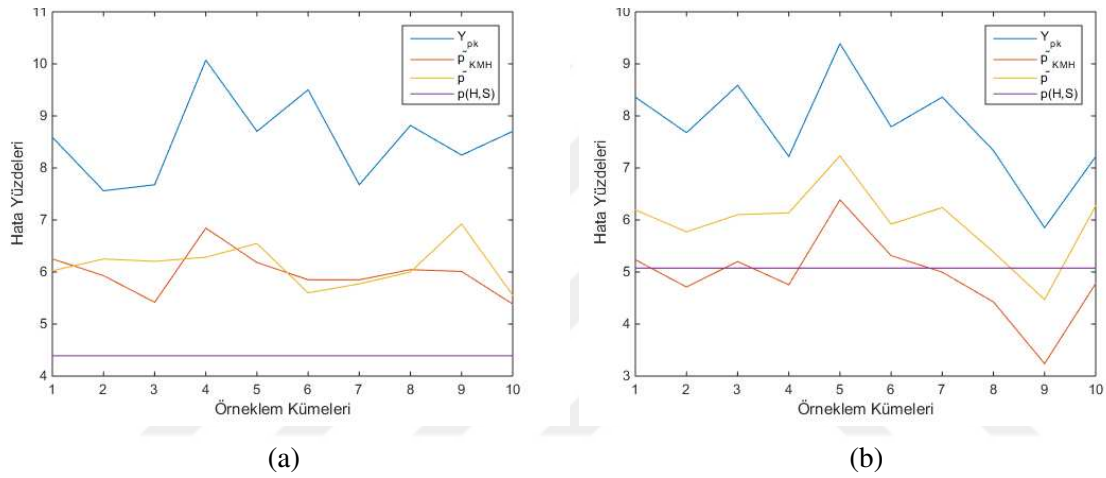
ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
1	0,9520	0,9315	0,9295	0,8382	%8,589	%6,251	%6,023	%4,391	0,9500	0,9226	0,9310	0,8322	%8,361	%5,236	%6,194	%5,075
2	0,9430	0,9287	0,9315	0,8382	%7,562	%5,931	%6,251	%4,391	0,9440	0,9180	0,9273	0,8322	%7,677	%4,711	%5,772	%5,075
3	0,9440	0,9242	0,9311	0,8382	%7,677	%5,418	%6,205	%4,391	0,9520	0,9223	0,9302	0,8322	%8,589	%5,201	%6,102	%5,075
4	0,9650	0,9367	0,9318	0,8382	%10,072	%6,844	%6,285	%4,391	0,9400	0,9184	0,9305	0,8322	%7,220	%4,756	%6,137	%5,075
5	0,9530	0,9309	0,9341	0,8382	%8,703	%6,182	%6,547	%4,391	0,9500	0,9327	0,9401	0,8322	%9,387	%6,388	%7,232	%5,075
6	0,9600	0,9280	0,9258	0,8382	%9,502	%5,851	%5,601	%4,391	0,9450	0,9233	0,9286	0,8322	%7,791	%5,315	%5,920	%5,075
7	0,9440	0,9280	0,9273	0,8382	%7,677	%5,851	%5,772	%4,391	0,9500	0,9205	0,9314	0,8322	%8,361	%4,996	%6,239	%5,075
8	0,9540	0,9297	0,9293	0,8382	%8,817	%6,045	%6,000	%4,391	0,9410	0,9155	0,9239	0,8322	%7,334	%4,426	%5,384	%5,075
9	0,9490	0,9294	0,9374	0,8382	%8,247	%6,011	%6,924	%4,391	0,9280	0,9051	0,9159	0,8322	%5,851	%3,239	%4,471	%5,075
10	0,9530	0,9239	0,9254	0,8382	%8,703	%5,384	%5,555	%4,391	0,9400	0,9186	0,9317	0,8322	%7,220	%4,779	%6,274	%5,075
ORTALAMA	0,9517	0,9291	0,9303	0,8382	%8,555	%5,977	%6,116	%4,391	0,9449	0,9197	0,9291	0,8322	%7,779	%4,905	%5,973	%5,075
VARYANS	0,0001	0,0000	0,0000	0,0000	%0,007	%0,002	%0,002	%0,000	0,0001	0,0000	0,0000	0,0000	%0,009	%0,006	%0,005	%0,000

Çizelge 4.10. Senaryo I için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 1000$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
ORTALAMA	0,9530	0,9319	0,9333	0,8382	%8,703	%6,297	%6,458	%4,391	0,9527	0,9254	0,9302	%7,779	%8,663	%5,555	%6,098	%5,075
VARYANS	0,0001	0,0001	0,0001	0,0000	%0,014	%0,013	%0,007	%0,000	0,0000	0,0000	0,0000	%0,009	%0,005	%0,003	%0,003	%0,000

Çizelge 4.9’da $n = 1000$ gözlemlili örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. Referans Bant Genişliği ile yapılan tahmin değerleri için $\hat{p}(H,S)$ yaklaşımının diğer yaklaşımlara göre daha iyi sonuç verdiği görülmektedir. Maksimum Düzleştirme Yöntemi ile bant genişliği ile yapılan tahminler sonucu elde edilen örnekleme kümeleri üzerinden hesaplanan hata yüzdeleri için ise $\tilde{p}_{KMH}$ yaklaşımının daha iyi sonuç verdiği söylenebilir.

Yüzde değerleri göre yöntemleri karşılaştırmak amacıyla çizdirilen çizgi grafikleri Şekil 4.5’te verilmiştir.



Şekil 4.5. Senaryo I için hata yüzdeleri çizgi grafiği ($n = 1000$, $\text{öks} = 10$)

Şekil 4.5’te verilen grafikler incelendiğinde ilk grafikteki sonuçlardan $\tilde{p}_{KMH}$, $\tilde{p}$ yaklaşımlarının hata yüzdeleri açısından benzer sonuçlar verdiği, diğer yaklaşımlara göre p ’yi daha iyi tahmin ettiği görülmektedir. Örneklem kümelerinin $n = 1000$ gözlem içerdiği göz önünde bulundurulduğunda dağılımın Normal dağıldığı varsayımı ile Referans Bant Genişliği tahmin yöntemi kullanıldığından ilk grafikte, $\tilde{p}$ ’nın p ’yi tahmin etmede başarılı olması ve gözlem sayısı büyüdükçe dağılımın normal dağılıma yakınsayarak klasik tahmin edici $\tilde{p}$ ’nin daha iyi sonuçlar vermesi olasıdır. Öyle ki ikinci grafik incelendiğinde Maksimum Düzleştirme Yöntemi ile tahmin edilen kernel yoğunluk tahminleri üzerinden çektirilen örneklem kümelerine $\tilde{p}_{KMH}$ yaklaşımı uygulandığında p ’yi daha iyi tahmin ettiği görülmektedir.

Çizelge 4.10’dan $n = 1000$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri

görülmektedir. 10 örneklem kümesi için elde edilen sonuçlardan farklı olarak her iki bant genişliği tahmin yönteminde de $\hat{p}(H,S)$ yaklaşımının daha iyi sonuçlar verdiği görülmektedir.

Örnekleme kümesi sayısı (öks) ve her bir örnekleme kümesindeki gözlem sayısı büyüdükçe kümeler arasında varyans değerinin küçüldüğü görülmektedir.

4.2. Senaryo II İçin Elde Edilen Sonuçlar

Senaryo II için Kernel yoğunluk tahmin yöntemine dayalı *MH* örnekleme ile elde edilen $n = 50, 100, 250, 500$ ve 1000 gözlemlili örneklem kümelerinden hesaplanan uygun ürün olasılıkları (Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$, $\hat{p}(H,S)$) kitle uygun ürün oranı p ile karşılaştırılarak, Hata yüzdeleri elde edilmiştir. Hesaplanan tahmin sonuçları $n = 50$ için Çizelge 4.11’de verilmiştir.

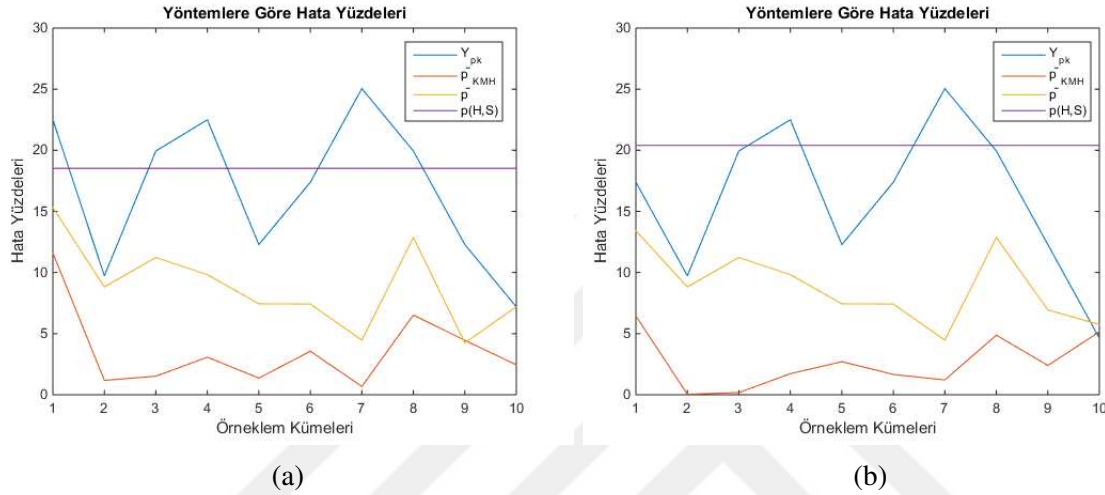
Çizelge 4.11. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 50$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
1	0,9600	0,8746	0,9037	0,6385	%22,496	%11,599	%15,312	%18,527	0,9200	0,8342	0,8889	0,6238	%17,392	%6,444	%13,424	%20,403
2	0,8600	0,7929	0,8528	0,6385	%9,736	%1,174	%8,817	%18,527	0,8600	0,7840	0,8528	0,6238	%9,736	%0,038	%8,817	%20,403
3	0,9400	0,7956	0,8717	0,6385	%19,944	%1,518	%11,229	%18,527	0,9400	0,7851	0,8717	0,6238	%19,944	%0,179	%11,229	%20,403
4	0,9600	0,8078	0,8607	0,6385	%22,496	%3,075	%9,825	%18,527	0,9600	0,7972	0,8607	0,6238	%22,496	%1,723	%9,825	%20,403
5	0,8800	0,7730	0,8420	0,6385	%12,288	%1,365	%7,439	%18,527	0,8800	0,7625	0,8420	0,6238	%12,288	%2,705	%7,439	%20,403
6	0,9200	0,8116	0,8419	0,6385	%17,392	%3,560	%7,426	%18,527	0,9200	0,7967	0,8419	0,6238	%17,392	%1,659	%7,426	%20,403
7	0,9800	0,7890	0,8187	0,6385	%25,048	%0,676	%4,466	%18,527	0,9800	0,7742	0,8187	0,6238	%25,048	%1,212	%4,466	%20,403
8	0,9400	0,8348	0,8846	0,6385	%19,944	%6,520	%12,875	%18,527	0,9400	0,8219	0,8846	0,6238	%19,944	%4,874	%12,875	%20,403
9	0,8800	0,7488	0,8170	0,6385	%12,288	%4,453	%4,249	%18,527	0,8800	0,7649	0,8381	0,6238	%12,288	%2,399	%6,941	%20,403
10	0,8400	0,7645	0,8401	0,6385	%7,184	%2,450	%7,197	%18,527	0,8200	0,7434	0,8287	0,6238	%4,632	%5,142	%5,742	%20,403
ORTALAMA	0,9160	0,7993	0,8533	0,6385	%16,881	%3,639	%8,884	%18,527	0,9100	0,7864	0,8528	0,6238	%16,116	%2,637	%8,818	%20,403
VARYANS	0,0023	0,0013	0,0008	0,0000	%0,373	%0,110	%0,124	%0,000	0,0024	0,0008	0,0005	0,0000	%0,394	%0,047	%0,089	%0,000

Çizelge 4.12. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 50$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
ORTALAMA	0,8980	0,7932	0,8521	0,6385	%16,465	%5,789	%9,588	%18,527	0,8960	0,7805	0,8522	0,6238	%15,955	%4,989	%9,530	%20,403
VARYANS	0,0068	0,0036	0,0022	0,0000	%0,488	%0,252	%0,194	%0,000	0,0062	0,0031	0,0024	0,0000	%0,487	%0,244	%0,238	%0,000

Çizelge 4.11’de, $n = 50$ gözlem sayısına sahip 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \varepsilon$) ortalama ile varyans değerleri görülmektedir. Ürün kitle parametresi p ’nin önerilen yaklaşım olan $\tilde{p}_{KMH}$ ile diğer üç yönteme göre daha iyi tahmin edildiği açıkça görülmektedir. $\tilde{p}_{KMH}$ yaklaşımının diğer yaklaşımlara göre p ’yi tahmin etmedeki başarısı, Şekil 4.6’da verilen çizgi grafiklerinden daha iyi anlaşılabilir.



Şekil 4.6. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 50$, $\text{öks} = 10$)

Şekil 4.6’da görülen her iki grafikte de $\tilde{p}_{KMH}$ yaklaşımının iyi sonuçlar verdiği ve hata yüzde seviyesinin diğerlerine göre çok aşağıda olduğu görülmektedir.

Çizelge 4.12’de $n = 50$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \varepsilon$) ortalama ile varyans değerleri görülmektedir. Bu değerler, Çizelge 4.11’de verilen 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \varepsilon$) ortalama ile varyans değerleri ile karşılaştırıldığında, örneklem küme sayısındaki artış sonucu kümeler arasında oluşan ortalama ve varyans değerlerinde çok fazla bir fark olmadığı söylenebilir.

$n = 100$ gözlem sayısına sahip örneklem kümeleri için sonuçlar Çizelge 4.13’da verilmiştir.

Çizelge 4.13. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 100$, $\ddot{o}ks = 10$)

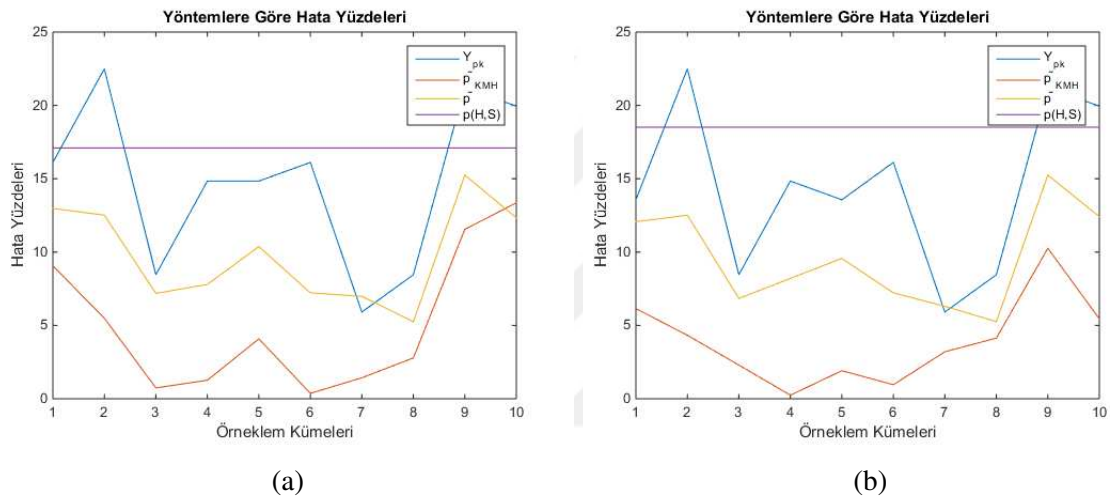
ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
1	0,9100	0,8547	0,8854	0,6497	%16,116	%9,060	%12,977	%17,098	0,8900	0,8319	0,8783	0,6386	%13,564	%6,150	%12,071	%18,515
2	0,9600	0,8267	0,8818	0,6497	%22,496	%5,487	%12,518	%17,098	0,9600	0,8176	0,8818	0,6386	%22,496	%4,326	%12,518	%18,515
3	0,8500	0,7779	0,8400	0,6497	%8,460	%0,740	%7,184	%17,098	0,8500	0,7657	0,8373	0,6386	%8,460	%2,297	%6,839	%18,515
4	0,9000	0,7936	0,8448	0,6497	%14,840	%1,263	%7,796	%17,098	0,9000	0,7856	0,8480	0,6386	%14,840	%0,242	%8,205	%18,515
5	0,9000	0,8157	0,8650	0,6497	%14,840	%4,083	%10,374	%17,098	0,8900	0,7987	0,8587	0,6386	%13,564	%1,914	%9,570	%18,515
6	0,9100	0,7866	0,8403	0,6497	%16,116	%0,370	%7,222	%17,098	0,9100	0,7762	0,8403	0,6386	%16,116	%0,957	%7,222	%18,515
7	0,8300	0,7725	0,8384	0,6497	%5,908	%1,429	%6,980	%17,098	0,8300	0,7586	0,8331	0,6386	%5,908	%3,203	%6,303	%18,515
8	0,8500	0,7619	0,8249	0,6497	%8,460	%2,782	%5,257	%17,098	0,8500	0,7513	0,8249	0,6386	%8,460	%4,134	%5,257	%18,515
9	0,9500	0,8742	0,9032	0,6497	%21,220	%11,548	%15,248	%17,098	0,9500	0,8641	0,9032	0,6386	%21,220	%10,259	%15,248	%18,515
10	0,9400	0,8884	0,8804	0,6497	%19,944	%13,360	%12,339	%17,098	0,9400	0,8264	0,8810	0,6386	%19,944	%5,449	%12,415	%18,515
ORTALAMA	0,9000	0,8152	0,8604	0,6497	%14,840	%5,012	%9,789	%17,098	0,8970	0,7976	0,8587	0,6386	%14,457	%3,893	%9,565	%18,515
VARYANS	0,0020	0,0020	0,0007	0,0000	%0,322	%0,224	%0,111	%0,000	0,0020	0,0013	0,0007	0,0000	%0,322	%0,086	%0,110	%0,000

Çizelge 4.14. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 100$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
ORTALAMA	0,9095	0,8256	0,8678	0,6497	%16,482	%6,517	%10,733	%17,098	0,9100	0,8122	0,8728	0,6386	%16,116	%5,076	%11,364	%18,515
VARYANS	0,0031	0,0018	0,0008	0,0000	%0,351	%0,153	%0,122	%0,000	0,0023	0,0013	0,0013	0,0000	%0,367	%0,087	%0,205	%0,000

Çizelge 4.13'te $n = 100$ gözlem sayısına sahip 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \varepsilon$) ortalama ile varyans değerleri görülmektedir. Ürün kitle parametresi p 'nin önerilen yaklaşım olan $\tilde{p}_{KMH}$ ile diğer üç yöntemle göre daha iyi tahmin edildiği açıkça görülmektedir. Öyle ki, diğer üç yöntemin hata seviyeleri $\%20$ 'lere yaklaşırken, $\tilde{p}_{KMH}$ 'nın p 'yi tahmin etmede ki hata yüzdeleri ortalama $\%3$ 'lerde seyretmektedir.

$\tilde{p}_{KMH}$ yaklaşımının diğer yaklaşımlara göre p 'yi tahmin etmedeki başarısı, Şekil 4.7'de verilen çizgi grafiklerinden daha iyi anlaşılabilir.



Şekil 4.7. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 100$, $\text{öks} = 10$)

Şekil 4.7'de her iki bant genişliği tahmin yöntemi için $\tilde{p}_{KMH}$ yaklaşımının iyi sonuçlar verdiği görülmektedir.

Çizelge 4.14'te $n = 100$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin $\tilde{p}_{KMH}$ yaklaşımı için iyi sonuçlar verdiği ve her iki bant genişliği tahmin yöntemi için de diğer yaklaşımlardan daha küçük hata seviyesi ile p 'yi tahmin ettiği görülmektedir.

$n = 250$ gözlem sayısına sahip örneklem kümeleri için sonuçlar Çizelge 4.15'de verilmiştir.

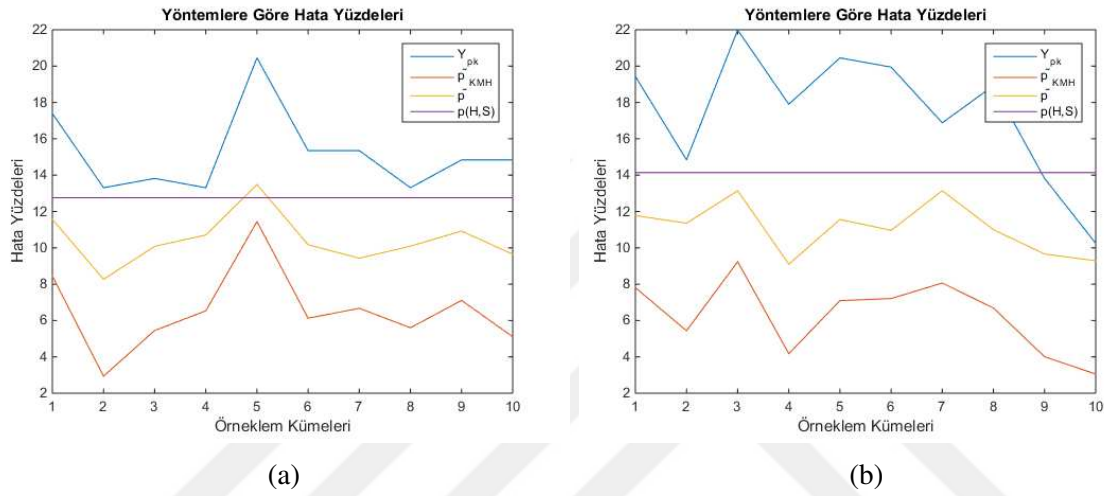
Çizelge 4.15. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 250$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
1	0,9200	0,8499	0,8744	0,6837	%17,392	%8,447	%11,573	%12,759	0,9360	0,8449	0,8760	0,6729	%19,433	%7,809	%11,777	%14,138
2	0,8880	0,8067	0,8485	0,6837	%13,309	%2,935	%8,268	%12,759	0,9000	0,8263	0,8727	0,6729	%14,840	%5,436	%11,356	%14,138
3	0,8920	0,8264	0,8627	0,6837	%13,819	%5,449	%10,080	%12,759	0,9560	0,8561	0,8867	0,6729	%21,985	%9,238	%13,143	%14,138
4	0,8880	0,8349	0,8676	0,6837	%13,309	%6,533	%10,706	%12,759	0,9240	0,8164	0,8550	0,6729	%17,902	%4,173	%9,098	%14,138
5	0,9440	0,8734	0,8894	0,6837	%20,454	%11,446	%13,487	%12,759	0,9440	0,8393	0,8743	0,6729	%20,454	%7,095	%11,561	%14,138
6	0,9040	0,8317	0,8634	0,6837	%15,350	%6,125	%10,170	%12,759	0,9400	0,8402	0,8696	0,6729	%19,944	%7,209	%10,961	%14,138
7	0,9040	0,8360	0,8576	0,6837	%15,350	%6,673	%9,430	%12,759	0,9160	0,8469	0,8867	0,6729	%16,881	%8,064	%13,143	%14,138
8	0,8880	0,8276	0,8628	0,6837	%13,309	%5,602	%10,093	%12,759	0,9320	0,8361	0,8699	0,6729	%18,923	%6,686	%10,999	%14,138
9	0,9000	0,8394	0,8693	0,6837	%14,840	%7,107	%10,923	%12,759	0,8920	0,8151	0,8594	0,6729	%13,819	%4,007	%9,659	%14,138
10	0,9000	0,8237	0,8593	0,6837	%14,840	%5,104	%9,647	%12,759	0,8640	0,8076	0,8565	0,6729	%10,246	%3,050	%9,289	%14,138
ORTALAMA	0,9028	0,8350	0,8655	0,6837	%15,197	%6,542	%10,438	%12,759	0,9204	0,8329	0,8707	0,6729	%17,443	%6,277	%11,099	%14,138
VARYANS	0,0003	0,0003	0,0001	0,0000	%0,050	%0,050	%0,020	%0,000	0,0008	0,0003	0,0001	0,0000	%0,127	%0,041	%0,021	%0,000

Çizelge 4.16. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 250$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
ORTALAMA	0,9120	0,8345	0,8663	0,6837	%16,371	%6,482	%10,539	%12,759	0,9104	0,8262	0,8666	0,6729	%16,167	%5,441	%10,574	%14,138
VARYANS	0,0006	0,0005	0,0003	0,0000	%0,094	%0,083	%0,046	%0,000	0,0006	0,0005	0,0003	0,0000	%0,094	%0,076	%0,042	%0,000

Çizelge 4.15'te $n = 250$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. Maksimum Düzleştirme Yöntemi ile elde edilen sonuçların Referans Bant Genişliği ile yapılan tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ürün kitle parametresi olan uygun ürün oranı p 'ye daha yakın sonuçlar verdiği görülmektedir. Ayrıca her iki tahmin yönteminde de $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelерinin diğer üç yönteme göre daha iyi sonuç verdiği söylenebilir.



Şekil 4.8. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 250$, $\text{öks} = 10$)

Şekil 4.8'de görülen grafikler Senaryo II'ye ait faz I verilerinden elde edilen ϵY_{pk} , $\epsilon \tilde{p}_{KMH}$, $\epsilon \tilde{p}$ ve $\epsilon \hat{p}(H,S)$ hata yüzdelерini göstermektedir. Grafikler incelendiğinde, her iki grafikte de $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelерinin diğer üç yönteme göre daha iyi sonuç verdiği görülmektedir. Her iki bant genişliği tahmin yöntemi kullanılarak elde edilen sonuçların benzer bir eğilim gösterdiği söylenebilir.

Çizelge 4.16'da Maksimum Düzleştirme Yöntemi ile elde edilen sonuçların Referans Bant Genişliği ile yapılan tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ve 10 örneklem kümesi için hesaplanan hata yüzderine benzer olarak ürün kitle parametresi olan uygun ürün oranı p 'ye daha yakın sonuçlar verdiği görülmektedir.

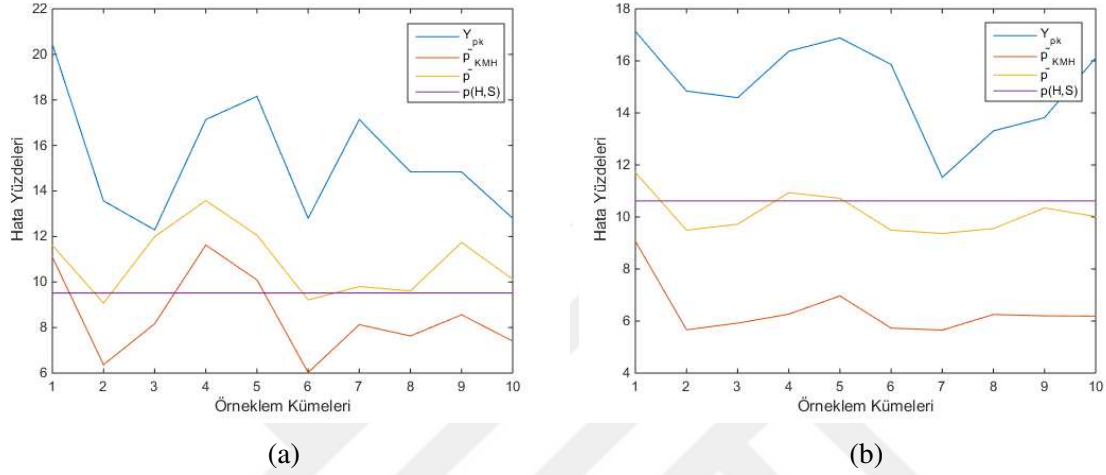
Çizelge 4.17. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 500$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
1	0,9440	0,8707	0,8747	0,7091	%20,454	%11,101	%11,612	%9,518	0,9180	0,8548	0,8755	0,7005	%17,137	%9,072	%11,714	%10,616
2	0,8900	0,8336	0,8548	0,7091	%13,564	%6,367	%9,072	%9,518	0,9000	0,8281	0,8581	0,7005	%14,840	%5,665	%9,493	%10,616
3	0,8800	0,8477	0,8777	0,7091	%12,288	%8,166	%11,994	%9,518	0,8980	0,8301	0,8599	0,7005	%14,585	%5,921	%9,723	%10,616
4	0,9180	0,8748	0,8901	0,7091	%17,137	%11,624	%13,577	%9,518	0,9120	0,8328	0,8694	0,7005	%16,371	%6,265	%10,935	%10,616
5	0,9260	0,8628	0,8782	0,7091	%18,157	%10,093	%12,058	%9,518	0,9160	0,8383	0,8677	0,7005	%16,881	%6,967	%10,718	%10,616
6	0,8840	0,8309	0,8559	0,7091	%12,798	%6,023	%9,213	%9,518	0,9080	0,8286	0,8581	0,7005	%15,861	%5,729	%9,493	%10,616
7	0,9180	0,8474	0,8605	0,7091	%17,137	%8,128	%9,800	%9,518	0,8740	0,8280	0,8571	0,7005	%11,522	%5,653	%9,366	%10,616
8	0,9000	0,8435	0,8590	0,7091	%14,840	%7,630	%9,608	%9,518	0,8880	0,8327	0,8586	0,7005	%13,309	%6,252	%9,557	%10,616
9	0,9000	0,8508	0,8757	0,7091	%14,840	%8,562	%11,739	%9,518	0,8920	0,8323	0,8648	0,7005	%13,819	%6,201	%10,348	%10,616
10	0,8840	0,8417	0,8630	0,7091	%12,798	%7,401	%10,119	%9,518	0,9100	0,8322	0,8622	0,7005	%16,116	%6,189	%10,017	%10,616
ORTALAMA	0,9044	0,8504	0,8690	0,7091	%15,401	%8,510	%10,879	%9,518	0,9016	0,8338	0,8631	0,7005	%15,044	%6,391	%10,136	%10,616
VARYANS	0,0005	0,0002	0,0001	0,0000	%0,074	%0,036	%0,023	%0,000	0,0002	0,0001	0,0000	0,0000	%0,032	%0,010	%0,006	%0,000

Çizelge 4.18. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 500$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
ORTALAMA	0,9090	0,8455	0,8675	0,7091	%15,988	%7,891	%10,687	%9,518	0,9072	0,8380	0,8675	0,7005	%15,759	%6,931	%10,692	%10,616
VARYANS	0,0003	0,0001	0,0001	0,0000	%0,043	%0,016	%0,012	%0,000	0,0003	0,0001	0,0001	0,0000	%0,046	%0,020	%0,014	%0,000

Çizelge 4.17’den $n = 500$ gözlem sayısına sahip 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) görülmektedir. burada her ne kadar örneklem kümelerinde $\tilde{p}_{KMH}$ yaklaşımı ile birlikte $\hat{p}(H,S)$ yaklaşımının da iyi sonuçlar verdiği gözlemlenmekle birlikte örneklem kümeleri ortalamaları incelendiğinde $\tilde{p}_{KMH}$ yaklaşımının her iki bant genişliği tahmin yöntemi için de diğer yaklaşımlardan iyi sonuç verdiği söylenebilir.



Şekil 4.9. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 500$, $\text{öks} = 10$)

Şekil 4.9’da yer alan çizgi grafiklerinde Çizelge 4.17’deki sonuçlar özetlenmiştir. Her iki tahmin yöntemi için de $\tilde{p}_{KMH}$ yaklaşımının iyi sonuçlar verdiği $n = 500$ gözlem sayısı için örneklem kümelerinden elde edilen hata yüzde değerleri arasındaki varyansın azaldığı özellikle Maksimum Düzleştirme Yöntemi ile elde edilen tahmin değerlerini gösteren ikinci grafikte açıkça oluşan düzgün eğilimden görülmektedir.

Çizelge 4.18’de $\tilde{p}_{KMH}$ yaklaşımına göre Maksimum Düzleştirme Yöntemi ile elde edilen sonuçların Referans Bant Genişliği ile yapılan tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ve 10 örneklem kümesi için hesaplanan hata yüzderine benzer olarak ürün kitle parametresi olan uygun ürün oranı p ’ye daha yakın sonuçlar verdiği görülmektedir.

$n = 1000$ gözlem sayısına sahip örneklem kümeleri için sonuçlar Çizelge 4.19’da verilmiştir.

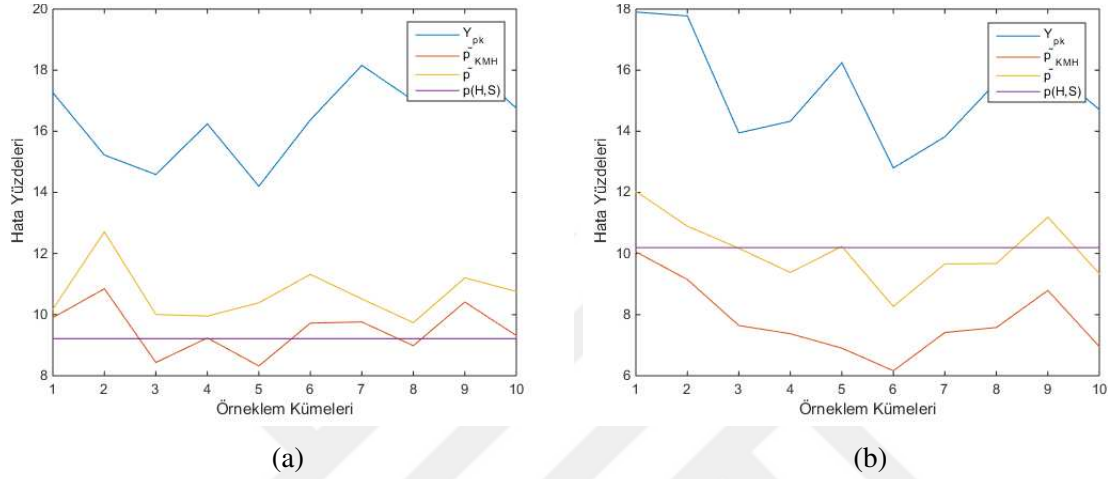
Çizelge 4.19. Senaryo II için tahmin değerleri ve hata yüzdeleri ($n = 1000$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
1	0,9190	0,8613	0,8635	0,7115	%17,264	%9,902	%10,182	%9,213	0,9240	0,8625	0,8780	0,7038	%17,902	%10,055	%12,033	%10,195
2	0,9030	0,8687	0,8833	0,7115	%15,223	%10,846	%12,709	%9,213	0,9230	0,8554	0,8691	0,7038	%17,775	%9,149	%10,897	%10,195
3	0,8980	0,8498	0,8621	0,7115	%14,585	%8,434	%10,004	%9,213	0,8930	0,8436	0,8634	0,7038	%13,947	%7,643	%10,170	%10,195
4	0,9110	0,8561	0,8617	0,7115	%16,243	%9,238	%9,953	%9,213	0,8960	0,8415	0,8572	0,7038	%14,329	%7,375	%9,379	%10,195
5	0,8950	0,8489	0,8651	0,7115	%14,202	%8,320	%10,387	%9,213	0,9110	0,8378	0,8639	0,7038	%16,243	%6,903	%10,234	%10,195
6	0,9120	0,8599	0,8724	0,7115	%16,371	%9,723	%11,318	%9,213	0,8840	0,8320	0,8485	0,7038	%12,798	%6,163	%8,268	%10,195
7	0,9260	0,8602	0,8661	0,7115	%18,157	%9,761	%10,514	%9,213	0,8920	0,8418	0,8594	0,7038	%13,819	%7,414	%9,659	%10,195
8	0,9170	0,8541	0,8600	0,7115	%17,009	%8,983	%9,736	%9,213	0,9060	0,8431	0,8595	0,7038	%15,605	%7,579	%9,672	%10,195
9	0,9270	0,8653	0,8715	0,7115	%18,285	%10,412	%11,203	%9,213	0,9100	0,8526	0,8714	0,7038	%16,116	%8,792	%11,191	%10,195
10	0,9150	0,8567	0,8680	0,7115	%16,754	%9,315	%10,757	%9,213	0,8990	0,8382	0,8569	0,7038	%14,712	%6,954	%9,340	%10,195
ORTALAMA	0,9123	0,8581	0,8674	0,7115	%16,409	%9,493	%10,676	%9,213	0,9038	0,8449	0,8627	0,7038	%15,325	%7,803	%10,084	%10,195
VARYANS	0,0001	0,0000	0,0000	0,0000	%0,019	%0,006	%0,008	%0,000	0,0002	0,0001	0,0001	0,0000	%0,029	%0,014	%0,012	%0,000

Çizelge 4.20. Senaryo II için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 1000$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
ORTALAMA	0,9071	0,8543	0,8661	0,7115	%15,746	%9,014	%10,520	%9,213	0,9065	0,8489	0,8672	0,7038	%15,669	%8,319	%10,656	%10,195
VARYANS	0,0001	0,0001	0,0000	0,0000	%0,017	%0,016	%0,008	%0,000	0,0001	0,0001	0,0001	0,0000	%0,016	%0,016	%0,009	%0,000

Çizelge 4.19’da $n = 1000$ ve 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelерinin ($\% \varepsilon$) ortalama ile varyans değeri görölmektedir. Burada Referans Bant Genişliği ile yapılan tahmin değeri ile yaklaşımlar üzerinden hesaplanan hata yüzdelерinin çoğunlukla $\hat{p}(H,S)$ yaklaşımı için iyi sonuçlar verdiği ve hata yüzde ortalamasının da yine $\hat{p}(H,S)$ için daha iyi olduğu görölmektedir. Ancak Maksimum Düzleştirme Yöntemine ile $\tilde{p}_{KMH}$ yaklaşımı daha iyi sonuçlar vermektedir.



Şekil 4.10. Senaryo II için hata yüzdeleri çizgi grafiği ($n = 1000$, $\text{öks} = 10$)

Şekil 4.10.(a) da $\hat{p}(H,S)$ ile $\tilde{p}_{KMH}$ den elde edilen hata yüzdelерinin çok yakın olduğu ve diğer iki yaklaşıma göre p 'yi daha iyi tahmin ettiği görölmektedir. Şekil 4.10.(b) de ise $\tilde{p}_{KMH}$ nin diğer üç yaklaşıma göre iyi sonuçlar verdiği 6. Örneklem kümesi için p 'nin en düşük hata yüzdesi ile tahmin edildiği görölmektedir.

Çizelge 4.20’de Maksimum Düzleştirme Yöntemi ile elde edilen sonuçların Referans Bant Genişliği ile yapılan tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ve 10 örneklem kümesi için hesaplanan hata yüzdesine benzer olarak ürün kitle parametresi olan uygun ürün oranı p 'ye daha yakın sonuçlar verdiği görölmektedir.

4.3. Senaryo III İçin Elde Edilen Sonuçlar

Senaryo III için Kernel yoğunluk tahmin yöntemine dayalı MH örnekleme ile elde edilen $n = 50, 100, 250, 500$ ve 1000 gözlemlili örneklem kümelerinden hesaplanan uygun ürün olasılıkları (Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$, $\hat{p}(H,S)$) ile kitle uygun ürün oranı p karşılaştırılarak, Hata yüzdeleri elde edilmiştir. Her bir gözlem sayısı için her iki bant genişliği matrisi kullanılarak hesaplanan tahmin sonuçları $n = 50$ için Çizelge 4.21’de verilmiştir.

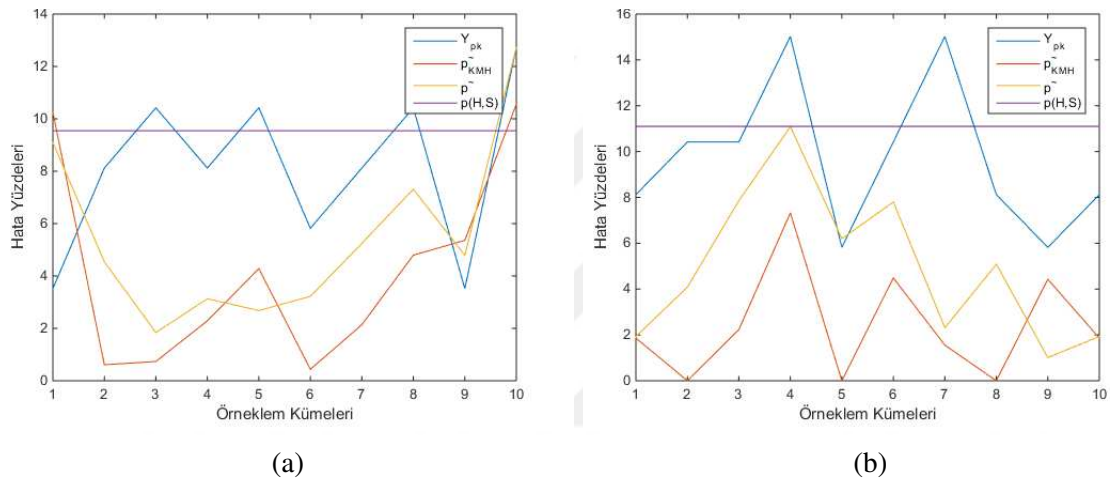
Çizelge 4.21. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 50$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
1	0,9000	0,7803	0,7904	0,7864	%3,520	%10,248	%9,087	%9,547	0,9400	0,8532	0,8862	0,7729	%8,121	%1,863	%1,932	%11,100
2	0,9400	0,8747	0,9089	0,7864	%8,121	%0,610	%4,543	%9,547	0,9600	0,8674	0,9049	0,7729	%10,421	%0,230	%4,083	%11,100
3	0,9600	0,8630	0,8854	0,7864	%10,421	%0,736	%1,840	%9,547	0,9600	0,8888	0,9377	0,7729	%10,421	%2,231	%7,856	%11,100
4	0,9400	0,8893	0,8966	0,7864	%8,121	%2,289	%3,129	%9,547	1,000	0,9331	0,9658	0,7729	%15,022	%7,327	%11,088	%11,100
5	0,9600	0,9066	0,8927	0,7864	%10,421	%4,279	%2,680	%9,547	0,9200	0,8731	0,9233	0,7729	%5,820	%0,426	%6,200	%11,100
6	0,9200	0,8656	0,8414	0,7864	%5,820	%0,437	%3,221	%9,547	0,9600	0,9084	0,9372	0,7729	%10,421	%4,486	%7,798	%11,100
7	0,9400	0,8880	0,9150	0,7864	%8,121	%2,139	%5,245	%9,547	1,000	0,8559	0,8895	0,7729	%15,022	%1,553	%2,312	%11,100
8	0,9600	0,9111	0,9330	0,7864	%10,421	%4,796	%7,315	%9,547	0,9400	0,8758	0,9136	0,7729	%8,121	%0,736	%5,084	%11,100
9	0,9000	0,8228	0,8278	0,7864	%3,520	%5,360	%4,785	%9,547	0,9200	0,8309	0,8782	0,7729	%5,820	%4,428	%1,012	%11,100
10	0,9800	0,9612	0,9800	0,7864	%12,721	%10,559	%12,721	%9,547	0,8400	0,7913	0,8293	0,7729	%8,121	%1,863	%1,932	%11,100
ORTALAMA	0,9400	0,8763	0,8871	0,7864	%8,121	%4,145	%5,457	%9,547	0,9440	0,8678	0,9066	0,7729	%9,731	%2,514	%4,930	%11,100
VARYANS	0,0007	0,0024	0,0030	0,0000	%0,094	%0,140	%0,113	%0,000	0,0021	0,0016	0,0015	0,0000	%0,106	%0,050	%0,108	%0,000

Çizelge 4.22. Senaryo III için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 50$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
ORTALAMA	0,9450	0,8806	0,9083	0,7864	%8,804	%3,260	%4,671	%9,547	0,9530	0,8692	0,9024	0,7729	%9,616	%4,323	%4,963	%11,100
VARYANS	0,0013	0,0010	0,0006	0,0000	%0,158	%0,038	%0,062	%0,000	0,0010	0,0021	0,0017	0,0000	%0,088	%0,112	%0,000	%0,000

Çizelge 4.21'den $n = 50$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. Örneklem kümeleri arasında oluşan varyans değerlerinin 0,000'a çok yakın olması nedeniyle tahmin ortalamaları üzerinden yaklaşımları karşılaştırmanın anlamlı olacağı söylenebilir. Örneklem kümeleri üzerinden elde edilen yaklaşım değerlerinin ortalamaları karşılaştırıldığında, Referans Bant Genişliği ve Maksimum Düzleştirme Yöntemi kullanılarak yapılan tahminler için en iyi sonucun önerilen $\tilde{p}_{KMH}$ yaklaşımı ile elde edildiği görülmektedir. Çizelge 4.21'de görülen hata yüzde değerlerinin karşılaştırma kolaylığı sağlaması açısından çizgi grafiği çizdirilerek Şekil 4.11'de verilmiştir.



Şekil 4.11. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 50$, $\text{öks} = 10$)

Şekil 4.11'de verilen grafikler incelendiğinde, her iki grafikte de $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelerinin diğer üç yönteme göre daha iyi sonuç verdiği görülmektedir.

Çizelge 4.22'de $n = 50$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri görülmektedir. Bu değerler, Çizelge 4.21'de verilen 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \epsilon$) ortalama ile varyans değerleri ile karşılaştırıldığında, örneklem küme sayısındaki artış sonucu kümeler arasında oluşan ortalama ve varyans değerlerinde çok fazla bir fark olmadığı görülmektedir.

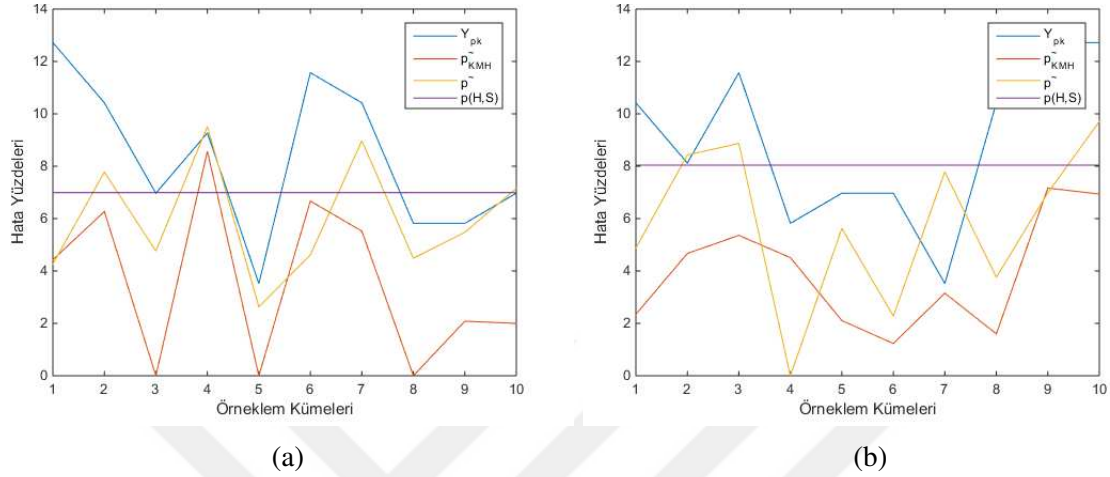
Çizelge 4.23. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 100$, $\delta ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
1	0,9800	0,9079	0,9065	0,8086	%12,721	%4,428	%4,267	%6,993	0,9600	0,8898	0,9118	0,7995	%10,421	%2,346	%4,877	%8,040
2	0,9600	0,9239	0,9370	0,8086	%10,421	%6,269	%7,775	%6,993	0,9400	0,9100	0,9427	0,7995	%8,121	%4,670	%8,431	%8,040
3	0,9300	0,8763	0,9108	0,8086	%6,970	%0,794	%4,762	%6,993	0,9700	0,9160	0,9465	0,7995	%11,571	%5,360	%8,868	%8,040
4	0,9500	0,9439	0,9521	0,8086	%9,271	%8,569	%9,512	%6,993	0,9200	0,8302	0,8686	0,7995	%5,820	%4,509	%0,092	%8,040
5	0,9000	0,8620	0,8923	0,8086	%3,520	%0,851	%2,634	%6,993	0,9300	0,8877	0,9183	0,7995	%6,970	%2,105	%5,625	%8,040
6	0,9700	0,9274	0,9095	0,8086	%11,571	%6,671	%4,612	%6,993	0,9300	0,8587	0,8892	0,7995	%6,970	%1,231	%2,277	%8,040
7	0,9600	0,9174	0,9474	0,8086	%10,421	%5,521	%8,972	%6,993	0,9000	0,8968	0,9371	0,7995	%3,520	%3,152	%7,787	%8,040
8	0,9200	0,8722	0,9084	0,8086	%5,820	%0,322	%4,486	%6,993	0,9600	0,8833	0,9021	0,7995	%10,421	%1,599	%3,761	%8,040
9	0,9200	0,8875	0,9170	0,8086	%5,820	%2,082	%5,475	%6,993	0,9800	0,9317	0,9301	0,7995	%12,721	%7,166	%6,982	%8,040
10	0,9300	0,8868	0,9316	0,8086	%6,970	%2,001	%7,154	%6,993	0,9800	0,9297	0,9538	0,7995	%12,721	%6,936	%9,708	%8,040
ORTALAMA	0,9420	0,9005	0,9213	0,8086	%8,351	%3,751	%5,965	%6,993	0,9470	0,8934	0,9200	0,7995	%8,926	%3,907	%5,841	%8,040
VARYANS	0,0007	0,0007	0,0004	0,0000	%0,088	%0,085	%0,051	%0,000	0,0007	0,0010	0,0007	0,0000	%0,097	%0,046	%0,097	%0,000

Çizelge 4.24. Senaryo III için tahmin değerleri ve hata yüzdelерinin ortalama ve varyansları ($n = 100$, $\delta ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
ORTALAMA	0,9400	0,8960	0,9178	0,8086	%8,121	%3,765	%5,568	%6,993	0,9500	0,8977	0,9241	0,7995	%9,271	%4,147	%6,308	%8,040
VARYANS	0,0009	0,0009	0,0007	0,0000	%0,113	%0,066	%0,090	%0,000	0,0004	0,0010	0,0007	0,0000	%0,058	%0,065	%0,092	%0,000

Çizelge 4.23'te $n = 100$ gözlem sayısına sahip 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \varepsilon$) ortalama ile varyans değerleri görülmektedir. Ürün kitle parametresi p 'nin önerilen yaklaşım olan $\tilde{p}_{KMH}$ ile diğer üç yöntemle göre daha iyi tahmin edildiği açıkça görülmektedir.



Şekil 4.12. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 100$, $\text{öks} = 10$)

Şekil 4.12'de verilen grafikler incelendiğinde, ilk grafikte görülen, Referans Bant Genişliğine göre hesaplanan hata yüzdeleri için 1. ve 6. Örneklem kümeleri dışındaki örneklem kümelerinde $\tilde{p}_{KMH}$ yaklaşımının iyi sonuç verdiği söylenebilir. İkinci grafikte Maksimum Düzleştirme Yönteminin 4. Örneklem kümesi dışındaki örneklem kümelerinde $\tilde{p}_{KMH}$ yaklaşımının daha iyi sonuç verdiği görülmektedir.

Çizelge 4.24'te $n = 100$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin 10 örneklem kümesine göre hesap edilen hata yüzdeleri ile benzer olduğu ve p 'nin önerilen yaklaşım olan $\tilde{p}_{KMH}$ ile diğer üç yöntemle göre daha iyi tahmin edildiği söylenebilir.

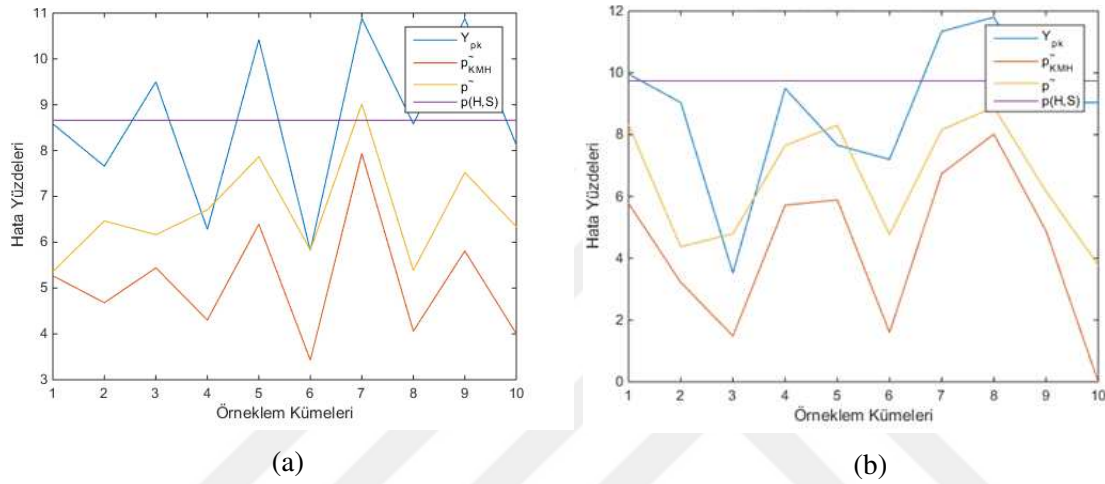
Çizelge 4.25. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 250$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
1	0,9440	0,9152	0,9160	0,7941	%8,581	%5,268	%5,360	%8,661	0,9560	0,9196	0,9417	0,7847	%9,961	%5,774	%8,316	%9,742
2	0,9360	0,9101	0,9256	0,7941	%7,660	%4,681	%6,464	%8,661	0,9480	0,8974	0,9074	0,7847	%9,041	%3,221	%4,371	%9,742
3	0,9520	0,9167	0,9230	0,7941	%9,501	%5,441	%6,165	%8,661	0,9000	0,8823	0,9111	0,7847	%3,520	%1,484	%4,796	%9,742
4	0,9240	0,9068	0,9277	0,7941	%6,280	%4,302	%6,706	%8,661	0,9520	0,9191	0,9359	0,7847	%9,501	%5,717	%7,649	%9,742
5	0,9600	0,9250	0,9378	0,7941	%10,421	%6,395	%7,867	%8,661	0,9360	0,9206	0,9416	0,7847	%7,660	%5,889	%8,305	%9,742
6	0,9200	0,8992	0,9201	0,7941	%5,820	%3,428	%5,832	%8,661	0,9320	0,8833	0,9108	0,7847	%7,200	%1,599	%4,762	%9,742
7	0,9640	0,9384	0,9478	0,7941	%10,881	%7,937	%9,018	%8,661	0,9680	0,9280	0,9403	0,7847	%11,341	%6,740	%8,155	%9,742
8	0,9440	0,9047	0,9162	0,7941	%8,581	%4,060	%5,383	%8,661	0,9720	0,9391	0,9467	0,7847	%11,801	%8,017	%8,891	%9,742
9	0,9640	0,9199	0,9348	0,7941	%10,881	%5,809	%7,522	%8,661	0,9480	0,9119	0,9230	0,7847	%9,041	%4,888	%6,165	%9,742
10	0,9400	0,9042	0,9245	0,7941	%8,121	%4,003	%6,338	%8,661	0,9480	0,8759	0,9022	0,7847	%9,041	%0,748	%3,773	%9,742
ORTALAMA	0,9448	0,9140	0,9274	0,7941	%8,673	%5,132	%6,666	%8,661	0,9460	0,9077	0,9261	0,7847	%8,811	%4,408	%6,518	%9,742
VARYANS	0,0002	0,0001	0,0001	0,0000	%0,032	%0,018	%0,014	%0,000	0,0004	0,0005	0,0003	0,0000	%0,055	%0,062	%0,038	%0,000

Çizelge 4.26. Senaryo III için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 250$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
ORTALAMA	0,9446	0,9113	0,9249	0,7941	%8,650	%4,818	%6,378	%8,661	0,9434	0,9046	0,9237	0,7847	%8,512	%4,052	%6,245	%9,742
VARYANS	0,0001	0,0002	0,0001	0,0000	%0,020	%0,020	%0,012	%0,000	0,0001	0,0002	0,0001	0,0000	%0,021	%0,013	%0,000	%0,000

Çizelge 4.25'te $n = 250$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \varepsilon$) görülmektedir. Maksimum Düzleştirme Yöntemi ile elde edilen sonuçların ortalama olarak Referans Bant Genişliği ile yapılan tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ürün kitle parametresi olan uygun ürün oranı p 'ye daha yakın sonuçlar verdiği görülmektedir. Ayrıca her iki tahmin yönteminde de $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelерinin diğer üç yönteme göre daha iyi sonuç verdiği söylenebilir.



Şekil 4.13. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 250$, $\text{öks} = 10$)

Şekil 4.13'deki grafikler incelendiğinde, her iki grafikte de $\tilde{p}_{KMH}$ yaklaşımı ile elde edilen hata yüzdelерinin diğer üç yönteme göre daha iyi sonuç verdiği görülmektedir.

Çizelge 4.26'da verilen hata yüzdelерinden, 10 örneklem kümesi için hesaplanan hata yüzderine benzer olarak $\tilde{p}_{KMH}$ yaklaşımını ürün kitle parametresi p 'yi daha iyi tahmin ettiği söylenebilir. Maksimum Düzleştirme Yöntemi ile elde edilen sonuçların Referans Bant Genişliği ile yapılan tahmin sonuçlarından daha az hata yüzdesine sahip olduğu görülmektedir.

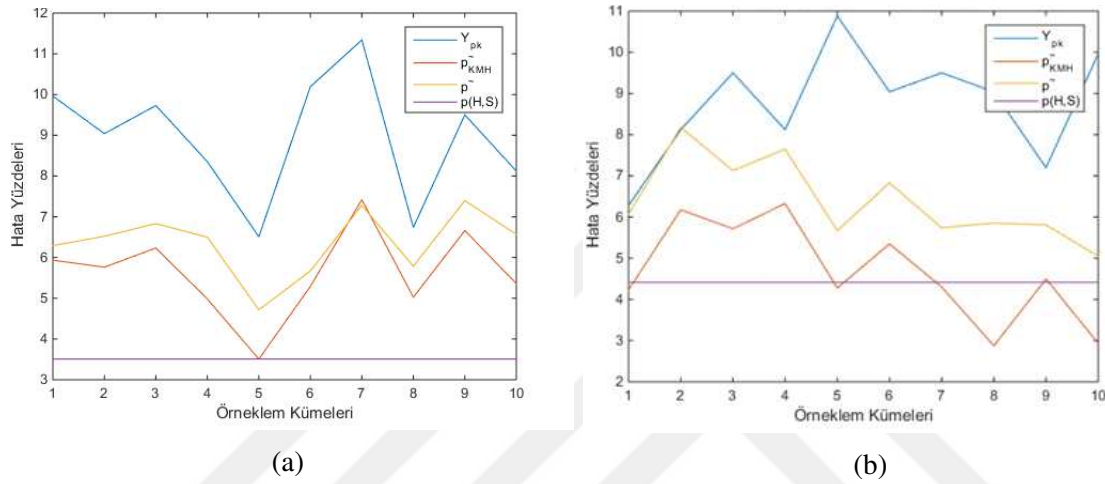
Çizelge 4.27. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 500$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
1	0,9560	0,9210	0,9241	0,8389	%9,961	%5,935	%6,292	%3,508	0,9240	0,9062	0,9222	0,8310	%6,280	%4,233	%6,073	%4,417
2	0,9480	0,9195	0,9261	0,8389	%9,041	%5,763	%6,522	%3,508	0,9400	0,9231	0,9405	0,8310	%8,121	%6,177	%8,178	%4,417
3	0,9540	0,9236	0,9288	0,8389	%9,731	%6,234	%6,832	%3,508	0,9520	0,9191	0,9314	0,8310	%9,501	%5,717	%7,131	%4,417
4	0,9420	0,9127	0,9259	0,8389	%8,351	%4,980	%6,499	%3,508	0,9400	0,9244	0,9359	0,8310	%8,121	%6,326	%7,649	%4,417
5	0,9260	0,8999	0,9104	0,8389	%6,510	%3,508	%4,716	%3,508	0,9640	0,9066	0,9187	0,8310	%10,881	%4,279	%5,671	%4,417
6	0,9580	0,9154	0,9187	0,8389	%10,191	%5,291	%5,671	%3,508	0,9480	0,9159	0,9288	0,8310	%9,041	%5,349	%6,832	%4,417
7	0,9680	0,9339	0,9327	0,8389	%11,341	%7,419	%7,281	%3,508	0,9520	0,9068	0,9193	0,8310	%9,501	%4,302	%5,740	%4,417
8	0,9280	0,9131	0,9197	0,8389	%6,740	%5,026	%5,786	%3,508	0,9480	0,8944	0,9203	0,8310	%9,041	%2,876	%5,855	%4,417
9	0,9520	0,9273	0,9337	0,8389	%9,501	%6,660	%7,396	%3,508	0,9320	0,9085	0,9199	0,8310	%7,200	%4,497	%5,809	%4,417
10	0,9400	0,9160	0,9265	0,8389	%8,121	%5,360	%6,568	%3,508	0,9560	0,8949	0,9133	0,8310	%9,961	%2,933	%5,049	%4,417
ORTALAMA	0,9472	0,9182	0,9247	0,8389	%8,949	%5,618	%6,356	%3,508	0,9456	0,9100	0,9250	0,8310	%8,765	%4,669	%6,399	%4,417
VARYANS	0,0002	0,0001	0,0000	0,0000	%0,023	%0,011	%0,006	%0,000	0,0001	0,0001	0,0001	0,0000	%0,018	%0,015	%0,010	%0,000

Çizelge 4.28. Senaryo III için tahmin değerleri ve hata yüzdelерinin ortalama ve varyansları ($n = 500$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
ORTALAMA	0,9466	0,9187	0,9281	0,8389	%8,880	%5,669	%6,755	%3,508	0,9461	0,9139	0,9279	0,8310	%8,822	%5,120	%6,734	%4,417
VARYANS	0,0001	0,0001	0,0001	0,0000	%0,018	%0,012	%0,009	%0,000	0,0001	0,0001	0,0001	0,0000	%0,016	%0,010	%0,008	%0,000

Çizelge 4.27’de $n = 500$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. Referans Bant Genişliği ve Maksimum Düzleştirme Yöntemi ile yapılan tahminlerde, $\hat{p}(H,S)$ yaklaşımının diğer üç yönteme göre daha iyi sonuç verdiği söylenebilir. Referans Bant Genişliği ile yapılan sonuçların Maksimum Bant Genişliği tahmin yöntemi ile elde edilen tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ürün kitle parametresi olan uygun ürün oranı p ’ye daha yakın sonuçlar verdiği görülmektedir.



Şekil 4.14. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 500$, $\text{öks} = 10$)

Şekil 4.14’te görülen grafikler Senaryo I’e ait faz I verilerinden elde edilen ϵY_{pk} , $\epsilon \tilde{p}_{KMH}$, $\epsilon \tilde{p}$ ve $\epsilon \hat{p}(H,S)$ hata yüzdelerini göstermektedir. İlk grafikte $\hat{p}(H,S)$ yaklaşımı ile p ’nin daha iyi tahmin edildiği görülmektedir. İkinci grafikte ise Maksimum Düzleştirme Yöntemi yardımıyla Kernel yoğunluk tahmini yapılarak aynı dağılımdan çektilen 10 tane örneklem kümesinden hesaplanmıştır. İkinci grafik incelendiğinde, örneklem kümelerinden hesaplanan tahmin sonuçlarının değişkenlik gösterdiği, 1.,5.,7.,8.,10. Örneklem kümelerinde $\tilde{p}_{KMH}$ yaklaşımı daha iyi sonuç verirken; diğerlerinde $\hat{p}(H,S)$ yaklaşımının başarılı olduğu görülmektedir.

Çizelge 4.28’deki 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinden $\hat{p}(H,S)$ ile daha iyi tahminler elde edildiği görülmektedir. Referans Bant Genişliği yöntemine göre $\hat{p}(H,S)$ yaklaşımının daha iyi sonuç verdiği görülmektedir. Maksimum Düzleştirme Yöntemi ile elde edilen sonuçlarda, 10 örneklem kümesi için hesaplanan hata yüzdelerinden farklı olarak $\tilde{p}$ yaklaşımının p ’ye daha yakın sonuçlar verdiği söylenebilir.

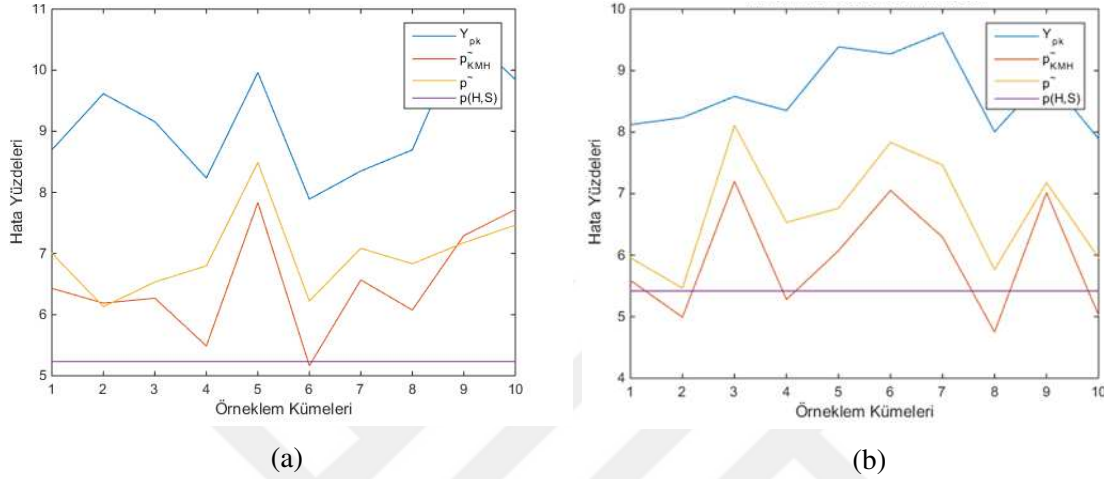
Çizelge 4.29. Senaryo III için tahmin değerleri ve hata yüzdeleri ($n = 1000$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
1	0,9450	0,9253	0,9303	0,8278	%8,696	%6,430	%7,005	%5,232	0,9400	0,9180	0,9212	0,8223	%8,121	%5,590	%5,958	%5,418
2	0,9530	0,9232	0,9227	0,8278	%9,616	%6,188	%6,131	%5,232	0,9410	0,9128	0,9169	0,8223	%8,236	%4,992	%5,464	%5,418
3	0,9490	0,9239	0,9262	0,8278	%9,156	%6,269	%6,533	%5,232	0,9440	0,9320	0,9399	0,8223	%8,581	%7,200	%8,109	%5,418
4	0,9410	0,9171	0,9285	0,8278	%8,236	%5,487	%6,798	%5,232	0,9420	0,9153	0,9262	0,8223	%8,351	%5,280	%6,533	%5,418
5	0,9560	0,9375	0,9432	0,8278	%9,961	%7,833	%8,489	%5,232	0,9510	0,9222	0,9282	0,8223	%9,386	%6,073	%6,763	%5,418
6	0,9380	0,9143	0,9235	0,8278	%7,890	%5,164	%6,223	%5,232	0,9500	0,9307	0,9375	0,8223	%9,271	%7,051	%7,833	%5,418
7	0,9420	0,9265	0,9310	0,8278	%8,351	%6,568	%7,085	%5,232	0,9530	0,9241	0,9343	0,8223	%9,616	%6,292	%7,465	%5,418
8	0,9450	0,9222	0,9288	0,8278	%8,696	%6,075	%6,832	%5,232	0,9390	0,9107	0,9195	0,8223	%8,006	%4,750	%5,763	%5,418
9	0,9620	0,9328	0,9318	0,8278	%10,651	%7,292	%7,177	%5,232	0,9470	0,9304	0,9318	0,8223	%8,926	%7,016	%7,177	%5,418
10	0,9550	0,9365	0,9343	0,8278	%9,846	%7,718	%7,465	%5,232	0,9380	0,9131	0,9212	0,8223	%7,890	%5,026	%5,958	%5,418
ORTALAMA	0,9486	0,9259	0,9300	0,8278	%9,110	%6,502	%6,974	%5,232	0,9445	0,9209	0,9277	0,8223	%8,638	%5,927	%6,702	%5,418
VARYANS	0,0001	0,0001	0,0000	0,0000	%0,008	%0,008	%0,005	%0,000	0,0000	0,0001	0,0001	0,0000	%0,004	%0,009	%0,008	%0,000

Çizelge 4.30. Senaryo III için tahmin değerleri ve hata yüzdelерinin ortalama ve varyansları ($n = 1000$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$	Y_{pk}	$\tilde{p}_{KMH}$	$\tilde{p}$	$\hat{p}(H, S)$	εY_{pk}	$\varepsilon \tilde{p}_{KMH}$	$\varepsilon \tilde{p}$	$\varepsilon \hat{p}(H, S)$
ORTALAMA	0,9467	0,9269	0,9321	0,8278	%8,885	%6,612	%7,208	%4,785	0,9466	0,9205	0,9283	0,8223	%8,874	%5,878	%6,772	%5,418
VARYANS	0,0001	0,0000	0,0000	0,0000	%0,009	%0,006	%0,003	%0,000	0,0001	0,0000	0,0000	0,0000	%0,004	%0,004	%0,000	%0,000

Çizelge 4.29 da, Referans Bant Genişliği ile yapılan tahminlerde 6. örneklem kümesi hariç, tüm örneklem kümelerinde en iyi tahminlerin $\hat{p}(H,S)$ yaklaşımı ile yapıldığı görülmektedir. Maksimum Düzleştirme Yöntemi ile yapılan tahminlerde ise $\tilde{p}_{KMH}$ ve $\hat{p}(H,S)$ yaklaşımların kitle parametresini tahmin etmedeki başarıları değişkenlik göstermektedir. Örneğin, $\tilde{\sigma}_{ks} = 2, 4, 8, 10$ için $\tilde{p}_{KMH}$ yaklaşımı daha iyi sonuçlar verdiği görülmektedir.



Şekil 4.15. Senaryo III için hata yüzdeleri çizgi grafiği ($n = 1000, \tilde{\sigma}_{ks} = 10$)

Şekil 4.15.(a) da $\hat{p}(H,S)$ nin diğer yaklaşımlara göre p 'yi daha iyi tahmin ettiği görülmektedir. Maksimum Düzleştirme Yöntemi ile yapılan tahminlerde ise $\tilde{p}_{KMH}$ ve $\hat{p}(H,S)$ nin kitle parametresini tahmin etmedeki başarıları değişkenlik göstermektedir. Şekil 4.15.(b) den $\tilde{\sigma}_{ks} = 2, 4, 8, 10$ için $\tilde{p}_{KMH}$ nin daha iyi sonuçlar verdiği görülmektedir.

Çizelge 4.30'daki 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelерinden $\hat{p}(H,S)$ ile daha iyi tahminler elde edildiği görülmektedir.

4.4. Senaryo IV İçin Elde Edilen Sonuçlar

Senaryo IV için Kernel yoğunluk tahmin yöntemine dayalı MH örnekleme ile elde edilen $n = 50, 100, 250, 500$ ve 1000 gözlemlili örneklem kümelerinden hesaplanan uygun ürün olasılıkları (Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$, $\hat{p}(H,S)$) kitle uygun ürün oranı p ile karşılaştırılarak, hata yüzdeleri elde edilmiştir. Kernel yoğunluk tahminleri yapılırken, MDY bant genişliği matrisi ve örneklem kümelerinin normal dağılım gösterdiği varsayımına dayalı RBG matrisi kullanılmıştır. Her bir gözlem sayısı için her iki bant genişliği matrisi kullanılarak hesaplanan tahmin sonuçları $n = 50$ için Çizelge 4.31'de verilmiştir.

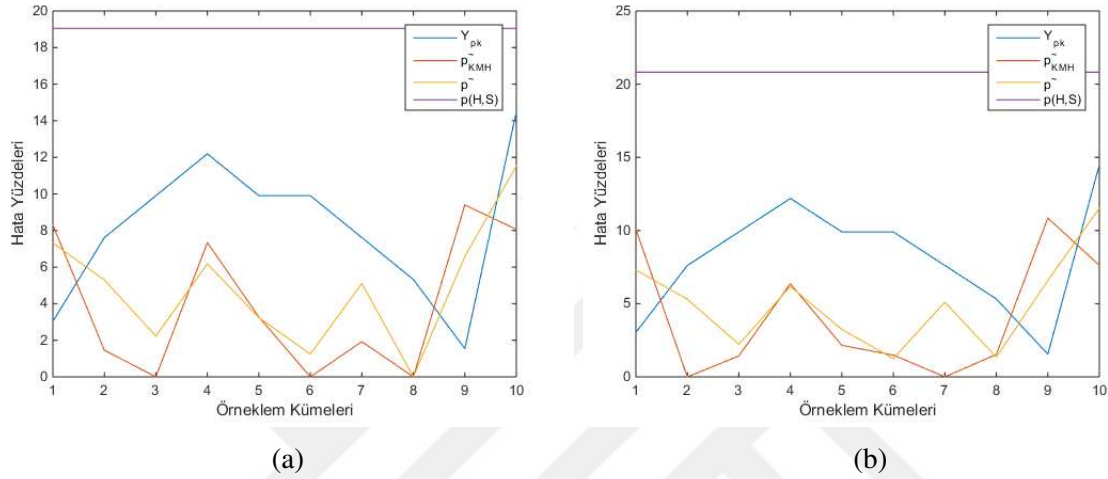
Çizelge 4.31. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 50$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
1	0,9000	0,8010	0,8097	0,7071	%3,034	%8,300	%7,304	%19,050	0,9000	0,7855	0,8097	0,6916	%3,034	%10,074	%7,304	%20,824
2	0,9400	0,8862	0,9198	0,7071	%7,613	%1,454	%5,301	%19,050	0,9400	0,8788	0,9198	0,6916	%7,613	%0,607	%5,301	%20,824
3	0,9600	0,8709	0,8929	0,7071	%9,903	%0,298	%2,221	%19,050	0,9600	0,8610	0,8929	0,6916	%9,903	%1,431	%2,221	%20,824
4	0,9800	0,9376	0,9275	0,7071	%12,192	%7,338	%6,182	%19,050	0,9800	0,9291	0,9275	0,6916	%12,192	%6,365	%6,182	%20,824
5	0,9600	0,9021	0,9018	0,7071	%9,903	%3,274	%3,240	%19,050	0,9600	0,8923	0,9018	0,6916	%9,903	%2,152	%3,240	%20,824
6	0,9600	0,8743	0,8626	0,7071	%9,903	%0,092	%1,248	%19,050	0,9600	0,8605	0,8626	0,6916	%9,903	%1,488	%1,248	%20,824
7	0,9400	0,8903	0,9181	0,7071	%7,613	%1,923	%5,106	%19,050	0,9400	0,8802	0,9181	0,6916	%7,613	%0,767	%5,106	%20,824
8	0,9200	0,8672	0,8812	0,7071	%5,323	%0,721	%0,882	%19,050	0,9200	0,8601	0,8853	0,6916	%5,323	%1,534	%1,351	%20,824
9	0,8600	0,7914	0,8160	0,7071	%1,546	%9,399	%6,583	%19,050	0,8600	0,7787	0,8160	0,6916	%1,546	%10,853	%6,583	%20,824
10	1,000	0,9439	0,9742	0,7071	%14,482	%8,060	%11,528	%19,050	1,0000	0,9401	0,9742	0,6916	%14,482	%7,624	%11,528	%20,824
ORTALAMA	0,9420	0,8765	0,8904	0,7071	%8,151	%4,086	%4,960	%19,050	0,9420	0,8666	0,8908	0,6916	%8,151	%4,290	%5,006	%20,824
VARYANS	0,0016	0,0025	0,0026	0,0000	%0,160	%0,140	%0,104	%0,000	0,0016	0,0027	0,0026	0,0000	%0,160	%0,162	%0,100	%0,000

Çizelge 4.32. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 50$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
ORTALAMA	0,9510	0,8850	0,9123	0,7071	%9,027	%3,719	%4,764	%19,050	0,9480	0,8731	0,9101	0,6916	%8,683	%3,146	%4,518	%20,824
VARYANS	0,0015	0,0014	0,0008	0,0000	%0,167	%0,058	%0,073	%0,000	0,0015	0,0014	0,0007	0,0000	%0,173	%0,083	%0,065	%0,000

Çizelge 4.31'den $n = 50$ gözlem sayısına sahip örneklem kümelerinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdeleri ($\% \epsilon$) görülmektedir. Referans Bant Genişliği ve Maksimum Düzleştirme Yöntemi kullanılarak yapılan tahminler için en iyi sonucun ortalamada önerilen $\tilde{p}_{KMH}$ yaklaşımı ile elde edildiği görülmektedir. Çizelge 4.31'de görülen hata yüzde değerlerinin karşılaştırma kolaylığı sağlaması açısından çizgi grafiği çizdirilerek Şekil 4.16'da verilmiştir.



Şekil 4.16. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 50, \text{öks} = 10$)

Şekil 4.16'daki verilen grafikler incelendiğinde ilk grafikteki sonuçlardan 1. ve 9. örneklem kümelerinden hesaplanan Y_{pk} yaklaşımı ile diğer yaklaşımlara göre p 'yi daha iyi tahmin ettiği görülmektedir. 2, 3, 6, 7. ve 8. örneklem kümelerinden hesaplanan $\tilde{p}_{KMH}$ yaklaşımı ile daha iyi sonuçlar elde edilmiştir. 4. ve 5. Örneklem kümelerinden hesaplanan $\tilde{p}$ yaklaşımı ile p daha iyi tahmin edilmiştir. İkinci grafikte görülen Maksimum Düzleştirme Yöntemi ile yapılan tahminlerde de ilk grafikteki gibi Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ yaklaşımların kitle parametresini tahmin etmedeki başarıları değişkenlik göstermektedir.

Çizelge 4.32'de $\tilde{p}_{KMH}$ yaklaşımına göre Maksimum Düzleştirme Yöntemi ile elde edilen sonuçların Referans Bant Genişliği ile yapılan tahmin sonuçlarından daha az hata yüzdesine sahip olduğu ve 10 örneklem kümesinde olduğu gibi $\tilde{p}_{KMH}$ yaklaşımı ile p daha iyi tahmin edildiği söylenebilir.

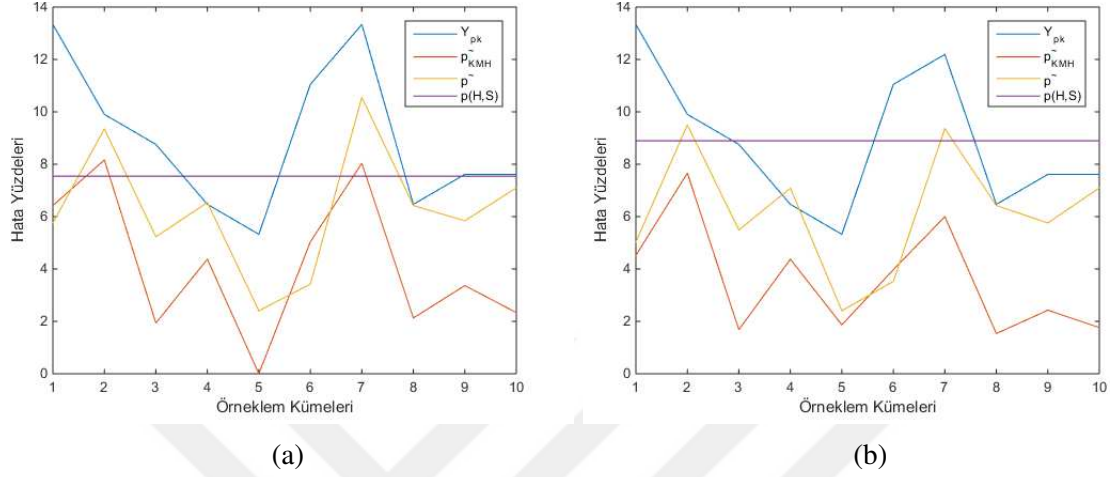
Çizelge 4.33. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 100$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
1	0,9900	0,9296	0,9238	0,8076	%13,337	%6,422	%5,758	%7,544	0,9900	0,9130	0,9174	0,7958	%13,337	%4,522	%5,026	%8,895
2	0,9600	0,9448	0,9552	0,8076	%9,903	%8,163	%9,353	%7,544	0,9600	0,9404	0,9565	0,7958	%9,903	%7,659	%9,502	%8,895
3	0,9500	0,8904	0,9192	0,8076	%8,758	%1,935	%5,232	%7,544	0,9500	0,8882	0,9214	0,7958	%8,758	%1,683	%5,488	%8,895
4	0,9300	0,9117	0,9306	0,8076	%6,468	%4,373	%6,537	%7,544	0,9300	0,9117	0,9354	0,7958	%6,468	%4,373	%7,086	%8,895
5	0,9200	0,8648	0,8945	0,8076	%5,323	%0,996	%2,404	%7,544	0,9200	0,8572	0,8945	0,7958	%5,323	%1,866	%2,404	%8,895
6	0,9700	0,9175	0,9034	0,8076	%11,048	%5,037	%3,423	%7,544	0,9700	0,9082	0,9043	0,7958	%11,048	%3,973	%3,526	%8,895
7	0,9900	0,9437	0,9657	0,8076	%13,337	%8,037	%10,555	%7,544	0,9800	0,9259	0,9553	0,7958	%12,192	%5,999	%9,365	%8,895
8	0,9300	0,8921	0,9296	0,8076	%6,468	%2,129	%6,422	%7,544	0,9300	0,8869	0,9296	0,7958	%6,468	%1,534	%6,422	%8,895
9	0,9400	0,9029	0,9245	0,8076	%7,613	%3,366	%5,839	%7,544	0,9400	0,8947	0,9238	0,7958	%7,613	%2,427	%5,758	%8,895
10	0,9400	0,8939	0,9355	0,8076	%7,613	%2,335	%7,098	%7,544	0,9400	0,8889	0,9355	0,7958	%7,613	%1,763	%7,098	%8,895
ORTALAMA	0,9520	0,9091	0,9282	0,8076	%8,987	%4,279	%6,262	%7,544	0,9510	0,9015	0,9274	0,7958	%8,872	%3,580	%6,168	%8,895
VARYANS	0,0006	0,0006	0,0005	0,0000	%0,081	%0,067	%0,059	%0,000	0,0005	0,0005	0,0004	0,0000	%0,071	%0,044	%0,051	%0,000

Çizelge 4.34. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 100$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
ORTALAMA	0,9405	0,8880	0,9115	0,8076	%7,670	%3,197	%4,361	%7,544	0,9405	0,8807	0,9119	0,7958	%7,670	%3,000	%4,391	%8,895
VARYANS	0,0007	0,0009	0,0007	0,0000	%0,096	%0,033	%0,085	%0,000	0,0007	0,0009	0,0006	0,0000	%0,096	%0,031	%0,085	%0,000

Çizelge 4.33'te $n = 100$ gözlem sayısına sahip 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin ($\% \varepsilon$) ortalama ile varyans değerleri görülmektedir. Ürün kitle parametresi p 'nin önerilen yaklaşım olan $\tilde{p}_{KMH}$ ile diğer üç yöntemle göre daha iyi tahmin edildiği açıkça görülmektedir.



Şekil 4.17. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 100$, $\text{öks} = 10$)

Şekil 4.17'de verilen grafikler incelendiğinde, Çizelge 4.33'te verilen sonuçların bir özeti görülmektedir. İlk grafikte Referans Bant Genişliği kullanılarak hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ yaklaşım değerlerinin hata yüzdelerinden, en küçük değerlerin $\tilde{p}_{KMH}$ yaklaşımı sonucu elde edildiği görülmektedir. Bununla beraber, 1. ve 6. örneklem kümeleri için $\tilde{p}$ yaklaşımının, 7. örneklem kümesi için de $\hat{p}(H,S)$ yaklaşımının daha iyi sonuç verdiği söylenebilir.

Çizelge 4.34 incelendiğinde, her iki yöntem için elde edilen sonuçlardan $\tilde{p}_{KMH}$ yaklaşımının $\tilde{p}$ 'yi tahmin etmede daha başarılı olduğu görülmektedir.

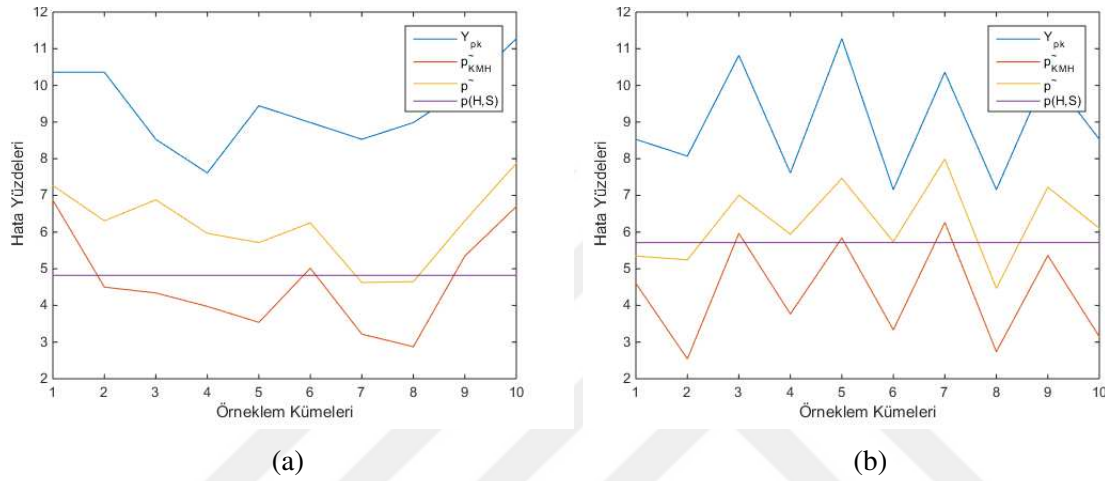
Çizelge 4.35. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 250$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
1	0,9640	0,9335	0,9370	0,8314	%10,361	%6,869	%7,270	%4,820	0,9480	0,9137	0,9202	0,8236	%8,529	%4,602	%5,346	%5,713
2	0,9640	0,9128	0,9286	0,8314	%10,361	%4,499	%6,308	%4,820	0,9440	0,8957	0,9193	0,8236	%8,071	%2,541	%5,243	%5,713
3	0,9480	0,9114	0,9336	0,8314	%8,529	%4,339	%6,880	%4,820	0,9680	0,9256	0,9347	0,8236	%10,819	%5,965	%7,006	%5,713
4	0,9400	0,9082	0,9256	0,8314	%7,613	%3,973	%5,965	%4,820	0,9400	0,9064	0,9254	0,8236	%7,613	%3,766	%5,942	%5,713
5	0,9560	0,9044	0,9234	0,8314	%9,445	%3,537	%5,713	%4,820	0,9720	0,9245	0,9387	0,8236	%11,276	%5,839	%7,464	%5,713
6	0,9520	0,9173	0,9281	0,8314	%8,987	%5,014	%6,251	%4,820	0,9360	0,9026	0,9236	0,8236	%7,155	%3,331	%5,736	%5,713
7	0,9480	0,9016	0,9139	0,8314	%8,529	%3,217	%4,625	%4,820	0,9640	0,9282	0,9433	0,8236	%10,361	%6,262	%7,991	%5,713
8	0,9520	0,8986	0,9141	0,8314	%8,987	%2,873	%4,648	%4,820	0,9360	0,8974	0,9125	0,8236	%7,155	%2,736	%4,465	%5,713
9	0,9600	0,9202	0,9286	0,8314	%9,903	%5,346	%6,308	%4,820	0,9640	0,9203	0,9366	0,8236	%10,361	%5,358	%7,224	%5,713
10	0,9720	0,9320	0,9423	0,8314	%11,276	%6,697	%7,876	%4,820	0,9480	0,9010	0,9268	0,8236	%8,529	%3,148	%6,102	%5,713
ORTALAMA	0,9556	0,9140	0,9275	0,8314	%9,399	%4,636	%6,184	%4,820	0,9520	0,9115	0,9281	0,8236	%8,987	%4,355	%6,252	%5,713
VARYANS	0,0001	0,0001	0,0001	0,0000	%0,012	%0,019	%0,011	%0,000	0,0002	0,0002	0,0001	0,0000	%0,025	%0,020	%0,013	%0,000

Çizelge 4.36. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 250$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
ORTALAMA	0,9490	0,9123	0,9233	0,8314	%8,643	%4,443	%5,695	%4,820	0,9472	0,9052	0,9224	0,8236	%8,437	%3,704	%5,598	%5,713
VARYANS	0,0003	0,0003	0,0002	0,0000	%0,043	%0,042	%0,030	%0,000	0,0003	0,0003	0,0002	0,0000	%0,043	%0,040	%0,031	%0,000

Çizelge 4.35. incelendiğinde, Referans Bant Genişliği ile yapılan tahminlerde örneklem kümelerinde en iyi tahminlerin $\tilde{p}_{KMH}$ ve $\hat{p}(H,S)$ yaklaşımları ile yapıldığı görülmektedir. Ortalamada $\tilde{p}_{KMH}$ yaklaşımının daha iyi sonuç verdiği söylenebilir. Maksimum Düzleştirme Yöntemi ile yapılan tahminlerde de $\tilde{p}_{KMH}$ ve $\hat{p}(H,S)$ yaklaşımların kitle parametresini tahmin etmedeki başarıları değişkenlik göstermektedir. Örneğin, 3,5 ve 7. Örneklem kümelerinde $\hat{p}(H,S)$ yaklaşımı daha iyi sonuçlar verdiği görülmektedir, ancak ortalamada $\tilde{p}_{KMH}$ yaklaşımı daha başarılıdır.



Şekil 4.18. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 250$, $\text{öks} = 10$)

Şekil 4.18'de verilen grafikler incelendiğinde, ilk grafikteki sonuçlardan 2, 3, 4, 5, 7 ve 8. Örneklem kümelerinde $\tilde{p}_{KMH}$ yaklaşımının diğer yaklaşımlara göre p 'yi daha iyi tahmin ettiği görülmektedir. Maksimum Düzleştirme Yöntemi ile yapılan tahminlerde ise 3,5 ve 7. Örneklem kümeleri için $\hat{p}(H,S)$ yaklaşımının daha iyi sonuçlar verdiği, ortalamada ise $\tilde{p}_{KMH}$ yaklaşımı daha iyi sonuçlar verdiği görülmektedir.

Çizelge 4.36'da $n = 250$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelerinin 10 örneklem kümesine göre hesap edilen hata yüzdeleri ile benzer olduğu ve p 'nin önerilen yaklaşım olan $\tilde{p}_{KMH}$ ile diğer üç yöntemle göre daha iyi tahmin edildiği söylenebilir.

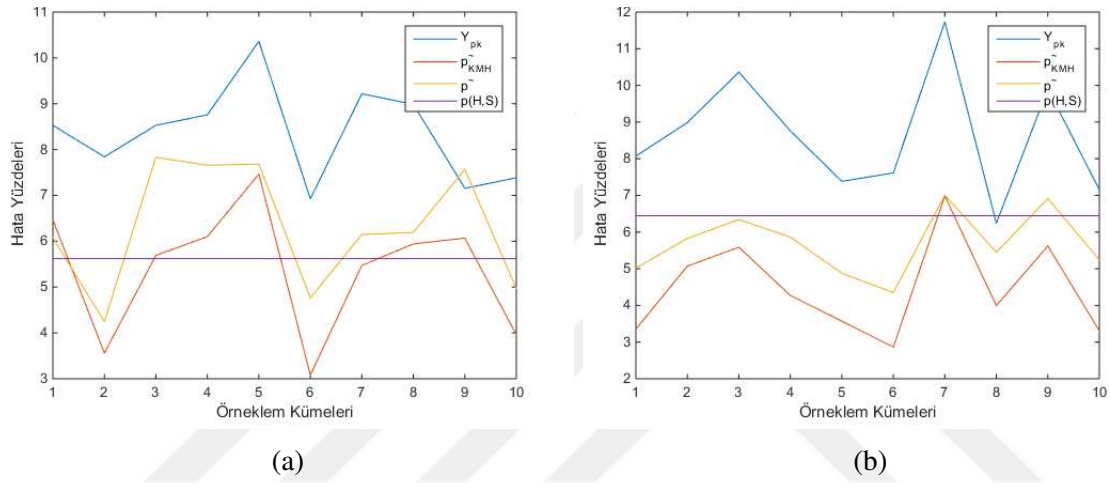
Çizelge 4.37. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 500$, $\ddot{o}ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
1	0,9480	0,9300	0,9266	0,8244	%8,529	%6,468	%6,079	%5,621	0,9440	0,9028	0,9173	0,8172	%8,071	%3,354	%5,014	%6,445
2	0,9420	0,9046	0,9106	0,8244	%7,842	%3,560	%4,247	%5,621	0,9520	0,9178	0,9244	0,8172	%8,987	%5,072	%5,827	%6,445
3	0,9480	0,9232	0,9419	0,8244	%8,529	%5,690	%7,831	%5,621	0,9640	0,9223	0,9289	0,8172	%10,361	%5,587	%6,342	%6,445
4	0,9500	0,9268	0,9404	0,8244	%8,758	%6,102	%7,659	%5,621	0,9500	0,9108	0,9247	0,8172	%8,758	%4,270	%5,861	%6,445
5	0,9640	0,9387	0,9406	0,8244	%10,361	%7,464	%7,682	%5,621	0,9380	0,9047	0,9161	0,8172	%7,384	%3,572	%4,877	%6,445
6	0,9340	0,9004	0,9151	0,8244	%6,926	%3,080	%4,762	%5,621	0,9400	0,8985	0,9115	0,8172	%7,613	%2,862	%4,350	%6,445
7	0,9540	0,9213	0,9272	0,8244	%9,216	%5,472	%6,148	%5,621	0,9760	0,9345	0,9347	0,8172	%11,734	%6,983	%7,006	%6,445
8	0,9520	0,9254	0,9276	0,8244	%8,987	%5,942	%6,193	%5,621	0,9280	0,9084	0,9211	0,8172	%6,239	%3,995	%5,449	%6,445
9	0,9360	0,9265	0,9396	0,8244	%7,155	%6,068	%7,567	%5,621	0,9600	0,9226	0,9339	0,8172	%9,903	%5,621	%6,915	%6,445
10	0,9380	0,9081	0,9169	0,8244	%7,384	%3,961	%4,969	%5,621	0,9360	0,9024	0,9193	0,8172	%7,155	%3,309	%5,243	%6,445
ORTALAMA	0,9466	0,9205	0,9287	0,8244	%8,369	%5,381	%6,314	%5,621	0,9488	0,9125	0,9232	0,8172	%8,621	%4,463	%5,688	%6,445
VARYANS	0,0001	0,0001	0,0001	0,0000	%0,011	%0,020	%0,018	%0,000	0,0002	0,0001	0,0001	0,0000	%0,028	%0,017	%0,008	%0,000

Çizelge 4.38. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 500$, $\ddot{o}ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	εY_{pk}	$\varepsilon \tilde{P}_{KMH}$	$\varepsilon \tilde{P}$	$\varepsilon \hat{P}(H, S)$
ORTALAMA	0,9513	0,9226	0,9282	0,8244	%8,907	%5,616	%6,262	%5,621	0,9499	0,9168	0,9277	0,8172	%8,746	%4,958	%6,202	%6,445
VARYANS	0,0001	0,0002	0,0001	0,0000	%0,013	%0,020	%0,019	%0,000	0,0001	0,0002	0,0002	0,0000	%0,014	%0,023	%0,020	%0,000

Çizelge 4.37'den $n = 500$ gözlem sayısına sahip 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelерinin ($\% \varepsilon$) görülmektedir. Burada Referans Bant Genişliği ile yapılan tahminlerden her ne kadar örneklem kümelerinde $\tilde{p}_{KMH}$ yaklaşımı ile birlikte $\hat{p}(H,S)$ yaklaşımının da iyi sonuçlar verdiği gözlemlenmiş olsa da örneklem kümeleri ortalamaları incelendiğinde $\tilde{p}_{KMH}$ yaklaşımının her iki bant genişliği tahmin yöntemi için de diğer yaklaşımlardan iyi sonuç verdiği görülmektedir. Maksimum Düzleştirme Yöntemi ile yapılan tahminlerde $\tilde{p}_{KMH}$ yaklaşımının daha iyi sonuçlar verdiği görülmektedir.



Şekil 4.19. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 500, öks = 10$)

Şekil 4.19'da görülen ilk grafikte 2, 6, 7 ve 10. Örneklem kümelerinde $\tilde{p}_{KMH}$ yaklaşımının diğer yaklaşımlara göre $\tilde{p}$ 'yi daha iyi tahmin ettiği görülmektedir. İkinci grafikte ise 7. Örneklem kümesi hariç, tüm örneklem kümeleri için $\tilde{p}_{KMH}$ yaklaşımının iyi sonuçlar verdiği ve hata yüzde seviyesinin diğerlerine göre aşağıda olduğu görülmektedir.

Çizelge 4.38'de $n = 500$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelерinin 10 örneklem kümesine göre hesap edilen hata yüzdeleri ile benzer olduğu ve $\tilde{p}$ 'nin önerilen yaklaşım olan $\tilde{p}_{KMH}$ ile diğer üç yöntemle göre daha iyi tahmin edildiği söylenebilir.

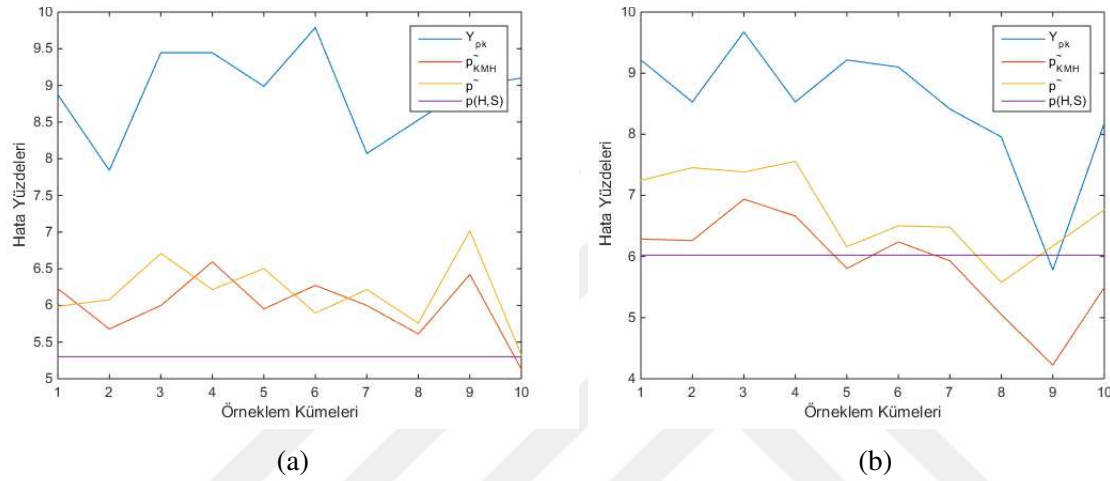
Çizelge 4.39. Senaryo IV için tahmin değerleri ve hata yüzdeleri ($n = 1000$, $\delta ks = 10$)

ÖRNEKLEM KÜMELERİ	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
1	0,9510	0,9279	0,9258	0,8272	%8,872	%6,228	%5,987	%5,301	0,9540	0,9284	0,9368	0,8209	%9,216	%6,285	%7,247	%6,022
2	0,9420	0,9231	0,9266	0,8272	%7,842	%5,678	%6,079	%5,301	0,9480	0,9282	0,9386	0,8209	%8,529	%6,262	%7,453	%6,022
3	0,9560	0,9259	0,9321	0,8272	%9,445	%5,999	%6,709	%5,301	0,9580	0,9341	0,9380	0,8209	%9,674	%6,938	%7,384	%6,022
4	0,9560	0,9311	0,9278	0,8272	%9,445	%6,594	%6,216	%5,301	0,9480	0,9317	0,9395	0,8209	%8,529	%6,663	%7,556	%6,022
5	0,9520	0,9255	0,9303	0,8272	%8,987	%5,953	%6,503	%5,301	0,9540	0,9242	0,9273	0,8209	%9,216	%5,804	%6,159	%6,022
6	0,9590	0,9283	0,9250	0,8272	%9,788	%6,274	%5,896	%5,301	0,9530	0,9280	0,9303	0,8209	%9,101	%6,239	%6,503	%6,022
7	0,9440	0,9259	0,9278	0,8272	%8,071	%5,999	%6,216	%5,301	0,9470	0,9253	0,9301	0,8209	%8,414	%5,930	%6,480	%6,022
8	0,9480	0,9225	0,9238	0,8272	%8,529	%5,610	%5,758	%5,301	0,9430	0,9176	0,9222	0,8209	%7,956	%5,049	%5,575	%6,022
9	0,9520	0,9296	0,9348	0,8272	%8,987	%6,422	%7,018	%5,301	0,9240	0,9104	0,9274	0,8209	%5,781	%4,224	%6,171	%6,022
10	0,9530	0,9183	0,9200	0,8272	%9,101	%5,129	%5,323	%5,301	0,9450	0,9215	0,9326	0,8209	%8,185	%5,495	%6,766	%6,022
ORTALAMA	0,9513	0,9258	0,9274	0,8272	%8,907	%5,989	%6,171	%5,301	0,9474	0,9249	0,9323	0,8209	%8,460	%5,889	%6,729	%6,022
VARYANS	0,0000	0,0000	0,0000	0,0000	%0,004	%0,002	%0,002	%0,000	0,0001	0,0000	0,0000	0,0000	%0,012	%0,006	%0,004	%0,000

Çizelge 4.40. Senaryo IV için tahmin değerleri ve hata yüzdelerinin ortalama ve varyansları ($n = 1000$, $\delta ks = 20$)

ORTALAMA DEĞERLER	REFERANS BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ								MAKSİMUM DÜZLEŞTİRME YÖNTEMİ BANT GENİŞLİĞİ İLE YAPILAN TAHMİN DEĞERLERİ VE HATA YÜZDELERİ							
	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$	Y_{pk}	$\tilde{P}_{KMH}$	$\tilde{P}$	$\hat{P}(H, S)$	ϵY_{pk}	$\epsilon \tilde{P}_{KMH}$	$\epsilon \tilde{P}$	$\epsilon \hat{P}(H, S)$
ORTALAMA	0,9510	0,9268	0,9276	0,8272	%8,872	%6,105	%6,199	%5,301	0,9496	0,9243	0,9313	0,8209	%8,706	%5,818	%6,615	%6,022
VARYANS	0,0000	0,0000	0,0000	0,0000	%0,005	%0,003	%0,004	%0,000	0,0001	0,0000	0,0000	0,0000	%0,007	%0,005	%0,005	%0,000

Çizelge 4.39’da $n = 1000$ gözlem sayısına sahip 10 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelерinin ($\% \varepsilon$) ortalama ile varyans değeri görülmektedir. Burada Referans Bant Genişliği ile yapılan tahmin değeri ile yaklaşımlar üzerinden hesaplanan hata yüzdelерinin 10. Örneklem kümesi dışında $\hat{p}(H,S)$ yaklaşımı için iyi sonuçlar verdiği ve hata yüzde ortalamasının da yine $\hat{p}(H,S)$ için daha iyi olduğu görülmektedir. Ancak Maksimum Düzleştirme Yöntemi kullanıldığında $\tilde{p}_{KMH}$ yaklaşımının ortalamada daha iyi sonuçlar verdiği görülmektedir.



Şekil 4.20. Senaryo IV için hata yüzdeleri çizgi grafiği ($n = 1000$, $\text{öks} = 10$)

Şekil 4.20’de ilk grafik için $\hat{p}(H,S)$ ile $\tilde{p}_{KMH}$ yaklaşımından elde edilen hata yüzdelерinin çok yakın olduğu ve $\hat{p}(H,S)$ yaklaşımının diğer iki yaklaşıma göre p ’yi daha iyi tahmin ettiği görülmektedir. İkinci grafikte ise 7. Örnekleme kadar $\hat{p}(H,S)$ ve $\tilde{p}_{KMH}$ yaklaşım değeri birbiriyle yakın olduğu ve $\hat{p}(H,S)$ diğer üç yaklaşıma göre iyi sonuçlar verdiği görülmektedir. Ancak $\tilde{p}_{KMH}$ yaklaşımının ortalamada p ’yi tahmin etmede daha başarılı olduğu ve 9. Örneklem kümesi ile p ’nin en düşük hata yüzdesi tahmin edildiği görülmektedir.

Çizelge 4.40’da $n = 1000$ gözlem sayısına sahip 20 örneklem kümesinden hesaplanan Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ve $\hat{p}(H,S)$ tahminleri ve hata yüzdelерinin 10 örneklem kümesine göre hesap edilen hata yüzdeleri ile benzer olduğu Referans Bant Genişliği ile yapılan tahmin değeri ile yaklaşımlar üzerinden hesaplanan hata yüzdelерinin $\hat{p}(H,S)$ yaklaşımı için daha iyi sonuçlar verdiği söylenebilir. Ancak Maksimum Düzleştirme Yönteminde $\tilde{p}_{KMH}$ yaklaşımın ortalamada daha iyi sonuçlar verdiği görülmektedir.



5. ÖRNEKLEM KÜMELERİNİN KERNEL TAHMİN GRAFİKLERİ

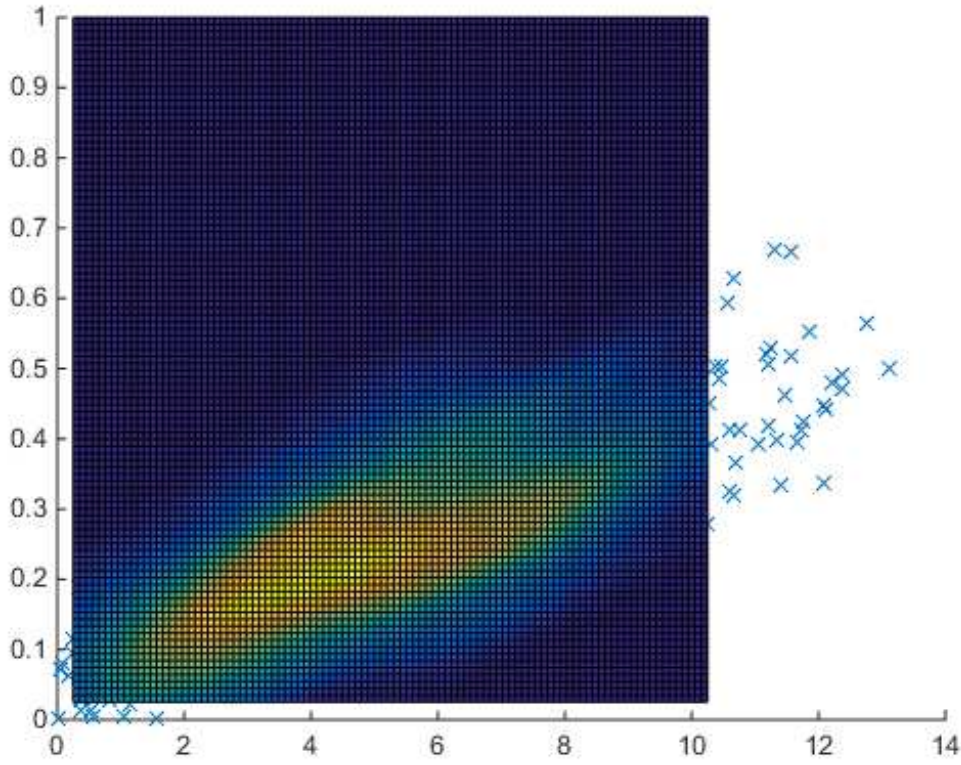
Çalışmanın 3. Bölümünde ayrıntılı olarak ele alınan bant genişliği matrisleri, Maksimum Düzleştirme Yöntemine göre elde edilmiş, ayrıca bant genişliği matrisinin dağılım bilgisine duyarlılığını değerlendirebilmek amacıyla ilk faz ürün veri setlerinin Normal dağılıma uygunluğu varsayımı ile Referans Bant Genişliği matris değerleri de tahmin edilerek Kernel tahmin grafikleri elde edilmiş olup Matlab programında yazılan kod dizileri Ek-1’de verilmiştir. Grafikler, tekrardan kaçınmak için sadece $n = 1000$ gözlemlili örneklem kümeleri için elde edilmiş ve yorumlanmıştır. Grafikler üzerinde görülen mavi çarpı (x) simgesi, $n = 1000$ gözlem sayısına sahip ilk ürün gözlem değerlerini; 100×100 hücreden oluşan koyu mavi kısım, süreç bölgesini ve renkli kısım ise gözlem değerlerinin spesifikasyon bölgesindeki yoğunluğunu göstermektedir. Spesifikasyon bölgesi dışına düşen gözlem noktaları (x) ile açıkça görülmektedir. Grafikler 4 farklı senaryo için elde edilmiş, bu grafikler oluşturulurken daha önce ayrıntılı olarak açıklanan bant genişliği matrisleri için iki farklı yöntem kullanılmış ve her bir grafik görsel anlamda daha iyi anlaşılabilmesi için hem iki boyutlu bir izdüşüm şeklinde, hem de spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünümü şeklinde çizdirilerek verilmeye çalışılmıştır. Bu nedenle, grafik sayısının fazla olması, her senaryo verisi üzerinden elde edilen örneklem kümelerinin grafiklerinin yorumu açısından benzer olması ve anlatım tekrarından kaçınmak amacıyla, her bir senaryo için ilk örneklem kümeleri üzerinden çizdirilen grafikler yorumlanmış, diğer örneklem kümelerine ait grafiklere çalışmada yer verilmemiştir.

5.1. Senaryo I İçin Tahmin Grafikleri

Senaryo I için çizdirilen grafikler, 1. örneklem kümesi için Şekil 5.1’den Şekil 5.4’e kadar verilen grafiklerden görülmektedir.

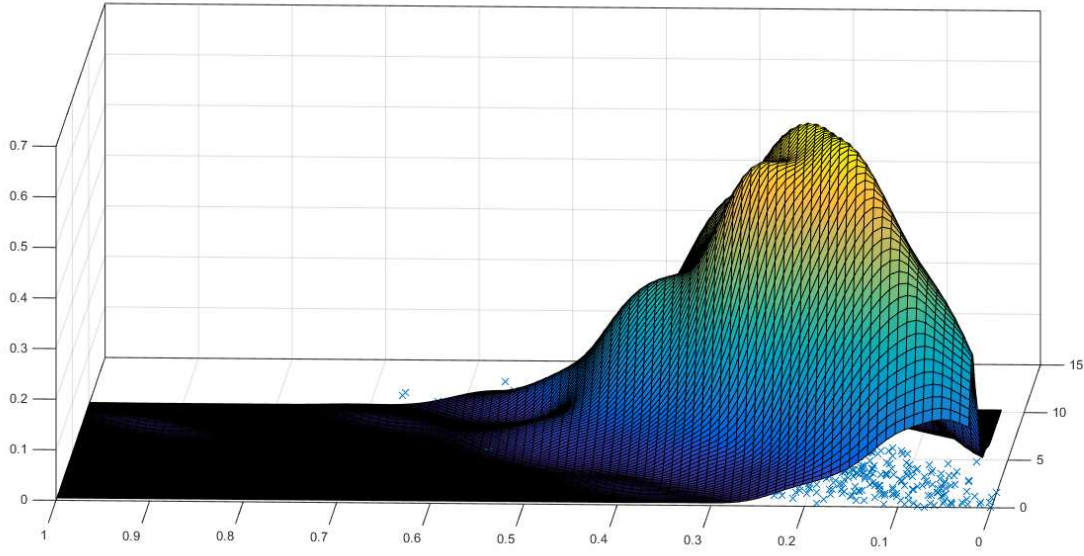
Referans Bant Genişliği ile Kernel yoğunluk tahmini

1. örneklem kümesi için referans bant genişliği matrisi kullanılarak elde edilen spesifikasyon bölgesine göre saçılım grafiği Şekil 5.1’de görülmektedir.



Şekil 5.1. Senaryo I, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (RBG)

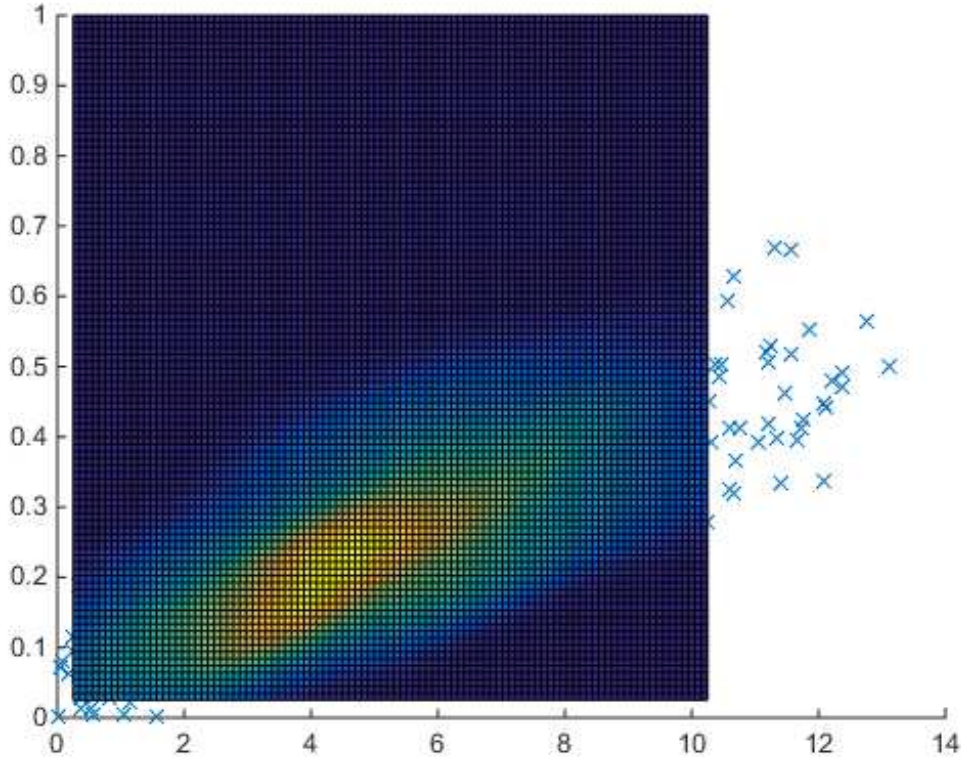
Sekil 5.1’de verilen grafikte, çalışmanın en başında belirlenen spesifikasyon sınırları içinde kalan bölgenin Kernel yoğunluk tahmin kontur sınırları, diğer bir ifade ile gözlem değerlerinin bant genişliği matrisine ve Kernel fonksiyon türüne göre düzgünleştirilerek tahmin edilen dağılımın şekli, iki boyutlu olarak görülmektedir. Ayrıca, spesifikasyon bölgesi dışına düşen gözlem noktaları (x) ile görülmektedir. Şekil 5.1’de verilen grafiğin farklı bir açıdan görüntüsü olan ve üç boyutlu uzayda Kernel yoğunluk tahmininin üst yüzey kesitinin görülmesine imkan veren dağılım tahmin grafiği Şekil 5.2’de verilmiştir.



Şekil 5.2. Senaryo I, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (RBG)

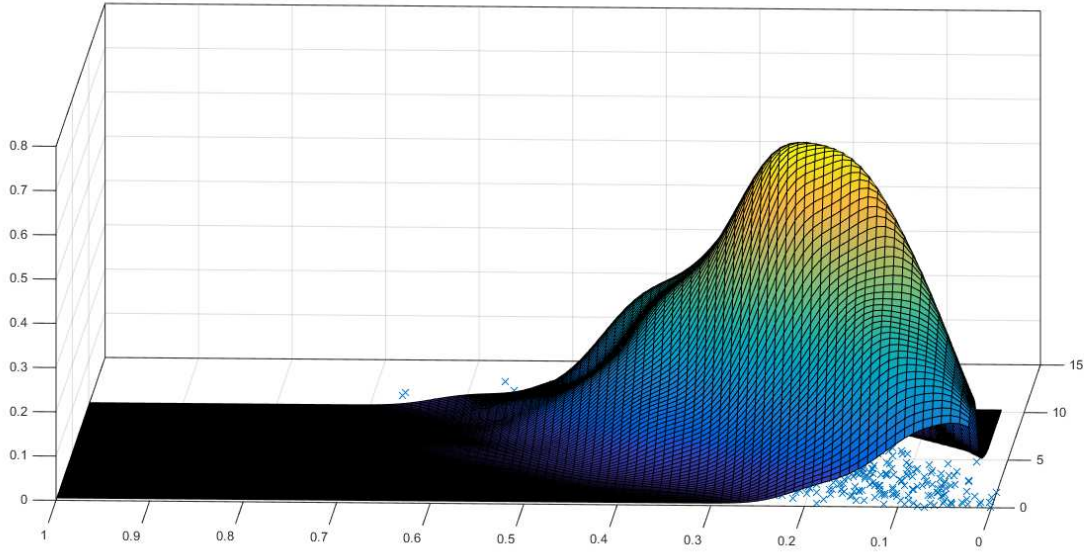
Şekil 5.2 ile verilen grafik, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünümünü vermekle birlikte, yine spesifikasyon bölgesi dışına düşen gözlemleri daha iyi bir şekilde görmemizi sağlamaktadır. 1. örneklem kümesi için referans bant genişliği matrisi kullanılarak elde edilen grafiklerle görsel açıdan farklılık oluşup oluşmadığını inceleyebilmek için Maksimum Düzleştirme Yöntemi bant genişliği ile elde edilen Kernel yoğunluk tahmin grafikleri Şekil 5.3 ve Şekil 5.4'te verilmiştir.

Maksimum Düzleştirme Yöntemi bant genişliği ile Kernel yoğunluk tahmin grafikleri



Şekil 5.3. Senaryo I, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (MDY)

Şekil 5.3'te 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği görülmektedir. Spesifikasyon bölgesi dışına düşen gözlem noktaları (x) ve tahmin kontur grafiği görülmektedir. Kernel tahmin grafiğinin farklı bir açıdan görüntüsü olan ve üç boyutlu uzayda Kernel yoğunluk tahmininin üst yüzey kesitinin görülmesine imkan veren dağılım tahmin grafiği Şekil 5.4'de verilmiştir.



Şekil 5.4. Senaryo I, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (MDY)

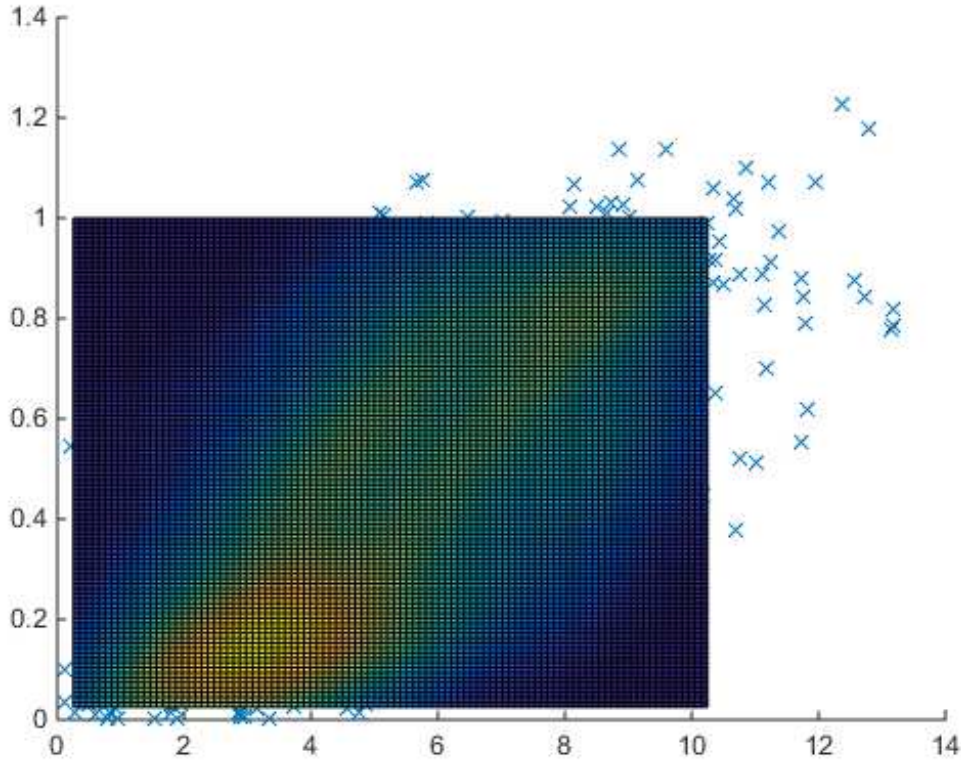
Şekil 5.4 ile verilen grafik, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünümünü vermekle birlikte, yine spesifikasyon bölgesi dışına düşen gözlemleri daha iyi bir şekilde görmemizi sağlamaktadır. 1. örneklem kümesi için elde edilen tahmin grafikleri hem Referans Bant Genişliği, hem de Maksimum Düzleştirme Yöntemi bant genişliği ile yapılan yoğunluk tahmin grafikleri açısından benzerlik göstermekte, ancak Maksimum Düzleştirme Yöntemi bant genişliği ile Kernel yoğunluk tahmin grafiklerinin daha pürüzsüz bir görüntü oluşturduğu, diğer bir ifade ile verilere maksimum düzgünleştirme uygulanarak, Kernel fonksiyon türüne daha yaklaşık bir yoğunluk tahmini elde etmemizi sağladığı söylenebilir.

5.2. Senaryo II İçin Tahmin Grafikleri

Senaryo II için çizdirilen grafikler, 1. örneklem kümesi için Şekil 5.5’den Şekil 5.8’e kadar verilen grafiklerden görülmektedir.

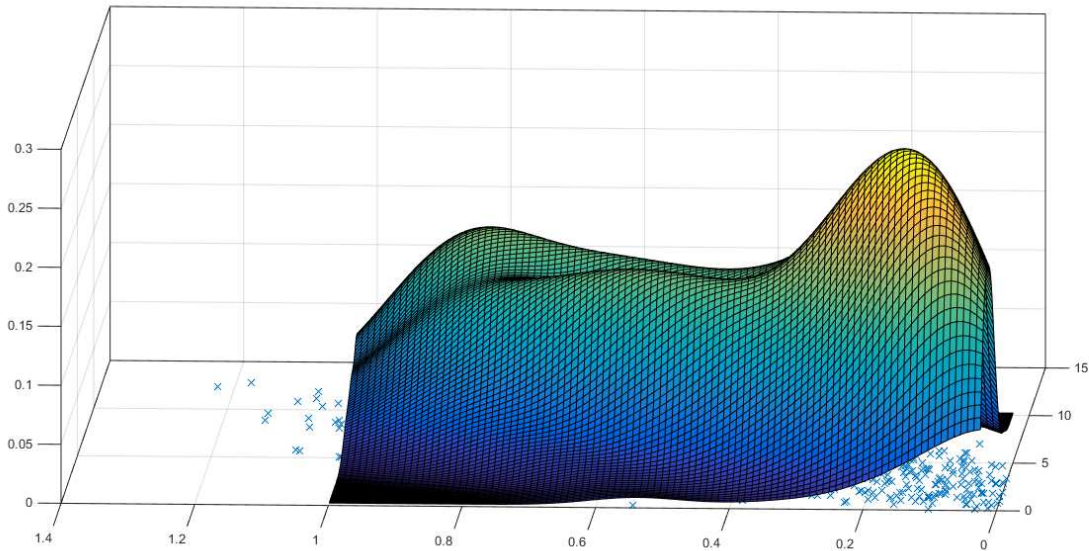
Referans bant genişliği ile Kernel yoğunluk tahmini

2. örneklem kümesi için Referans Bant Genişliği matrisi kullanılarak elde edilen spesifikasyon bölgesine göre saçılım grafiği Şekil 5.5’de görülmektedir.



Şekil 5.5. Senaryo II, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (RBG)

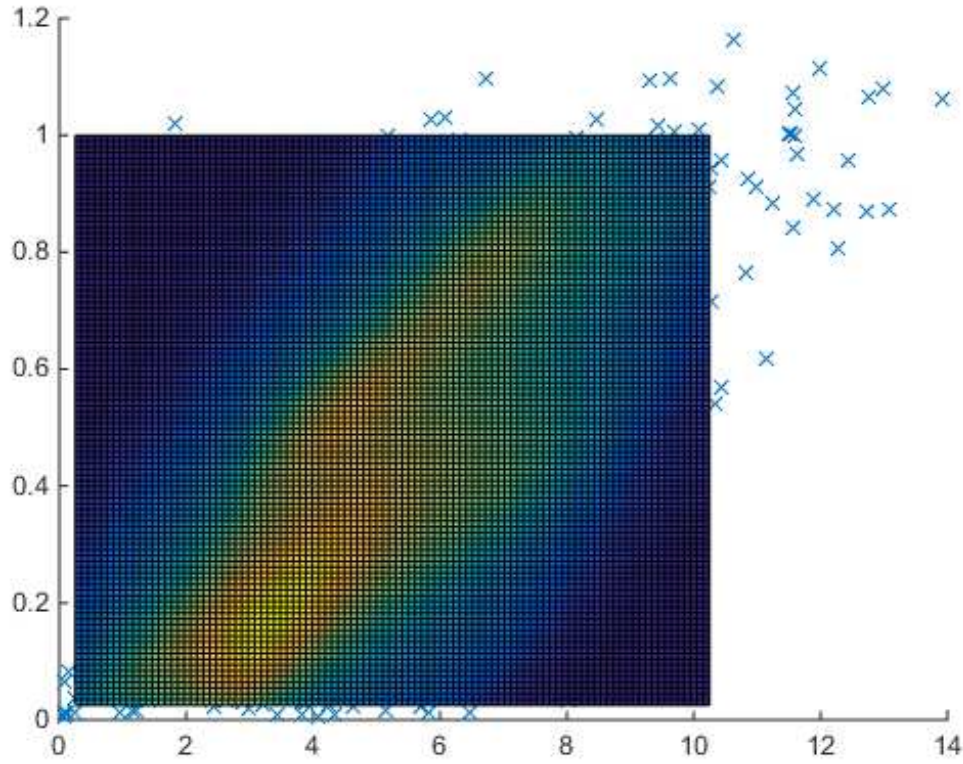
Şekil 5.5’de verilen grafikte, spesifikasyon bölgesi dışına düşen gözlem noktaları (x) ile görülmektedir. Şekil 5.5’de verilen grafiğin farklı bir açıdan görüntüsü olan ve üç boyutlu uzayda Kernel yoğunluk tahmininin üst yüzey kesitinin görülmesine imkan veren dağılım tahmin grafiği Şekil 5.6’da verilmiştir.



Şekil 5.6. Senaryo II, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (RBG)

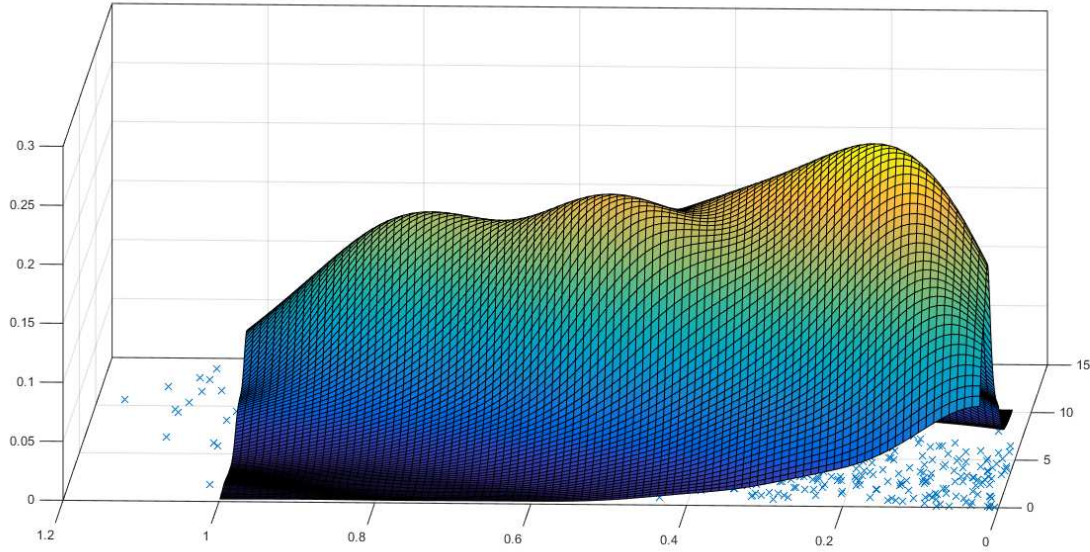
Şekil 5.6 ile verilen grafik, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünümünü vermekle birlikte, yine spesifikasyon bölgesi dışına düşen gözlemleri daha iyi bir şekilde görmemizi sağlamaktadır. 1. örneklem kümesi için referans bant genişliği matrisi kullanılarak elde edilen grafiklerle görsel açıdan farklılık oluşup oluşmadığını inceleyebilmek için Maksimum Düzleştirme Yöntemi bant genişliği ile elde edilen Kernel yoğunluk tahmin grafikleri Şekil 5.7 ve Şekil 5.8’da verilmiştir.

Maksimum düzleştirme yöntemi bant genişliği ile Kernel yoğunluk tahmin grafikleri



Şekil 5.7. Senaryo II, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (MDY)

Şekil 5.7’de 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği görülmektedir. Spesifikasyon bölgesi dışına düşen gözlem noktaları (x) ve tahmin kontur grafiği görülmektedir. Kernel tahmin grafiğinin farklı bir açıdan görüntüsü olan ve üç boyutlu uzayda Kernel yoğunluk tahmininin üst yüzey kesitinin görülmesine imkan veren dağılım tahmin grafiği Şekil 5.9’da verilmiştir.



Şekil 5.8. Senaryo II, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (MDY)

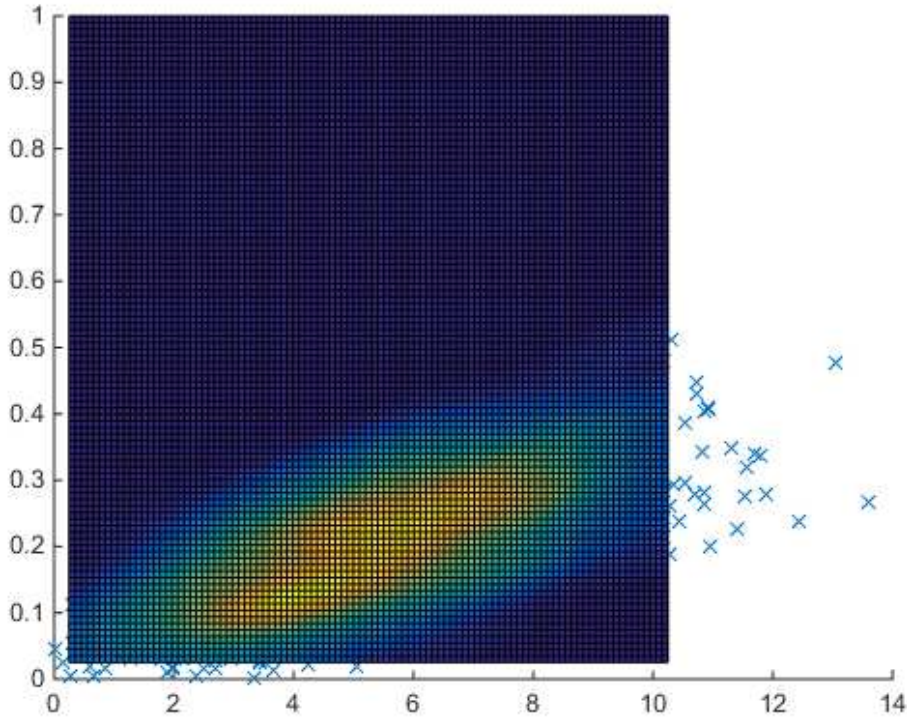
Şekil 5.8 ile verilen grafik, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünümünü vermekle birlikte, yine spesifikasyon bölgesi dışına düşen gözlemleri daha iyi bir şekilde görmemizi sağlamaktadır. 1. örneklem kümesi için elde edilen tahmin grafikleri hem Referans Bant Genişliği, hem de Maksimum Düzleştirme Yöntemi bant genişliği ile yapılan yoğunluk tahmin grafikleri açısından benzerlik göstermekte, ancak Maksimum Düzleştirme Yöntemi bant genişliği ile Kernel yoğunluk tahmin grafiklerinin daha pürüzsüz bir görüntü oluşturduğu, diğer bir ifade ile verilere maksimum düzgünleştirme sağlayarak, Kernel fonksiyon türüne daha yaklaşık bir yoğunluk tahmini elde etmemizi sağladığı görülmektedir.

5.3. Senaryo III İçin Tahmin Grafikleri

Senaryo III için çizdirilen grafikler, 10 farklı örneklem kümeleri için Şekil 5.9'dan Şekil 5.12'ye kadar verilen grafiklerden görülmektedir.

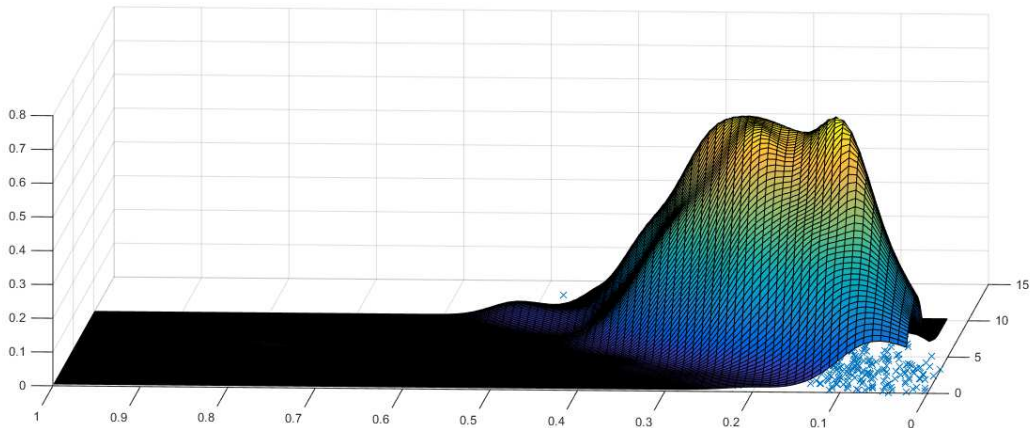
Referans bant genişliği ile Kernel yoğunluk tahmini

1. örneklem kümesi için referans bant genişliği matrisi kullanılarak elde edilen spesifikasyon bölgesine göre saçılım grafiği Şekil 5.9'da görülmektedir.

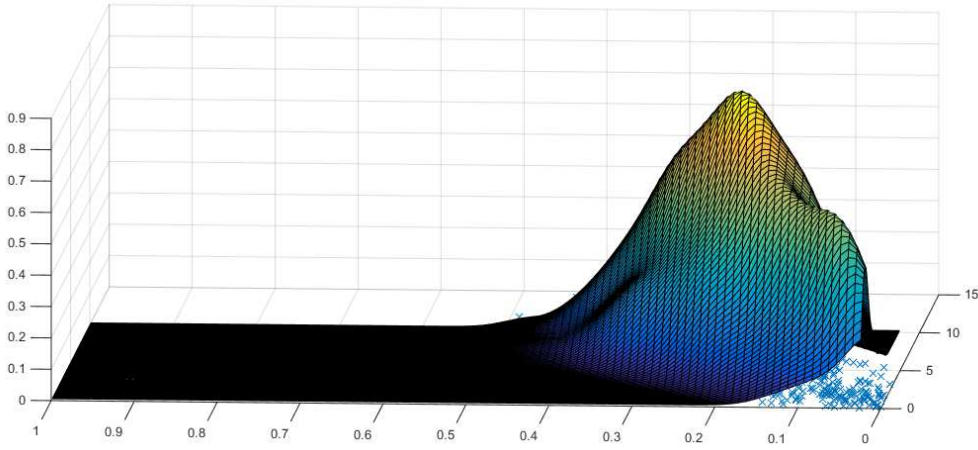


Şekil 5.9. Senaryo III, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (RBG)

Şekil 5.9’da verilen grafikte, çalışmanın en başında belirlenen spesifikasyon sınırları içinde kalan bölgenin Kernel yoğunluk tahmin kontur sınırları, diğer bir ifade ile gözlem değerlerinin bant genişliği matrisine ve Kernel fonksiyon türüne göre düzgünleştirilerek tahmin edilen dağılımın şekli iki boyutlu olarak görülmektedir. Ayrıca, spesifikasyon bölgesi dışına düşen gözlem noktaları (x) ile görülmektedir. Şekil 5.9’da verilen grafiğin farklı bir açıdan görüntüsü olan ve üç boyutlu uzayda Kernel yoğunluk tahmininin üst yüzey kesitinin görülmesine imkan veren dağılım tahmin grafiği Şekil 5.10’da verilmiştir.



Şekil 5.10. Senaryo III, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (RBG)



Şekil 5.12. Senaryo III, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (MDY)

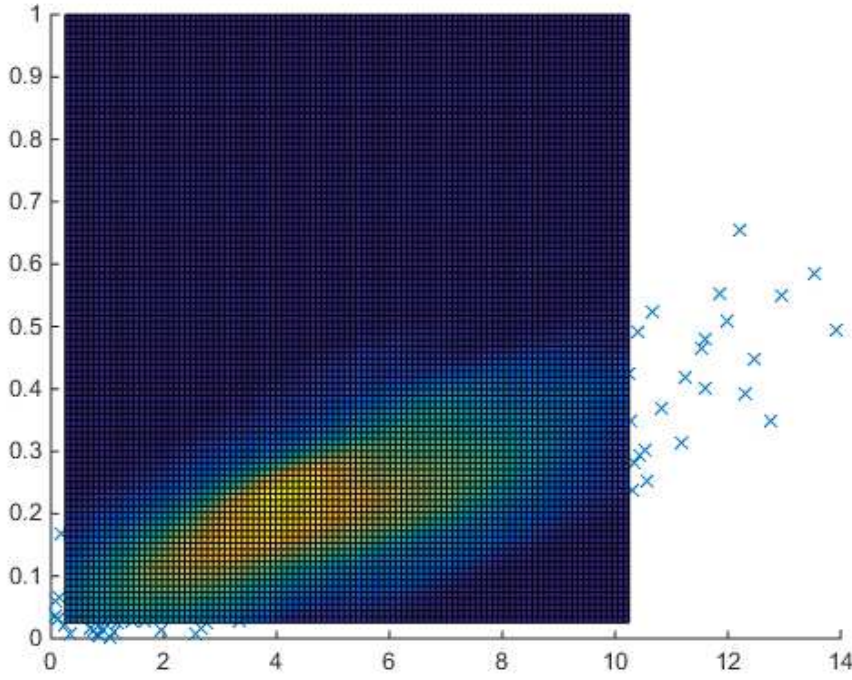
Şekil 5.12 ile verilen grafik, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünümünü vermekte ve spesifikasyon bölgesi dışına düşen gözlemleri daha iyi bir şekilde görmemizi sağlamaktadır. 1. örneklem kümesi için elde edilen tahmin grafikleri hem Referans Bant Genişliği, hem de Maksimum Düzleştirme Yöntemi bant genişliği ile yapılan yoğunluk tahmin grafikleri açısından benzerlik göstermekte, ancak Maksimum Düzleştirme Yöntemi bant genişliği ile Kernel yoğunluk tahmin grafiklerinin daha pürüzsüz bir görüntü oluşturduğu, diğer bir ifade ile verilere maksimum düzgünleştirme sağlayarak, Kernel fonksiyon türüne daha yaklaşık bir yoğunluk tahmini elde etmemizi sağladığı görülmektedir.

5.4. Senaryo IV İçin Tahmin Grafikleri

Senaryo IV için çizdirilen grafikler, 10 farklı örneklem kümeleri için Şekil 5.13'den Şekil 5.16'ya kadar verilen grafiklerden görülmektedir.

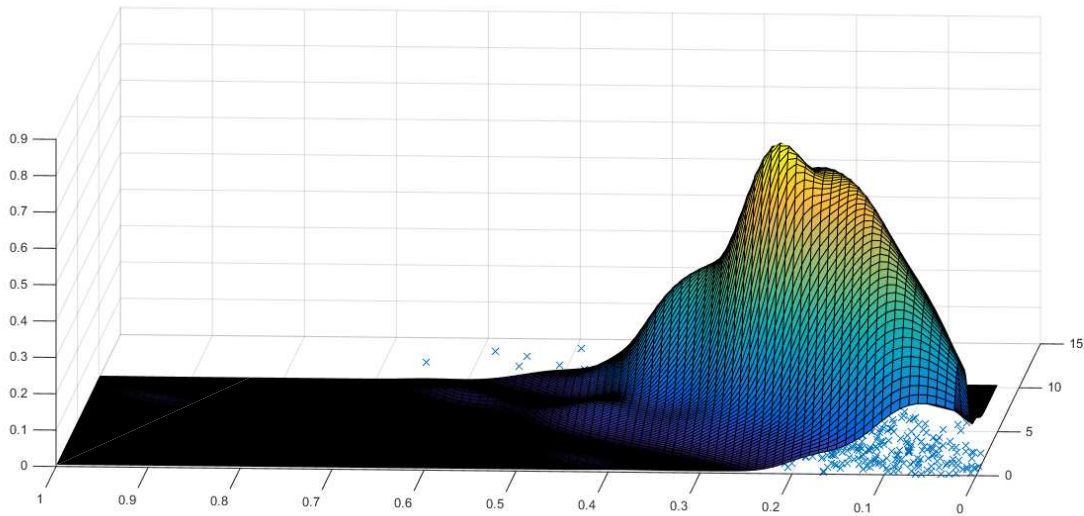
Referans bant genişliği ile Kernel yoğunluk tahmini

1. örneklem kümesi için Referans Bant Genişliği matrisi kullanılarak elde edilen spesifikasyon bölgesine göre saçılım grafiği Şekil 5.13'de görülmektedir.

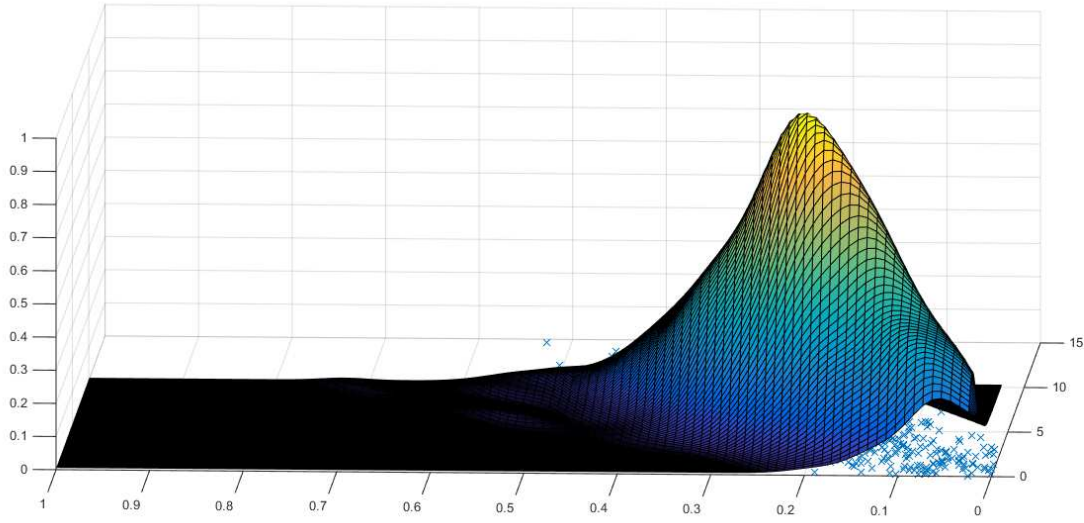


Şekil 5.13. Senaryo IV, 1. örneklem kümesinin spesifikasyon bölgesine göre saçılım grafiği (RBG)

Şekil 5.13’de verilen grafikte, gözlem değerlerinin bant genişliği matrisine ve Kernel fonksiyon türüne göre düzgünleştirilerek tahmin edilen dağılımın şekli iki boyutlu olarak görülmektedir. Ayrıca, spesifikasyon bölgesi dışına düşen gözlem noktaları (x) ile görülmektedir. Şekil 5.13’de verilen grafiğin farklı bir açıdan görüntüsü olan ve üç boyutlu uzayda Kernel yoğunluk tahmininin üst yüzey kesitinin görülmesine imkan veren dağılım tahmin grafiği Şekil 5.14’de verilmiştir.



Şekil 5.14. Senaryo IV, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (RBG)



Şekil 5.16. Senaryo IV, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünüm grafiği (MDY)

Şekil 5.16 ile verilen grafik, 1. örneklem kümesinin spesifikasyon bölgesine göre konumlanmasının yüzeysel açıdan görünümünü vermekle birlikte, yine spesifikasyon bölgesi dışına düşen gözlemleri daha iyi bir şekilde görmemizi sağlamaktadır. 1. örneklem kümesi için elde edilen tahmin grafikleri hem Referans Bant Genişliği, hem de Maksimum Düzleştirme Yöntemi bant genişliği ile yapılan yoğunluk tahmin grafikleri açısından benzerlik göstermektedir. Maksimum Düzleştirme Yöntemi bant genişliği ile Kernel yoğunluk tahmin grafiklerinin daha pürüzsüz bir görüntü sağlamaktadır.

6. SONUÇ VE ÖNERİLER

Üretimin devamlılığı açısından, üretilen mal ve hizmetin müşteri isteklerine ve beklentilerine cevap vermesinin sağlanması, üretim aşamasında uygun olmayan ürün sayısının azaltılması ve böylece üretim maliyetlerinin düşürülebileceği göz önünde bulundurulduğunda, süreç yeterlilik analizinin fabrika süreç problemlerinde ne denli önemli rol oynadığı anlaşılabılır.

Sürecin yeterli olduğuna karar vermede, süreç dağılımının Normal dağılıma uygun olması, klasik yeterlilik indekslerin doğru tahmin edilmesinde önemli bir rol oynamaktadır.

Üretim sürecinin Normal dağılıma uymadığı durumda, klasik yeterlilik indeksleri ile doğru tahminler elde edilemediği, yapılan literatür çalışmasıyla, varsayım bozulmaları sonucu doğru tahminlere ulaşmak için yazarlar tarafından çok çeşitli yaklaşımlar geliştirildiği görülmektedir. Örneğin, tek değişkenli üretim süreçleri için, normallik varsayımının sağlanmadığı durumlarda kullanılabilecek yaklaşımlardan ilki, normal dağılıma uygun olmayan veriyi normal dağılıma dönüştürmek ve daha sonra yeterlilik indeksleri için önerilen klasik yaklaşımları kullanmaktır. Bu amaçla en çok kullanılan dönüşüm teknikleri, Johnson (1949) ve Box-Cox (1964) Güç dönüşüm teknikleridir. İkinci yaklaşım ise, dağılımın bilinmediği durumlarda veriyi bir dağılıma uydurmak ve daha sonra normal olmayan yüzdelikler kullanılarak süreç yeterlilik indekslerini tahmin etmektir. Bu amaçla en çok kullanılan yöntemlere örnek olarak, Clements (1989) yüzde metodu ve Burr yüzde metodu verilebilir.

Çok değişkenli üretim süreçleri için normallik varsayımının sağlanmadığı durumlarda süreç yeterlilik indekslerinin hesaplanması, spesifikasyon bölgesi ve değişkenler arasındaki ilişki durumu da göz önünde bulundurulduğunda, daha güç ve karmaşık olabilmektedir. Normal Dağılıma uygun olmayan çok değişkenli süreç yeterlilik analizinde, yeterlilik indekslerini, az hata ile tahmin edebilmek ancak, çalışılan verinin dağılım bilgisine uygun tahmin yöntemleri kullanmakla mümkündür. Süreç dağılımı, bilinen çok değişkenli dağılımlara uygun olabileceği gibi, herhangi bilinen bir dağılıma uygun olmayan bir süreç de olabilir. Süreç dağılımının bilinmediği durumlarda, yoğunluk tahmini için kullanılan en yaygın yöntem Kernel Yoğunluk Tahmin yöntemidir. Kernel Yoğunluk Tahmin yöntemi, X rastgele

değişkenin olasılık yoğunluk fonksiyonu $f(X)$ 'i tahmin etmek için kullanılan parametrik olmayan bir tahmin yöntemidir. Genellikle süreçten elde edilen gözlem verisinin dağılımı bilinmediği ya da bilinen bir dağılıma uygun olmadığı durumlarda gözleme dayalı olarak Kernel yoğunluk tahminleri elde edilebilir.

Süreç dağılımının Normal dağılıma yakınsaması için gerekli örneklem büyüklüğü ile çalışmak ve bilinen Merkezi Limit Teoreminden yola çıkarak, klasik yeterlilik indekslerini kullanmak da mümkündür. Ancak, Normal Dağılım yaklaşımı için gerekli örnek sayısının büyük olması halinde, örneklem çekme işlemi maliyetli ya da imkansız olabilmektedir. Bu aşamada örnekleme yöntemlerinden faydalanmak kaçınılmazdır. Bu nedenle, bu tez çalışmasında çarpık dağılıma sahip süreçler dikkate alınarak t-Copula yöntemi ile üretilen Senaryo verilerine, Metropolis-Hastings (MH) örnekleme uygulanmıştır.

MH örnekleme kullanılarak süreç dağılımı ile aynı dağılım özelliklerine sahip örneklem büyüklüğü $n = 50, 100, 250, 500$ ve 1000 olan 10 'a ve 20 'şer tane farklı örneklem kümesi elde edilmiştir. Daha önce tanımlarına ve özelliklerine ayrıntılı şekilde yer verilen çok değişkenli yeterlilik indeksler Y_{pk} , $\tilde{p}_{KMH}$, $\tilde{p}$ ile tez çalışmasında önerilen $\tilde{p}_{KMH}$ indeksinin, uygun ürün oranı p 'ye ilişkin tahminlerde, örneklem büyüklüğünden (n) ve örnekleme kümesi sayısından nasıl etkilendiği incelenmiştir.

Sonuçlar özetlenecek olursa, Senaryo I'de örneklem büyüklüğü $n = 50$ için elde edilen hata yüzdeleri, Referans Bant Genişliği (RBG) ve Maksimum Düzleştirme Yöntemi (MDY) ile yapılan ortalama tahminler için benzerlik göstermekle birlikte $\tilde{p}_{KMH}$ yaklaşımı ile uygun ürün oranı p 'nin daha az hata yüzdesi ile tahmin edildiği söylenebilir. 10 tane örneklem kümesinden hesaplanan RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%3,799$, $\varepsilon Y_{pk} = \%8,970$, $\varepsilon\tilde{p} = \%4,814$ ve $\varepsilon\hat{p}(H,S) = \%10,095$ olarak elde edilmiştir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%3,436$, $\varepsilon Y_{pk} = \%8,742$, $\varepsilon\tilde{p} = \%4,756$ ve $\varepsilon\hat{p}(H,S) = \%11,566$ olarak hesaplanmıştır. Her iki yöntemde de $n = 50$ büyüklüğü için hata yüzdeleri incelendiğinde, $\tilde{p}_{KMH}$ 'nin diğer yöntemlere göre daha düşük bir hata yüzdesine sahip olduğu görülmektedir. 20 örneklem kümesinden elde edilen sonuçlar için benzer tahminler elde edilmiştir.

$n = 100$ için RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%3,579$, $\varepsilon Y_{pk} = \%8,475$,

$\varepsilon\tilde{p} = \%5,541$ ve $\varepsilon\hat{p}(H,S) = \%6,331$ olarak elde edilmiştir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%3,259$, $\varepsilon Y_{pk} = \%8,019$, $\varepsilon\tilde{p} = \%5,303$ ve $\varepsilon\hat{p}(H,S) = \%7,608$ olarak elde edilmiştir. Her iki yöntemde de $n = 100$ büyüklüğü için hata yüzdeleri incelendiğinde, $\tilde{p}_{KMH}$ 'nin diğer yöntemlere göre daha düşük bir hata yüzdesi ile tahmin sağladığı söylenebilir. 20 örneklem kümesinden hesaplanan tahmin sonuçları benzerdir.

$n = 250$ için RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%5,841$, $\varepsilon Y_{pk} = \%9,737$, $\varepsilon\tilde{p} = \%7,261$ ve $\varepsilon\hat{p}(H,S) = \%8,281$ olarak bulunmuştur. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%3,926$, $\varepsilon Y_{pk} = \%8,133$, $\varepsilon\tilde{p} = \%6,124$ ve $\varepsilon\hat{p}(H,S) = \%9,296$ olarak elde edilmiştir. $n = 250$ için $\tilde{p}_{KMH}$ indeksinin her iki yöntemde de düşük hata yüzdesine sahip olduğu görülmekle birlikte MDY'ne göre daha iyi sonuç verdiği söylenebilir. 20 örneklem kümesinden hesaplanan tahminler için de aynı yorumlar yapılabilir.

$n = 500$ için RBG ile yapılan tahminlerden hesaplanan ortalama hata yüzdeleri $\varepsilon\tilde{p}_{KMH} = \%5,517$, $\varepsilon Y_{pk} = \%8,498$, $\varepsilon\tilde{p} = \%6,239$ ve $\varepsilon\hat{p}(H,S) = \%3,239$ olarak bulunmuştur. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%5,050$, $\varepsilon Y_{pk} = \%8,977$, $\tilde{p} = \%6,110$ ve $\varepsilon\hat{p}(H,S) = \%4,136$ olarak elde edilmiştir. $n = 500$ için $\hat{p}(H,S)$ indeksinin her iki yöntemde de düşük hata yüzdesine sahip olduğu, $\tilde{p}_{KMH}$ indeksinin MDY'ne göre $\hat{p}(H,S)$ yaklaşımına yakın sonuçlar verdiği söylenebilir. 20 örneklem kümesinden hesaplanan tahminler için de $\hat{p}(H,S)$ yaklaşımının p 'yi tahmin etmede daha başarılı olduğu söylenebilir.

$n = 1000$ için RBG ile yapılan tahminlerden hesaplanan ortalama hata yüzdeleri $\varepsilon\tilde{p}_{KMH} = \%5,977$, $\varepsilon Y_{pk} = \%8,555$, $\varepsilon\tilde{p} = \%6,116$ ve $\varepsilon\hat{p}(H,S) = \%4,391$ olarak bulunmuştur. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%4,905$, $\varepsilon Y_{pk} = \%7,779$, $\varepsilon\tilde{p} = \%5,973$ ve $\varepsilon\hat{p}(H,S) = \%5,075$ olarak elde edilmiştir. $n = 1000$ için $\hat{p}(H,S)$ indeksinin RBG ile yapılan tahminlerde daha düşük hata yüzdesine sahip olduğu, $\tilde{p}_{KMH}$ indeksinin ise MDY ile p 'yi daha iyi tahmin ettiği görülse de $\hat{p}(H,S)$ yaklaşımına yakın sonuçlar verdiği söylenebilir. 20 örneklem kümesinden hesaplanan tahminlerde ise $\hat{p}(H,S)$ yaklaşımının her iki yöntem için p 'yi tahmin etmede daha başarılı olduğu söylenebilir.

Senaryo II için sonuçlar özetlenecek olursa, örneklem büyüklüğü $n = 50$ için elde edilen

hata yüzdeleri, RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%3,639$, $\varepsilon Y_{pk} = \%16,881$, $\varepsilon\tilde{p} = \%8,884$ ve $\varepsilon\hat{p}(H,S) = \%18,527$ olarak elde edilmiştir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%2,637$, $\varepsilon Y_{pk} = \%16,116$, $\varepsilon\tilde{p} = \%8,818$ ve $\varepsilon\hat{p}(H,S) = \%20,403$ olarak hesaplanmıştır. Her iki yöntemde de $n = 50$ büyüklüğü için hata yüzdeleri incelendiğinde, $\tilde{p}_{KMH}$ 'nin diğer yöntemlere göre çok düşük bir hata yüzdesine sahip olduğu görülmektedir. 20 örneklem kümesinden elde edilen sonuçlar için benzer tahminler elde edilmiştir.

$n = 100$ için RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%5,012$, $\varepsilon Y_{pk} = \%14,840$, $\varepsilon\tilde{p} = \%9,789$ ve $\varepsilon\hat{p}(H,S) = \%17,098$ olarak elde edilmiştir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%3,893$, $\varepsilon Y_{pk} = \%14,457$, $\varepsilon\tilde{p} = \%9,565$ ve $\varepsilon\hat{p}(H,S) = \%18,515$ olarak elde edilmiştir. Her iki yöntemde de $n = 100$ büyüklüğü için hata yüzdeleri incelendiğinde, $\tilde{p}_{KMH}$ 'in diğer yaklaşımlara göre daha düşük bir hata yüzdesi ile tahmin sağladığı söylenebilir. 20 örneklem kümesinden hesaplanan tahmin sonuçları aynı şekilde yorumlanabilir.

$n = 250$ için RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%6,542$, $\varepsilon Y_{pk} = \%15,197$, $\varepsilon\tilde{p} = \%10,438$ ve $\varepsilon\hat{p}(H,S) = \%12,759$ olarak bulunmuştur. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%6,277$, $\varepsilon Y_{pk} = \%17,443$, $\varepsilon\tilde{p} = \%11,099$ ve $\varepsilon\hat{p}(H,S) = \%14,138$ olarak elde edilmiştir. $n = 250$ için $\tilde{p}_{KMH}$ indeksinin her iki yöntemde de düşük hata yüzdesine sahip olduğu görülmektedir. 20 örneklem kümesinden hesaplanan tahminler incelendiğinde, yine $\tilde{p}_{KMH}$ indeksinin diğer yöntemlere göre daha iyi sonuç verdiği söylenebilir.

$n = 500$ için RBG ile yapılan tahminlerden hesaplanan ortalama hata yüzdeleri $\varepsilon\tilde{p}_{KMH} = \%8,510$, $\varepsilon Y_{pk} = \%15,401$, $\varepsilon\tilde{p} = \%10,879$ ve $\varepsilon\hat{p}(H,S) = \%9,518$ olarak hesaplanmıştır. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%6,391$, $\varepsilon Y_{pk} = \%15,044$, $\varepsilon\tilde{p} = \%10,136$ ve $\varepsilon\hat{p}(H,S) = \%10,616$ olarak elde edilmiştir. $n = 500$ için $\tilde{p}_{KMH}$ indeksinin her iki yöntemde de düşük hata yüzdesine sahip olduğu görülmektedir. 20 örneklem kümesinden hesaplanan tahminler için de $\tilde{p}_{KMH}$ yaklaşımının p 'yi tahmin etmede daha başarılı olduğu söylenebilir.

$n = 1000$ için RBG ile yapılan tahminlerden hesaplanan ortalama hata yüzdeleri

$\varepsilon\tilde{p}_{KMH} = \%9,493$, $\varepsilon Y_{pk} = \%16,409$, $\varepsilon\tilde{p} = \%10,676$ ve $\varepsilon\hat{p}(H,S) = \%9,213$ olarak bulunmuştur. Hata yüzdeleri incelendiğinde, $\hat{p}(H,S)$ yaklaşımının daha iyi sonuç verdiği görülse de $\tilde{p}_{KMH}$ yaklaşımının hata yüzdesi ile çok yakın olduğu görülmektedir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%7,803$, $\varepsilon Y_{pk} = \%15,325$, $\varepsilon\tilde{p} = \%10,084$ ve $\varepsilon\hat{p}(H,S) = \%10,195$ olarak elde edilmiştir. $n = 1000$ için $\hat{p}(H,S)$ indeksinin RBG ile yapılan tahminlerde daha düşük hata yüzdesine sahip olduğu, $\tilde{p}_{KMH}$ indeksinin ise MDY ile p' 'yi daha iyi tahmin ettiği görülmektedir. 20 örneklem kümesinden hesaplanan tahminlerde ise $\tilde{p}_{KMH}$ yaklaşımının her iki yöntem için p' 'yi tahmin etmede daha başarılı olduğu söylenebilir.

Senaryo III için sonuçlar özetlenecek olursa, örneklem büyüklüğü $n = 50$ için elde edilen hata yüzdeleri, RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%4,145$, $\varepsilon Y_{pk} = \%8,121$, $\varepsilon\tilde{p} = \%5,457$ ve $\varepsilon\hat{p}(H,S) = \%9,547$ olarak elde edilmiştir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%2,514$, $\varepsilon Y_{pk} = \%9,731$, $\varepsilon\tilde{p} = \%4,930$ ve $\varepsilon\hat{p}(H,S) = \%11,100$ olarak hesaplanmıştır. Her iki yöntemde de $n = 50$ büyüklüğü için hata yüzdeleri incelendiğinde, $\tilde{p}_{KMH}$ 'nin diğer yöntemlere göre çok düşük bir hata yüzdesine sahip olduğu görülmektedir. 20 örneklem kümesinden elde edilen sonuçlar için benzerdir. $\tilde{p}_{KMH}$ diğer yöntemlere daha iyi sonuç vermektedir.

$n = 100$ için RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%3,751$, $\varepsilon Y_{pk} = \%8,751$, $\varepsilon\tilde{p} = \%5,965$ ve $\varepsilon\hat{p}(H,S) = \%6,993$ olarak elde edilmiştir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%3,907$, $\varepsilon Y_{pk} = \%8,926$, $\varepsilon\tilde{p} = \%5,841$ ve $\varepsilon\hat{p}(H,S) = \%8,040$ olarak elde edilmiştir. Her iki yöntemde de $n = 100$ büyüklüğü için hata yüzdeleri incelendiğinde, $\tilde{p}_{KMH}$ yaklaşımının hata yüzdesinin diğer yaklaşımlara göre daha düşük olduğu söylenebilir. 20 örneklem kümesinden hesaplanan tahmin sonuçları aynı şekilde yorumlanabilir.

$n = 250$ için RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%5,132$, $\varepsilon Y_{pk} = \%8,673$, $\varepsilon\tilde{p} = \%6,666$ ve $\varepsilon\hat{p}(H,S) = \%8,661$ olarak bulunmuştur. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%4,408$, $\varepsilon Y_{pk} = \%8,811$, $\varepsilon\tilde{p} = \%6,518$ ve $\varepsilon\hat{p}(H,S) = \%9,742$ olarak elde edilmiştir. $n = 250$ için $\tilde{p}_{KMH}$ indeksinin her iki yöntemde de düşük hata yüzdesine sahip olduğu görülmektedir. 20 örneklem kümesinden hesaplanan

tahminler incelendiğinde, yine $\tilde{p}_{KMH}$ indeksinin diğer yöntemlere göre daha iyi sonuç verdiği söylenebilir.

$n = 500$ için RBG ile yapılan tahminlerden hesaplanan ortalama hata yüzdeleri $\epsilon\tilde{p}_{KMH} = \%5,618$, $\epsilon Y_{pk} = \%8,949$, $\epsilon\tilde{p} = \%6,356$ ve $\epsilon\hat{p}(H,S) = \%3,508$ olarak hesaplanmıştır. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\epsilon\tilde{p}_{KMH} = \%4,669$, $\epsilon Y_{pk} = \%8,765$, $\epsilon\tilde{p} = \%6,399$ ve $\epsilon\hat{p}(H,S) = \%4,417$ olarak elde edilmiştir. $n = 500$ için $\hat{p}(H,S)$ yaklaşımının her iki yöntemde de düşük hata yüzdesine sahip olduğu görülmektedir. MDY için $\hat{p}(H,S)$ ve $\tilde{p}_{KMH}$ 'den elde edilen sonuçların yakın olduğu görülmektedir. 20 örneklem kümesinden hesaplanan tahminler için referans bant genişliği tahmini ile hesaplanan $\hat{p}(H,S)$ yaklaşımı daha düşük hata yüzdesine sahipken; MDY de $\tilde{p}$ yaklaşımının p' 'yi tahmin etmede daha başarılı olduğu görülmüştür.

$n = 1000$ için RBG ile yapılan tahminlerden hesaplanan ortalama hata yüzdeleri $\epsilon\tilde{p}_{KMH} = \%6,502$, $\epsilon Y_{pk} = \%9,110$, $\epsilon\tilde{p} = \%6,974$ ve $\epsilon\hat{p}(H,S) = \%5,232$ olarak bulunmuştur. Hata yüzdeleri incelendiğinde, $\hat{p}(H,S)$ yaklaşımının daha iyi sonuç verdiği görülmektedir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\epsilon\tilde{p}_{KMH} = \%5,937$, $\epsilon Y_{pk} = \%8,638$, $\epsilon\tilde{p} = \%6,702$ ve $\epsilon\hat{p}(H,S) = \%5,418$ olarak elde edilmiştir. $n = 1000$ için $\hat{p}(H,S)$ indeksinin her iki yöntem ile p' 'yi diğer yaklaşımlara göre daha iyi tahmin ettiği görülmektedir. MDY için $\hat{p}(H,S)$ ve $\tilde{p}_{KMH}$ 'den elde edilen sonuçların yakın olduğu söylenebilir. 20 örneklem kümesinden hesaplanan tahminlerde ise $\hat{p}(H,S)$ yaklaşımının her iki yöntem için p' 'yi tahmin etmede daha başarılı olduğu söylenebilir.

Son olarak, Senaryo IV için elde edilen sonuçlar özetlenecek olursa, örneklem büyüklüğü $n = 50$ için hesaplanan hata yüzdeleri, RBG ile yapılan tahminler için ortalama $\epsilon\tilde{p}_{KMH} = \%4,086$, $\epsilon Y_{pk} = \%8,151$, $\epsilon\tilde{p} = \%4,960$ ve $\epsilon\hat{p}(H,S) = \%19,050$ olarak elde edilmiştir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\epsilon\tilde{p}_{KMH} = \%4,290$, $\epsilon Y_{pk} = \%8,151$, $\epsilon\tilde{p} = \%5,006$ ve $\epsilon\hat{p}(H,S) = \%20,824$ olarak hesaplanmıştır. Her iki yöntemde de $n = 50$ büyüklüğü için hata yüzdeleri incelendiğinde, $\tilde{p}_{KMH}$ 'nin diğer yöntemlere göre özellikle, Y_{pk} ve $\hat{p}(H,S)$ yaklaşımlarından çok düşük bir hata yüzdesine sahip olduğu görülmektedir. 20 örneklem kümesinden elde edilen sonuçlar için benzerdir. $\tilde{p}_{KMH}$ 'nin diğer yöntemlere daha iyi sonuç verdiği görülmüştür.

$n = 100$ için RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%4,279$, $\varepsilon Y_{pk} = \%8,987$, $\varepsilon\tilde{p} = \%6,262$ ve $\varepsilon\hat{p}(H,S) = \%7,544$ olarak elde edilmiştir. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%3,580$, $\varepsilon Y_{pk} = \%8,872$, $\varepsilon\tilde{p} = \%6,168$ ve $\varepsilon\hat{p}(H,S) = \%8,895$ olarak elde edilmiştir. Her iki yöntemde de $n = 100$ büyüklüğü için hata yüzdeleri incelendiğinde, $\tilde{p}_{KMH}$ yaklaşımının diğer yaklaşımlara göre daha düşük bir hata yüzdesine sahip olduğu kolaylıkla görülmektedir. 20 örneklem kümesinden hesaplanan tahmin sonuçları aynı şekilde yorumlanabilir.

$n = 250$ için RBG ile yapılan tahminler için ortalamada $\varepsilon\tilde{p}_{KMH} = \%4,636$, $\varepsilon Y_{pk} = \%9,399$, $\varepsilon\tilde{p} = \%6,184$ ve $\varepsilon\hat{p}(H,S) = \%4,820$ olarak bulunmuştur. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\tilde{p}_{KMH} = \%4,355$, $Y_{pk} = \%8,987$, $\varepsilon\tilde{p} = \%6,252$ ve $\varepsilon\hat{p}(H,S) = \%5,713$ olarak elde edilmiştir. $n = 250$ için $\tilde{p}_{KMH}$ indeksinin her iki yöntemde de düşük hata yüzdesine sahip olduğu görülmektedir. $\tilde{p}_{KMH}$ indeksinin RBG ile yapılan tahmin sonuçlarında $\hat{p}(H,S)$ yaklaşımına yakın sonuçlar verdiği görülmektedir. 20 örneklem kümesinden hesaplanan tahminler incelendiğinde, yine $\tilde{p}_{KMH}$ indeksinin diğer yöntemlere göre daha iyi sonuç verdiği söylenebilir.

$n = 500$ için RBG ile yapılan tahminlerden hesaplanan ortalama hata yüzdeleri $\varepsilon\tilde{p}_{KMH} = \%5,381$, $\varepsilon Y_{pk} = \%8,369$, $\varepsilon\tilde{p} = \%6,314$ ve $\varepsilon\hat{p}(H,S) = \%5,621$ olarak hesaplanmıştır. MDY ile yapılan tahminlerde ise ortalama hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%4,463$, $\varepsilon Y_{pk} = \%8,621$, $\varepsilon\tilde{p} = \%5,688$ ve $\varepsilon\hat{p}(H,S) = \%6,445$ olarak elde edilmiştir. $n = 500$ için $\tilde{p}_{KMH}$ yaklaşımının her iki yöntemde de düşük hata yüzdesine sahip olduğu görülmektedir. RBG ile yapılan tahminler için $\hat{p}(H,S)$ ve $\tilde{p}_{KMH}$ 'den elde edilen sonuçların yakın olduğu görülmektedir. 20 örneklem kümesinden hesaplanan tahminler için referans bant genişliği tahmini ile hesaplanan $\tilde{p}_{KMH}$ yaklaşımının daha düşük hata yüzdesine sahip olduğu görülmüştür.

$n = 1000$ için RBG ile yapılan tahminlerden hesaplanan ortalama hata yüzdeleri $\varepsilon\tilde{p}_{KMH} = \%5,989$, $\varepsilon Y_{pk} = \%8,907$, $\varepsilon\tilde{p} = \%6,171$ ve $\varepsilon\hat{p}(H,S) = \%5,301$ olarak bulunmuştur. Hata yüzdeleri incelendiğinde, her ne kadar $\hat{p}(H,S)$ yaklaşımının daha iyi sonuç verdiği görülse de, $\tilde{p}_{KMH}$ yaklaşım sonuçları ile yakın sonuç verdiği görülmektedir. MDY ile yapılan tahminlerde ise hata yüzdeleri, $\varepsilon\tilde{p}_{KMH} = \%5,889$, $\varepsilon Y_{pk} = \%8,460$, $\varepsilon\tilde{p} = \%6,729$ ve $\varepsilon\hat{p}(H,S) = \%6,022$ olarak elde edilmiştir. $n = 1000$ için MDY ile $\tilde{p}_{KMH}$

indeksinin p 'yi diğer yaklaşımlara göre daha iyi tahmin ettiği görülmektedir. 20 örneklem kümesinden hesaplanan tahminlerde ise $\hat{p}(H,S)$ yaklaşımının RBG ile yapılan tahminler için p 'yi tahmin etmede daha başarılı olduğu, MDY ile $\tilde{p}_{KMH}$ indeksinin p 'yi daha iyi tahmin ettiği görülmektedir.

Sonuç olarak, kurgulanan 4 farklı üretim senaryosundan hesaplanan hata yüzde değerleri karşılaştırıldığında, kitle uygun ürün oranı p 'yi, en iyi $\tilde{p}_{KMH}$ indeksinin tahmin ettiği, önerilen $\tilde{p}_{KMH}$ indeksinin, örneklem kümesi sayısına (10 ve 20 örneklem kümesi için) duyarlı olmadığı gözlenmiştir. Ancak, büyük çaplı örneklem için ($n > 500$), $\hat{p}(H,S)$ indeksinin $\tilde{p}_{KMH}$ indeksine ve diğer yaklaşımlara göre kitle parametresini daha az hata ile tahmin ettiği görülmüştür. $\tilde{p}_{KMH}$ indeksinin, çekilen örnek sayısına duyarlı olduğu, örneklem büyüklüğü arttıkça ($n > 500$), hata yüzdelerinin $\hat{p}(H,S)$ indeksi için hesaplanan hata yüzdeleri ile yakın sonuçlar verdiği görülmüştür.

Örneklem kümeleri üzerinden Referans Bant Genişliği matrisi ve Maksimum Düzleştirme Yöntemi ile yapılan tahmin sonuçları incelendiğinde, bant genişliği matris yöntemi bakımından, $\tilde{p}_{KMH}$ indeksinin örneklem büyüklüğüne duyarlı olmadığı ve hata yüzdelerinin her iki bant genişliği yöntemi için diğer indeks hesaplamalarına göre daha iyi sonuç verdiği gözlenmiştir.

Beklenildiği gibi, örneklem büyüklüğünün az olması ve süreç parametre değerinin küçük çaplı örneklem ile daha iyi tahmin edilebilmesi bu tez çalışmasında amaçlanan bir sonuçtur. $\tilde{p}_{KMH}$ indeksinin diğer yöntemlere göre p 'yi tahmin etmede başarılı olduğu görülmektedir. Şüphesiz ki tahmin sonuçlarını iyileştirmek, süreç dağılımına daha uygun bant genişliği yöntemleri seçimi ile mümkün olacaktır. Bunun için sürecin dağılımı iyi yorumlanmalı, önerilen $\tilde{p}_{KMH}$ indeksinin uygulamada kullanımı için süreç dağılımının iyi bir şekilde tahmin edilmesi gerekmektedir.

Bu tez çalışmasında, çok değişkenli süreçler için süreç yeterliliğine karar vermede kullanılacak olan yeni bir yöntem önerilmiştir. KMH yöntemi olarak adlandırılan bu yöntemin literatürdeki diğer yeterlilik indekslerine göre başarısını ölçmek amacıyla $\tilde{p}_{KMH}$ indeksi tasarlanmıştır. Çalışmada kurgulanan senaryolar üzerinden süreç dağılımı Kernel Yoğunluk Tahmin yöntemi ile tahmin edilmeye çalışılmıştır. Kernel Yoğunluk Tahmini ile

süreç dağılımının en iyi şekilde tahmin edilmesi Kernel fonksiyon türüne ve bant genişliği parametresine bağlıdır. Bu nedenle doğru tahminlere ulaşılabilmek için veriye uygun Kernel fonksiyon türünün belirlenmesi ve bant genişliği parametresinin doğru şekilde tahmin edilmesi gereklidir. Çalışmada, bant genişliği parametresinin tahmini için kullanılan yöntemlere alternatif olarak farklı yöntemlerin kullanılması da mümkündür. Çalışmada kurgulanan senaryolarda baz alınan Gamma ve Beta marjinal dağılımlarından farklı olarak Normal dağılıma uygun olmayan diğer çarpık dağılımlar da kullanılabilir. Ayrıca marjinal dağılımlar arasındaki ilişki derecesinin çeşitlendirilmesi yapılacak ileriki çalışmaları zenginleştirecektir.





KAYNAKLAR

1. Abbasi, B. and Niaki, S. T. (2010). Estimating process capability indices of multivariate nonnormal process. *International Journal of Advanced Manufacturing Technology*, 50(5-8), 823-830.
2. Ahmad, S., Abdollahian M., and Zeephongsekul, P. (2008). Process capability estimation for non-normal quality characteristics: A Comparison of Clements, Burr And Box-Cox Methods. *Anziam Journal*, 49, 642-665.
3. Bissell, A. F. (1990). How Reliable is Your Capability Index? *Journal of the Royal Statistical Society*, 39(3), 331- 340.
4. Box, G., Cox, D. R. (1964). An analysis of transformation. *Journal of the Royal Statistical Society, Series B*, 26, 211–243.
5. Boyles, R. A. (1991). The Taguchi capability index. *Journal of Quality Technology*, 23, 17–26.
6. Burr, I. W. (1942). Cumulative frequency functions. *The Annals of Mathematical Statistics*, 13(2), 215-232.
7. Castagliola, P. (1996). Evaluation of non-normal process capability indices using Burr's distributions. *Quality Engineering*, 8(4), 587-593.
8. Castagliola, P., Castellanos, J. V. G. (2005). Capability indices dedicated to the two quality characteristics case. *Quality Technology and Quantitative Management*, 2, 201–220.
9. Chan, L. K., Cheng, S. W. and Spiring, F. A. (1988). A new measure of process capability: Cpm. *Journal of Quality Technology*, 20(3), 162-175.
10. Chen, H. A. (1994). Multivariate process capability index over a rectangular solid tolerance zone. *Statistica Sinica*, 4, 749-758.
11. Chou, Y. M. and Owen, D. B. (1989). On the distributions of the estimated process capability indices. *Communication in Statistics-Theory and Methods*, 18, 4549-4560.
12. Chou, Y. M., Polansky, A. K. and Mason, R. L. (1998). Transforming nonnormal data to normality in statistical process control. *Journal of Quality Technology*, 30(2), 133–141.
13. Clements, J. A. (1989). Process capability calculations for nonnormal distributions. *Qual Prog*, 22, 95–100.
14. Deleryd M. (1999). A pragmatic view on process capability studies. *International Journal of Production Economics*, 58(3), 319–330.
15. Farnum, N. R. (1996). *Using Johnson curves to describe non-normal process data*. *Quality Engineering*, 9(2), 329-336.
16. Finley, J. C. (1992). What is capability? Or what is Cp and Cpk. *ASQC Quality Congress Transactions*, Nashville, 186–191.
17. Franklin L. A. (1999). Sample size determination for lower confidence limits for

- estimating process capability indices. *Computers & Industrial Engineering*, 36(3), 603–614.
18. Franklin, L. A. and Wasserman, G. (1992). Bootstrap lower confidence limits for capability indices. *Journal of Quality Technology*, 24 (4), 196-210.
 19. Goethals, P. L., Cho, B.R. (2011). The development of a target-focused process capability index with multiple characteristics. *Quality and Reliability Engineering International*, 27(3), 297-311.
 20. Gonzalez, I and Sanchez, I. (2009). Capability indices and nonconforming proportion in univariate and multivariate processes. *The International Journal of Advanced Manufacturing Technology*, 44(9–10), 1036–1050.
 21. Gunter, B. H. (1989). *The use and abuse of Cpk*. Parts 1-4, Quality Progress, 22(1), 72-73; 3, 108-109; 5, 79-80; 7, 86-87.
 22. Heavlin, W. D. (1988). Statistical properties of capability indices. *Technical Report No.320*, Sunnyvale, CA: Tech. Library.
 23. Hitchcock, D. B. (2003). A history of the Metropolis-Hastings Algorithm. *The American Statistician*, 57, 254-257.
 24. Horová, I., Koláček, J. and Zelinka, J. (2012). Kernel smoothing in MATLAB: Theory and practice of kernel smoothing. *Singapore: World Scientific*, 244.
 25. Hosseinifard, S. Z., Abbasi, B., Ahmad, S. and Abdollahian, M. (2009). A transformation technique to estimate the process capability index for non-normal processes. *International Journal of Advanced Manufacturing Technology*, 40, 512-517.
 26. Hsiang, T. C. and Taguchi. G. (1985). A tutorial on quality control and assurance – The Taguchi methods. Las Vegas, Nevada: *ASA Annual Meeting*.
 27. Huang, W., Pahwa, A. and Kong, Z. (2012). Kernel density estimation and metropolis hastings sampling in process capability analysis of unknown distributions. *ASME Manufacturing Science and Engineering Conference*, University of Notre Dame, Notre Dame.
 28. Hubele, N., Shahriari, H. and Cheng, C. (1991). A bivariate process capability vector. In J Keats, D. Montgomery (Eds.), *Statistical process control in manufacturing marcel dekker*. New York, 299–310.
 29. Işığışok, E. (2012). *Toplam Kalite Yönetimi Bakış Açısıyla İstatistiksel Kalite Kontrol*. (2. Baskı). Bursa: Ezgi Kitapevi, 255.
 30. İnternet: Duong, T. (2005). Bandwidth selectors for multivariate kernel density estimation. URL: <http://www.webcitation.org/query?url=http%3A%2F%2Fmvstat.net%2Ftduong%2Fresearch%2Fpublications%2Fduong-2005-thesis.pdf&date=2018-06-29> Son Erişim Tarihi: 05.05.2018.
 31. İnternet: Steyvers, Mark. Computational Statistics with Matlab. mark.steyvers@uci.edu. 2018-12-04. URL: http://www.webcitation.org/query?url=http%3A%2F%2Fpsiexp.ss.uci.edu%2Fresearch%2FteachingP205C%2F205C_old.pdf&date=2018-12-04, Son Erişim Tarihi: 2018-12-04.

32. İnternet: URL: <http://www.webcitation.org/query?url=http%3A%2F%2Fwww.mathworks.com%2Fhelp%2Fstats%2Fexamples%2Fsimulating-dependent-random-variables-using-copulas.html&date=2018-12-04>, Son Erişim Tarihi: 2018-12-04.
33. Johnson, N. L. (1949). Systems of frequency curves generated by methods of translation. *Biometrika*, 36, 149-176.
34. Kane, V. E. (1986). Process capability indices. *Journal Quality Technology*, 18, 41-52.
35. Kocherlakota S., Kocherlakota K. (1991). Process capability indices: Bivariate normal distribution. *Communications in Statistics - Theory and Methods*, 20, 2529-2547.
36. Kotz, S. and Johnson, N. L. (1993). *Process capability indices*. London: Chapman & Hall, 212.
37. Kotz, S. & Johnson, N. L. (1999). Delicate relations among the basic process capability indices Cp, Cpk, Cpm and their modifications. *Communications Statistics - Theory and Method*, 28(3), 849-861
38. Kurekova, E. (2001). Measurement process capability-trends and approaches. *Measurement Science Review*, 1(1), 43-46.
39. Kushler, R.H. and Hurley, P. (1992). Confidence bounds for capability indices. *Journal of Quality Technology*, 24(4), 188-195.
40. Liu, P.H. and Chen, F.L (2006). Process capability analysis of non-normal process data using the Burr XII distribution. *International Journal Advanced Manufacturing and Technology*, 27, 975-984.
41. Martinez, W. L. And Martinez, A. R. (2002). *Computational statistics handbook with MATLAB*. New York: Chapman & Hall.
42. Montgomery, D. C. (2009). *Introduction to Statistical Quality Control (6rd Edition)*. New York, USA: John Wiley and Son.
43. Nagata, Y. and Nagahata, H. (1992). Approximation formulas for the confidence intervals of process capability indices. *Reports of Statistical Application Research, Japanese Union of Scientists and Engineers*, 39(3), 15-29.
44. Nagata, Y. and Nagahata, H. (1994). Approximation formulas for the lower confidence intervals of process capability indices. *Okayama Economic Review*, 25, 301-314.
45. Niaki, S. T. A. and Abbasi, B. (2007). Skewness reduction approach in multi-attribute process monitoring. *Journal of Communications in Statistics - Theory and Methods*, 36(12), 2313 -2325.
46. Owen, D. B. (1965). A special case of bivariate non-central t-distribution. *Biometrika*, 52(3/4), 437-446.
47. Pan, J. N. and Wu, S. L. (1997). Process capability analysis for non-normal relay test data. *Microelectronics and Reliability*, 37(3), 421-428.
48. Pan, J. N., Lee, C. Y. (2010). New capability indices for evaluating the performance of multivariate manufacturing processes. *Quality and Reliability Engineering International*, 26, 3-15.

49. Pearn, W. L. and Chen, K. S. (1995). Estimating process capability indices for non-normal Pearsonian populations. *Quality and Reliability Engineering International*, 11(5), 386-388.
50. Pearn, W. L. and Chen, K. S. (1997). Capability indices for non-normal distributions with an application in electrolytic capacitor manufacturing. *Microelectron Reliability*, 37(12), 1853-1858.
51. Pearn, W. L. and Chen, K. S. (1998). New generalization of process capability index Cpk. *Journal of Applied Statistics*, 25, 801-810.
52. Pearn, W. L. and Kotz, S. (1994). Application of Clements' method for calculating second and third generation process capability indices for Non-normal personian populations. *Quality Engineering*, 7, 139-145.
53. Pearn, W. L. and Kotz, S. (2006). Encyclopedia and Handbook of Process Capability Indices: A Comprehensive Exposition of Quality Control Measure. Singapore: *World Scientific*.
54. Pearn, W. L., Kotz, S. and Johnson, N. L. (1992). Distributional and inferential properties of process capability indices. *Journal of Quality Technology*, 24(4), 216-33.
55. Pearn, W. L. and Lin G. H. (2002). Estimated incapability index: Reliability and decision making with sample information. *Quality and Reliability Engineering International*, 18(2), 141-147.
56. Pearn, W. L., Lin, G. H. and Chen, K. S. (1998). Distributional and inferential properties of the process accuracy and process precision indices. *Communications in Statistics - Theory and Methods*, 27, 985-1000.
57. Pearn, W. L., Wang, F. K. and Yen, C. H. (2007). Multivariate capability indices: Distributional and inferential properties. *Journal of Applied Statistics*, 34(8), 941-962.
58. Polansky, A. M. (2001). A smooth nonparametric approach to multivariate process capability. *Technometrics*, 43(2), 199-211.
59. Polansky, A. M., Chou, Y. M. and Mason, R. L. (1998). Estimating process capability indices for a truncated distribution. *Quality Engineering*, 11(2), 257-265.
60. Shinde, R. L. and Khadse, K. G. (2009). Multivariate process capability using principal component analysis. *Quality and Reliability Engineering International*, 25(1), 69-77.
61. Silverman, B. W. (1986). *Density Estimation For Statistics And Data Analysis*. Chapman and Hall, London.
62. Somerville, S. and Montgomery, D. (1996). Process capability indices and non-normal distributions. *Quality Engineering*, 19(2), 305-316.
63. Spiring, F. (2010). Determining and Assessing Process Capability for Engineers and Manufacturing. New York: *Nova Science Publishers*.
64. Stoumbos, Z. G. (2002). Process capability indices: Overview and extensions. *Nonlinear Analysis: Real World Applications*, 3, 191-210.
65. Sullivan, L. P. (1984). Reducing variability: A new approach to quality. *Quality Progr.*, 17, 15-21.

66. Sullivan, L. P. (1985). Letters. *Quality Progr.*, 18, 7-8.
67. Taam, W., Subbaiah, P. and Liddy, J. W. (1993). A note on multivariate capability indices. *Journal of Applied Statistics*, 20, 339-351.
68. Terrell, G.R. (1990). The maximal smoothing principle in density estimation. *Journal of the American Statistical Association*, 85, 470-477.
69. Vännman, K. (1995). A unified approach to capability indices. *Statistica Sinica*, 5, 805-820.
70. Wang, F. K. and Chen, J. C. (1998). Capability index using principal components analysis. *Quality Engineering*, 11, 21-27.
71. Wang, C. H. (2005). Constructing multivariate process capability indices for short-run production. *International Journal of Advanced Manufacturing Technology*, 26, 1306-1311.
72. Wu, H. H., Wang, J. S, Liu, T. L. (1998). Discussions of the Clements-based process capability indices. In: *Proceedings of the 1998 CIIE National Conference*. 561-566.
73. Zhang, N. F., Stenback, G. A. and Wardrop, D. M. (1990). Interval estimation of process capability index Cpk, *Communications in Statistics-Theory and Methods*, 19, 4455-4470.
74. Zimmer, L. S. and Hubele, N. F. (1997). Quantiles of the sampling distribution of Cpm. *Quality Engineering*, 10, 309-329.
75. Zimmer, L. S., Hubele, N. F. and Zimmer, W. J. (2001). Confidence intervals and sample size determination for Cpm. *Quality and Reliability Engineering International*, 17, 51-68.





EKLER

EK-1. t-Copula yöntemi kullanılarak veri üretme

```

function y= Multivariatedata(n,rho,k,theta,alfa,beta)
Z=mvnrnd([0 0], [1 rho; rho 1], n);
U=normcdf(Z);
X=[gaminv(U(:,1),k,theta) betainv(U(:,2),alfa,beta)]
[n1,ctr1]=hist(X(:,1),20);
[n2,ctr2]=hist(X(:,2),20);
subplot(2,2,2);plot(X(:,1),X(:,2),'.');axis([0 20 -2 2]);
h1=gca;
title('Gamma ve Beta Dagilimlarina Bagli Olarak Uretilen
      Verinin Grafigi');
subplot(2,2,4);
bar(ctr1,-n1,1);
axis([0 20 -max(n1)*1.1 0]);
axis('off');
h2 = gca;
subplot(2,2,1);
barh(ctr2,-n2,1);
axis([-max(n2)*1.1 0 -2 2]);
axis('off');
h3 = gca;
h1.Position = [0.35 0.35 0.55 0.55];
h2.Position = [.35 .1 .55 .15];
h3.Position = [.1 .35 .15 .55];
colormap([.8 .8 1]);
B=X
mean(B)
save('...mat','B')
end
% Kaynak: http://www.mathworks.com/help/stats/examples/
simulating-dependent-random-variables-using-copulas.html

```

EK-1. (devam) MH örnekleme ile hedef dağılımdan örneklem çekilmesi

```

close all

%% Metropolis Hastings Orneklemesi
T =1000; % Maksimum iterasyon sayisi
thetamin = [0.25 0.025]; thetamax = [10.25 1]; % theta1 ve
    theta2 icin minimum ve maksimum sinirlar
theta = zeros( 2, T ); % orneklem depolanma alanı
% seed=1; rand( 'state' , seed ); randn('state',seed );
% rastgele orneklemi sabit tutmak icin tanimli seed
mu = [6.0014 0.2500]; sigma = [11.9958 0.3432; 0.3432
    0.0209];

%% orneklemeye baslanir
N=20; A=zeros((N*2),T);
for j=1:2:(N*2)
    t = 1;
    theta(1,1)= abs(unifrnd( thetamin(1) , thetamax(1) ));%
        theta1 icin baslangic deger
    theta(2,1)= abs(unifrnd( thetamin(2) , thetamax(2) ));%
        theta2 icin baslangic deger
    for i=1:N;
        while t < T % T ornekleme ulasilana kadar devam
            edilir
            t= t + 1;
            % theta icin yeni bir deger onerilir
            theta_star=abs(mvnrnd( mu, sigma));
            theta_star1=theta_star(1,1); theta_star2=
                theta_star(1,2);
            theta1=theta(1,t-1); theta2=theta(2,t-1);
            pratio= mernel(theta_star1, theta_star2)/kernel(
                theta1, theta2);
            alpha = min( [ 1 pratio ] ); % kabul orani
            hesaplanir
            u= rand; % U[ 0 1 ] dagilimdan rastgele bir sayi
            cekilir
            if u < alpha; % oneri kabul edilecek mi?
                theta(:,t)=theta_star; % onerilen deger theta
                    'nin yeni degeri olur
            else
                theta(:,t)=theta(:,t-1);% theta icin eski
                    deger gecerli olur
            end
        end
    end
    theta;
    A(j:j+1,:)=A(j:j+1,:)+ theta;
end
B=(A(:,1:T))'; save('...mat','B'); nonconforming=0;

```

EK-1. (devam) MH örnekleme ile hedef dağılımdan örneklem çekilmesi

```
for i=1:2:(2*N-1);  
    for j=1:(T);  
        if ((0.2505<= B(j,i)&& B(j,i)<= 10.2495))  
            &&((0.02505<= B(j,i+1) && B(j,i+1)<= 0.9999));  
            nonconforming=nonconforming+0;  
        else  
            nonconforming=nonconforming+1;  
        end  
    end  
end  
nonconforming  
ypk=1-(nonconforming/(N*(T)))  
% Kaynak: http://psiexp.ss.uci.edu/research/teachingP205C/205C\_old.pdf
```

EK-1. (devam) Gaussian Kernel Fonksiyonu için yazılan kod dizisi

```

function z=kernel(theta1, theta2)
close all
load('...mat'); % Veri dosyasinin yuklenmesi
n=numel(B(:,1));
% Ilk faz veri icin elde edilen bant genisligi matrisi
H=[1.2843 0.0357;0.0357 0.0021]; M=[theta1 theta2];
D= repmat(M, n, 1); d=2;
count=1; C=zeros(n,d);
for j=1:1
    for i=1:n
        a(i,:) = D(j,:)-B(i,:);
    end
    C(:,count)=C(:,count)+a(:,1);
    C(:,count+1)=C(:,count+1)+a(:,2);
    count=count+2;
end
i1=sqrt(inv(det(H)))*(1/(2*pi))*exp(-0.5*((C)*(inv(H))*(C)'))
;
z=(1/n)*trace(i1);
end

```

EK-1. (devam) Y_{pk} değerinin hesaplanması

```
function out = ypk
clc
load('...mat')
n=10; X=B;
nonconforming=0;
for j=1:n;
    if ((0.2505<= X(j,1)&& X(j,1)<=10.2495))&&((0.02505<= X(j,2) && X(j,2)<= 0.9999));
        nonconforming=nonconforming+0;
    else
        nonconforming=nonconforming+1;
    end
end
nonconforming
conforming=n-nonconforming
out=(conforming/(n))
```

EK-1. (devam) $\tilde{p}_{KMH}$ değerin hesaplanması

```
function out = trapezoidal_rule_double_integral(x, y, mat)
x=linspace(0.2505,10.2495,100)';
y=linspace(0.02505,0.9999,100);
load ('X.mat');
load ('Y.mat');
mat=fullkern;
out=trapz(y,trapz(x,mat,1),2);
end
```

```
function [A] = fullkern(X,Y)
clear; close all
load('...mat');
n=numel(B(:,1));
% H=msp(B',NaN,NaN,NaN,'gaus')
H=[ (4/(n*4))^(1/3) ]*cov(B)
F=B;
load ('X.mat'); load ('Y.mat');
m=numel(X(:,1)); k=numel(Y(1,:));
A=zeros(m,k);
for j=1:k;
    for i=1:m;
        M=[X(i,j),Y(i,j)];
        C=repmat(M, n, 1);
        z=C-F;
        i1=sqrt(inv(det(H)))*(1/(2*pi))*exp(-0.5*((z)*(inv(H)
            )*(z)'));
        t=(1/n)*trace(i1);
        A(i,j)=A(i,j)+t;
    end
end
clearvars F X Y
end
```

EK-1. (devam) $\tilde{p}$ değerin hesaplanması

```

function t= integral(X1,X2)
clc; clear; close all;
load('...mat')
X1=B(:,1)
X2=B(:,2)
C=cov(X1,X2)
mu_u=mean(X1)
mu_k=mean(X2)
% C = [sigma_u^2 ro*sigma_u*sigma_k;... ro*sigma_u*sigma_k
       sigma_k]
mu = [mu_u;mu_k];
fun =@(X1,X2) ((1/sqrt(det(C)*(2*pi)^2))*exp(-0.5*transpose
([X1;X2]-mu)*inv(C)*([X1;X2]-mu)))
p = integral2(@(X1,X2)arrayfun(fun,X1,X2)
,0.2505,10.2495,0.02505,0.9999)
q=1-p;
end

```

EK-1. (devam) Çizgi grafiklerinin elde edilmesi

```
function çizgigrafik
a = xlsread('grafikx.xlsx');
save 'grafikx.mat.'
y=[a(:,1) a(:,2) a(:,3) a(:,4)];
x=(1:10);
plot(x,y, 'DisplayName', 'y');
hold on
legend('Y_p_k', 'p^~_K_M_H', 'p^~', 'p(H,S) ')
xlabel('Orneklem Kumeleri')
ylabel('Hata Yuzdeleri')
end
```

EK-1. (devam) Spesifikasyon bölgesine göre saçılım grafiğinin elde edilmesi

```
function createcontour(xdata1, ydata1, zdata1)
%CREATECONTOUR(XDATA1, YDATA1, ZDATA1)
% XDATA1: contour x
% YDATA1: contour y
% ZDATA1: contour z
filename1='gridX.xlsx';
xdata1=xlsread(filename1);
filename2='gridY.xlsx';
ydata1=xlsread(filename2);
%filename3='ZDATA1.xlsx'
%zdata1=xlsread(filename3);
zdata1=kerncontour;
hold all
%contour(xdata1,ydata1,zdata1)
createplot
surf(xdata1,ydata1,reshape(zdata1,100,100))
end
```

EK-1. (devam)

```

function [zdata1] = kerncontour(X,Y)
clear; close all;
load('...mat'); n=numel(B(:,1));
%H=msp(B',NaN,NaN,NaN,'gaus');
H=[(4/(n*4))^(1/3)]*cov(B);
F=B; load('X.mat'); load('Y.mat');
m=numel(X(:,1)); k=numel(Y(1,:));
zdata1= zeros(m,k);
for j=1:k;
    for i=1:m;
        M=[X(i,j),Y(i,j)];
        C= repmat(M, n, 1);
        z=C-F;
        i1=sqrt(inv(det(H)))*(1/(2*pi))*exp(-0.5*((z)*(inv(H)
            )*(z)'));
        t=(1/n)*trace(i1);
        zdata1(i,j)=zdata1(i,j)+t;
    end
end
zdata1
clearvars F X Y
xlswrite('ZDATA1.xlsx',zdata1);
end

```

```

function createplot(a1, b1)
% CREATEPLOT1(X1, Y1)
% X1: vector of x data
% Y1: vector of y data
% Auto-generated by MATLAB on 25-Jul-2016 17:55:19
% filename='N11.xlsx'
% a=xlsread(filename)
load('...mat'); a=B; a1=a(:,1); b1=a(:,2);
% Create plot
plot(a1,b1,'Marker','x','LineStyle','none');

```

ÖZGEÇMİŞ

Kişisel Bilgiler

Soyadı, Adı : DARENDE ŞİMŞEK, Burcu
 Uyuğu : T.C.
 Doğum tarihi ve yeri : 31.05.1983, Bandırma
 Medeni hâli : Evli
 e-mail : bdarende@gmail.com



Eğitim

Derece	Eğitim Birimi	Mezuniyet Tarihi
Doktora	Gazi Üniversitesi /İstatistik	Hâlen
Yüksek Lisans	Hacettepe Üniversitesi /İstatistik	2009
Lisans	Hacettepe Üniversitesi /İstatistik	2006

İş Deneyimi

Yıl	Yer	Görev
2013 – Hâlen	Sosyal Güvenlik Kurumu	Denetmen
2008 – 2013	Başkent Üni. İstatistik ve Bilg. Bil. Bölümü	Araştırma Görevlisi

Yabancı Dil

İngilizce

Yayınlar

- Şimşek, B. D. ve Yeniay, Ö. (2012). Optimization of Locating Casualty Collection Points in a Possible Istanbul Earthquake. *Anadolu University Journal of Science and Technology–A Applied Sciences and Engineering*, 13(1), 1-11.
- Yeloglu, H. O. ve Şimşek, B. D., (2011). Örgütsel Söylem, Toplumsal Bağlam ve Kurumsal Sosyal Sorumluluk Projelerinin Eşbiçimliliği: Türkiye Üzerine Genel Bir Değerlendirme. 19. *Ulusal Yönetim ve Organizasyon Kongresi Bildiriler Kitabı*, 176-179.

Hobiler

Yüzme, Müzik, Tenis



GAZİ GELECEKTİR...