



**SINIFLANDIRMADA KULLANILAN VERİ MADENCİLİĞİ
YÖNTEMLERİNİN PERFORMANSLARININ VERİ SETİ
ÖZELLİKLERİNE GÖRE KARŞILAŞTIRILMASI**

Görkem Ceyhan

**DOKTORA TEZİ
EĞİTİM BİLİMLERİ ANABİLİM DALI**

**GAZİ ÜNİVERSİTESİ
EĞİTİM BİLİMLERİ ENSTİTÜSÜ**

TEMMUZ, 2020

TELİF HAKKI VE TEZ FOTOKOPİ İZİN FORMU

Bu tezin tüm hakları saklıdır. Kaynak göstermek koşuluyla tezin teslim tarihinden itibaren 12 (on iki) ay sonra tezden fotokopi çekilebilir.

YAZARIN

Adı : Görkem
Soyadı : CEYHAN
Bölümü : Eğitimde Ölçme ve Değerlendirme
İmza :
Teslim tarihi :

TEZİN

Türkçe Adı : Sınıflandırmada Kullanılan Veri Madenciliği Yöntemlerinin Performanslarının Veri Seti Özelliklerine Göre Karşılaştırılması
İngilizce Adı : Comparison of Performance of Data Mining Methods Used for Classification in Terms of Data Characteristics

ETİK İLKELERE UYGUNLUK BEYANI

Tez yazma sürecinde bilimsel ve etik ilkelere uyduğumu, yararlandığım tüm kaynakları kaynak gösterme ilkelerine uygun olarak kaynakçada belirttiğimi ve bu bölümler dışındaki tüm ifadelerin şahsıma ait olduğunu beyan ederim.

Yazar Adı Soyadı: Gökem CEYHAN

İmza:

JÜRİ ONAY SAYFASI

Görkem CEYHAN tarafından hazırlanan “Sınıflandırmada Kullanılan Veri Madenciliği Yöntemlerinin Performanslarının Veri Seti Özelliklerine Göre Karşılaştırılması” adlı tez çalışması aşağıdaki jüri tarafından oy birliği ile Gazi Üniversitesi Eğitim Bilimleri Anabilim Dalı’nda Doktora Tezi olarak kabul edilmiştir.

Danışman: Prof. Dr. İsmail KARAKAYA

(Eğitimde Ölçme ve Değerlendirme, Gazi Üniversitesi)

Başkan: Prof. Dr. Nuri DOĞAN

(Eğitimde Ölçme ve Değerlendirme, Hacettepe Üniversitesi)

Üye: Prof. Dr. Hakan Yavuz ATAR

(Eğitimde Ölçme ve Değerlendirme, Gazi Üniversitesi)

Üye: Prof. Dr. Mehtap ÇAKAN

(Eğitimde Ölçme ve Değerlendirme, Gazi Üniversitesi)

Üye: Prof. Dr. Şener Büyükoztürk

(Eğitimde Ölçme ve Değerlendirme, Hasan Kalyoncu Üniversitesi)

Tez Savunma Tarihi: 09/06/2020

Bu tezin Eğitim Bilimleri Anabilim Dalı’nda Doktora tezi olması için şartları yerine getirdiğini onaylıyorum.

Prof. Dr. Selma YEL

Eğitim Bilimleri Enstitüsü Müdürü

TEŞEKKÜR

Doktora eğitimim boyunca akademik ve bilimsel gelişimimde katkıları olan, tez çalışma sürecinde beni destekleyen, motive eden, çalışmamın tüm aşamalarında bana yol gösteren, tecrübelerini paylaşan, akademik çalışma hayatı dışında da yaşadığım olumlu veya olumsuz durumlarda desteğini esirgemeyen, çalışma disiplini ve ciddiyeti ile örnek aldığım değerli hocam ve tez danışmanım Prof Dr. İsmail KARAKAYA'ya en içten dileklerimle teşekkür ederim.

Doktora eğitim sürecine adım attığım ilk andan bu güne kadar kendisinden çok şey öğrendiğim, ufkumu ve zihnimi açan, birçok akademik çalışmada yer almamı sağlayan, akademik hayata dair bitmeyen enerjisi, pozitif bakış açısı ve insanlara karşı ilişkileri ile benim için rol model olan Prof. Dr. Şener BÜYÜKÖZTÜRK'e sonsuz teşekkür ederim.

Tez çalışmam boyunca sık sık fikirlerine başvurduğum ve her defasında değerli vaktini bana ayırıp yardımcı olan verdiği geri bildirimler ve düzeltmelerle tezimin tüm aşamalarında bana büyük destek sunan, akademik bakış açısı ve ciddiyeti ile örnek aldığım Prof. Dr. Hakan Yavuz ATAR'a çok teşekkür ederim.

Tez savunma jürimde bulunmasından mutluluk duyduğum kıymetli hocalarım Prof. Dr. Mehtap ÇAKAN ve Prof. Dr. Nuri DOĞAN'a tezime sundukları değerli katkılardan ötürü çok teşekkür ederim.

Tez çalışma konumda bana fikir veren, zorlandığım noktalarda bana destek olan ve bilgilerini paylaşan, akademik yaşantı dışında da her zaman desteği ile yanımda olan, akademik hayatı ve çalışmayı ilk kez yanında tecrübe ettiğim saygıdeğer hocam Prof. Dr. Murat KAYRI'ye çok teşekkür ederim.

Doktora eğitimim süresince değerli bilgilerini paylaşan, sorularıma vakit ayıran ve akademik hayatta desteklerini hissettiren, Prof. Dr. Şeref TAN, Prof. Dr. Bayram ÇETİN, Prof. Dr. Nilüfer KAHRAMAN, Doç. Dr. Ayfer SAYIN'a ve Gazi Üniversitesi Eğitimde Ölçme ve Değerlendirme anabilim dalındaki tüm mesai arkadaşlarıma çok teşekkür ederim.

Fikir ve destekleri ile tez sürecimin tüm aşamalarında katkılar sunan meslektaşım ve arkadaşım Arş. Gör. Ayşenur ERDEMİR'e çok teşekkür ederim.

Doktora sürecimin başından sonuna kadar karşılaştığım her durumda destekleriyle yanımda olan, akademik ve sosyal hayatımda birlikte olmaktan mutlu olduğum, kişilikleri ve dostluklarıyla bana güç veren değerli dostlarım, Arş. Gör. Ergün Cihat ÇORBACI, Arş. Gör. Dr. Fuat ELKONCA, Arş. Gör. Dr. Hikmet ŞEVGİN, Dr. Öğr. Üyesi Mahmut Sami KOYUNCU, Dr. Öğr. Üyesi Mehmet ŞATA, Dr. Öğr. Üyesi Sami PEKTAŞ ve Dr. Sibel ADA'ya çok teşekkür ederim.

İyi ve kötü günümde her zaman yanımda olan, destekleri ve tecrübeleri ile bana güç veren, doktora eğitimim süresince yaşadığım birçok sıkıntıda yanımda olan kendi aileme ve eşimin ailesine çok teşekkür ederim.

Akademik hayata adım attığım ilk andan bugüne kadar gösterdiği olağanüstü anlayış ve emekle birçok sorumlulukta yükümü hafifleterek bana daima rahat bir çalışma imkânı sunan eşim Arş. Gör. Dr. Sümeyra CEYHAN'a ve vakitlerinden çaldığım çocukluk dönemlerinde birçok güzel anını paylaşamadığım kızım Reyhan Sude ve oğlum İbrahim Eymen'e bana olan sabırlarından dolayı teşekkür ederim.

**SINIFLANDIRMADA KULLANILAN VERİ MADENCİLİĞİ
YÖNTEMLERİNİN PERFORMANSLARININ VERİ SETİ
ÖZELLİKLERİNE GÖRE KARŞILAŞTIRILMASI
(Doktora Tezi)**

**Görkem Ceyhan
GAZİ ÜNİVERSİTESİ
EĞİTİM BİLİMLERİ ENSTİTÜSÜ
Temmuz, 2020**

ÖZ

Bu çalışmanın amacı, PISA (2015) fen başarıları puanlarına göre Yapay Sinir Ağları, Rastgele Orman Algoritması, Destek Vektör Makinesi, Sınıflandırma ve Regresyon Ağaçları ve Lojistik Regresyon yöntemlerinin sınıflandırma performanslarının bağımlı değişkenin kategori sayısı, bağımsız değişken sayısı ve örneklem büyüklüğü açısından incelenmesidir. Araştırmada PISA (2015) uygulamasına katılan 15 yaş grubundaki öğrencilere ait veriler arasından bütün ülkelere uygulanmayan anketlere ilişkin değişkenlere bağlı olarak ilgili öğrencilerin veri setinden çıkarılması ile geriye kalan 169326 öğrenciye ait veri kullanılmıştır. Sınıflama modellerinde kullanılacak bağımsız değişkenler belirlenirken bağımlı değişkenin sürekli puanlarına göre hesaplanan korelasyon analizi ve VIF değerlerinden yararlanılmıştır. Elde edilen sonuçlar doğrultusunda bağımsız değişken olarak seçilen değişkenler matematik başarı puanı, ekonomik, sosyal ve kültürel statü indeksi, epistemolojik inançlar, evde bulunan kültürel eşyalar, feni sevmeye, evdeki eğitim kaynakları, evde sahip olunan eşyalar, çevresel farkındalık, öğrenme süresi ve fen özyeterlik inancıdır. Yöntemlerin performansları öncelikle bağımlı değişkenin kategori sayısına göre incelenmiştir. Bunun için bağımlı değişkenin 2, 3 ve 6 adet sınıfa sahip olduğu durumlara ait sınıfların öğrenci yüzdeleri göz önünde tutularak, rastgele seçim yöntemiyle veri setinden büyüklüğü 5000 olan 25'er adet çalışma grubu seçilmiştir. Ardından 10 bağımsız değişkenin yer aldığı sınıflandırma modellerine YSA, RO, DVM, SVRA ve LR yöntemleri uygulanmıştır. Sonuçlar, bütün yöntemlerin sınıflandırma performanslarının bağımlı

değişkenin sınıf sayısının azalması durumunda artış gösterdiğini göstermektedir. Yöntemlerin performansları bağımsız değişken sayısının değişimine göre incelenirken, yöntemlerin en iyi sınıflandırma performansını gösterdiği bağımlı değişkenin 2 kategoriye sahip olma durumu ele alınmıştır. Sırasıyla 10, 7 ve 4 bağımsız değişkenin yer aldığı sınıflandırma modellerine YSA, RO, DVM, SVRA ve LR yöntemleri uygulanmıştır. Bütün yöntemlerin sınıflandırma performanslarının bağımsız değişkenin sayısına göre anlamlı bir değişim göstermediği tespit edilmiştir. Ardından fen başarı puanı ile yüksek ilişkiye sahip olan matematik başarı puanı değişkeni modellerden çıkarılmış ve bağımsız değişkenin 9, 6 ve 3 olması koşulunda sınıflandırma modellerine YSA, RO, DVM, SVRA ve LR yöntemleri tekrar uygulanmıştır. Elde edilen sonuçlara göre bağımlı değişkenle yüksek korelasyona sahip matematik başarı puanı değişkeninin modellerde yer almadığı durumda bağımsız değişken sayısı arttıkça bütün yöntemlerin sınıflandırma performansının da artış gösterdiği tespit edilmiştir. Yöntemlerin performansları örneklem büyüklüğüne göre incelenirken, yöntemlerin en iyi sınıflandırma performansını gösterdiği bağımlı değişkenin 2 kategorili ve 10 bağımsız değişkenin yer aldığı sınıflandırma modelleri ele alınmıştır. Veri setinden 100, 250, 500, 1000, 2500 ve 5000 örneklem büyüklüğünün her biri için rastgele seçim yöntemiyle 25'er adet çalışma grubu oluşturulmuştur. Sonuçlar, bütün yöntemlerin sınıflandırma performanslarının örneklem büyüklüğüne göre değiştiğini göstermektedir. YSA, DVM ve LR yöntemlerinin 500 ve daha fazla örneklem büyüklüğünde 100 ve 250 örneklem büyüklüğüne göre daha yüksek ve birbirine benzer değerler ürettiği dolayısıyla daha iyi bir sınıflama performansı sergilediğini sonucuna ulaşılmıştır. RO ve SVRA yöntemlerinin ise 1000 ve üzeri örneklem büyüklüklerinde 100, 250 ve 500 örneklem büyüklüğüne göre daha yüksek bir sınıflama performansına sahip olduğu tespit edilmiştir. Öte yandan bütün koşullar altında yöntemlerin birbirlerine göre performansları da karşılaştırılmış ve elde edilen sonuçlar doğrultusunda YSA, DVM ve LR yöntemlerinin sınıflandırma performanslarının RO ve SVRA yöntemlerine göre daha iyi olduğu sonucuna ulaşılmıştır. Ayrıca YSA, DVM ve LR yöntemlerinin birbirlerine göre benzer RO yönteminin ise SVRA yöntemine göre daha iyi sınıflandırma performansı gösterdiği belirlenmiştir.

Anahtar Kelimeler : Yapay Sinir Ağları, Rastgele Orman, Destek Vektör Makinesi, Sınıflandırma ve Regresyon Ağaçları, Lojistik Regresyon, PISA 2015.

Sayfa Adedi : xix+215

Danışman : Prof. Dr. İsmail KARAKAYA

**COMPARISON OF PERFORMANCE OF DATA MINING METHODS
USED FOR CLASSIFICATION IN TERMS OF DATA
CHARACTERISTICS**

(Ph. D Thesis)

Görkem Ceyhan

GAZI UNIVERSITY

GRADUATE SCHOOL OF EDUCATIONAL SCIENCES

July, 2020

ABSTRACT

This study aims to examine the classification performances of Artificial Neural Networks, Random Forest Algorithm, Support Vector Machine, Classification, and Regression Trees and Logistic Regression methods according to PISA (2015) science achievement scores in terms of the number of classes of the dependent variable, number of independent variables and sample size. The population of the research is all 15-year-old students who participated in PISA (2015) application. The target universe consists of the remaining 169326 students with the exclusion of the students from the data set, depending on the variables related to the questionnaires that are not applied to all countries. While determining the independent variables to be used in the classification models, the correlation analysis in which the continuous scores of the dependent variable were used, and VIF values were taken into consideration. In line with the results, the variables chosen as independent variables obtained are mathematics achievement score, economic, social and cultural status index, epistemological beliefs, cultural possessions at home, enjoyment of science, home educational resources, home possessions, environmental awareness, learning time, and science self-efficacy. The performances of the methods were first examined according to the number of classes of the dependent variable. For this purpose, considering the student percentages of the classes in which the dependent variable has 2, 3, and 6 classes, 25 samples were selected from the target universe, the size of which was 5000 by the random selection method. Afterward, AAN, RF, SVM, CART, and LR methods were applied to the

classification models with ten independent variables. The results show that the classification performance of all methods increases when the number of classes of the dependent variable decreases. The performances of the methods according to the number of independent variables were examined in the case that the dependent variable was in 2 categories, in which the methods showed the best classification performance. ANN, RF, SVM, CART, and LR methods were applied to classification models with 10, 7, and 4 independent variables, respectively. It was found that the classification performances of all methods did not show a significant change according to the number of independent variables. Then, the mathematics achievement score variable, which has a high correlation with the science achievement score, was removed from the models and ANN, RF, DVM, CART, and LR methods were applied to the classification models provided that the independent variable was 9, 6, and 3. According to the results, when the math achievement score with a high correlation with the dependent variable is not included in the models, as the number of independent variables increases, the classification performance of all methods also increases. While the performances of the methods were examined according to the sample size, the classification models with the two-class dependent variable and ten independent variables, in which the methods showed the best classification performance, were tested. Twenty-five samples were created using the random selection method for each sample size of 100, 250, 500, 1000, 2500, and 5000 from the target universe. Results showed that the classification performance of all methods varies according to the sample size. It can be said that the ANN, SVM, and LR methods produced higher values compared to 100 and 250 sample sizes in 500 and more sample sizes; therefore, they showed a better classification performance. Besides, RF and CART methods produced higher values compared to 100, 250, and 500 sample sizes in 1000 and more sample sizes. On the other hand, under all conditions, the performances of the methods were compared, and it was concluded that the classification performance of ANN, SVM, and LR methods were better than those of RF and CART methods. Also, it was found that ANN, DVM, and LR methods were performed similarly to each other, and the RF method showed better classification performance than the CART method.

Key Words : Artificial Neural Networks, Random Forest Algorithm, Support Vector Machine, Classification and Regression Trees, Logistic Regression, PISA 2015

Page Number : xix+215

Supervisor : Prof. Dr. İsmail KARAKAYA

İÇİNDEKİLER

TELİF HAKKI VE TEZ FOTOKOPİ İZİN FORMU	i
ETİK İLKELERE UYGUNLUK BEYANI.....	ii
JÜRİ ONAY SAYFASI.....	iii
TEŞEKKÜR.....	iv
ÖZ	vi
ABSTRACT	viii
İÇİNDEKİLER.....	x
TABLolar LİSTESİ.....	xv
ŞEKİLLER LİSTESİ.....	xvii
SİMGELER VE KISALTMALAR LİSTESİ.....	xix
BÖLÜM I	1
GİRİŞ.....	1
1.1. Problem Durumu	1
1.2. Araştırmanın Amacı	12
1.3. Araştırmanın Önemi	12
1.4. Sınırlılıklar	15
BÖLÜM II.....	17
KURAMSAL TEMELLER.....	17

2.1. Veri Madenciliği	17
2.2. Veri Madenciliği Süreci	22
2.3. Veri Madenciliği Modelleri	27
2.3.1. Kümeleme	28
2.3.2. Birliktelik Kuralları	30
2.4. Sınıflandırma ve Regresyon Ağaçları.....	32
2.4.1. Bölünme Kriterleri.....	36
2.4.1.1. Gini Algoritması	36
2.4.1.2. Twoing Algoritması	37
2.4.1.3. Rastgele Orman Algoritması	39
2.4.1.3.1. Rastgele Orman Algoritmasının İşleyişi	42
2.5. Yapay Sinir Ağları	45
2.5.1. Yapay Sinir Ağının Yapısı.....	49
2.5.1.1. Biyolojik Sinir Hücresi.....	49
2.5.1.2. Yapay Sinir Hücresi.....	50
2.5.2. Yapay Sinir Ağlarının Sınıflandırılması	54
2.5.2.1. Öğrenme Türlerine Göre YSA	54
2.5.2.1.1. Denetimli (Danışmanlı) Öğrenme.....	55
2.5.2.1.2. Danışmansız Öğrenme	56
2.5.2.1.3. Destekleyici Öğrenme	58
2.5.2.2. Öğrenme Zamanına Göre Yapay Sinir Ağları.....	59
2.5.2.3. Katman Sayısına Göre Yapay Sinir Ağı Modelleri.....	60
2.5.2.3.1. Tek Katmanlı Ağlar	60

2.5.2.3.2. Çok Katmanlı Ağlar	61
2.5.2.4. Topolojik Yapısına Göre Yapay Sinir Ağı modelleri.....	63
2.5.2.4.1. Geri Beslemeli (Tekrarlayan) Ağlar	63
2.5.2.4.2. İleri Beslemeli Ağlar	65
2.5.2.5. Çok Katmanlı Algılayıcı (Multilayer Perception; MLP).....	66
2.6. Destek Vektör Makinesi	68
2.6.1. Destek Vektör Makinelerinde Sınıflandırma	73
2.6.1.1. Sert Marjın Lineer Sınıflandırma (Hard Margin)	73
2.6.1.2. Soft Marjın Lineer Sınıflandırma (Doğrusal Olarak Ayrılama)	77
2.6.1.3. Lineer Olmayan Sınıflandırma	80
2.6.2. Çok Sınıflı Destek Vektör Makinesi	81
2.6.2.1. Bire-Karşı-Hepsi (One-Against-All)	81
2.6.2.2. İkili Sınıflama/Bire-Karşı-Bir (Pairwise Classification/One Versus One).....	82
2.7. Lojistik Regresyon	83
2.7.1. Lojistik Regresyon Modelleri ve Matematiksel Yapıları	84
2.8. İlgili Araştırmalar	87
BÖLÜM III	98
YÖNTEM.....	98
3.1. Araştırmanın Türü.....	98
3.2. Çalışma Grubu	98
3.3. Veri Analizi.....	102

3.3.1. Değişken Seçimi.....	102
3.3.2. Koşullar.....	104
3.3.3. Sınıflama Yöntemlerinin Uygulanması.....	107
3.3.4. Karşılaştırma testleri	107
3.3.5. Değerlendirme Kriterleri	109
BÖLÜM IV	114
BULGULAR	114
4.1. Bağımlı Değişken Kategorisi Sayısının 6, 3 ve 2 Olması Durumuna Göre Elde Edilen Bulgular.....	114
4.2. Bağımsız Değişken Sayısının 10, 7 ve 4 Olması Durumuna Göre Elde Edilen Bulgular.....	122
4.3. Bağımsız Değişken Sayısının 9, 6 ve 3 Olması Durumuna Göre Elde Edilen Bulgular.....	127
4.4. Örneklem Büyüklüğünün Değişimine Göre Elde Bulgular	134
BÖLÜM 5.....	145
TARTIŞMA, SONUÇ VE ÖNERİLER	145
5.1. Tartışma	145
5.2. Sonuç	154
5.3. Öneriler	157
5.3.1. Uygulayıcılara Yönelik Öneriler.....	157
5.3.2. Araştırmacılara Yönelik Öneriler	159
KAYNAKLAR.....	161
EKLER.....	186

EK 1. Değişkenlere ait betimsel istatistikler	187
EK 2. Korelasyon katsayısı 0,20'den küçük olan değişkenler	188
EK 3. Bağımlı değişken koşuluna göre değerlendirme kriterlerine ait betimsel istatistikler	189
EK 4. Bağımsız değişken sayısının 10, 7 ve 4 olması koşuluna göre değerlendirme kriterlerine ait betimsel istatistikler	191
EK 5. Bağımsız değişken sayısının 9, 6 ve 3 olması koşuluna göre değerlendirme kriterlerine ait betimsel istatistikler	194
EK 6. Örneklem büyüklüğü koşuluna göre değerlendirme kriterlerine ait betimsel istatistikler	197
EK 7. Bağımlı değişkenin kategori sayısının 6, 3 ve 2 olması durumunda yöntemlerin değerlendirme kriterlerinin değerlerine ait grafikler.....	204
EK 8. Bağımsız değişken sayısının 10, 7 ve 4 olması durumunda yöntemlerin değerlendirme kriterlerinin değerlerine ait grafikler.....	206
EK 9. Bağımsız değişken sayısının 9, 6 ve 3 olması durumunda yöntemlerin değerlendirme kriterlerinin değerlerine ait grafikler.....	209
EK 10. Örneklem büyüklüğüne göre yöntemlerin değerlendirme kriterlerinin değerlerine ait grafikler.....	212
EK 11. Etik Kurul Onayı.....	215

TABLÖLAR LİSTESİ

Tablo 1. Kümeleme Analizinde Başlıca Kullanılan Algoritmalar	30
Tablo 2. YSA'da Kullanılan Bazı Toplam Fonksiyonları	51
Tablo 3. YSA'da Kullanılan Bazı Aktivasyon Fonksiyonları	52
Tablo 4. Literatürde Yer Alan Bazı Kernel Fonksiyonları	81
Tablo 5. PISA 2015 Öğrenci Verisinde Yer Alan Öğrencilerin Ülkelere Göre Dağılımı....	99
Tablo 6. Fen Başarı Puanlarına Ait Ortalama Standart Puanları ile Regresyon Faktör Puanları İçin Hesaplanan Korelasyon Analizi Sonuçları	100
Tablo 7. Bağımlı Değişkenin 6, 3 ve 2 Kategori Olduğu Durumlara Ait Frekans ve Yüzdelere	101
Tablo 8. Bağımlı Değişken ile Bağımsız Değişkenler Arasında Hesaplanan Korelasyon Katsayıları.....	103
Tablo 9. Bağımlı Değişkenin Kategori Sayısının 6, 3 ve 2 Olması Durumlarında Üretilen Veri Setleri	105
Tablo 10. Bağımsız Değişkenin 4, 7 ve 10 ile 3, 6 ve 9 Olması Durumlarında Üretilen Veri Setleri	106
Tablo 11. Örneklem Büyüklüklerinin 100, 250, 500, 1000, 2500 ve 5000 Olması Durumlarında Üretilen Veri Setleri	107
Tablo 12. Sınıflama İki Düzey Olduğunda Kullanılan Confusion Matrisi	109
Tablo 13. Çoklu Sınıflama Olduğunda Kullanılan Hata Matrisi	110

Tablo 14. Yöntemlerin Sınıflandırma Performanslarının Bağımlı Değişkenin Kategori Sayısının Değişimine Göre Değerlendirilmesine Yönelik Uygulanan Kruskal Wallis Testi	115
Tablo 15. Bağımlı Değişkenin Kategori Sayısının 6,3 ve 2 Olduğu Durumlarda Yöntemlerin Sınıflandırma Performanslarının Birbirlerine Göre Değerlendirilmesine Yönelik Uygulanan Friedman Testi Sonuçları	119
Tablo 16. Yöntemlerin Sınıflandırma Performanslarının Bağımsız Değişken Sayısının Değişimine Göre Değerlendirilmesine Yönelik Uygulanan Friedman Testi Sonuçları	123
Tablo 17. Her Bir Bağımsız Değişken Sayısında Yöntemlerin Sınıflandırma Performanslarının Birbirlerine Göre Karşılaştırılmasına Yönelik Uygulanan Friedman Testi Sonuçları	124
Tablo 18. Yöntemlerin Sınıflandırma Performanslarının Kendi İçinde Bağımsız Değişken Sayısının Değişimine Göre Değerlendirilmesine Yönelik Uygulanan Friedman Testi Sonuçları	128
Tablo 19. Her Bir Bağımsız Değişken Sayısında Yöntemlerin Sınıflandırma Performanslarının Birbirlerine Göre Karşılaştırılmasına Yönelik Uygulanan Friedman Testi Sonuçları	131
Tablo 20. Yöntemlerin Sınıflandırma Performanslarının Kendi İçinde Örneklem Büyüklüğünün Değişimine Göre Değerlendirilmesine Yönelik Uygulanan Kruskal Wallis Testi Sonuçları	135
Tablo 21. Her Bir Örneklem Büyüklüğünde Yöntemlerin Sınıflandırma Performanslarının Birbirlerine Göre Karşılaştırılmasına Yönelik Uygulanan Friedman Testi Sonuçları	139

ŞEKİLLER LİSTESİ

Şekil 1. Veri tabanı ve veri yönetimi gelişim süreci.	18
Şekil 2. Veri madenciliği süreci Fayyad vd. (1996).....	23
Şekil 3. Veri madenciliği süreci.	24
Şekil 4. CRISP-DM süreci.	26
Şekil 5. Örnek karar ağacı.	35
Şekil 6. RO yönteminin işleyişine yönelik aşamalar.....	43
Şekil 7. Bir biyolojik nöronun yapısı.	49
Şekil 8. Bir biyolojik ve yapay sinir hücrelerinin elemanlarına ait karşılaştırma.	53
Şekil 9. Yapay bir nöronun işleme yetenekleri.	53
Şekil 10. Danışmanlı öğrenme biçimini gösteren bir diyagram.	56
Şekil 11. Danışmansız öğrenme biçimini gösteren bir diyagram.....	57
Şekil 12. Destekleyici öğrenme biçimini gösteren bir diyagram.	58
Şekil 13. Basit ileri beslemeli ağ modeli.	60
Şekil 14. Tek katmanlı YSA modellerinin matematiksel formu.	61
Şekil 15. Basit ileri beslemeli ağ modeli.	62
Şekil 16. Geri beslemeli ağların işleyiş modeli.....	64
Şekil 17. İleri beslemeli ağların işleyiş modeli.	65
Şekil 18. Çok katmanlı ileri beslemeli geri yayımlı bir YSA modeli.	68

<i>Şekil 19.</i> DVM'nın işleyiş modeli.....	72
<i>Şekil 20.</i> Veri setini birbirinden ayıran üç olası hiperdüzleme ait gösterimi.	73
<i>Şekil 21.</i> İki boyutlu bir hiperdüzlem modeli.....	74
<i>Şekil 22.</i> Lineer sınıflandırma modeli.	75
<i>Şekil 23.</i> Verinin doğrusal olarak ayrılmama durumuna ait model.	78
<i>Şekil 24.</i> Bağımlı değişkenin kategori sayısı koşulu için her bir yönteme ait doğruluk değerleri.	122
<i>Şekil 25.</i> Bağımlı değişkenin kategori sayısı koşulu için her bir yönteme ait F ölçütü değerleri.	122
<i>Şekil 26.</i> Bağımsız değişken sayısının 10, 7 ve 4 olması durumunda her bir yönteme ait doğruluk değerleri.....	126
<i>Şekil 27.</i> Bağımsız değişken sayısının 10, 7 ve 4 olması durumunda her bir yönteme ait F ölçütü değerleri.	127
<i>Şekil 28.</i> Bağımsız değişken sayısının 9, 6 ve 3 olması durumunda her bir yönteme ait doğruluk değerleri.....	133
<i>Şekil 29.</i> Bağımsız değişken sayısının 9, 6 ve 3 olması durumunda her bir yönteme ait F ölçütü değerleri.	134
<i>Şekil 30.</i> Farklı örneklem büyüklükleri koşulunda her bir yönteme ait doğruluk değerleri.	143
<i>Şekil 31.</i> Farklı örneklem büyüklükleri koşulunda her bir yönteme ait F ölçütü değerleri.	144

SİMGELER VE KISALTMALAR LİSTESİ

MEB	Milli Eğitim Bakanlığı
VM	Veri Madenciliği
EVM	Eğitimde Veri Madenciliği
PISA	Programme for International Student Assessment
YSA	Yapay Sinir Ağları
RO	Rastgele Orman
DVM	Destek Vektör Makinası
SVRA	Sınıflandırma ve Regreson Ağaçları
LR	Lojistik Regresyon
%	Yüzde

BÖLÜM I

GİRİŞ

Bu bölümde problem durumu ortaya konulmaya çalışılmış, araştırmanın amacı ve önemi hakkında bilgiler verilmiştir.

1.1. Problem Durumu

1950'lerin sonu ve 1960'ların başında gündeme gelen elektronik verilerin yönetimi, başlangıçta dosyaların işlenmesi için gerekli olan donanım ve yazılım bakımından oldukça temel düzeyde ve pahalıydı (Shim vd. 2002). Günümüzde ise, teknolojinin ve bilgi sistemlerinin gelişimi ile pazarlama, bankacılık, telekomünikasyon, sağlık ve eğitim gibi birçok alanda elde edilen elektronik veriler sayısal ortamlarda saklanabilir ve kolayca erişilebilir hale gelmiştir (Fakı & Öztayşı, 2015). Bu bağlamda ortaya çıkan ve her gün artış gösteren büyük veriler, organizasyonları verilerin etkin bir biçimde nasıl analiz edilebileceği, karar destek aşamasında nasıl kullanılabileceği ve faydalanılabileceği gibi problemler ile karşı karşıya bırakmıştır (Westphal & Blaxton 1998). Böyle büyük veriler olduğu durumlarda, geleneksel veri inceleme araçlarının yetersiz kaldığının düşünülmesi, yeni yöntem ve teknolojilerin geliştirildiği “Veri Tabanlarında Bilgi Keşfi” sürecini ortaya çıkarmıştır. Veriden elde edilebilecek gerekli bütün bilginin alınabilmesini ifade eden bu süreç içindeki en önemli bölüm, modellerin kurulduğu ve değerlendirildiği veri madenciliği sürecidir (Akpınar, 2000). Veri Madenciliği sürecinde, veri analiz araçlarının kullanılarak veride yer alan örüntüleri ve aynı zamanda ilişkileri ortaya çıkarmak ve yapılan tahminlerin geçerli olarak elde edilmesi amaçlanmaktadır (Crows, 1999).

Veri madenciliği sürecindeki problemlerin çözümünde izlenmesi gereken adımları, Veri Madenciliği Sürecinde Alanlar Arası Standartlar (Cross Industry Standard Process for Data Mining; CRISP-DM) süreci olarak tanımlanmıştır. Bu adımlar, işin anlaşılması, verinin anlaşılması, veri hazırlama, modelleme ve değerlendirmedir (SPSS, 2000, s.13-14). Veri madenciliği sürecinde veri setinin yapısına ve analizin amacına uygun olan yöntemlerin seçiminin yapıldığı modelleme basamağı, tanımlayıcı (descriptive) modeller ve tahmin edici (predictive) modeller olmak üzere ikiye ayrılmaktadır (Albayrak & Koltan-Yılmaz, 2009; Şimşek-Gürsoy, 2012). Tanımlayıcı modellerde, verilerdeki örüntüler ortaya çıkarılarak bu örüntülerin karar verme sürecinde kullanılması sağlanmaktadır. Tahmin edici modellerde ise, gerçek durumları bilinen veriler kullanılarak bir model geliştirilir ve gerçek durumu bilinmeyen veriler için gerçek durumun tahmininde bu modelin kullanılması amaçlanmaktadır (Akpınar, 2000). Veri madenciliğinde kullanılan modeller, işlevlerine göre üç ana başlık altında ifade edilebilir: Sınıflama ve Regresyon (Classification and Regression), Kümeleme (Clustering), Birliktelik Kuralları (Association Rules). Bu modellerden Sınıflama ve Regresyon modelleri tahmin edici; Kümeleme ve Birliktelik Kuralları ise tanımlayıcı modeller olarak ifade edilebilir (Akpınar, 2000; Albayrak & Koltan-Yılmaz, 2009).

Sınıflama ve Regresyon modellerinin kendi içerisindeki temel farklılığı, bağımlı değişkenin kategorik veya sürekli olup olmadığıdır. Veri setinde yer alan bağımlı değişken sürekli veri özelliğindeyse regresyon, kategorik bir bağımlı değişken varsa sınıflama olarak ele alınır. Ancak regresyon içinde yer alan lojistik regresyon tekniği ile kategorik bir bağımlı değişkenin de yordanması söz konusudur. Bu durumda sınıflama ve regresyonun giderek birbirine yaklaşmakta olduğu ve her ikisinde de aynı tekniklerden faydalanılmasının mümkün olduğu belirtilmektedir (Kıran, 2010). Sınıflama, verinin içerdiği ortak özelliklerine göre ayrıştırılması işlemi veya gerçek sınıfı bilinen halihazırdaki verilerden yararlanarak, gerçek sınıfı bilinmeyen verilerin bir sınıfa atanıp atanamayacağını tahmin etmek eğer atanacaksa hangi sınıf olduğunu belirlemek şeklinde tanımlanmaktadır (Albayrak & Koltan-Yılmaz, 2009; Atılğan, 2011; Emel & Taşkın, 2005; Şimşek-Gürsoy,

2009). Sınıflama problemi grupların ayırt edici niteliklerinin belirlendiği veya bunlar kullanılarak oluşturulan modeller ile bireylerin uygun gruplara atandığı istatistiksel bir karar verme sürecidir (Burmaoğlu, Oktay & Özen, 2009). Sınıflama analizleri fen ve sosyal bilimlerde olduğu gibi eğitim bilimleri alanında da oldukça sıklıkla kullanılan veri madenciliği yöntemidir.

Eğitimde veri madenciliğinin kullanıldığı çalışmalar “Eğitimsel Veri Madenciliği (EVM) (Educational Data Mining)” çalışmaları olarak adlandırılmaktadır. Eğitimsel veri madenciliği topluluğu, Eğitimsel Veri Madenciliği’ni, bir eğitim ortamından gelen veri türlerini keşfetmek için yöntemlerin geliştirildiği ve bu verilerin öğrencileri ve öğrenme ortamlarını daha iyi anlamak için kullanıldığı bir disiplin olarak tanımlamıştır (Ihantola vd. 2015). Romero ve Ventura (2010) EVM’nin, verilerin analizi ve görselleştirilmesi, eğitimcilere geri dönüt sağlanması, öğrencilere tavsiyelerde bulunulması, öğrenci performansının yordanması, öğrenci becerilerinin modellemesi, istenmeyen öğrenci davranışlarının tespit edilmesi, öğrencilerin sınıflandırılması, sosyal ağ analizi, kavram haritaları geliştirilmesi, eğitim yazılımlarının geliştirilmesi, planlama ve programlamanın yapılması gibi amaçlar için kullanıldığını belirtmiştir. Bu alandaki ilgili literatürün derlendiği bazı çalışmalar bulunmaktadır (Baker & Yacef, 2009; Pena-Ayala, 2014; Romero & Ventura, 2007). EVM yöntemleri eğitimsel verideki çok düzeyli hiyerarşiyi kullanması açısından genel veri madenciliği literatüründeki yöntemlerden farklıdır. EVM’de psikometri literatüründeki yöntemler, makine öğrenmesi ve veri madenciliği literatüründeki yöntemlerle entegre edilir (Baker, 2010). Lewis’e göre (1986, s.11) psikometri, belirli bir test maddesine verilen yanıtlar ile bireyin varsayılan özellikleri arasındaki ilişkinin incelenmesini istatistiksel çıkarım problemi olarak ele almaktadır. Psikometrik paradigma, bir yüzyıldan daha fazla bir süredir eğitimsel değerlendirmelerde baskın bir analiz çerçevesi oluşturmuştur. EVM ise bu çerçevenin ötesine geçen öğrenme ve değerlendirme yollarının geliştirilmesini amaçlamaktadır (Mislevy, Behrens, Dicerbo & Levy, 2012). Dolayısıyla EVM ve psikometri alanları arasındaki ilişkinin öne çıkarılması için çalışmalar yapılması gerektiğinden söz edilebilir.

EVM alanında yapılan çalışmalar ilerledikçe, araştırmacılar geçerliliği arttıran yöntemler geliştirmekle daha fazla ilgilenmekte ve teorik yapı geçerliliğini (bir modelin özelliklerinin amaçlanan yapıyı yansıtırma derecesi) artırma amacına yönelik çalışmalar yapmaktadırlar (Ocumpaugh, Baker, Gowda, N. Heffernan & Heffernan, 2014). Deneysel ve inceleme çalışmaları artmaya devam ettikçe, konuların çoğunluğu veri madenciliği tekniklerinin karşılaştırılmasını ve birleştirilmesini içermektedir (Whitley, 2018, s.7). Ancak, EVM literatüründe birçok önemli çalışma olmasına rağmen özellikle hesaplama olarak zor algoritmaları gerektiren büyük veri setleri için denetimli öğrenme yöntemleri ile ilgili araştırma eksikliği bulunmaktadır (Kılıç-Depren, Aşkın & Öz, 2016; Sinharay, 2016). Ayrıca, Aksu ve Doğan (2019) da eğitim alanında sınıflama amacıyla veri madenciliğinin kullanıldığı az sayıda çalışmanın olduğunu belirtmektedir. Bu durumda, alana katkı sunması açısından psikometrik özelliklerle ilişkili olarak daha fazla EVM çalışmasına ihtiyaç duyulduğu söylenebilir.

Eğitim alanında yapılan ölçme ve değerlendirme süreci sonunda geçti-kaldı, öğrencinin başarı düzeyinin ne olduğu ya da ilgilenilen bir davranışa sahip olup olmadığına ilişkin kararlar alınabilmektedir. Eğitim bilimleri açısından önemli görülen bu kararların alınması, bir şekilde öğrencilerin ilgilenilen durum ya da düzey ile ilgili olarak hangi sınıfa dahil olduklarının incelenmesidir (Başokçu, 2011). Eğitim bilimleri alanında ölçme araçlarından elde edilen sonuçlar, öğrenciler hakkında karar vermek amacıyla kullanılabilmektedir. Bu kararların doğruluğu ölçme araç veya yöntemlerinin psikometrik özelliklerine bağlıdır. Bu psikometrik özelliklerden biri de geçerliliklerdir. Psikolojik testlerin geçerliliğini belirlemek ve dolayısıyla bu testlerin amaçlarına ne derece hizmet ettiğini belirlemek önemlidir (Kelecioğlu & Şahin, 2014). En genel tanımıyla geçerlilik, bir testin ölçmek istediği özelliği diğer başka değişkenlerle karıştırmadan ölçebilme derecesidir (Thorndike & Hagen, 1961; Turgut & Baykul, 2013). Kane'e (2013) göre geçerlilik çalışmaları test puanlarının kullanılması ve yorumlanması temeline dayanmaktadır. Erkuş'a göre ise (2003) ölçeklerin asıl amacı karar vermektir ve süreç ve işlemler açısından geçerlilik ikiye ayrılır. Bunlardan ilki, ölçeğin hangi özelliği ölçmeyi amaçlıyorsa sadece onu ölçüp ölçmediğinin irdelenmesi

yani ölçme geçerliğidir. İkincisi ise ölçek geliştirildikten sonra bireyler hakkında verilecek kararların tutarlığı olup karar geçerliğidir. Karar geçerliliği, sınıflama ve sıralama geçerliliği kavramlarını içinde barındırmaktadır. Seçme, tanı koyma ve yerleştirme amacı ile kullanılan ölçme araçları vasıtasıyla alınan sınıflama kararlarının doğruluğu sınıflama geçerliliği olarak ifade edilmektedir. Murphy ve Davidshofer (2005) da bir test uygulandıktan sonra verilecek kararların doğruluğunun, test puanlarının geçerliliği ile ilişkili olduğunu ifade etmektedir. Ayrıca Turgut ve Baykul (1992) da ölçülen yapıdan bağımsız bir şekilde ölçme sonunda alınan kararların çoğunun sınıflama düzeyinde olduğunu ifade etmektedir. Buradan hareketle sınıflama geçerliliği kavramının ve bu bağlamda kanıt toplamanın önemli olduğu söylenebilir. Eğitimde de alınan kararlar çoğunlukla geçti/kaldı, kabul edildi/edilmedi ya da öğrenci özelliklerine yönelik olarak sahip olup olmadığına ilişkin olabilmektedir. Bu durumda sınıflama geçerliliği kanıtına sahip olmayan test puanları ile karar verildiğinde haksız kararların alınmasının olası olduğu ve bunun önüne geçilmesi gerektiği düşünülmektedir.

Eğitimde yapılan sınıflama analizleri sonucunda öğrenciler hakkında doğru kararların alınabilmesi için ölçümlerin geçerlilik ve güvenilirliklerine ilişkin kanıtlar elde edilmesi önemli görülmektedir (Çelikten & Çakan, 2019). Yapılan sınıflama işlemleri sonucunda elde edilen sınıflama doğruluğu ve sınıflama tutarlılığı indeksleri, sırasıyla sınıflama geçerliliği ve güvenilirliği olarak ele alınmaktadır (Barnett & Macmann, 1992; Lee, Hanson & Brennan, 2000). Bu doğrultuda, hesaplanan bu indeksler yoluyla, ölçümler sonucunda alınan kararların geçerliliği ve güvenilirliğine ilişkin incelemeler yapılmasının gerekli olduğu ifade edilmektedir (Çelikten & Çakan, 2019). Aksu ve Doğan (2018) da veri madenciliği yöntemlerini kullandıkları sınıflama çalışmalarında doğru sınıflama sayısı, doğru sınıflama oranı, Kappa istatistiği, ortalama mutlak hata, karekök hata, göreceli mutlak hata ve göreceli karekök hata değerlerini güvenilirlik ölçütü olarak, doğru pozitif oranı, yanlış pozitif oranı, duyarlık, geri çağırma, F-ölçütü, Matthew korelasyon katsayısı, ROC eğrisi altında kalan alan ve Duyarlık-Geri çağırma eğrisi altında kalan alan değerleri ise geçerlilik ölçütleri

olarak almışlardır. Dolayısıyla bu değerlerin yüksek olması daha geçerli ve güvenilir sonuçlar elde edilmesi anlamına gelmektedir.

Sınıflamada geçerli ve güvenilir sonuçlar elde edilmesi için önemli olan bazı hususlar bulunmaktadır. Holden ve Kelly (2010) sınıflamada kullanılan bağımsız değişkenlerin doğru kullanılmasının sınıflama performansı ile ilişkili olduğunu belirtmiştir. Öte yandan, sınıflamaya ait düzey sayısının da sınıflama doğruluğu ve tutarlılığını etkilediği ifade edilmektedir (Ercikan & Julian, 2002; Lathrop & Cheng, 2014). Ayrıca, sınıflamada kullanılan istatistiksel yöntemler ve bu işlemlerde yapılan yanlışlıklar hatalı sınıflama yapılmasına ve dolayısıyla geçerliliğin düşmesine sebep olabilmektedir (Holden & Kelly, 2010). Agarwal, Pandey ve Tiwari (2012) de çalışmalarında, tanımlanmış bir problem için en iyi sınıflandırıcıyı seçmek amacıyla sınıflandırma yöntemlerinin karşılaştırmalı analizlerin yapılması gerektiğini vurgulamıştır. Eğitim bilimleri alanında sınıflama uygulamalarına sıklıkla yer verildiği ve bu uygulamalar sonucunda alınan kararların eğitim ve öğretim sürecinde yer alan birçok paydaşı doğrudan etkilediği bilinmektedir. Dolayısıyla, sınıflama analizinden elde edilen sonuçların doğruluğunu etkileyen faktörler arasında yer alan veri seti özelliklerinin ve kullanılan yöntemlerin performanslarının sınıflama geçerliliğinde ve güvenilirliğinde büyük bir önem taşıdığı söylenebilir.

Son zamanlarda, farklı senaryolardan elde edilen veri setlerinin analizi oldukça önemli olmuştur. Bu veri setleri, büyüklüğüne, veri formatının ve kaynağının çeşitliliğine göre tanımlanabilmektedir. Farklı yapıdaki veri setlerinin, geleneksel veri madenciliği algoritmalarında nadiren ele alındığı ve gerçekleştirilen analizlerde önemli bir faktör olmalarına rağmen az sayıdaki çalışmada, veri madenciliği algoritmalarının performansının hangi koşullarda arttığı ya da azaldığının incelendiği belirtilmektedir (Kwon & Sim, 2013). Veri setinin özellikleri, sınıflama algoritmasının performansının belirlenmesindeki kilit noktalardan birisidir. Örneğin Raudys ve Pikelis (1980) çalışmalarında, sınıflama hatasının örneklem büyüklüğü, değişken sayısı ve sınıflama türü ile ilgili olduğunu ifade etmişlerdir. Performans değerlendirme kriterleri de veri setinin özelliklerine göre birbirinden farklı sonuçlar gösterebilmektedir (Dhurandhar & Dobra, 2009). Bu özelliklerden biri olan

örneklem büyüklüğü ile sınıflama doğruluğu arasında ilişki bulunmaktadır (Okamoto, 1963). Ayrıca, örneklem büyüklüğü küçük olduğunda her bir veri madenciliği yöntemi farklı performans gösterebilmektedir (Morgan, Dougherty, Hilchie, & Carey, 2003). Bağımlı değişkenin sınıf sayısının ikili ya da çoklu olup olmadığı ve buna bağlı olarak verinin gruplardaki dengesiz dağılımı da performansı etkileyen bir diğer özellik olabilmektedir (Minaei-Bidgoli, Kashy, Kortemeyer & Punch, 2003; Li & Hung, 2009; Weng & Poon, 2006). Öte yandan, modelde ele alınan bağımsız değişken sayısının da yöntemlerin performansları üzerinde etkili olduğu ifade edilmektedir (Mannila, 2000). Yapılan bazı çalışmalarda, bağımsız değişken sayısının artmasının sınıflama doğruluğunu arttırdığı bulunmuştur (Dolgun, 2014; Kumar & Chong, 2018). Öte yandan, veri madenciliği yöntemlerinde ele alınan değişken seçiminin asıl amacı, en az değişken seti kullanarak en yüksek doğruluğu elde edebilmektir (Dash & Liu, 1997). Değişken seçimi, önemli değişkenleri tanımlamaya, gereksiz ya da fazla olan değişkenleri çıkarmaya ve bu sayede sınıflama performansını arttırmaya yardımcı olmaktadır (Anwar, Ghany & Elmahdy, 2015; Chandrashekar & Şahin, 2014; Chizi, Rokach & Maimon, 2009) Yani, bağımsız değişkenin sayısından ziyade, bağımsız değişkenin niteliğinin önemli olduğu da vurgulanmaktadır (Dağ, Sayın, Yenidoğan, Albayrak & Acar, 2012). Yu ve Liu (2003) bağımsız değişkenlerin birbirleri ve bağımlı değişken ile olan ilişkilerinin modellerin performansını etkilediğini belirtmiştir. Bu bağlamda araştırma kapsamında veri seti özellikleri olarak ele alınan bağımlı değişken kategori sayısı, bağımsız değişken sayısı ve örneklem büyüklüğü koşullarına yönelik olarak sınıflandırma analizlerinin performanslarının incelenmesi bir gereklilik oluşturmaktadır.

Sınıflama analizinden elde edilen sonuçların doğruluğunu etkileyen faktörlerden birinin de kullanılan yöntemlerin performansları olduğu söylenebilir. Çünkü psikometristler, ele aldıkları argümanların niceliğini belirlemek için verilerdeki ilgili formlara ve kalıplara dayanarak istatistiksel modeller kullanırlar (Mislevy vd. 2012). Dolayısıyla bu istatistiksel modellemelerde kullanılan yöntemlerin performanslarının, ele alınan psikometrik özellikleri yani geçerlilik ve güvenilirliği de etkileyen önemli unsurlardan biri olduğu söylenebilir.

Geleneksel istatistiksel yöntemler, tüm parametrik istatistiksel teknikler ve bunların parametrik olmayan versiyonlarını kapsayacak şekilde tanımlandığında, önceden tasarlanmış belli bir plan dahilinde yürütülen çalışmalardan elde edilen verilerin işlenmesinde iyi performans göstermektedirler. Ancak son yıllarda, elektronik veri toplama ve veri tabanı teknolojisi, geleneksel yöntemlerden farklı şekilde veri toplanmasına olanak sağlamaktadır (Wegman, 1988; Hand, Mannila & Smyth, 2001, s.17). Bu durum sonucunda ortaya çıkan büyük veri setlerinin analizinde geleneksel istatistiksel teknikler birçok soruna karşı çözüm üretmekte yetersiz kalabilmektedir. Geleneksel yöntemler, önceden tasarlanmış belli bir plana göre elde edilmeyen veri setlerinde bulunan gürültülü veriler ve düzensizlikler (hatalı veri, eksik veri, uç değer) ile başa çıkmak için yeterli değildir. Ayrıca bu yöntemler genellikle normallik, doğrusallık ve bağımsızlık gibi güçlü model varsayımlarının ihlaline karşı dayanıklı değildir (Hand vd. 2001, s.17). Öte yandan, geleneksel yöntemler temel olarak bir veya birkaç bağımlı değişkenin değişkenliği ve bu değişkenliğe bağlı araştırma problemleri ile ilgilidir; ancak daha karmaşık bir veri setinde, veri yapısı ve değişkenler arasındaki fonksiyonel ilişkiler hakkında fikir edinirken daha fazla araştırma problemi ortaya çıkabilir. Bu nedenle, istatistiksel çıkarımdan yapısal çıkarıma kayan bir durum söz konusudur (Wegman, 1988). Wegman (2000) de geleneksel istatistiklerin büyük veri setlerinin analizindeki dezavantajlarını vurgulayarak, veri toplamanın doğası ve yollarındaki değişikliklerin veri analizinde yeni araçlar ve teknikler gerektirdiğini savunmuştur. Bu nedenle, araştırmacıların geleneksel yaklaşımlarla sınırlandırılmaması ve yeni sorunları çözmek için aktif olarak yeni yöntemler aramaları önerilmektedir. Özellikle çok sayıda değişkeni olan büyük ölçekli verilerle çalışırken, araştırmacıların gizli örüntüleri keşfetmek ve daha önce bilinmeyen değişken ilişkilerini ortaya koymak için doğru yöntemleri aramaları gerektiği ifade edilmektedir (Xu, 2003, s.23). Bu bağlamda eğitim alanında sınıflamaya dönük alınan kararlarda daha geçerli ve güvenilir sonuçlara ulaşabilmek için yeni yöntemlerin kullanılması ve bu yöntemlerin karşılıklı olarak incelenmesi gerektiği söylenebilir.

Bu çalışmada ele alınan yöntemler, Yapay Sinir Ağları (YSA; Artificial Neural Networks), Rastgele Orman (RO; Random Forest), Destek Vetör Makinası (DVM; Support Vector Machine), Sınıflama ve Regresyon Ağaçları (SVRA; Classification and Regression Trees) ve Lojistik Regresyon'dur (LR; Logistic Regression). Bu yöntemler eğitimde veri madenciliğinde tahminleme modellerinde yani sınıflama ve regresyonda en popüler olan yöntemlerdir (Baker, 2010). Bu yöntemlerin tercih edilmesindeki gerekçeler herbir yöntemin sahip olduğu belirgin özellikler doğrultusunda ortaya konmaya çalışılmıştır. SVRA, 1984 yılında Breiman, Friedman, Olshen ve Stone tarafından sınıflandırmaya yönelik analizler yapmak amacıyla geliştirilmiş parametrik olmayan istatistiksel bir yöntemdir (Yohannes & Webb, 1999, s.4). Karar ağacı modellerinden biri olan SVRA, ağaç oluşumundaki bölünmelerde yakınsamaya dayalı sayısal yöntemleri kullanır. Hızlıdır, eksik değerlerle baş edebilir, büyük veri setlerini girişte kolayca yönetebilir, hatayı minimize etmek için kendi içinde bir takım işlemler gerçekleştirebilir (Köse, 2018, s.238-239). SVRA'nın diğer bir avantajı parametrik istatistiklerin varsayımlarına tabi olmaması ve belirli bir sonucu tahmin etmek için çok sayıdaki bağımsız değişken arasından en önemli olanları seçebilmesidir (Chudzinska & Baralkiewicz, 2011). 2001 yılında Breiman tarafından geliştirilen RO yöntemi ise, sınıflama değişkenlerini belirlemek, değişkenler arasındaki ilişkiyi tespit etmek ve kümeleme analizi yapmak amacıyla kullanılır (Breiman & Cutler, t.y; Gislason, Benediktsson & Sveinsson, 2006). RO yönteminde verilerin herhangi bir işleminden geçirilmesine gerek duyulmaz, diğer yöntemlerin aksine ihtiyaç duyulan tek şey ağaç ve değişken sayısı parametreleridir (Yan, 2017, s.31). Aşırı uyuma karşı da dayanıklı olan bu yöntem, büyük veri tabanlarında hızlı ve verimli bir performans sergiler (Breiman ve Cutler, t.y).

YSA, ilk kez 1943 yılında McCulloch ve Pitts tarafından ortaya konan ve biyolojik sinir ağlarını temel alan bir modeldir (Akpınar, 2017, 309-310). YSA, öğrenme, ilişkilendirme, sınıflandırma, genelleme, özellik belirleme ve optimizasyon gibi bir çok konuda iyi performans göstermektedir (Öztemel, 2016). YSA'nın bazı özellikleri şunlardır: Eksik veya normal dağılım göstermeyen verilerle çalışabilmesi, hata toleransına sahip olması, gürültülü

veri setinden etkilenmemesi, üzerinde çalışılan problem hakkında detaylı bilgiye ihtiyaç duymaması ve daha önce görmedikleri örnekler hakkında bilgiler üretmesidir (Ateş, 2008; Özçalıcı & Ayriçay, 2016; Öztemel, 2016). Çalışmada ele alınan bir diğer yöntem DVM ise, 1992 yılında Boser, Guyon ve Vapnik tarafından geliştirilmiştir (Akpınar, 2017, s.291). Fen bilimlerinde çokça çalışılmıştır ve verimli uygulamaları ortaya konmuş olan bu yöntem son zamanlarda sosyal bilimler alanında da kullanılmaktadır (Attewell, Monaghan & Kwong, 2015, s.147). Çözüm üretirken eğitim verilerinin küçük bir alt kümesini kullanmasından ötürü hesaplama kolaylığına sahip olan bu yöntem, aynı zamanda çok boyutlu durumlarda çıkan zorlukların çözümünü de kolaylaştırmaktadır (Cristianini & Shawe-Taylor, 2000). DVM yönteminin avantajları, güvenilirliğinin yüksek olması, ezber öğrenme yeteneğinin olması ve lineer olmayan sınıflamalarda da yüksek başarıya sahip olmasıdır (Akpınar, 2017, s.291).

Çalışmada ele alınan diğer bir yöntem olan LR analizinin temelini, Berkson tarafından 1944, 1953 ve 1955 yıllarında yapılan çalışmalar oluşturmaktadır (Çokluk, Şekercioğlu & Büyüköztürk, 2012, s.57). LR bir veya daha fazla bağımsız değişken ile bir ikili veya çoklu sınıfa sahip bağımlı değişken arasındaki ilişkiyi inceler (Mohamed, 2015, s.21). LR'nin bağımlı ve bağımsız değişkenlerin normal dağılım göstermesi, bağımlı değişken ile bağımsız değişken/ler arasında doğrusal ilişkinin olması ve grupların eşit varyansa sahip olması gibi varsayımlara ihtiyacı yoktur. Ayrıca LR modelleri diğer yöntemlere göre daha esnek olup sürekli, kesikli veya her ikisinin de bulunduğu karma değişken gruplarıyla çözüm üretebilecek yapıya sahiptir (Tabachnick, Fidell & Ullman, 2013, s.439). Görüldüğü üzere sınıflandırma yöntemlerinin sahip oldukları özellikler ve avantajlar farklılık göstermektedir. Dolayısıyla farklı koşullar altında sınıflandırma analizlerinin performanslarının incelenmesi ve hangisinin daha etkin bir çözüm sunduğunun belirlenmesi gerekmektedir. Aynı zamanda çok sayıda yöntemin bulunması nedeniyle hangilerinin daha etkin performans gösterdiğinin belirlenmesi sonuçların güvenilirliğini ve geçerliliğini etkilemesi bakımından önemli görülmektedir (Souza, Matwin & Japkowicz, 2002)

Eğitim bilimleri alanında öğrencilerin sınıflandırılmasına yönelik olarak ölçme ve değerlendirme süreci sonunda alınan kararlar büyük önem taşımaktadır. Çünkü alınan bu kararlar, öğrencilerin eksikliklerinin belirlenmesini, geri dönüt verilmesini, uygulanan eğitim öğretim programlarının biçimlendirilmesini, belirli bir performansa yönelik olarak öğrencilerin yeterli olup olmadığının tespit edilmesini sağlamakta ve birçok aşamada eğitim ve öğretim sürecinde yer alan paydaşlara bilgi vermektedir. Dolayısıyla eğitim ve öğretim sürecini ve bu süreçte yer alan paydaşların aldığı kararları etkileyen bu bilgilerin mümkün olduğu kadar gerçek durumu yansıtmaları önemlidir. Yani ölçme ve değerlendirme sürecindeki tüm adımların mümkün olduğu kadar hatalardan arınık olması gerekmektedir. Bu bağlamda öğrencilerin sınıflandırılmasına yönelik verilen kararların doğruluğunda ve geçerliğinde elde edilen ölçümlere uygulanan analiz yöntemlerinin performansları da büyük önem taşımaktadır. Çünkü uygulanan analiz yöntemlerinin performansı ne kadar iyi olursa ulaşılan sonuçlara karışan hata miktarı daha az olur ve dolayısıyla kararların doğruluğu, yorumlanabilirliği ve geçerliği bu duruma bağlı olarak daha yüksek olur. Analiz yöntemlerinin performansı, kullanılan yöntemin yapısı ve veriye ait özelliklerle doğrudan ilişkilidir. Dolayısıyla sınıflama analizinde kullanılacak yöntemlere karar verilirken hem yöntemlere ait özelliklerin hem de veri setlerine ait özelliklerin göz önünde bulundurulması gerekmektedir. Buna yönelik olarak araştırmacıların ve uygulayıcıların sınıflama analizi yaparken veri setine ait özellikler arasında yer alan sınıf sayısı yani bağımlı değişkenin kategori sayısı, bağımsız değişken sayısı ve örneklem büyüklüğünü ve yöntemlerin bu koşullardaki performansını göz önünde bulundurması gerekmektedir. Çünkü veri setinin sahip olduğu özellikler ve yöntemlerin performansı sınıflandırma analizinden elde edilen sonuçlara göre alınacak kararların doğruluğunu, yorumlanabilirliğini ve geçerliğini doğrudan etkileyecektir. Bu bağlamda eğitim bilimleri alanında sınıflama ve regresyon analizlerinde çoğunlukla kullanılan yöntemler arasında yer alan YSA, RO, DVM, SVRA ve LR yöntemlerinin bağımlı değişkenin kategori sayısına, bağımsız değişken sayısına ve örneklem büyüklüğüne bağlı olarak sınıflama performanslarının incelenmesi bu çalışmanın problemini oluşturmaktadır.

1.2. Araştırmanın Amacı

Bu araştırmanın amacı, sınıflandırma analizlerinde kullanılan YSA, RO, DVM, SVRA ve LR yöntemlerinin performanslarını, bağımlı değişkenin kategorisi sayısı, bağımsız değişken sayısı ve örneklem büyüklüğüne göre PISA 2015 verilerini kullanarak karşılaştırmaktır. Bu amaç doğrultusunda aşağıdaki araştırma sorularına cevap aranmıştır.

1. Öğrencilerin sınıflandırılmasında kullanılan YSA, RO, DVM, SVRA ve LR yöntemlerinin performansları bağımlı değişkenin kategori sayısının değişimine göre farklılaşmakta mıdır?
2. Öğrencilerin sınıflandırılmasında kullanılan YSA, RO, DVM, SVRA ve LR yöntemlerinin performansları bağımsız değişken sayısının değişimine göre farklılaşmakta mıdır?
3. Öğrencilerin sınıflandırılmasında kullanılan YSA, RO, DVM, SVRA ve LR yöntemlerinin performansları örneklem büyüklüğünün değişimine göre farklılaşmakta mıdır?
4. Öğrencilerin sınıflandırılmasında YSA, RO, DVM, SVRA ve LR yöntemlerinin performansları birbirlerine göre farklılaşmakta mıdır?

1.3. Araştırmanın Önemi

Eğitim bilimleri alanında, öğrenciler veya eğitim öğretim sürecindeki diğer paydaşların belirli bir özelliğe yönelik yeterli olup olmaması, belirlenen sınıflardan hangisinde yer alacağına karar verilmesi, sınıf düzeyini etkileyen değişkenlerin belirlenmesi ve bu değişkenler arasındaki ilişkilerin incelenmesi ve bunlara benzer bir çok problemde sınıflama uygulamalarına rastlanmaktadır. Örneğin, sınıflama zorbalığa uğrama durumu, bağımlılık, motivasyon, tutum ve hazırbulunuşluk gibi birçok öğrenci özelliklerinin kapsamlı analizi ve geçti/kaldı, yeterlilik, muafiyet gibi iki veya daha fazla çıktıya sahip değerlendirme uygulamalarında tahminde bulunmak için kullanılabilir. Eğitim bilimleri alanında sınıflamaya dönük birçok uygulamanın bulunması ve bu uygulamalar sonucunda alınan

kararlar doğrultusunda eğitim-öğretim sürecinde yer alan paydaşların doğrudan etkilenmesi alınan bu kararların doğruluğunu, yorumlanabilirliğini, kullanılabilirliğini ve geçerliğini önemli kılmaktadır. Dolayısıyla ulusal ve uluslararası alanda eğitim ve öğretim sürecine yönelik yapılan sınıflama uygulamalarının mümkün olduğunca hatalardan arınık ve gerçek durumu yansıtacak düzeyde olması büyük önem taşımaktadır. Buradan hareketle sınıflamaya dönük uygulamalar için gerçekleştirilen ölçme ve değerlendirme sürecinin tüm aşamalarının olabildiğince hassas ve hatasız bir biçimde yürütülmesi gerektiği söylenebilir. Örneğin, yetenek sınavı ile öğrenci alan bir yükseköğretim kurumunda kabul edildi/edilmedi biçiminde iki sınıfa sahip olarak alınan kararda hata yapıldığında öğrenciler hak kaybına uğrayabilir. Ortaya çıkan bu olumsuz durum uygulayıcılardan, ölçme ortamından, ölçme araçlarından ve bunun gibi birçok durumdan kaynaklandığı gibi veri setine ait özelliklerden ve sınıflamada kullanılan yöntemlerin performansından da kaynaklanabilir. Bu bağlamda yapılan bu araştırma ile sınıflama uygulamalarının hangi durumlarda daha az hataya sahip olacağı ortaya konulmaya çalışılacaktır ve araştırma bu yönü ile önemli görülmektedir.

Sınıflama uygulamaları çerçevesinde yürütülen ölçme ve değerlendirme sürecinin önemli aşamalarından biri de yapılan uygulamalar sonucunda elde edilen ölçümlerin değerlendirilmesidir. Şüphesiz teknolojinin gelişimi ile paralel olarak büyük bir ivmeyle artış gösteren veri yığınları içerisinde sağlıklı sonuçlara ulaşmak gün geçtikçe daha zorlaşmaktadır. Veri setlerindeki değişken sayısının fazla olması, örneklem büyüklüğünün geniş olması, değişkenler arasındaki örüntü yapılarının farklılaşması, sınıflarda yer alan bireylerin farklı özelliklere sahip olması, kayıp veri ve çoklu bağlantılılık gibi bir çok analize ait varsayımların sağlanmasının zor olması gibi bir çok karmaşık durum bulunmaktadır. Bu nedenle PISA, TIMMS ve ABİDE gibi geniş bir örneklemde ve hatta veri seti özelliklerinden dolayı daha küçük örneklemde bile gerçekleştirilen sınıflamaya dönük uygulamalar sonucunda karar alınmasını zorlaşmaktadır. Ayrıca COVID-19 salgını süresince ve sonrası online eğitimin kullanımının yaygın hale gelmesi ve ilerleyen süreçlerde de yaşanabilecek bu durumların ortaya çıkması veri setlerinden elde edilecek sınıflamaya dönük alınacak kararlarda karmaşık yapılarla başedebilen veri madenciliğine dayalı yöntemlerin kullanımını

ön plana çıkarmaktadır. Dolayısıyla ölçümlere ait veri setlerinin taşıdığı özelliklerin ve bu veri setlerinden sonuca ulaşmak için kullanılan analiz yöntemlerinin, sınıflama uygulamaları sonucunda elde edilecek kararların doğruluğunu, yorumlanabilirliğini, kullanılabilirliğini ve geçerliğini doğrudan etkilediği söylenebilir. Bu bağlamda gerek eğitim bilimleri gerekse farklı disiplinlerde yapılan araştırmalar sonucunda sınıflamaya dönük alınacak kararların daha doğru biçimde gerçekleştirilebilmesi için, veri setlerinin özelliklerine ve kullanılan analiz yöntemlerine yönelik olarak çeşitli çözüm önerileri gündeme gelmektedir. Bu bağlamda önerilen çözümler sonucunda ortaya çıkan “Sınıflamaya dönük uygulamalarda hangi öneri daha etkin ve doğrudur?” sorusunun cevabı büyük önem taşımaktadır.

Eğitim bilimleri alanında sınıflamaya yönelik yapılan uygulamalarda kullanılan veri setlerinin farklılaştığı birçok faktör vardır. Bu faktörler arasında yer alan bağımlı değişkenin kategori sayısı, bağımsız değişken sayısı ve örneklem büyüklüğü değişimine göre sınıflama performansının etkilenip etkilenmemesi durumunun incelenmesi uygulayıcılara bir çok durum hakkında ön bilgi verecektir. Sınıflama, ölçme ve değerlendirme süreci sonunda elde edilen ölçümlerin analiz edildiği bir uygulama basamağı olmaktan ziyade tüm sürecin içinde yer alan ve süreçle birlikte ilerlemesi gereken uygulama basamaklarını içermektedir. Bu nedenle sürecin en başından itibaren mümkün olduğu kadar iyi bir sınıflama performansı elde etmek için bağımlı değişkenin kategori sayısının belirlenmesinde, sınıflama için hangi bağımsız değişkenlerin önemli olduğunun ve kaç tanesinin modelde yer alması gerektiğinin tespit edilmesinde ve çalışmada kullanılacak en düşük örneklem büyüklüğüne karar vermede, yapılan bu araştırma sonucunda elde edilen bulguların büyük katkı sağlayacağı öngörülmektedir. Elde edilen bulgular çerçevesinde, benzer koşullar altında gerçekleştirilecek sınıflama uygulamalarında uygulayıcıların değişken özelliklerini belirlerken öncelikle daha iyi bir performansta sınıflama uygulaması gerçekleştireceği, verileri elde ederken ve sınıflamaya dönük modeller kurarken daha az vakit haracayacağı ve bireyler hakkında daha doğru kararlar elde edeceği düşünülmektedir. Araştırma bu yönüyle sağlayacağı avantajlar açısından önemli görülmektedir.

Eđitim bilimleri alanında sınıflama yapmak amacıyla kullanılan bir ok yntem ve bu yntemlere ait farklı algoritmalar bulunmaktadır. Bu algoritmaların bir ok eşidi vardır ve her biri farklı parametrelere gre alışır. Ayrıca algoritmaların amalarına gre farklılık gstermesi, kullandıkları veri kaynaklarının ve veri trlerinin farklılaşması, modellerde yer alan deęiřkenlerin farklı zelliklere sahip olması gibi bir ok nedenden dolayı farklı sonular elde edilebilir. Dolayısıyla, hangi algoritmanın daha yansız ve doęru sonular rettięi, hangi algoritmanın hangi kořullarda iyi performans sergiledięi hususlarına ynelik edinilecek bilgiler uygulamaların verimini arttıracak, harcanan zamanı azaltacak ve daha doęru sonuların elde edilmesini saęlayacaktır. Bu durumlar gznne alındıęında yapılan bu arařtırma erevesinde algoritmaların farklı kořullara gre genel performanslarının karřılařtırılarak incelenmesinmli grlmektedir.

1.4. Sınırlılıklar

Bu arařtırmada baęımlı deęiřkenin kategori sayısının deęiřimine gre yntemlerin performansı incelenirken, baęımlı deęiřkenin 6, 3 ve 2 kategorili olması durumu ele alınmıřtır. PISA 2015 verilerine gre hazırlanan teknik rapordan hareket edilerek orijinal 6 kategoriye ek olarak 3 ve 2 kategorili yapılar oluřturulabilmiř ancak 4 ve 5 iin uygun bir belirleme metodu seilememiřtir. Baęımlı deęiřkenin srekli halinden yola ıkarak bařka yntemlerle, srekli deęiřken kategorik hale getirilmeye alıřılmıř ancak bu durumda da PISA 2015 teknik raporunda 6 dzey iin belirlenen ęrenci daęılımında ok farklı bir duruma ulařılmıřtır. Bu durum arařtırmada kullanılacak veri setini gerek durumdan uzaklařtıracadıęı iin PISA 2015 teknik raporuna sadık kalınarak 6 kategorili yapı yalnızca 2 ve 3 kategorili hale getirilmiřtir. Dolayısıyla arařtırma baęımlı deęiřkenin kategori sayısının 2, 3 ve 6 olarak alınmıř ve aradaki kategori sayıları gndeme alınmadıęı iin bu ynyle arařtırma aısından bir sınırlılık oluřmuřtur.

Arařtırmada modelde yer alın baęımsız deęiřkenler seilirken mmkn olduęun kadar korelasyon katsayısı birbirine yakın olan deęiřkenler kullanılmaya alıřılmıřtır. Ancak gerek veri setine dayalı olarak alıřıldıęı iin modele alınan veya kořullara gre modelden

ıkarılan deęiřkenlerin korelasyon katsayıları arasında kk farklılıklar bulunmaktadır. Bu bakımdan sınırlıdır.

BÖLÜM II

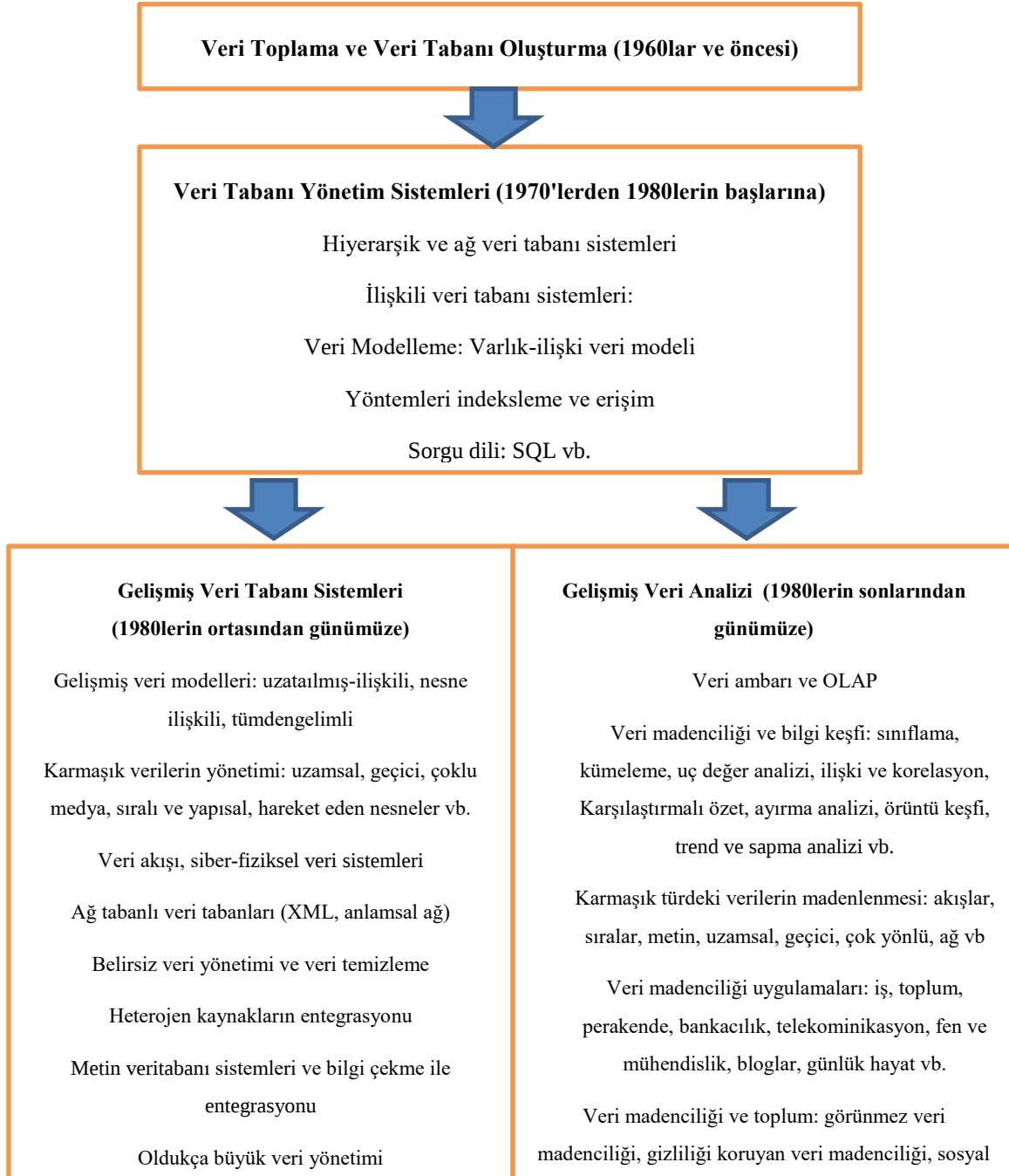
KURAMSAL TEMELLER

Bu bölümde araştırmanın konusu ve problem durumu ile ilgili kuramsal temeller sunulmuştur. Bu kapsamda veri madenciliği, veri madenciliği modelleri, Sınıflandırma ve Regresyon Ağaçları (SVRA), Rastgele Orman (RO), Yapay Sinir Ağları (YSA), Destek Vektör Makinası (DVM) ile Lojistik Regresyon (LR) yöntemlerine ilişkin bilgilere ve literatürde yer alan çalışmalara yer verilmiştir.

2.1. Veri Madenciliği

Günümüzde teknolojinin gelişmesi, veri tabanlarının artışı ve bilgi teknolojilerinin yaygınlaşan kullanımından dolayı, organizasyonlar karmaşık veri tabanları ve büyük verilerle karşılaşmaktadır (Ersöz, 2019, s.13). Bu tür veriler üzerinde çözümleme yaparak yararlı bilgiye ulaşmak amacıyla ihtiyaç duyulan çözümleme yöntemlerini kapsayan kavram veri madenciliğidir (Özkan, 2016, s.12). Bu kavramın temelini 1960'larda istatistikçilerin analiz öncesinde kötü verilerin tespit edilip ayıklanması amacıyla kullandığı Veri Avlama (Data Fishing) veya Veri Tarama (Data Dredging) terimleri oluşturmaktadır (Pektaş, 2013, s.99). 1960'lı yıllardan beri, veritabanı ve bilgi teknolojisi, ilkel dosya işleme sistemlerinden gelişmiş ve güçlü veritabanı sistemlerine kadar sistematik olarak gelişmiştir. 1970'lerden bu yana veri tabanı sistemlerinde yapılan araştırma ve geliştirmeler, önceki hiyerarşik ve ağ veri tabanı sistemlerinden ilişkisel veri tabanı sistemlerine, veri modelleme araçlarına, indeksleme ve erişim yöntemlerine doğru ilerleme göstermiştir. Veri tabanı ve veri yönetimi

gelişim süreci Şekil 1’de yer almaktadır (Han, & Kamber, 2001, s.3; Han, Pei, & Kamber, 2012, s.3).



Şekil 1. Veri tabanı ve veri yönetimi gelişim süreci.

Veri madenciliği, bilgi teknolojisinin doğal evrimi sonucu görülebilir (Han, Pei, & Kamber, 2012, s.2.). Teknoloji ile paralel olarak her geçen gün gelişmekte ve kullanım alanı yayılmakta olduğu için, veri madenciliği kavramını tek bir tanım ile ifade etmek mümkün

değildir (Silahtaroglu, 2016, s.12). Bu bağlamda literatür incelendiğinde veri madenciliğinin birçok tanımına rastlamak mümkündür. Köse'ye (2018) göre veri madenciliği “Büyük verilerde yer alan ilginç, yararlı ve keşfedilmemiş örüntüler ile ilişkilerin tespit edilmesi işlemidir” (s.48). Özkan'a (2016) göre ise “Büyük ölçekli veriler arasından değerli olan bir bilgiyi elde etme işidir” (s.12). Ersöz'e (2019) göre veri madenciliği gelişmiş teknoloji ve iş deneyimleri ile birlikte kullanılarak veri yığınlarından anlamlı bilgilerin ortaya çıkarıldığı interaktif ve iteraktif bir süreçtir (s.14). Veri madenciliği, hızlı bir biçimde gelişimini devam ettirmekte olan bir disiplindir. Temel amacı büyük veri tabanlarında yer alan faydalı bilgilerden faydalanarak, kullanışlı model ve örüntüleri ortaya çıkaran araştırmaya yönelik istatistiklerin bir dalı olarak tanımlanabilir (Zhi, Yan & Yan, 2016). Veri madenciliği, kabul edilebilir hesaplama verimliliği sınırlamaları altında, veriler üzerinde belirli bir model sayımı (veya model) üreten veri analizi ve keşif algoritmalarının uygulanmasından oluşan bir adımdır. Başka bir ifadeyle veriden desen çıkarmak için özel algoritmaların uygulanmasıdır (Fayyad, Piatetsky-Shapiro & Smyth, 1996). Veri madenciliği, kullanılan veri setini araştırmacılar için anlaşılabilir ve yararlı yeni yöntemlerle özetlemek ve beklenmeyen ilişkileri keşfetmek amacıyla genellikle büyük gözlemlere ait veri setlerinin analiz edilmesidir (Hand, Mannila & Smyth, 2001, s.6). Veri madenciliği, büyük veri tabanlarından bilgi elde etme sorununun çözümü için makine öğrenmesi, örüntü tanıma, istatistik, veri tabanı ve görselleştirme tekniklerini bir araya getiren disiplinlerarası bir alandır (Cabena, Hadjinian, Stadler, Verhees & Zanasi, 1998).

Literatürde yer alan tanımlar incelendiğinde veri madenciliğinin, çeşitli disiplinleri kapsayan disiplinlerarası bir araştırma alanı olduğu görülmektedir (Hand, Mannila & Smyth, 2001, s.7). Fayyad vd. (1996)'ne göre veri madenciliği, makine öğrenmesi, istatistik ve veritabanı alanları ile yakından ilişkilidir. Veri madenciliği ve istatistik, benzer yöntemleri kullanarak veriden bilgi elde eden iki disiplindir (Köse, 2018, s.49). Veri madenciliğinde yer alan yöntemler temelde istatistik ve yapay zekânın uzantısı olan makine öğreniminden beslenmektedir (Akpınar, 2017, s.54). Veri madenciliğinde yer alan uygulamalar, veri yığınlarındaki ilişkilerin, örüntülerin ve eğilimlerin ortaya çıkarılması ihtiyacına yönelik

olarak istatistiksel analiz ve makine öğrenimi tekniklerinin birlikte kullanımı sonucunda ortaya çıkmıştır (Ersöz, 2019, s.16).

Veri madenciliği ve istatistik, veriden bilgi öğrenilmesi, verinin bilgiye dönüştürülmesi, verilerin anlamlandırılması, belirsizliklerin ortadan kaldırılması, bir olayı etkileyen değişkenlerin belirlenmesi ve üretilen modeller aracılığıyla ileriye dönük tahminler yapılması hususlarında yakın ilişki içinde olup benzerlikler göstermektedir. Veri madenciliği büyük miktardaki veriden bilgi çıkarmak ve yapıyı açıklamak için veri analizi kullanması bakımından istatistik ile iç içedir (Ersöz, 2019, s.22, s.24).

Önceki dönemlerde evren üzerinde çalışma zorluğu yaşanırken, veri tabanı yönetim sistemlerinin ve teknolojinin gelişmesi sayesinde büyük miktardaki verilerin transferi kolaylaşmıştır. Bu durum ise veri madenciliği ve istatistik arasındaki farkın daha net ortaya çıkmasına neden olmuştur. Veri madenciliği ve istatistik arasındaki en önemli fark üzerinde çalışılan verinin kapsam ve miktarıdır. İstatistikte örneklem ile çalışma yapılırken, veri madenciliğinde ana kitle (evren) ile çalışmalar yapılır (Köse, 2018, s.49). Veri madenciliği istatistiksel uygulamalara çok benzer ancak istatistikte yer alan uygulamalar düzenlenmiş ve çoğunlukla özet verilerle çalışırken, veri madenciliğinde yer alan uygulamalar büyük miktardaki veri ve çok fazla sayıdaki değişken ile çalışır (Özkan, 2016, s.12).

Zhao ve Luan (2006) veri madenciliği ve istatistik arasındaki farklılıkları aşağıdaki biçimde sıralamıştır:

Teorinin rolü. İstatistik ile teori birlikte hareket eden bir ilişki yapısına sahiptir. Teori gözlemler ve olaylar hakkında bir çerçeve ve bakış açısı sağlar. Bu nedenle teorinin yardımı olmadan ortaya çıkan gözlem ve olaylar tam olarak anlamlandırılmayabilir. İstatistik doğrulayıcı bir süreçtir ve istatistiksel analiz önceki bilgilere dayalı olan bir teoriyle başlar ve bu teorinin reddedilmesi veya onaylanması için kanıtlar sunar. Veri madenciliği teorinin doğrulanmasına odaklanmaz. Ancak bu durum bilgisayarın aranan örüntü ve tahminleri kendiliğinden bulacağı anlamına gelmez. Bu durum veri madenciliği hakkındaki ön

yargılardan biridir. Fakat veri madenciliği araştırmacıdan gelecek açık talimatlara ihtiyaç duyar.

Değişkenler arasındaki varsayımlar. Veri madenciliği istatistiğe göre değişkenler arasındaki ilişkiler hakkındaki varsayımlar bakımından daha az sınırlandırılmıştır. Bundan dolayı başka bir biçimde elde edilemeyen sonuçların ortaya çıkması için daha geniş bir alana sahiptir. Veri madenciliği bireysel düzeyde tahmin doğruluğuna odaklanır ve modelin açıklayıcı gücünden fazla endişe etmez. Veri madenciliği kolektif nedenleri yorumlamaya veya izole etmeye çalışmak yerine her bir analiz birimi ile ilgilenir. Bu nedenle değişkenler arasındaki yüksek korelasyon sonucunda ortaya çıkan çoklu bağlantılılık (multicollinearity) gibi istatistikte büyük bir problem olarak görülen durumlar, veri madenciliğinde kullanılan çeşitli yöntemlerle giderilebilir.

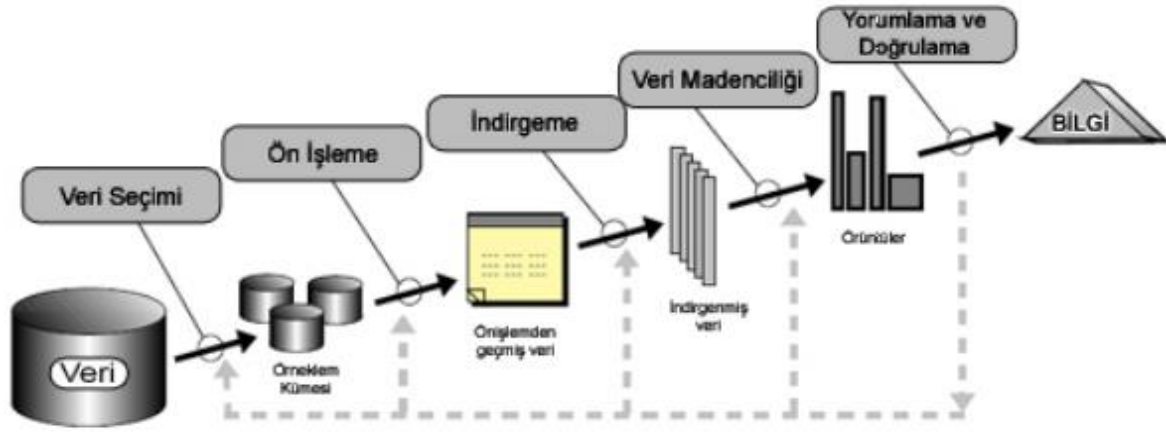
Genellenebilirlik. İstatistikte araştırmacının ilgisinin örneklem olması nadir rastlanan bir durumdur. Araştırmacılar evren hakkında bilgi edinebildiği miktardaki örneklemle ilgilenir. İstatistik bireysellikle başa çıkmak yerine konular arasındaki ortaklıkları keşfetmek için tasarlanmıştır. Veri madenciliğinin amacı, daha ayrıntılı, özel ve yerel bilgileri bir araya getirmektir ve çok geniş bir gruba genellenebilecek geniş çaplı bir açıklama yapılması aşırı derecede önemsizdir (Luan, 2002).

Hipotez testleri. Veri madenciliğinde bir teori veya hipotez ile başlanmaz ve sonuçların bir evrene genelleştirilmesi önemli sayılmaz bu nedenle hipotez testleri özel bir anlam ifade etmez. Veri madenciliği büyük miktarda veri ile çalışma imkânı verdiği için çoğunlukla evren, çalışma evrenidir. Dolayısıyla istatistikte çıkarımların doğruluğunu belirlemek amacıyla kullanılan anlamlılık düzeyi geçerliğini yitirir. Sosyal bilimlerdeki araştırmacılar, genellikle klinik deney türünde bir analiz yapma lüksüne sahip değildir. Böylece, hipotez testi en kesin anlamıyla ciddi biçimde kısıtlanır. Daha detaylı olarak inceleme imkânı tanıdığı için veri madenciliği bireysel konular veya vakalar için puanlar üreterek daha bireysel ve ayrıntılı bulgular sağlamada avantajlıdır.

Veri madenciliği ile ilişkili diğer önemli bir alan makine öğrenmesidir. Makine öğrenmesi, istatistiksel yapay zekâ algoritmaları kullanarak veri hakkında sonuç çıkarmaya, değerlendirme yapmaya ve karar vermeye imkân verir (Ersöz, 2019, s.14). Makine öğrenmesi, bir makinenin kendi kendine en iyi performansı göstermesi amacıyla, verilerin kümelenmesini ve sınıflandırılmasını sağlayan öğrenmeye dayalı tekniklerin yer aldığı, yapay zekâ ile ilişkili bir çalışma alanıdır (Balaban & Kartal, 2018, s. 2). Veri madenciliği ve makine öğrenmesi çoğunlukla birbirleri yerine kullanılan iki kavramdır. Bu iki disiplinde ortak olan en önemli kavram öğrenme olup, öğrenmeye dayalı olan birçok yöntem hem veri madenciliği hem de makine öğrenmesi alanlarında kullanılır (Köse, 2018, s.205). Öğrenme, yazılımların giriş olarak verilen veri setinden algoritmalar yardımıyla sonuçlar ve kurallar çıkarmasıdır (Akpınar, 2017, s.50). Makine öğrenmesinde veri setinin özelliği bakımına göre ele alınan algoritmaların seçiminde iki kavramından bahsedilir. Bunlar danışmanlı (denetimli) ve danışmansız (denetimsiz) öğrenmedir. Eğer ele alınan veri setinde çıktı değerleri hakkında bir tanımlama mevcut ise denetimli öğrenme aksi durumda denetimsiz öğrenmedir. Denetimli öğrenmede temel amaç sınıflandırmayı sağlayan bir model oluşturmak ve aynı şartlarda yeni verilere yönelik tahminler yapmaktır. Denetimsiz öğrenmede amaç belirli özelliklerdeki benzerliklerine göre birbirine yakın örnekleri kümelemek veya model oluşturmada ilişkiler dikkate alınarak birliktelik analizleri ile çözüm oluşturmaktır (Balaban ve Kartal, 2018, s. 3).

2.2. Veri Madenciliği Süreci

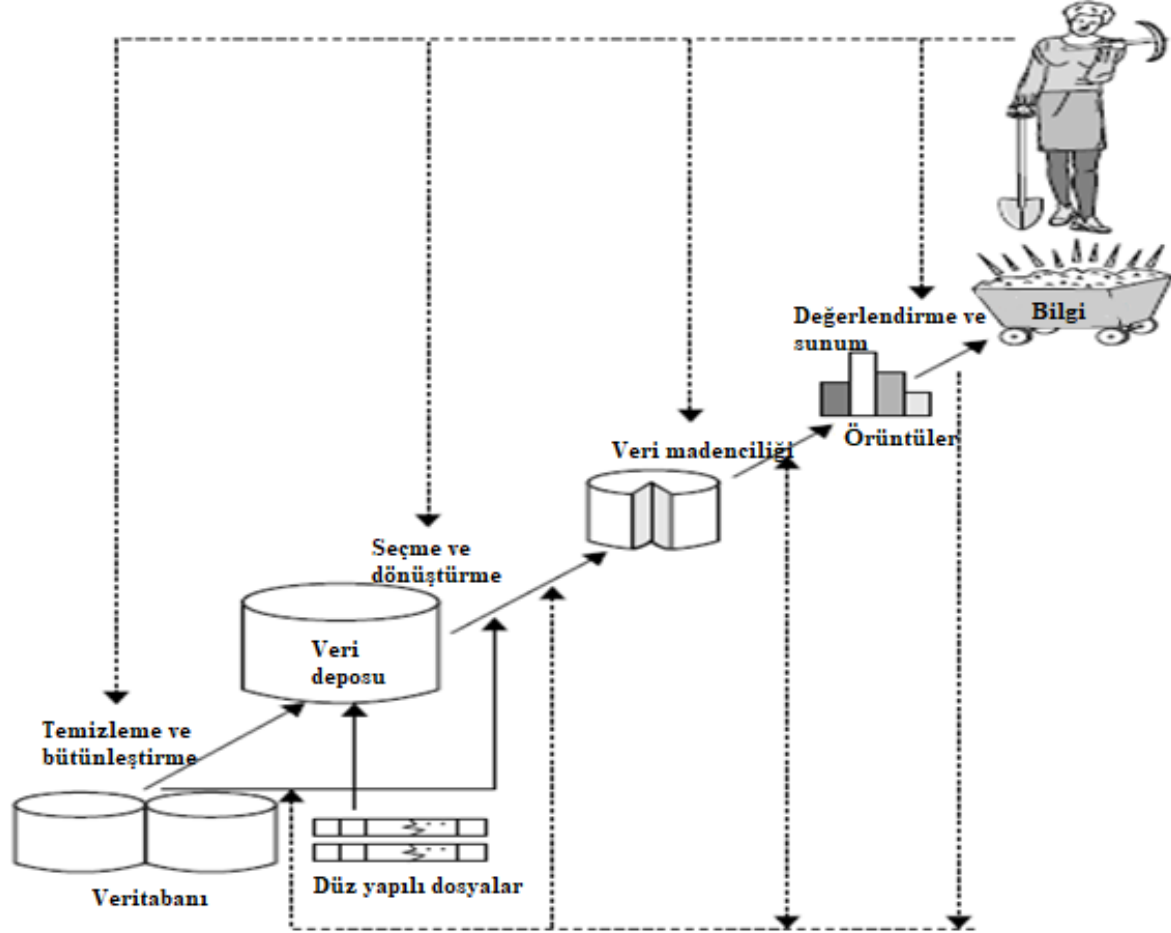
Literatür incelendiğinde veri madenciliği ve veri tabanlarında bilgi keşfi konusunda araştırmacılar tarafından çeşitli süreçler oluşturulmuştur. Bunlardan biri Fayyad vd. (1996) tarafından seçim, ön işlem, dönüştürme, veri madenciliği ve yorumlama olmak üzere beş ana, toplamda ise dokuz adımdan meydana gelen akademik bir süreçtir (Akpınar, 2017, s.76). Fayyad vd. (1996)'ne göre oluşturulan veri madenciliği süreci Şekil 2'de yer almaktadır.



Şekil 2. Veri madenciliği süreci Fayyad vd. (1996).

Fayyad vd. (1996)'ne göre veri madenciliği sürecinin ilk adımında, sürecin amacı belirlenerek uygulama alanı ve önceki bilgiler hakkında bir anlayış geliştirilmelidir. İkinci adımda, üzerinde araştırmanın yapılacağı değişkenlere ve örnekleme odaklanma imkânı veren bir veri seti seçilmelidir. Üçüncü adımda, veri temizlemesi ve ön işleme yapılır. Bunun için kayıp veriler için uygun stratejiler belirlenir, gürültüyü önlemek amacıyla çalışmalar yapılır ve zaman serilerinde bilinen değişimler hesaba katılır. Dördüncü adımda, çalışmanın amacına bağlı olarak verileri temsil etmek amacıyla kullanışlı değişkenleri temsil eden özellikler belirlenerek veri azaltılır. Beşinci adımda, birinci adıma uygun olacak veri madenciliği yöntemi belirlenir. Altıncı adımda, keşfedici analiz, model ve hipotezler belirlenir. Bu amaçla veri örüntülerini araştırmak için kullanılacak veri madenciliği metot ve algoritmaları seçilir. Bu süreç, hangi modellerin ve parametrelerin uygun olacağına karar vermeyi içerir (örneğin, kategorik verilerin modelleri, sürekli verilerin modellerinden farklıdır). Yedinci adım, veri madenciliğidir; sınıflandırma kuralları veya ağaçları, regresyon ve kümeleme dâhil olmak üzere belirli bir temsil biçimindeki kullanılacak ilgili modeller araştırılır. Kullanıcı, önceki adımları doğru şekilde uygulayarak veri madenciliği adımına önemli ölçüde yardımcı olabilir. Sekizinci adımda, ortaya çıkan sonuçlar yorumlanır. Bu adımda aynı zamanda ortaya çıkan model veya modellerin görselleştirilmesi yapılır. Daha fazla sonuç elde etmek için 1 ve 7 arasındaki adımlara dönüş yapılabilir. Dokuzuncu adımda, keşfedilen bilgi ya doğrudan kullanılır ya ilgili birimlere rapor edilir ya da sonraki çalışmalar

için başka bir sisteme dâhil edilir. Bu süreçte daha öncesinde ortaya çıkan sonuçlarla uyumsuz olan durumlar kontrol edilir ve çözümler sunulur. Han ve Kamber (2012, s.7) tarafından yedi adımda oluşturulan veri madenciliği süreci Şekil 3’te yer almaktadır.



Şekil 3. Veri madenciliği süreci.

Han ve Kamber’e (2012, s.7-8) göre bilgi keşif işlemi aşağıdaki adımların tekrarlamalı bir sırası olarak gösterilmektedir:

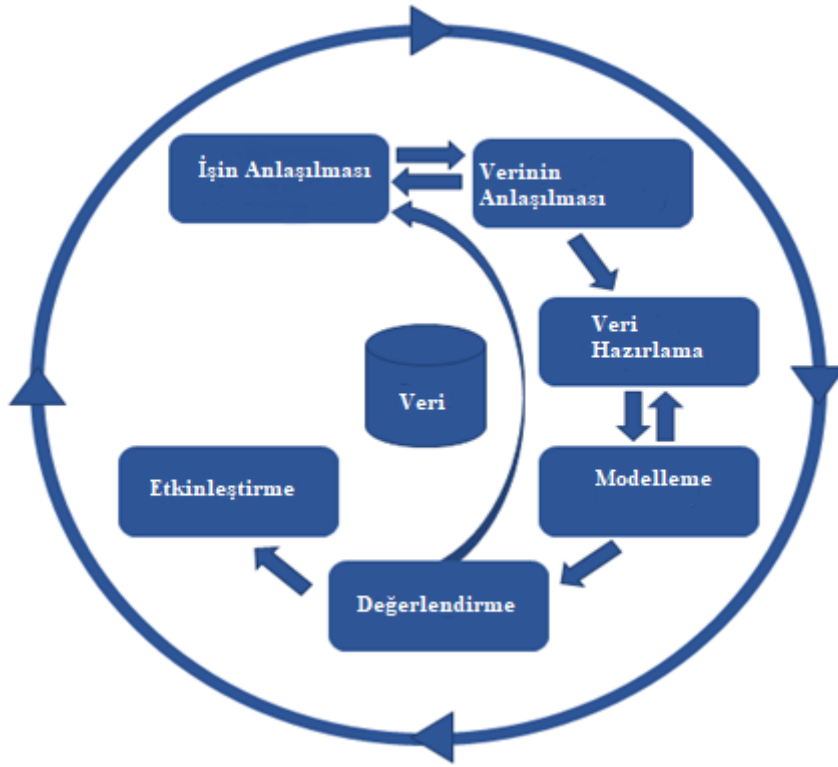
- ✓ Veri temizleme: Verilen örneklerdeki gürültülü ve tutarsız verilerin silinmesi.
- ✓ Veri entegrasyonu: Birden fazla veri kaynağının birleştirilmesi.
- ✓ Veri seçimi: Veri tabanından analizle ilgili olan verilerin seçilmesi.
- ✓ Veri dönüşümü: Verilerin özet veya toplama işlemleri gerçekleştirilerek madencilğe uygun formlara dönüştürülmesi ve birleştirilmesi. Bazen veri dönüşümü ve birleştirme, özellikle veri depolama durumunda, veri seçim işleminden önce gerçekleştirilir.

Bütünlüğünden ödün vermeden orijinal verilerin daha küçük bir gösterimini elde etmek için veri azaltma da yapılabilir.

- ✓ Veri madenciliği: Veri kalıplarını çıkarmak için yöntemlerin uygulanması.
- ✓ Örüntü değerlendirmesi: Örüntülerden elde edilen bilgilerin, gerçek ölçütlere dayanan bilgilerle karşılaştırılması.
- ✓ Bilgi sunumu: Elde edilen bilgiyi kullanıcılara sunmak için görselleştirme ve bilgi temsil tekniklerinin kullanılması.

Adım 1 ila 4, verilerin madencilik için hazırlandığı farklı veri ön işleme biçimleridir. Veri madenciliği adımı, kullanıcı veya bir bilgi tabanı ile etkileşime girebilir. İlgili modeller kullanıcıya sunulabilir veya bilgi tabanında yeni bilgiler olarak saklanabilir.

SPSS (2000, s.13), veri madenciliği sürecindeki problemlerin çözümünde izlenmesi gereken adımları, CRISP-DM süreci ile aşamalandırmıştır. Bu süreç modeli, bir veri madenciliği çalışmasının işleyiş döngüsüne genel bir bakış sunar ve çalışmada yer alan aşamaları, görevleri ve bu görevler arasındaki ilişkileri içerir. Süreç, altı aşamadan oluşmaktadır ve aşamalar arasındaki geçiş sırası esnekler. Aşamalar arasında bir önceki aşamada yer alan görevin sonucuna bağlı olarak ileri veya geri hareket etmek mümkündür. Şekil 4, CRISP-DM sürecinin aşamalarını göstermektedir.



Şekil 4. CRISP-DM süreci.

Şekil 4'te yer alan dış daire veri madenciliğinin döngüsel niteliğini sembolize etmektedir. Aşamalar arasında yer alan oklar ise en önemli ve en sık karşılaşılan bağımlı durumları göstermektedir. CRISP-DM sürecinde yer alan aşamaları aşağıdaki biçimde açıklanmaktadır (SPSS, 2000, s.13-14).

İşin Anlaşılması: Bu aşamada çalışmanın amaçlarının ve gereksinimleri organizasyonun bakış açısı doğrultusunda belirlenerek elde edilen bu bilgi doğrultusunda, veri madenciliğinin problem tanımına ve hedefine ulaşmak için bir ön plan hazırlanır.

Verinin Anlaşılması: İlk başta verinin elde edilmesi ile başlar ardından verilere aşina olma, veri kalitesi problemlerini belirleme, verilere yönelik ilk bakış açısını belirleme veya gizli bilgilerin ortaya çıkarılması için hipotezler oluşturmak amacıyla ilgili altkümeleri saptama faaliyetlerle devam eder.

Veri Hazırlama: Bu aşama ham veriden başlayarak nihai veri seti elde edilene kadar geçen süreç içindeki bütün işlemleri kapsar. Veri hazırlama işlemlerinin önceden belirlenmiş bir

sıra olmadan birçok kez yapılması muhtemeldir. Bu işlemler, tablo, kayıt, öznitelik (değişken) seçimi ve modelleme araçları için verinin dönüştürülmesi ve temizlenmesidir.

Modelleme: Çeşitli modelleme yöntemleri seçilir ve uygulanır. Genel olarak aynı veri madenciliği problemine yönelik birden fazla yöntem vardır. Bazı yöntemlerin uygulanabilmesi için veri formunun özel gereksinimleri olabilir. Bundan dolayı, çoğu zaman modelleme aşamasından veri hazırlama aşamasına geri dönmek gerekebilir.

Değerlendirme: Bir önceki aşamada elde edilen model veya modeller veri analizi açısından yüksek kaliteye sahip olduğu düşünülerek oluşturulmuştur. Değerlendirme aşamasında, oluşturulan bu model veya modeller ilgili birimlerle paylaşılmadan önce uygulanan işlemlerin daha ayrıntılı bir biçimde değerlendirilmesi yapılır. Bu değerlendirme, modelin çalışmanın hedeflerine uygun bir biçimde ulaştığından emin olmak için önemlidir. Bu aşamada anahtar amaç, yeterince düşünülmemiş bazı önemli noktaların olup olmadığını belirlemektir. Bu aşama sonunda, veri madenciliği sonuçlarının kullanımı konusunda bir karara varılır.

Veri madenciliği süreci, bir çözüme ulaşılsa bile sona ermez. Sürecin aşamalarından ve ulaşılan çözümden elde edilen bilgiler yeni araştırma sorularını ortaya çıkarır ve daha sonra ele alınacak veri madenciliği süreçlerinde önceki bu deneyimlerden faydalanılır.

2.3. Veri Madenciliği Modelleri

Veri madenciliği modelleri tahminleyici ve tanımlayıcı olmak üzere iki gruba ayrılır (Akpınar, 2000; Bounsaythip & Rinta-Runsala, 2001, s.10; Dunham, 2003, s.4;). Bir veri madenciliği çalışmasının amacı genel olarak tanımlayıcı veya tahminleyici bir model oluşturmaktır (Jain & Srivastava, 2013). Tahminleyici model, geçmiş veya farklı verilerden elde edilen sonuçları kullanarak verilerin değerleri hakkında bir tahmin yapar (Dunham, 2003, s.5). Tahminleyici modelin amacı bir değişkenin genellikle gelecekteki bilinmeyen bir değerinin tahmin edilmesini sağlamaktır (Jain & Srivastava, 2013). Bunun için sonuçları bilinen veriler kullanılarak geliştirilen modellerden yararlanılarak sonuçları bilinmeyen veri

kümeleri için tahminde bulunur (Akpınar, 2000). Veri madenciliği uygulamalarında bağımlı değişken kategorik ise veri madenciliği uygulaması sınıflandırma olarak adlandırılırken bağımlı değişkenin sürekli olması durumunda regresyon olarak adlandırılır (Jain & Srivastava, 2013).

Tanımlayıcı model kısa olarak bir veri setinin temel özelliklerini sunar. Bu modeller temel olarak veri göstergelerinin bir özeti olup, veri setinin önemli yönlerini incelemeyi sağlar. Tanımlayıcı modeller veri setindeki örüntüleri keşfeder ancak bu örüntülerin yorumlanmasını araştırmacıya bırakır (Jain & Srivastava, 2013). Bu modellerde temel amaç karar vermeye rehberlik etmede kullanılabilecek örüntüleri ortaya çıkarmaktır (Akpınar, 2000). Tanımlayıcı model, tahminleyici modelden farklı olarak ele alınan verideki özellikleri incelemeye yarar, yeni özellikler hakkında yordama yapmaz sadece verilerdeki örüntüleri ve ilişkileri tanımlar (Dunham, 2003, s.5).

Veri madenciliğinde yaygın olarak, Kümeleme ile Birliktelik Kuralları (Association Rules) ve Ardışık Zamanlı Örüntüler (Sequential Patterns) tanımlayıcı modeller, Sınıflama ve Regresyon ise tahminleyici modeller olarak sınıflandırılmaktadır (Bounsathip & Rintarunsala, 2001, s.10).

2.3.1. Kümeleme

Kümeleme analizi hakkında ilk fikirler 1932 yılında Harold Edson Driver ve Alfred Louis Kreber tarafından ortaya atılmış, 1938 yılında Joseph Zubin'in çalışması ile gelişmeye başlamıştır. Kümeleme analizi yönteminin anlatıldığı ilk eser Robert Tryon tarafından 1939 yılında yayımlanan Kümeleme Analizi'dir (Cluster Analysis). Sokal ve Sneath'ın (1963) çalışmalarından sonra ise kümeleme analizi daha yaygın bir kullanım alanına sahip olmuştur (Akpınar, 2017, s.361). Günümüzde biyoloji, zooloji, kimya, yer bilimleri, tıp, mühendislik, işletme ve ekonomi ve sosyal bilimleri alanlarında çok sayıda kümeleme analizi ile yapılan çalışmalara rastlamak mümkündür (Gore, 2000, s.298).

Kümeleme analizinde genel olarak aşağıdaki sorulara cevap aranır (Köse, 2018, s.126):

- ✓ Belirli $(N_1, N_2, N_3, \dots, N_n)$ niteliklerine göre birbiri ile benzer olan varlıklar nelerdir?
- ✓ Belirli $(N_1, N_2, N_3, \dots, N_n)$ niteliklerine göre elimizdeki varlıkları önceden belirlenen K tane farklı kümeye ayıracak olursak, hangi varlıklar hangi kümede yer alır?
- ✓ Belirli $(N_1, N_2, N_3, \dots, N_n)$ niteliklerine göre elimizdeki varlıkları kümeleyecek olursak birbirinden farklı kaç küme meydana gelir.
- ✓ Belirli (X, Y, Z) niteliklerine göre elimizdeki varlıkları kümeleyecek olursak, kümelerin koordinat düzleminin şekli nasıl olur?
- ✓ Belirli (X, Y, Z) niteliklerine göre elimizdeki varlıkları kümeleyecek olursak, kümelerin koordinat düzlemindeki dağılımı nasıl olur?

Tanımlayıcı bir veri madenciliği modeli olan kümeleme analizi, ön bilgilerimizin kısıtlı olduğu veya muhtemel benzerliklerin araştırmacı tarafından fark edilemediği durumlarda, belirli özellikler bakımından hangi varlıkların aynı kümelerde olacağını tespit etmede kullanılır (Köse, 2018, s.125). Kümeleme analizinde amaç çok boyutlu veriler içinden doğal kümeleri bulmak ve belirli bir özellik bakımından birbirine benzeyen varlıkları aynı kümeye benzemeyenleri ise farklı bir kümeye yerleştirmektir (Ersöz, 2019, s.77). Kümeleme analizi de bir sınıflama yaklaşımıdır, ancak önceden belirlenmiş sınıflar yani bağımlı değişkene ait y değerleri bulunmadığından dolayı sınıf oluşturma işlemi ön planda olup varlıkların araştırma kapsamında ele alınan bağımsız değişken değerlerine göre sınıflandırılması yapılmaktadır (Akpınar, 2017, s.74). Kümeleme analizi veri setini benzer gruplara bölme işlemi olup verileri kümelerle modeller ve verilerin daha az küme seti ile temsil edilmesini sağlar. Bu durum verilerin sadeleştirilmesini imkân verir ancak ince ayrıntıların kaybolmasına neden olur. Verilerin modellenmesi kümelemeyi matematik, istatistik, veri analizi ve makine öğrenmesi alanlarına dayanan önemli bir perspektife yerleştirir. Makine öğrenmesi bakımından kümeler gizli örüntülere karşılık gelir ve ortaya çıkan sistem bir veri kavramını temsil eder (Berkhin, 2006, s.1). Kümeleme analizi makine öğrenmesi alanında denetimsiz öğrenme yöntemlerinden olup bağımlı ve bağımsız değişkenler arasındaki

birliktelik ve örüntüler tanımlanır (Ersöz 2019, s,77). Kümeleme analizinde bağımlı ile bağımsız değişkenler arasında bir ilişki kurulması söz konusu değildir, varlıkların küme içine homojen kümeler arasında ise heterojen bir yapıya sahip olması beklenir (Akpınar, 2017, s.360).

Kümeleme analizi matematiksel olarak aşağıdaki biçimde tanımlanabilir

$D = \{X_1, X_2, X_3, \dots, X_n\}$ ve $n = 1, 2, 3, \dots, m$ olacak biçimde bir veri tabanı olsun ve her X_n bir veriyi temsil etsin. $X = \{x_1, x_2, x_3, \dots, x_n\}$ ve $i = 1, 2, 3, \dots, m$ olacak biçimde her x_n ad, soyad, yaş, vb. bir değişken olsun. Kümeleme analizinin amacı; D veri tabanını, j adet kümeye bölmek ve $K \subseteq D$ koşulunu elde etmektir (Silahtaroglu, 2016, s.155).

Kümeleme analizinde kullanılan algoritmaları “Klasik Yaklaşımlar” ve “Çağdaş Yaklaşımlar” olmak üzere iki ana başlık altında toplayabiliriz. Kümeleme analizinde başlıca kullanılan algoritmalar Tablo 1’de yer almaktadır (Akpınar, 2017, s.374).

Tablo 1

Kümeleme Analizinde Başlıca Kullanılan Algoritmalar

Klasik Yaklaşımlar	Çağdaş Yaklaşımlar
❖ Tekli bağlantı En yakın komşu yöntemi	❖ Hiyerarşik temelli küme analizi BIRCH, CURE, ROCK, Chameleon
❖ Tam bağlantı En uzak komşu yöntemi	❖ Bölümleyici küme analizi k-ortalamlar, k-medoidler (PAM, CLARA, CLARANS), k-prototipler, k-modları, k-medyan, k-ortalamlar+
❖ Aritmetik ortalamlı bağlantı Gruplar arası bağlantı	❖ Yoğunluk temelli kümeleme analizi DBSCAN, OPTICS, DENCLUE
❖ Merkezi bağlantı	❖ Izgara temelli algoritmalar STING, CLIQUE
❖ Medyan bağlantı	❖ Alt uzay arama algoritmaları PROCLUS, SUBCLUE
❖ Ward bağlantısı	

2.3.2. Birliktelik Kuralları

Temelleri pazar sepeti analizine dayanan, bir dizi öge içinden iki veya daha fazlasının eşzamanlı veya farklı zamanlarda meydana gelmesinin analiz edildiği ve birliktelik kurallarının ortaya konulduğu, veri madenciliği yöntemlerine alanyazında örüntü madenciliği de denir (Akpınar, 2017, s.201). Sepet analizi olarak da bilinen birliktelik kuralları, belirli bir sonucu bir dizi koşulla ilişkilendirmeye çalışarak olayların birlikte

meydana gelme kurallarına yönelik olasılıklar üretir (Ersöz, 2019, s.85, Özkan, 2016, s.36-217). Birliktelik kurallarında genel olarak aşağıdaki sorulara cevap aranır (Köse, 2018, s.127).

- ✓ Tekrar eden bir olayda yer alan öğelerin her biri için, diğer öğelerle aynı olayda belli bir doğruluk çerçevesinde bir arada olma durumlarını temsil eden kurallar var mıdır?
- ✓ Bir olayda belirli bir öğenin yer aldığı biliniyorsa, bu öğe ile birlikte başka hangi öğelerin yer alabileceğini belirli bir doğruluk çerçevesinde tahmin edebilir miyiz?

Tanımlayıcı bir model olan birliktelik kuralları, veriler arasındaki ilişkileri açığa çıkarır. Pazarlama, reklam, envanter kontrolü, perakendecilik ve telekomünikasyon gibi birçok alanda kullanılır. Birliktelik kurallarında açığa çıkarılan ilişkiler herhangi bir nedensellik veya korelasyonu temsil etmez sadece öğelerin ortak kullanımını yani birlikte meydana gelişini ifade eder (Dunham, 2003, s. 8, s.164). Yani aslında birliktelik kuralları sık sık beraber gözlenen olaylar arasındaki probabilistik korelasyondur (Silahtaroglu, 2016, s.53).

Birliktelik kuralları matematiksel formu; $I = I_1, I_2, I_3, \dots, I_m$; ikili değişkenlerin (özelliklerin) bir kümesi ve T işlemlere ait bir veri tabanı olsun. Her bir t işlemi, I_k değişkeninin gerçekleşmesi durumunda $t[k]=1$ aksi halde $t[k]=0$ olacak biçimde ikili vektörlerle temsil edilsin. Her bir işlem için veri tabanında bir değişken grubu vardır. X , I 'daki bazı değişkenlerin kümesi olsun. Eğer X 'deki tüm I_k 'lar için $t[k]=1$ ise, t işleminin X 'i karşıladığını söyleyebiliriz. Bir birliktelik kuralı olarak, $X \Rightarrow I_j$ formunun dolaylı anlatımı kast edilmektedir ve burada X , I 'daki bazı değişkenlerinin bir kümesidir ve I_j , I 'da yer alıp X 'te bulunmayan tek bir değişkendir (Agrawal, Imielinski & Swami, 1993). Yani temel olarak bir birliktelik kuralı ile şu ifade edilmektedir: $X \subset I$, $I_j \subset I$, $I_j \notin X$ ve $X \Rightarrow I_j$ olmak üzere; $X \Rightarrow I_j$ kuralının T için var olduğunun söylenebilmesi için belirli bir güven seviyesinin mevcut olması gerekir. Başka bir ifadeyle T içindeki tüm X 'lerin ne kadarının I_j 'yi sağladığı %c değeriyle ifade edilir. Bu durumda $0 \leq c \leq 1$ güven seviyesinde birliktelik kuralı $X \Rightarrow I_j | c$ biçiminde gösterilir (Silahtaroglu, 2016, s.140). Birliktelik kurallarında kullanılan bazı algoritmalar Öncül (apriori), sıklık örüntü gelişim algoritması (frequent patterns growth

algorithm) ve eşdeğerlik sınıf dönüşümü (equivalence class transformation; Eclat) algoritmalarıdır (Akpınar, 2017, s.209).

2.4. Sınıflandırma ve Regresyon Ağaçları

Sınıflandırma ve Regresyon Ağaçları, literatürde SVRA (Classification and Regression Trees; CART) olarak kısaltılır. SVRA algoritması, Morgan ve Sonquist (1963) tarafından ilk karar ağacı yaklaşımı olarak sunulan Otomatik Etkileşim Çıkarma (Automatic Interaction Detection; AID) algoritmasından bu yana geliştirilen ağaç tabanlı sınıflandırma ve tahmin yöntemlerinin tümünün bir birleşimini temsil eder (Da Rosa, Veiga & Medeiros, 2003, s.2). SVRA, ilk olarak 1984 yılında Breiman, Friedman, Olshen ve Stone tarafından tanımlanmış bir karar ağacı algoritmasıdır. SVRA, kategorik ve sürekli değişkenlerden faydalanılarak sınıflandırmaya yönelik analizler yapmak amacıyla geliştirilmiş parametrik olmayan istatistiksel bir yöntemdir. Bağımlı değişkenin kategorik olması durumunda sınıflandırma ağacı oluştururken, bağımlı değişkenin sürekli olması durumunda regresyon ağacı oluşturur (Yohannes & Webb, 1999, s.4). SVRA, regresyon ağacı modellerini klasik regresyon yöntemlerinin parametrik olmayan önemli bir alternatifi haline dönüştürmesinden dolayı birçok araştırmacıyı karar ağaç tekniklerini bilinen istatistiksel yöntemlerle birleştiren hibrit modelleme stratejileri oluşturmaya yöneltmiştir. Bu bağlamda Segal'in (1992) boylamsal verilere yönelik analizlerini, Ahn'ın (1996) risk analizini ve Cooper'ın (1998) zaman serileri analizi örnek olarak verilebilir (Da Rosa, vd., 2003, s.2).

SVRA ve bu tür sınıflandırma yöntemlerinin temel amacı belirli özelliklere ait mevcut olan veri setlerinden faydalanılarak gelecekteki gözlem veya gözlemlerin sınıfının tahmin edilmesini sağlamaktır (Yohannes & Webb, 1999, s.4). Literatür incelendiğinde karar ağaçlarının, borsa, tıp, arge, turizm, banka ve risk yönetimi alanlarında kullanıldığı söylenebilir (Sayıcı, 2013, s. 43).

SVRA'nin lojistik regresyon ve geleneksel lineer yöntemlere göre avantajı, parametrik ve doğrusal olmayan bir yöntem olarak parametrik istatistiklerin varsayımlarına tabi olmamasıdır. Bu nedenle, SVRA, bağımsız değişkenler arasındaki altta yatan ilişkilerden

etkilenmeden, belirli bir sonucu tahmin etmek için en önemli olanları çok sayıdaki bağımsız değişken arasından seçebilir (Chudzinska & Baralkiewicz, 2011). Öte yandan geleneksel istatistiksel analizler için kayıp verilerin bulunması bir sorun teşkil ederken SVRA yöntemi için bir sorun teşkil etmez. Çünkü SVRA birincil değişkenin bölünmesini taklit eden vekil değişken kullanarak bu sorunun üstesinden gelebilecek özelliğe sahiptir. Ağacın oluşumu sırasında bazı durumlar için kayıp veriler varsa SVRA bu durumları göz ardı etmez. Bunun yerine bu durumları birincil bölme değişkenine benzeyen vekil değişkene göre sınıflandırır. Vekil değişken birincil bölünmeden farklı bir kesme noktasına sahip olabilir, ancak vekil bölünmenin sol ve sağ düğümlere gönderdiği kişi sayısı birincil bölünmeye çok yakın olmalıdır (Yohannes & Webb, 1999, s.18). Aynı zamanda SVRA, değişkenler arasında doğrusal olmayan ilişkilerin ve çoklu bağlantının mevcut olduğu durumlarla başa çıkabilen esnek ve sapmasız bir analitik yöntem olup, karmaşık verilerin analizi için kullanışlıdır (De'ath & Fabricius, 2000). SVRA yöntemiyle kurulan modellerde bağımsız değişkenlerin kategorik, sıralamalı veya sürekli olması konusunda herhangi bir kısıtlama olmayıp, bağımsız değişkenlerin olası tüm kombinasyonlarını modele kattığı için mümkün olabilecek en doğru sınıflamaya ulaşma şansı yüksektir. Bağımsız değişkenlerin tüm kombinasyonlarını modele alması nedeniyle modele esneklik ve daha geniş bir bakış açısı kazandırırken aynı zamanda bağımlı değişkeni önemli düzeyde etkileyen bağımsız değişkenlerin belirlenmesine de imkân verir (Orekeci-Temel, 2004, s.8). Ayrıca karar ağaçları sahip olduğu görsel özelliğinden dolayı kolay anlaşılabilir ve etkin yorum yapabilme imkânı verir (Sayıcı, 2013, s. 43).

SVRA algoritması, ağaç yapısını önce örneklemin kendisinden başlayarak, örneklem ve örneklemin alt kümelerini ikili biçimde tekrarlı bölünmeleri ile oluşturur (Breimann, Friedman, Stone & Olshen, 1984, s.20). Bu tekrarlı bölünmelerden homojen alt gruplar oluşur ve ağaç bu şekilde büyümeye devam eder (Köktürk, 2012, s.51). Genel olarak karar ağaçları, düğüm, dal ve yaprak olmak üzere üç temel bileşenden meydana gelir ve düğümler bağımsız değişkenleri temsil eder (Atasever, 2011, s. 21). SVRA kök düğümünden itibaren her bir sonraki seviyesinde daha homojen düğümlerin elde edildiği ikili karar ağacı üreten

bir tekniktir (Akpınar, 2017, s.267, Güler, 2017, s.30). Dolayısıyla her alt sınıf için bir alt sınıf oluşturulması nedeniyle her defasında yalnızca iki alt sınıf meydana gelir (Dunham, 2003, s. 102; Yohannes & Webb, 1999, s.4-6). Genel olarak karar ağacı algoritmalarında öncelikle kök düğümünden başlanılarak veri iki alt düğüme ve daha sonra her bir alt düğüm tekrar iki alt düğümlerine bölünerek durdurma kuralı olmadan ağaç en yüksek büyüklüğüne ulaşmaya kadar devam eder. Daha sonra en uygun ağaç yapısına ulaşmak için kök düğüme doğru budama işlemi yapılır. Budama işlemi yapılırken ağacın toplam performansına katkısı en az olan bölünmeler ortadan kaldırılır (Kuzey, 2012, s.82). Karar ağaçlarının matematiksel olarak açılımı Dunham (2003, s.93) tarafından aşağıdaki biçimde tanımlanmıştır.

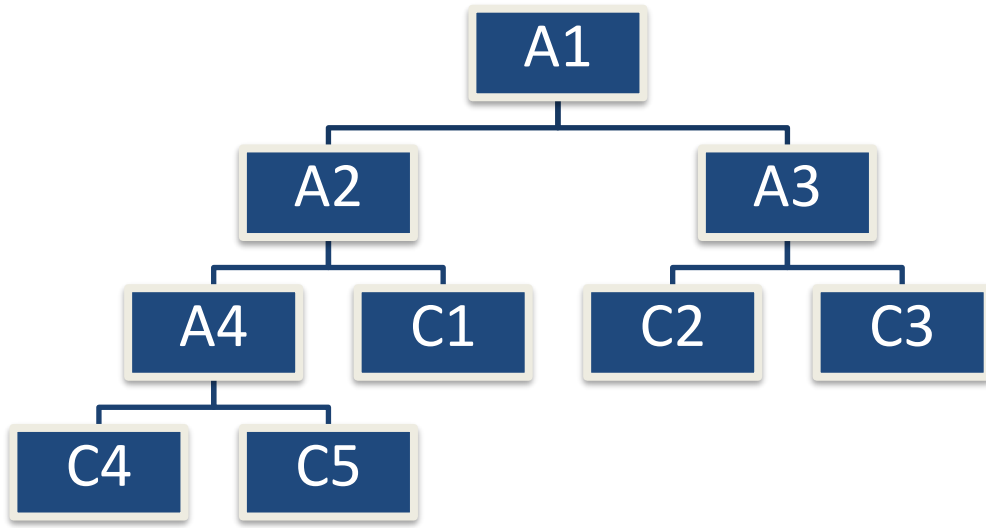
$t_i=(t_{i1},\dots,t_{ih})$ eşitliğinin sağlandığı ve $(A_1, A_2, \dots,A_h)$ özelliklerini içeren veritabanı şemasını içeren bir $D= (t_1,\dots,t_n)$ veri tabanı olsun. Ayrıca $C= (C_1, \dots, C_m)$ sınıflarını içeren bir set verilsin. Karar ağacı (KA) ya da sınıflama ağacı, aşağıdaki özelliklere sahip D ile ilişkili bir ağaçtır:

- Her bir düğüm A_i değişkeni ile etiketlenmiştir.
- Her yaprak bir C_j sınıfı ile etiketlenmiştir.

Karar ağaçlarını kullanarak sınıflandırma problemlerinin çözümü iki aşamalı bir süreçtir:

1. Karar ağacı atama: Eğitim verisini kullanarak bir karar ağacı oluşturun.
2. Her bir $t_i \in D$ için, DT 'yi sınıfı belirlemek üzere kullanın.

Yukarıdaki biçimde tanımlanan bir karar ağacının şekilsel olarak gösterimi Şekil 5'te yer almaktadır (Silahtaroglu, 2016, s.69).



Şekil 5. Örnek karar ağacı.

Şekil 5’te yer alan örnek karar ağacını incelendiğinde, A1 kök düğüm olup kendisinden sonra A2 ve A3 düğümlerinin geldiği iki dala bölünmüştür. C1, C2 ve C3 ise bölünmelerin sonlandığı yapraklar başka bir ifadeyle sınıflardır. A4 ise A2 düğümünün alt düğümü olup kendi içinde iki dala bölünerek C4 ve C5 sınıflarını oluşturmuştur. A1 düğümü veri setinde yer alan tüm bireyleri içermektedir ve $A1=A2UA3$ ’tür. Diğer düğüm ve sınıflar ise, $A2=A4UC1$, $A3=C2UC3$ ve $A4=C4UC5$ biçimindedir.

SVRA algoritmasında iki kategoriye sahip kategorik bir bağımsız değişken için, her kategorinin bir sınıf olarak tanımlandığı sadece bir bölünme mevcuttur. Eğer bağımsız değişken 2’den fazla olacak biçimde k tane kategoriye sahipse bu durumda boş küme hariç $2^{k-1}-1$ tane farklı bölünme mümkündür. Sürekli bağımsız değişkenin olması durumunda ise, bölme işlemi belirlenen bazı değerlerden daha küçük ve daha büyük değerlere göre oluşturulur. Bu nedenle yalnızca sürekli değişkenlerin sıralaması bir bölünmeyi belirler ve belirli bir u değeri için u-1 tane olası bölünmeden bahsedilebilir (De'ath, & Fabricius, 2000). SVRA, ID3 algoritmasına benzer olarak en iyi bölünme niteliğini ve kriterini seçmek için entropi kullanır ve hesaplanan entropi değerine göre en iyi bölünme noktası olarak belirlenen yerden bölünme işlemi yapılır (Dunham, 2003, s. 102). Ancak SVRA algoritmasında en iyi bölünme kriterinin belirlenmesinde kullanılan formül C4.5 ve ID3 karar ağacında kullanılan formülden farklılık göstermektedir (Silahtaroglu, 2016, s. 82).

2.4.1. Bölünme Kriterleri

SVRA algoritmasında eğer bağımlı değişken kategorik ise yani sınıflama ağacı modeli kullanılacaksa bölünme kriteri olarak iki farklı yaklaşıma göre elde edilen Twoing ve Gini algoritmaları kullanılır (Breiman, vd., 1984, s.103).

2.4.1.1. Gini Algoritması

Gini algoritmasının temel çalışma prensibi en geniş sınıfı diğer sınıflardan ayırmaktır (Orekeci-Temel, 2004, s.15). Gini ölçütü, frekans dağılımında yer alan değerler arasındaki eşitsizliğin ölçüsü olup, değeri 0 ile 1 arasında değişir ve mümkün olabildiği kadar küçük olması istenilen bir durumdur (Özkan, 2016, s.111). Matematiksek olarak Gini ölçütü Breiman vd. (1984, s.103) tarafından aşağıdaki biçimde tanımlanmıştır.

$$i(t) = \sum_{i \neq j} p(i | t) p(j | t)$$

$$i(t) = \sum_{i \neq j} p(j | t)(1 - p(j | t))$$

$$i(t) = 1 - \sum_{i \neq j} (p(j | t))^2$$

Eğer bağımlı değişkenin kategori sayısı yani sınıf sayısı 2 ise,

$$i(t) = 2p(1 | t)p(2 | t)'dir$$

Gini algoritması, Twoing algoritmasına benzer olarak bağımsız değişkenin değerlerinin sol ve sağda olmak üzere ikiye bölünmesi temeline dayanmaktadır (Dondurmacı, 2011, s.75). Özkan (2016, s.112) sağ ve sol olarak ikiye bölünmeler sonucunda her bir bölünme için ayrı ayrı Gini ölçütü hesaplanmasına ve elde edilen sonuçlar karşılaştırılmasına yönelik olarak gerçekleştirilen işlem adımlarını aşağıdaki gibi sıralamıştır:

1. Öncelikle bağımsız değişken sağ ve sol olarak ikiye bölünerek bu bölünmelere karşılık gelen sınıf değerleri bir arada gruplandırılır.

2. Her bir bağımsız değişken için $Gini_{sol}$ ve $Gini_{sağ}$ değerleri hesaplanır. Bunun için k =sınıf sayısı, T =bir düğümde yer alan örneklerin sayısı, $|T_{sol}|$ =solda kalan örneklerin sayısı, $|T_{sağ}|$ =sağda kalan örneklerin sayısı, L_i =sol tarafta i . kategorideki örnek sayısı, R_i =sağ tarafta i . kategorideki örnek sayısı, olmak üzere;

$$Gini_{sol} = 1 - \sum_{i=1}^k \left(\frac{L_i}{|T_{sol}|} \right)^2 \text{ 'dir}$$

$$Gini_{sağ} = 1 - \sum_{i=1}^k \left(\frac{R_i}{|T_{sağ}|} \right)^2 \text{ 'dir}$$

3. Her bir j bağımsız değişkeni için, eğitim kümesindeki satır sayısı n olmak üzere; $Gini_j$ değerleri hesaplanır:

$$Gini_j = \frac{1}{n} (|T_{sol}| Gini_{sol} + |T_{sağ}| Gini_{sağ})$$

4. Her bir j bağımsız değişkeni için hesaplanan $Gini_j$ değerleri arasında en küçük olanı seçilerek bölünmeler bu bağımsız değişken üzerinden gerçekleştirilir.
5. Bir sonraki bölünme için ilk basamağa dönülerek işleme devam edilir.

2.4.1.2. Twoing Algoritması

Twoing algoritmasının temel çalışma prensibi düğümde yer alan örneklerin %50'sini içeren ve birbirine benzemeyen sınıfları diğer sınıflardan ayırmaktır (Orekeci-Temel, 2004, s.17). Twoing ölçütü, her adım için üzerinde çalışılan örneklemin iki parçaya ayrılması (C_1 ve C_2 olacak biçimde iki aday üst sınıf) temeline dayanır (Duda, Hart, & Stork, 2012, s.401). Twoing ölçütüne göre elde edilen değerlerden büyük olanı tercih edilir (Larose, 2005, s.110). Matematiksek olarak Twoing ölçütü Breiman vd. (1984, s.104) tarafından şöyle tanımlanmıştır: $C = \{1, 2, \dots, J\}$ sınıfları belirtmek üzere, $C_1 = \{j_1, \dots, j_n\}$, $C_2 = C - C_1$ olacak biçimde her düğümdeki sınıflar iki üst sınıfa ayrılınsın ve birinci sınıfta yer alan örnekler C_1 'e, ikinci sınıfta yer alan örnekler C_2 'ye yerleştirilir.

Daha sonra düğümün belirli bir s bölünmesi için, $\Delta_i(s, t)$ değerini bağımlı değişkenin 2 kategorili yani sınıf sayısının 2 olması durumu gibi hesaplanır.

Aslında $\Delta_i(s, t)$ değeri, C_1 'in seçimine bağlıdır, dolayısıyla $\Delta_i(s, t, C_1)$ biçiminde gösterilebilir. Daha sonra $\Delta_i(s, t, C_1)$ değerini en büyük yapan $s^*(C_1)$ bölünmesi tespit edilir ve son olarak $\Delta_i(s^*(C_1), t, C_1)$ değerini en büyük yapan üst sınıf bulunur.

Sağ ve sol olarak ikiye bölünmeler sonucunda her bir bölünme için ayrı ayrı Twoing ölçütü hesaplanmasına ve elde edilen sonuçlar karşılaştırılmasına yönelik olarak gerçekleştirilen işlem adımlarını aşağıdaki gibi sıralanabilir:

1. Öncelikle bağımsız değişken değerleri gözönüne alınarak örneklem sağ ($t_{sağ}$) ve sol (t_{sol}) olarak ikiye bölünür (Özkan, 2016, s.93).
2. Her bir aday bölünme için P_{sol} , $P(j | t_{sol})$, $P_{sağ}$, ve $P(j | t_{sağ})$ olasılıkları hesaplanır (Larose, 2005, s.110).

$t_{sol}=t$ düğümünün sol alt düğümü ve $t_{sağ}=t$ düğümünün sağ alt düğümü olmak üzere;

$$P_{sol} = \frac{t'_{sol}daki\ kayıt\ sayısı}{Örneklem}, P(j | t_{sol}) = \frac{t'_{sol}daki\ kayıtların\ j\ sınıfı\ sayısı}{Örneklem},$$

$$P_{sağ} = \frac{t'_{sağ}daki\ kayıt\ sayısı}{Örneklem}, P(j | t_{sağ}) = \frac{t'_{sağ}daki\ kayıtların\ j\ sınıfı\ sayısı}{Örneklem},$$

3. $\Phi(s | t)$, t düğümündeki s bölünmesinin uyum ölçüsü değeri hesaplanır (Larose, 2005, s.110).

$$\Phi(s | t) = 2P_{sol}P_{sağ} \sum_{j=1}^{classes} |P(j | t_{sol}) - P(j | t_{sağ})|$$

4. En büyük $\Phi(s | t)$ seçilir ve bu değer ait olduğu aday bölünme dallanmanın yapılacağı satırı gösterir. Dallanma yapıldıktan sonra bu bölünmeye yönelik karar ağacı çizilir. Daha sonra ilk basamağa dönülerek alt düğüm için aynı işlemler uygulanmaya başlanır (Özkan, 2016, s.94).

Breiman vd. (1984, s.109, s.111) tarafından yapılan uygulamalar sonucunda, genel olarak Gini ve Twoing algoritmasına göre yapılan bölünmeler arasında benzerliklerin olduğu ancak

farklılık olan durumlarda genel anlamda Gini bölünmelerinin daha iyi olduğu sonucuna ulaşılmıştır. Ayrıca Twoing algoritmasının sınıf sayısının az olduğu durumlarda, Gini algoritmasının ise orta düzeyde sınıf sayısına sahip değişkenlerin olduğu durumlarda kullanılabileceği ifade edilmiştir.

2.4.1.3. Rastgele Orman Algoritması

Rastgele Orman algoritması, veri setinden bootstrap tekniği ile örneklem seçilerek birbirinden bağımsız bir biçimde karar ağaçlarının oluşturulduğu torbalama (Bagging; Breiman, 1996) ve tüm değişkenler arasından rastgele seçilen az sayıdaki değişkenden karar ağacındaki her düğümünde en iyi bölünmeyi oluşturan değişkenin seçildiği rastgele altuzay (Random Subspace; Ho, 1998) yöntemlerinin birleşimidir (Akman, 2010, s.33). RO ilk olarak 2001 yılında Breiman tarafından tanımlanmış bir topluluk algoritmasıdır (Gislason vd., 2006). RO yönteminin temelinde karar ağaçları vardır ve birbirinden bağımsız bir biçimde oluşturulan karar ağaçlarının bir araya gelmesiyle karar ormanları elde edilerek, karar ağaçları tarafından yapılmış olan tahminlerin birleşimi ile tek bir tahmin yapılması amaçlanır (Atasever, 2011, s.31). RO yönteminde üretilen her bir ağaç diğer ağaçlardan etkilenmeden meydana gelir (Ercire, 2019, s.27). Popüler bir topluluk yöntemi olarak RO, belirli sayıda tam olarak yetiştirilmiş SVRA karar ağacını göz önünde bulundurarak, aşırı uyumu (overfitting) önlemek için veri alt kümelerinden öğrenmek üzere tasarlanmıştır. Hem sınıflandırma hem de regresyon için kullanılırlar (Bhalla, 2014). RO yönteminde karar ağaçları oluşturulurken bölünme kriteri olarak, sınıf değişkenlerinin uygunluğunu ölçen Gini algoritmasını kullanır (Pal, 2005). RO yöntemi sınıflandırmada kullanılacak değişkenlerin belirlenmesi, değişkenler arasındaki etkileşimlerin tespit edilmesi, kayıp verilerin belirlenmesi, uç değerlerin belirlenmesi ve kümeleme analizi yapılması amacıyla kullanılır (Breiman & Cutler, t.y). Veri setinde çok sayıda değişkenin olması durumunda, önem derecelerinden yola çıkarak değişken sayısının azaltılması ve modelin indirgenmesi ile daha etkin tahminlerin yapılmasında kullanışlıdır (Akman, 2010, s.33).

Son yıllarda RO parametrik olmayan regresyon için genetik, bioinformatik, klinik tıp ve psikoloji gibi birçok bilimsel alanda popüler ve yaygın bir yöntem haline gelmiştir. RO yönteminin genel olarak yüksek tahmin doğruluğuna sahip olduğu ve çok boyutlu problemler için uygulanabilir olduğu söylenebilir (Strobl & Zeileis, 2008). RO yöntemi bağımlı değişkenin ikili veya çoklu sınıfa sahip olması durumunda ve bağımsız değişkenlere yönelik herhangi bir sınır olmadan sınıflamalı, sıralamalı ve sürekli değişkenler için kullanılabilir. Modelde yer alan değişkenleri önem sırasına göre sıralayabilir (Korkem, 2013, s. 12). Büyük veri tabanlarında verimli bir biçimde çalışma performansı sergiler. Değişken silinmesine gerek duymadan binlerce giriş değişkenini işleyebilir. Orman üretimi devam ettikçe genelleme hatasının yansız bir kestirimini üretir. Eksik verilerin tahmini için etkili bir yöntemdir ve verilerin büyük bir kısmı eksik olduğu durumlarda bile doğruluğunu koruma eğilimi gösterir. Sınıflara ait örneklem dağılımının dengesiz olduğu veri kümelerinde hatayı dengeleyebilir. Oluşturulan ormanlar ileride kullanılmak üzere saklanabilir. Değişkenler ve sınıflandırma arasındaki ilişki hakkında bilgi veren prototipler hesaplanabilir. Değişken etkileşimlerini tespit etmek için deneysel bir yöntem sunar. RO yöntemi aşırı uyuma karşı dayanıklıdır ve ne kadar ağaç üretilirse üretilsin hızlıdır (Breiman & Cutler t.y). Torbalama ve rastgele değişken seçiminin birbiri ardına uygulandığı RO yönteminde, ormanda yer alan ağaçlar mevcut eğitim setinden elde edilen her yeni eğitim seti için rastgele değişken seçilerek oluşturulduğundan dolayı budamaya ihtiyaç duyulmaz (Breiman, 2001). Bir başka ifadeyle oluşturulan bir ağaç her zaman yeni eğitim seti üzerinde seçilen değişkenlerin kombinasyonları kullanılarak maksimum derinliğe kadar büyüme gerçekleştirir ve budanmaz. Bu özellik RO algoritmasının Quinlan (1993) tarafından önerilen diğer karar ağacı yöntemlerine göre en büyük avantajıdır (Pal, 2005). Çünkü Pal ve Mather'a (2003) göre budama yöntemlerinin seçimi ağaç tabanlı sınıflandırıcıların performansını etkilemektedir. Öte yandan kullanıcı tarafından belirlenen sayıda ağacın üretildiği RO yönteminde, eğer sınıflandırma işlemi yapılıyorsa eğitilmiş ağaçların moduna regresyon işlemi yapılıyorsa eğitilmiş ağaçların ortalamasına göre sonuç elde edilir ve RO yöntemi sürekli değişkenlere göre kategorik değişkenlerde özellikle daha başarılıdır (Bhalla,

2014). RO yöntemi uygulandığı bakımından kolaydır (Akman, 2010, s.34). Diğer sınıflandırıcıların aksine RO yönteminin uygulanabilmesi için yalnızca ağaç ve değişken sayısı parametrelerine ihtiyaç duyulmaktadır. Ayrıca bazı algoritmaları kullanmadan önce, doygunluğu önlemek için verileri normalleştirmeye ihtiyaç duyulduğundan dolayı verilerin ön işlemlerden geçirilmesi gerekirken RO yönteminde böyle bir işlem gerekmez (Yan, 2017, s.31).

RO yöntemi çok sayıda ağaca sahip olduğu için ağaçlar şekilsel olarak görülemez ve belli bir güven aralığı hesaplanmasını engeller. Bu nedenle yapılan sınıflandırma işleminin güven aralığına yönelik bir değerden bahsetmek mümkün değildir. Öte yandan RO algoritmasında yer alan bootstrap tekniğinin yapılan sınıflandırmayı genellemesinden dolayı zaten güven aralığına ihtiyaç duyulmaz (Korkem, 2013, s. 13). RO yönteminin önemli bir dezavantajı sınıflandırmadaki her bir değişkenin etkisini tam olarak açıklamak (Yan, 2017, s.31) ve ormandaki tüm ağaçların yapısının araştırılmasının mümkün olmamasından dolayı bağımsız değişkenler arasındaki ilişkilerin yorumlanması (Prasad, Iverson & Liaw, 2006) kolay değildir. RO yöntemi ile hipotez testi yapmak zordur (Hallet, Fan, Su, Levine & Nunn, 2014). RO yönteminin performansı diğer karmaşık sınıflandırıcılardan özellikle fazla sayıda ağaçların gerektiği gradyan arttırıcı ağaçlardan daha kötüdür. (Yan, 2017, s.32).

RO yönteminin matematiksel olarak açılımı Breiman (2001) tarafından aşağıdaki biçimde tanımlanmıştır (A. Cutler, Cutler & Stevens, 2012, s.163)

RO temel öğrenciler olarak $h_j(X, Q_j)$ ağaçlarını kullanır. Eğitim verisi için, $\mathcal{D} = \{(x_1, y_1), \dots, (x_N, y_N)\}$ eşitliği söz konusudur ve bu eşitlikte, $x_i = (x_{i,1}, \dots, x_{i,p})^T$ 'dir ve p yordayıcıları, y_i yanıtları temsil eder. Uyumlu hale getirilmiş (fitted) ağaç ise $\hat{h}_j(x, \theta_j, \mathcal{D})$ ile gösterilir. Bu Breiman tarafından ileri sürülen orijinal formülleştirme olmasına rağmen, θ_j rastgele bileşeni seçkisizliği enjekte etmek üzere doğrudan değil dolaylı olarak uygulamaya konulur. İlk olarak, her bir ağaç orijinal veriden bağımsız önyükleme (bootstrap) örnekleme yolu ile oluşturulur. Bootstrap örneklemedeki randomizasyon, θ_j 'nin bir parçasını verir. İkinci olarak, bir düğüm ayrıldığında, en iyi ayırım tüm p tane yordayıcı yerine rastgele

seçilen m tane yordayıcı değişkenin alt setinde (her bir düğümde bağımsız olarak) bulunur. Yordayıcıları örneklemek için kullanılan randomizasyon θ_j 'nin kalan bölümlerini verir. Ağaçlar budama olmadan büyür. İlk olarak, Breiman şunu önermiştir: uç noktalar kusursuz olana kadar ya da her bir uç noktadaki daha önceden belirlenen veri noktası sayısından çok az daha fazla oluncaya kadar ağaçları büyütmeyi önermiştir. Son zamanlarda, uç noktaların maksimum sayısını kontrol altına almayı önermektedir.

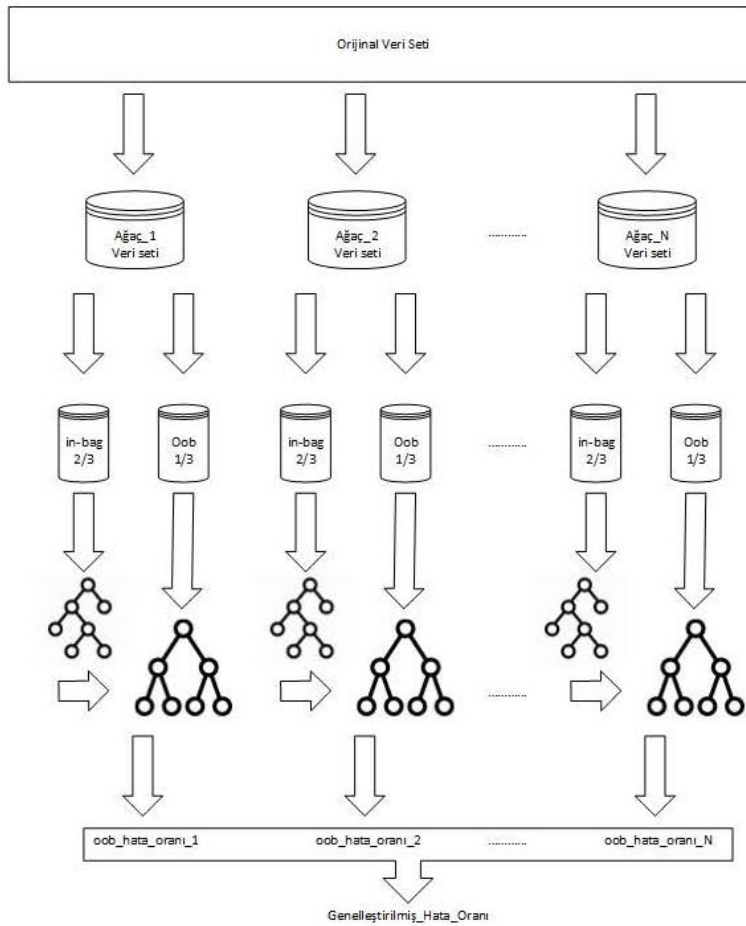
2.4.1.3.1. Rastgele Orman Algoritmasının İşleyişi

Basit yapıda bir sınıflandırma yöntemi olan RO algoritmasında belirli (kullanıcı tarafından) sayıdaki ağaçlardan oluşan orman aşağıdaki biçimde meydana gelir (Breiman, 2004):

- Eğitim verisi, araştırma yapılan veri setinden bootstrap örneklem tekniği ile elde edilir (Breiman, 2004). RO yönteminde model oluşturulurken seçilen verinin bir dahaki sefere yerine konulması ve rastgele seçim yapılması şartıyla veri setinin $2/3$ 'ü öğrenme veri seti (inBag), $1/3$ 'ü ise test veri seti (Out-Of-Bag) olarak ayrılır (Bhalla, 2014). Bootstrap tekniği ile oluşturulması düşünülen karar ağacı sayısı kadar örneklem oluşturulur ve her örneklem için eğitim ve test verisi ayrılır (Akman, 2010, s.34).
- m tane bağımsız değişken için, her bir düğümde m tane değişken arasında en iyi bölünmeyi sağlayan ve rastgele seçilen belirli bir m sayısı kadar değişken mevcuttur. m değeri orman oluşurken sabit tutulur (Breiman, 2004). m , sınıflandırma için tüm bağımsız değişkenlerin sayısının karekökü iken regresyon için m 'nin üçe bölünmesiyle elde edilir (Bhalla, 2014). Bootstrap örneklemlerinin her biri için budama yapılmamış sınıflandırma veya regresyon ağaçları oluşturulurken; her bir düğümde, tüm değişkenler arasında en iyi bölünmeyi seçmek yerine değişkenlerin rastgele örneklemlerinin (m tane) arasındaki en iyi bölünme seçilir (Liaw & Wiener, 2002). Ormanda yer alacak karar ağaçları SVRA yöntemi ile oluşturulur. Bunun için tahmin değişkeni belirlendikten sonra Gini indeksi ile en iyi dallanma kriteri indeksi hesaplanır ve düğüm iki dala ayrılır (Akman, 2010, s.35).

- Regresyonda bir x test vektörü her bir ağaç indirildiğinde bulunduğu düğümdeki y değerlerinin ortalama değerine atanır. Bu değerlerin ormanda yer alan bütün ağaçlardaki ortalaması x için öngörülen değerdir. Sınıflandırma için öngörülen değer ise orman oylarının çoğunluğunu (mod) alan sınıftır (Breiman, 2004). Başka bir ifadeyle n tane ağacın tahminleri birleştirilerek yeni veri setine göre tahmin yapılır. Bu birleştirme işlemi için regresyonda ortalama ve sınıflamada oy çoğunluğu (mod) kullanılır (Liaw & Wiener, 2002).

RO yönteminin işleyişine yönelik aşamalar Şekil 6'da yer almaktadır (Ercire, 2019, s.28).



Şekil 6. RO yönteminin işleyişine yönelik aşamalar.

RO yönteminde, bootstrap tekniği ile seçilen örneklem ve her düğümde değişkenler arasından belirlenen sayıda seçilen değişkenler yöntemin performansı üzerinde önemli etkiye sahiptir (Atasever, 2011, s.32). RO yönteminin işleyişinde kullanıcı tarafından

tanımlanan iki parametre büyük önem taşımaktadır. Bunlar ağaç oluşturmak için her bir düğümde kullanılan değişken ve oluşturulacak ağaç sayısıdır. Her bir düğüm oluşturulurken yalnızca seçilen değişkenler incelenir. RO yönteminde üretilcek ağaç sayısı kullanıcı tarafından belirlenen herhangi bir değer olabilir (Pal, 2005).

RO yönteminde orman hata oranını etkileyen iki önemli unsur vardır. Bunlardan biri ağaçlar arasındaki korelasyondur. Eğer iki ağaç arasındaki korelasyon yüksek ise bu durum ormandaki hata oranını artırır. Orman hata oranını etkileyen diğer bir unsur ise ağacın gücüdür. Hata oranı düşük olan bir ağaç güçlü bir sınıflandırıcıdır. Dolayısıyla her bir ağacın hata oranını azaltmak başka bir ifadeyle ağacı güçlü hale getirmek orman hata oranını azaltır. Analizde kullanılacak değişken sayısı (m) azaldıkça hem korelasyon hem de gücü azaltır. Artırmak ise her ikisini de artırır. Değişken sayısının en düşük ve en yüksek değerlerinin arasında optimal bir m değer aralığı vardır. Bu aralık test verisi hata oranı kullanılarak hesaplanabilir (Breiman & Cutler, t.y). Verilerden bir önyükleme örnekleme alındığında, bazı gözlemler önyükleme örnekleminde yer almaz. Bunlara eğitim veri seti (out of bag) denir ve genelleme hatası ile değişken önemini tahmin etmek için son derece kullanışlıdır (A. Cutler vd., 2012, s.164). Eğitim veri seti hata oranını hesaplamak için öncelikle veri seti, önyükleme aşamasında ağaç oluşturmak için kullanılan (inbag) ve kullanılmayan (out of bag) veri olmak üzere iki gruba ayrılır. Ardından kullanılmayan veri ile ağaç test edilerek hata tahmini yapılır. Daha sonra bireysel ağaçlardan elde edilen eğitim veri seti tahminleri birleştirilerek ormanın eğitim veri seti hata oranı hesaplanır (Akman, 2010, s.36). Regresyon yapılırken bu hata oranı $MSE_{oob} = \frac{1}{N} \sum_{i=1}^N (y_i - \hat{f}_{oob}(x_i))^2$ formülü ile hesaplanırken, sınıflama yapıldığında, $E_{oob} = \frac{1}{N} \sum_{i=1}^N I(y_i \neq \hat{f}_{oob}(x_i))$ formülü ile hesaplanır. Her iki formülde $\hat{f}_{oob}(x_i)$, i . gözlem için tahmin edilen eğitim veri setidir (A. Cutler vd., 2012, s.165).

2.5. Yapay Sinir Ağları

Yapay sinir ağları, biyolojik organizmalardaki beyin ve sinir sistemi çalışmalarından esinlenerek oluşturulan bir hesaplama yapısıdır (Karunanithi, Whitley & Malaiya, 1992; V. Rao & Rao, 1998, s.14). Biyolojik sistemlerin hesaplama stilini temel alan YSA, paralel olarak hareket eden basit ve birbirine bağlı işlem birimlerinin (nöronların) bir topluluğu olarak görülebilir. Bu ayrı birimler değeri veya ağırlığı olan tek yönlü bağlantılar (veya sinapslar) kullanarak birbirleriyle iletişim kurarlar (Yuhas & Ansari, 2012, s. 2). YSA genellikle çok sayıda nörondan meydana gelen, karmaşık biçimde birbirine bağlanan ve katmalar halinde organize edilen doğrusal olan ve olmayan hesaplama elemanlarından oluşur (Sarle, 1994).

YSA kavramı ile ilgili ilk temel çalışma 1943 yılında McCulloch ve Pitt tarafından yapılmıştır. Daha önceki dönemlerde beyin fonksiyonları hakkında yapılan çalışmalar olsa da günümüzdeki YSA kavramına en yakın olanı bu çalışmadır. Daha sonra 1949 yılında Donald Olding Hebb günümüzde de YSA mimarilerinde kullanılan ve birçok öğrenme kuralının temelini oluşturan, yapay hücrelerden oluşan bir yapay sinir ağının değerini değiştirebilen *Hebbian Öğrenme Kuralını* geliştirmiştir (Öztemel, 2016, s.37). 1969 yılında Marvin Minsky ve Seymour Papert tarafından *Algılayıcı (perceptron) Öğrenme Kuramına* yönelik yapılan eleştiriler, bilim dünyasındaki araştırmacıların YSA'na yönelik ilgisini azaltmış ve 1980'li yıllara kadar çok az sayıda çalışmanın yapılmasına neden olmuştur. 1982 yılında John Hopfield basit bir analog devre modeli ile Hopfield ağlarındaki hesaplama yeteneklerini kullanarak bazı problemlere çözüm getirmiştir. Ardından James McClelland, David Rumelhart ve Francis Crick'in yer aldığı PDP (Parallel Distributed Processing) araştırma grubunun yaptığı çalışmalar YSA ağlarının gelişiminde önemli bir rol oynamış ve Amerika Fizik Topluluğu (The American Physical Society) tarafından 1985 yılında ilk defa YSA araştırmalarının yer aldığı bir konferans düzenlenmiştir (Akpınar, 2017, s.323). 1990'lı yıllardan sonra ise YSA ile çok sayıda araştırma yapılmış ve günümüze kadar gelişimini göstermeye devam etmiştir (Öztemel, 2016, s.41).

Sinir ağıları genelleme yapabilmek için büyük bir potansiyele sahiptir, çünkü bileşenlerin hesaplamaları büyük ölçüde birbirinden bağımsızdır. Bir sinir ağı, nöronlar arasındaki bağlantı örüntüsü (ağ mimarisi olarak adlandırılır) ve bağlantılardaki ağırlıkları belirleme yöntemi (eğitim veya öğrenme algoritması olarak adlandırılır) ile dizayn edilmiştir. Ağırlıklar veri bazında ayarlanır (Wu & McLarty, 2012, s.11). Başka bir ifadeyle, sinir ağıları örneklerden öğrenir (Wu & McLarty, 2012, s.11; Zurada, 1992, s.2). Bu sayede eğitim verilerinin ötesinde bir miktar genelleme yeteneği sergiler (Wu & McLarty, 2012, s.11). Bu özellik, bu tür hesaplama modellerini, çözülecek sorunun çok az veya eksik anlaşıldığı, ancak eğitim verilerinin kolayca bulunduğu uygulama alanlarında çok cazip hale getirir (Boznar, Lesjak & Malke, 1993). Sinirbilimden esinlenmiş olmalarına rağmen, YSA'nın çoğunluğu parametrik olmayan örüntü sınıflandırıcıları, kümeleme algoritması, doğrusal olmayan filtreler ve istatistiksel regresyon modelleri gibi geleneksel istatistiksel yöntemlerle yakın ilişkilidir (Sarle, 1994; Wu & McLarty, 2012, s.11). Yapılan çalışmalar ele alındığında YSA'nın veri setinde doğrusal olmayan, çok boyutlu, gürültülü ve kayıp verilerin olduğu ve çözülmesi istenilen probleme yönelik matematiksel bir modelin veya algoritmanın bulunmadığı durumlarda yaygın bir biçimde kullanıldığı görülmektedir (Öztemel, 2016, s.36). Günümüzde, çeşitli endüstriyel tesislerin kontrolü dâhil olmak üzere, sinir ağlarının endüstriyel uygulamalara çok sayıda uygulaması vardır. Bunlardan bazıları örüntü tanıma (retina özellik çıkarma, parmak izi tanıma), iş tahmini (enerji talebi), risk kontrolü (kredi kartı risk kontrolü) (Taylor, 2002, s.87), konuşma sentezi ve tanıma, insan ve karmaşık fiziksel sistemler arasındaki uyarlanabilir arayüzler, fonksiyon yaklaşımı, görüntü sıkıştırma, kümeleme, optimizasyon, doğrusal olmayan sistem modellemesi ve kontrolü (Wu & McLarty, 2012, s.11), probablistik fonksiyon kestirimi, zaman serileri analizi, veri sıkıştırma, sinyal filtreleme, doğrusal olmayan modelleme ve sınıflandırmadır (Öztemel, 2016, s.36). YSA'nın sağladığı bazı avantaj ve dezavantajları aşağıda özetlenmektedir:

Doğrusal olmama: Yapay bir nöron doğrusal veya doğrusal olmayabilir. Doğrusal olmayan nöronların ara bağlantısından oluşan bir sinir ağı, doğrusal değildir. Ayrıca, doğrusal olmama, ağın her tarafına dağılması açısından özel bir türdür (Haykin, 2009, s.3). Bu özellik

YSA modelleri için doğrusal olmayan gerçek dünya problemlerine yönelik öğrenmeler için veri üretimi yapabilme yeteneği kazandırır (Kantardzic, 2011, s.200).

Adapte olabilirlik: Sinir ağları sinaptik ağırlıklarını çevredeki değişikliklere uyarlamak için yerleşik bir kapasiteye sahiptir. Özellikle, belirli bir ortamda çalışmak üzere eğitilmiş bir sinir ağı, çalışma ortam koşullarındaki küçük değişikliklerle başa çıkmak için kolayca yeniden eğitilebilir. Dahası, durağan olmayan bir ortamda (yani istatistiklerin zamanla değiştiği bir ortamda) çalıştığında, bir sinir ağı sinaptik ağırlıklarını gerçek zamanlı olarak değiştirmek için tasarlanabilir. Desen sınıflandırması, sinyal işleme ve kontrol uygulamaları, ağın uyarlanabilir kabiliyeti ile birleştiğinde, onu uyarlanabilir model sınıflandırması, uyarlanabilir sinyal işleme ve uyarlanabilir kontrol için yararlı bir araç haline getirir. Genel bir kural olarak, bir sistemi ne kadar uyarlanabilir hale getirirsek, sistemin sabit kalmasını sağlayan zaman, sistemin durağan olmayan bir ortamda çalışması gerektiğinde performansı ne kadar sağlam olur (Haykin, 2009, s.3).

Bağlamsal bilgi: Bilgi, bir sinir ağının yapısı ve aktivasyon durumu ile temsil edilir. Ağdaki her nöron, ağdaki diğer tüm nöronların küresel aktivitesinden potansiyel olarak etkilenir. Sonuç olarak, bağlamsal bilgiler doğal olarak bir sinir ağı tarafından ele alınır (Haykin, 2009, s.4).

Kanıt Cevabı: Örüntü sınıflandırması bağlamında, bir sinir ağı sadece hangi özel örüntünün seçileceği hakkında değil, aynı zamanda alınan karara olan güven hakkında da bilgi sağlamak üzere tasarlanabilir. Bu son bilgiler, ortaya çıkmaları durumunda belirsiz örüntüleri reddetmek için kullanılabilir ve böylece ağın sınıflandırma performansını iyileştirir (Haykin, 2009, s.4).

Hata Toleransı: Uygulanan bir sinir ağı, performansının olumsuz çalışma koşulları altında azalmasına karşın doğası gereği hataya dayanıklı veya güçlü hesaplama yeteneğine sahiptir (Haykin, 2009, s.4).

VLSI Uygulanabilirliği: Bir sinir ağının büyük ölçüde paralel doğası, belirli görevlerin hesaplanması için potansiyel olarak hızlı olmasını sağlar. Bu aynı özellik, çok büyük ölçekli

entegre (VLSI) teknolojisini kullanarak bir sinir ağını uygulama için çok uygun hale getirir (Haykin, 2009, s.4).

Analiz ve Tasarım Tekdüzeligi: Temel olarak, sinir ağları bilgi işlemcileri olarak evrensellikten hoşlanırlar. Bunu, sinir ağlarının uygulanmasını içeren tüm alanlarda aynı gösterimin kullanıldığı anlamında söyleriz. Bu özellik kendini farklı şekillerde gösterir:

- Nöronlar, şu veya bu şekilde, tüm sinir ağlarında ortak olan bir bileşeni temsil eder.
- Bu ortak nokta, sinir ağlarının farklı uygulamalarında teorileri ve öğrenme algoritmalarını paylaşmayı mümkün kılar.
- Modüler ağlar, modüllerin sorunsuz entegrasyonu ile kurulabilir (Haykin, 2009, s.4).

Nörobiyolojik Analoji: Sinir ağının tasarımı, hataya dayanıklı paralel işlemenin sadece fiziksel olarak değil, aynı zamanda hızlı ve güçlü olduğunun kanıtı olan beyinle kıyaslanarak motive edilir (Haykin, 2009, s.4).

Örneklerden Öğrenme: Bir YSA, bir dizi eğitim veya öğrenme örneği uygulayarak ara bağlantı ağırlıklarını değiştirir. Bir öğrenme sürecinin nihai etkileri, bir ağın ayarlanmış parametreleridir (parametreler yerleşik modelin ana bileşenleri aracılığıyla dağıtılır) ve mevcut sorun için örtük olarak saklanan bilgileri temsil eder (Kantardzic, 2011, s.200).

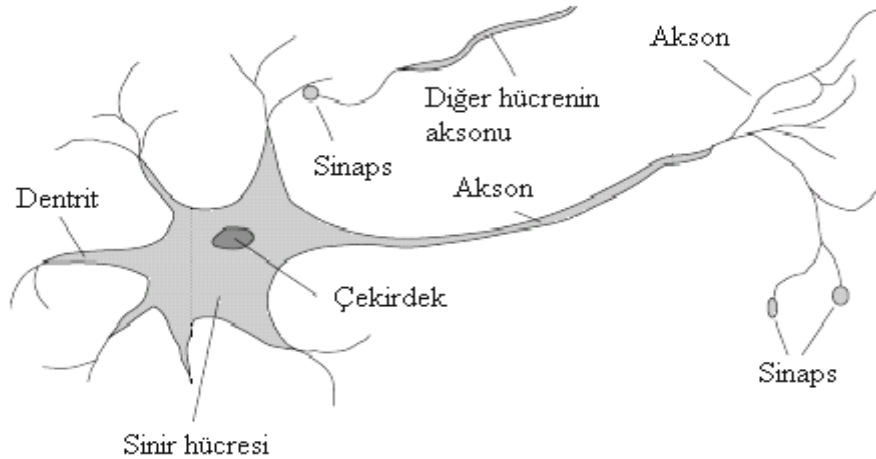
Birçok YSA modeli iyi bilinen istatistiksel modellerle benzer veya özdeş olsa da, YSA literatüründeki terminoloji, istatistik literatüründeki terminolojiden oldukça farklıdır. Bunlardan bazıları değişken/özellik, bağımsız değişken/girdi, tahmin değeri/çıktı, bağımlı değişken/hedef veya eğitim değeri, artık/hata, kestirim/eğitim, öğrenme, adaptasyon, parametre tahmini/sinaptik ağırlıklar, tahmin ölçütü/hata işlevi, maliyet işlevi veya Lyapunov işlevi, gözlem/desen veya eğitim çifti, etkileşim/üst düzey nöronlar, dönüşüm/fonksiyonel bağlantı, regresyon ve ayırt edici analiz/denetimli öğrenme, veri azaltma/denetimsiz öğrenme, kodlama veya otomatik ilişkilendirme, kümeleme analizi/rekabetçi öğrenmedir (Sarle, 1994).

2.5.1. Yapay Sinir Ağının Yapısı

Yapay sinir ağları biyolojik beyin ve sinir sistemini temel alarak yapılandırıldığı için biyolojik sinir yapısıyla karşılıklı olarak değerlendirilmelidir. Bu nedenle öncelikle biyolojik sinir yapısını tanıtmak ve daha sonra yapay sinir hücrelerini açıklamak gerekir.

2.5.1.1. *Biyolojik Sinir Hücresi*

İnsan beyni tarafından gerçekleştirilen bilgi işleme, düşünme ve öğrenme gibi uygun işlevleri üretmek için paralel olarak çalışan biyolojik işleme bileşenleri tarafından gerçekleştirilir. Merkezi sinir sisteminin temel hücresi nörondur ve rolü, belirli çalışma koşulları altında dürtüleri (fiziksel-kimyasal reaksiyonlardan kaynaklanan elektrik uyarıları) yürütmek için kullanılır (Da Silva, Spatti, Flauzino, Liboni & dos Reis Alves, 2017, s.9). Bir nöron, çoklu dentritler ve bir akson olmak üzere iki tip uzantıya sahip basit hücrenin bir uzantısı olup normal bir hücrenin tüm özelliklerine sahiptir. Nöronlar, dendrit, soma (hücre gövdesi), akson ve sinapstan oluşan dört bileşenli bir yapıya sahiptir (Du & Swamy, 2013, s.4). Bir biyolojik nöronun yapısı Şekil 7’de yer almaktadır (Canan, 2006, s.41).



Şekil 7. Bir biyolojik nöronun yapısı.

Şekil 7’ye göre bir nöronun işleyişi, soma ve dendritleri aracılığıyla diğer nöronlardan sinyaller alır, bunları birleştirir ve aksonu aracılığıyla diğer nöronlara çıkış sinyalleri

gönderir. Dendritler birkaç komşu nörondan sinyal alır ve bunları hücre gövdesine geçirir ve burada işlenir ve elde edilen sinyal bir akson üzerinden aktarılır (Du & Swamy, 2013, s.4). Dendritlerin temel amacı, sürekli olarak, diğer birkaç nörondan veya çevre ile temas eden bazı nöronların (duyu nöronları olarak da bilinir) dış ortamdan uyarıcıları elde etmektir. Hücre gövdesi, nöronun aksonu boyunca bir elektrik sinyalini tetikleyip tetikleyemeyeceğini gösteren bir aktivasyon potansiyeli üretmek için dendritlerden gelen tüm bilgilerin işlenmesinden sorumludur. Akson, görevi elektriksel uyarıları diğer nöronlara bağlayan veya doğrudan kas dokusuna (efferent nöronlar) bağlı nöronlara yönlendirmek olan tek bir uzantıdan oluşur. Aksonun son kısmı ise sinaptik terminaller olarak adlandırılan dallardan oluşur. Sinapslar, elektrik sinyallerinin belirli bir nörondan diğer nöronların dendritlerine aktarılmasını sağlayan bağlantılardır. Sinaptik birleşme noktasını oluşturan nöronlar arasında fiziksel bir temas olmadığından dolayı elemanları bir nörondan diğerine ağırlıklandırarak iletir (Da Silva, vd. 2017, s.9).

2.5.1.2. Yapay Sinir Hücresi

Biyolojik sinir hücresini temel alan yapay bir sinir hücresi, girdi, ağırlık, toplama fonksiyonu, aktivasyon fonksiyonu ve çıktı olmak üzere beş bileşenden meydana gelir. Girdi, yapay sinir hücresine dış dünyadan, kendisinden veya başka hücrelerden gelen bilgilerdir. Ağırlık, girdilerin önemini ve hücre üzerindeki etkisini gösteren değişken veya sabit değerlerdir. Bu değerlerin büyük veya küçük olması önem sırası üzerinde bir etkiye sahip değildir. Ağırlığı küçük olan bir girdi büyük öneme sahip olabilir (Öztemel, 2006, s.50). Toplam fonksiyonu nöronun ilgili sinaptik kuvvetlerine göre ağırlıklandırılmış giriş sinyallerini toplar (Haykin, 2009, s.10). YSA’da kullanılan bazı toplam fonksiyonları Tablo 2’de gösterilmektedir (Çelik, 2018).

Tablo 2

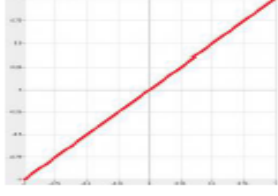
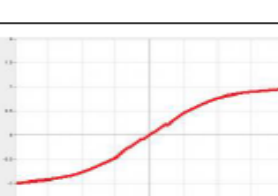
YSA’da Kullanılan Bazı Toplam Fonksiyonları

Fonksiyonun adı	Fonksiyon	Açıklama
Ağırlıklandırılmış toplam	$NET = \sum x_i \cdot w_i$	Girdiler ve ağırlıklı değerler çarpılır. Hesaplanan değerler toplanır.
Çarpma	$NET = \prod x_i \cdot w_i$	Girdiler ve ağırlıklı değerler çarpılır. Hesaplanan değerler çarpılır.
Maksimum	$NET = \text{Max}(x_i \cdot w_i)$	Girdiler ve ağırlıklı değerler çarpılır. En yüksek hesaplanan değer alınır.
Minimum	$NET = \text{Min}(x_i \cdot w_i)$	Girdiler ve ağırlıklı değerler çarpılır. En düşük hesaplanan değer alınır.
Artan toplam	$NET_k = NET_{k-1} + \sum x_i \cdot w_i$	Ağırlıklandırılmış toplam hesaplanır. Daha önce hesaplanan ağırlıklandırılmış toplam hesaplanır.

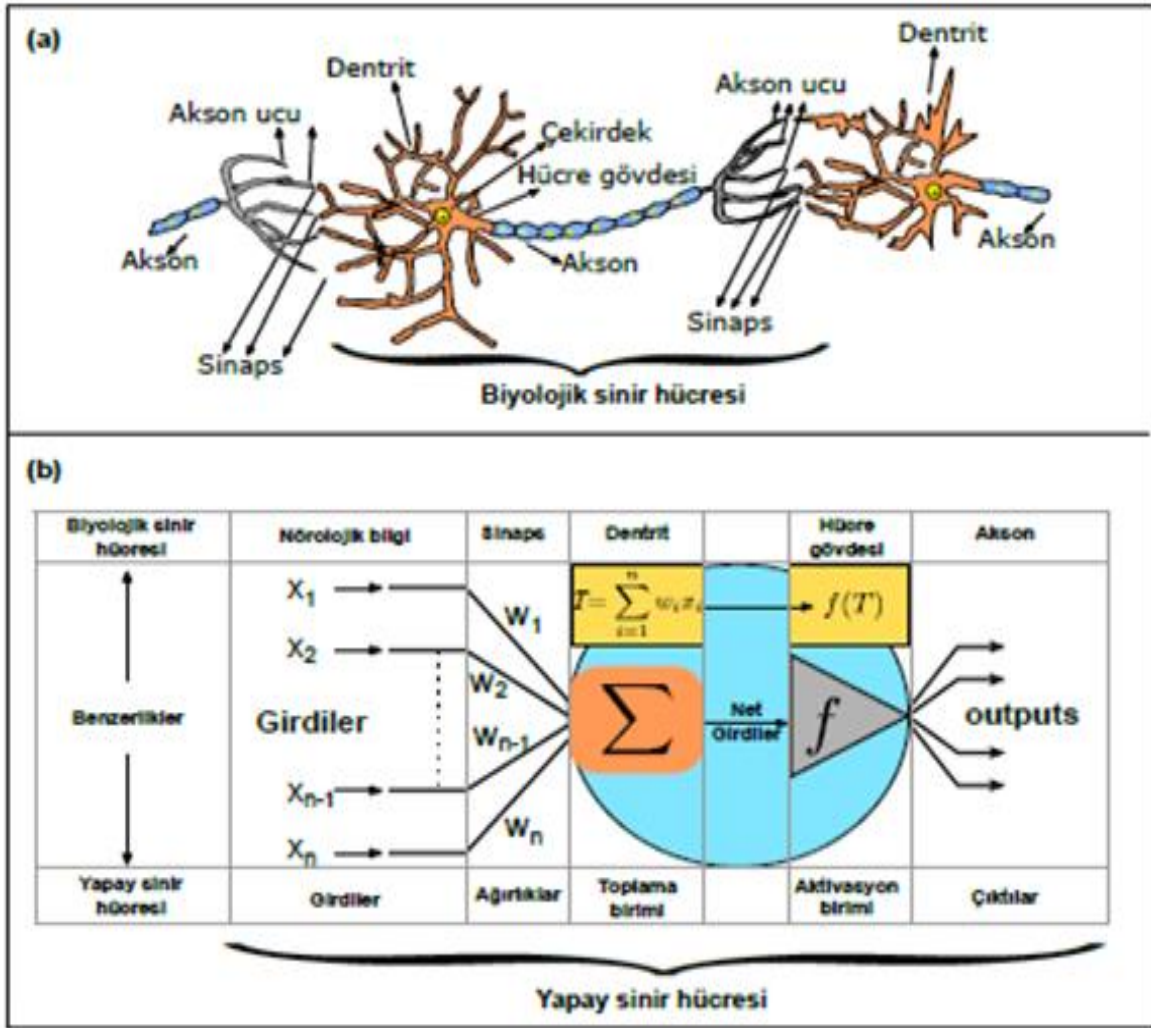
Toplam fonksiyonu bir hücredeki net girdi değeridir. Birçok toplam fonksiyonu mevcuttur ancak çoğunlukla ağırlıklı toplam fonksiyonu kullanılır. Girdilerin kendi ağırlıkları ile çarpımlarının toplam değeri olarak hesaplanır (Öztemel, 2006, s.50). Aktivasyon fonksiyonu, bir nöronun çıktı aralığını sonlu bir değerle sınırlar (Haykin, 2009, s.10). YSA’da kullanılan bazı aktivasyon fonksiyonları aşağıdaki gibidir (Kantardzic, 2011, s.203).

Tablo 3

YSA’da Kullanılan Bazı Aktivasyon Fonksiyonları

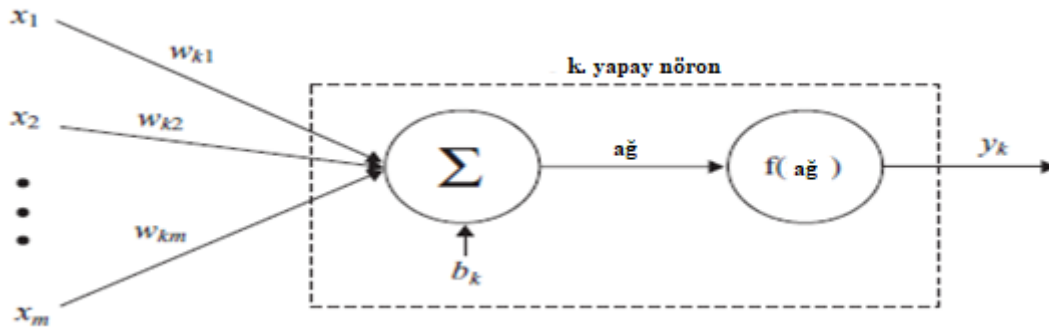
Doğrusal (Lineer) Aktivasyon Fonksiyonu		$F(Net) = A * Net$ (A sabit bir sayı)
Adım (Step) Aktivasyon Fonksiyonu		$F(Net) = \begin{cases} 1 & \text{if } Net > \text{Eşik Değer} \\ 0 & \text{if } Net \leq \text{Eşik Değer} \end{cases}$
Sigmoid Aktivasyon Fonksiyonu		$F(Net) = \frac{1}{1 + e^{-Net}}$
Tanjant Hiperbolik Aktivasyon Fonksiyonu		$F(Net) = \frac{e^{Net} + e^{-Net}}{e^{Net} - e^{-Net}}$
Eşik Değer Fonksiyonu		$F(Net) = \begin{cases} 0 & \text{if } Net \leq 0 \\ Net & \text{if } 0 < Net < 1 \\ 1 & \text{if } Net \geq 1 \end{cases}$
Sinüs Aktivasyon Fonksiyonu		$F(Net) = \sin(Net)$

Aktivasyon fonksiyonu, toplam fonksiyonundan elde edilen net girdileri işleyerek yapay sinir hücresindeki her bir girdiye karşılık gelen çıktı değerini belirler. Çok sayıda aktivasyon fonksiyonu bulunur ancak çoğunlukla özellikle de çok katmanlı YSA modellerinde sigmoid fonksiyonu kullanılır. Yapay sinir hücresi içindeki tüm elemanlara aynı aktivasyon fonksiyonunun uygulanması zorunlu değildir (Öztemel, 2006, s.50). Bir biyolojik ve yapay sinir hücrelerinin elemanlarına ait karşılaştırma Şekil 8’de yer almaktadır (Yeşilkanat, 2016, s.32).



Şekil 8. Bir biyolojik ve yapay sinir hücrelerinin elemanlarına ait karşılaştırma.

Matematiksel olarak, yapay bir nöron, doğal bir nöronun soyut bir modelidir ve işleme yetenekleri Şekil 9’da yer almaktadır (Kantardzic, 2011, s.202).



Şekil 9. Yapay bir nöronun işleme yetenekleri.

Her bir x_i girdisi karşılık gelen w_{ki} ağırlığı ile çarpılır, k bir YSA'daki verilen bir nöronun indeksidir. Ağırlıklar, doğal bir nörondaki biyolojik sinaptik güçleri simüle eder. $i=1, \dots, m$ için $x_i w_{ki}$ 'lerin toplamının ağırlıklı toplamı, YSA literatüründeki “ağ (net)” olarak ifade edilir:

$$net_k = x_1 w_{k1} + x_2 w_{k2} + \dots + x_m w_{km} + b_k$$

Benimsenen notasyonu ile $w_{k0} = b_k$ ve $x_0 = 1$ ön tanımlı girdi kullanılarak, ağ toplamının yeni bir versiyonu şu şekildedir:

$$net_k = x_0 w_{k0} + x_1 w_{k1} + x_2 w_{k2} + \dots + x_m w_{km} = \sum_{i=0}^m x_i w_{ki}$$

Aynı toplam, iki m -boyutlu vektörlerin skalar toplamı olarak vektör halinde ifade edilebilir:

$$net_k = X \cdot W$$

Burada, $X = \{x_0, x_1, x_2, \dots, x_m\}$, $W = \{w_{k0}, w_{k1}, w_{k2}, \dots, w_{km}\}$ ile ifade edilir. Son olarak, bir yapay nöron y_i çıktısını net_k değerinin belirli bir fonksiyonu olarak hesaplar: $y_k = f(net_k)$

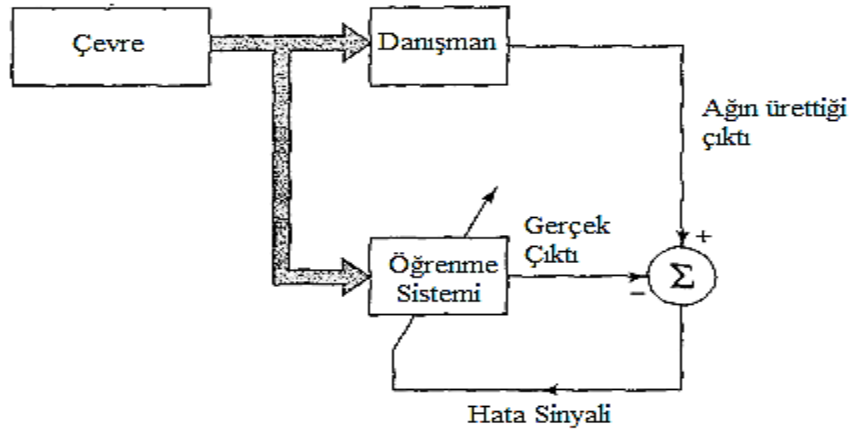
2.5.2. Yapay Sinir Ağlarının Sınıflandırılması

2.5.2.1. Öğrenme Türlerine Göre YSA

YSA'da öğrenme bağlantılardaki ağırlıkların, etkinliklerin veya aktarım işlevlerinin değişimi ile meydana gelir. Bunun sonucunda eğitilen ağlar sınıflama veya şekil tanıma gibi birçok işlemi gerçekleştirebilir (Elmas, 2007, s. 87). Ağ yapısında bir düzeltme sağlamak için belirli bir yanıt belirtilebilir veya belirtilmeyebilir. Girişler / çıkışlar hakkındaki bilgiler bilinmediği veya önsel olduğu zaman öğrenme gereklidir aksi halde önceden bir ağ tasarımı gerçekleştirilemez (Zurada, 1992, s.56). YSA örneklerden öğrenen bir sisteme sahiptir ve bu sistem kullanılan öğrenme algoritmasına göre değişim göstermektedir. Genel olarak YSA modelleri için danışmanlı (öğretmenli), danışmansız (öğretmensiz) ve destekleyici olmak üzere üç öğrenme stratejisi vardır (Öztemel, 2016, s.25).

2.5.2.1.1. Denetimli (Danışmanlı) Öğrenme

Denetimli öğrenme, bilinen girdi/çıkı örneklerinden bilinmeyen bir bağımlı değeri tahmin etmek için kullanılır. Sınıflama ve regresyon, bu tür tümevarımsal öğrenme tarafından desteklenen uygulamalardır. Denetimli öğrenme bir danışmanın/öğretmenin-uygunluk fonksiyonunun veya önerilen modeli tahmin etmek için başka bir dış yöntemin varlığını varsayar. “Denetimli” terimi, eğitim verisi için çıktı değerlerinin bilindiğini gösterir (Kantardzic, 2011, s.99). Sinir yapıları işlemi ve sonucunu temsil eden girdi / çıktı verileri için oluşan bir değer tablosu üzerinden öğrenilen sistem hakkında “hipotez” oluşturacaktır. Bu durumda, denetimli öğrenmenin uygulanması yalnızca söz konusu değer tablosunun kullanılabilirliğine ve her örnek için doğru yanıtın ne olduğunu öğreten danışmana/öğretmene bağlıdır. Ağın sinaptik ağırlıkları ve eşikleri, ayarlama prosedüründeki bu farkı kullanarak, üretilen çıktılar arasındaki uyumsuzluğu denetleyen öğrenme algoritmasının, kendisi tarafından yürütülen karşılaştırmalı sonuçları sürekli olarak ayarlanır. Bu tutarsızlık, çözümlerin genelleştirilmesi amaçları dikkate alınarak kabul edilebilir bir değer aralığında olduğunda ağ “eğitilmiş” kabul edilir (Da silva vd., 2017, s.26). Başka bir ifadeyle ağa giriş ve çıkış değerleri birlikte verilerek, ağın istenen çıktıları oluşturabilmesi için ağırlıklarını güncellemesi sağlanır. Ardından ağın bütün çıktıları ile beklenen çıktıları arasındaki fark hesaplanarak hata payı elde edilir. Son olarak bu hata payına göre ağın çıktıları, ağın yeni ağırlıkları ve her eleman için düşen ağırlık payı yeniden düzenlenir (Çayıroğlu, 2015, s.3). Şekil 10, danışmanlı öğrenme biçimini gösteren bir diyagramı göstermektedir (Haykin, 1999, s.63).



Şekil 10. Danışmanlı öğrenme biçimini gösteren bir diyagram.

Kavramsal olarak, öğretmeni çevre hakkında bilgi sahibi olarak düşünebiliriz, bu bilgi bir dizi girdi-çıkı örneğiyle temsil edilir. Bununla birlikte, çevre, ilgili sinir ağı tarafından bilinmemektedir. Şimdi, öğretmenin ve sinir ağının her ikisinin de ortamdan çizilen bir eğitim vektörüne (yani örnek) maruz kaldığını varsayalım. Yerleşik bilgi sayesinde, öğretmen sinir ağına o eğitim vektörü için istenen bir tepki verebilir. Aslında, istenen yanıt sinir ağı tarafından gerçekleştirilecek optimum eylemi temsil eder. Ağ parametreleri, eğitim vektörünün ve hata sinyalinin birleşik etkisi altında ayarlanır. Hata sinyali, istenen yanıt ile ağın gerçek yanıtı arasındaki fark olarak tanımlanır. Bu ayarlama, nihayetinde sinir ağını öğretmeni taklit etmek amacıyla adım adım yinelemeli olarak gerçekleştirilir; benzemeye çalışmanın bazı istatistiksel anlamda optimum olduğu varsayılır. Bu şekilde, öğretmen için mevcut ortam bilgisi mümkün olduğunca eksiksiz bir şekilde eğitim yoluyla sinir ağına aktarılır. Bu duruma ulaşıldığında, öğretmenden vazgeçebilir ve sinir ağının çevreyle tamamen başa çıkmasına izin verebiliriz (Haykin, 1999, s.63).

2.5.2.1.2. Danışmansız Öğrenme

Gözetimsiz öğrenme şeması altında, bir öğrenme sistemine sadece girdi değerleri olan örnekler verilir ve öğrenme sürecinde çıktı hakkında bir fikir yoktur. Gözetimsiz öğrenme öğretmeni ortadan kaldırır ve öğrencinin modeli kendi başına oluşturmasını ve değerlendirmesini gerektirir. Denetimsiz öğrenmenin amacı girdi verilerinde “doğal” yapıyı

keşfetmektir. Biyolojik sistemlerde algı, denetimsiz tekniklerle öğrenilen bir görevdir (Kantardzic, 2011, s.100). Denetimli öğrenmeden farklı olarak, denetimsiz öğrenmeye dayalı bir algoritmanın uygulanması, istenen ilgili çıktılar hakkında herhangi bir bilgi gerektirmez. Dolayısıyla, tüm örnek kümesini oluşturan elementler arasında benzerlikler sunan alt kümeleri (veya kümeleri) tanımlayan mevcut özellikler olduğunda ağın kendini organize etmesi gerekir. Öğrenme algoritması, bu kümeleri ağın içinde yansıtmak için ağın sinaptik ağırlıklarını ve eşiklerini ayarlar. Alternatif olarak, ağ tasarımcısı, sorun hakkındaki bilgisini (a priori) kullanarak bu olası kümelerin maksimum miktarını belirleyebilir (Da silva vd., 2017, s.26). Başka bir ifadeyle, danışmansız öğrenmede ağa öğrenme sırasında sadece örnek girdiler verilirken herhangi bir çıktı bilgisi verilmez. Girdilere göre ağ bağlantı ağırlıklarını aynı özellikte olan örnekleri ayırabilecek biçimde düzenler ve her bir örneği kendi arasında sınıflandıracak şekilde kendi kurallarını oluşturur (Çayıroğlu, 2015, s.3). Şekil 11, danışmansız öğrenme biçimini gösteren bir diyagramı göstermektedir (Haykin, 1999, s.65).

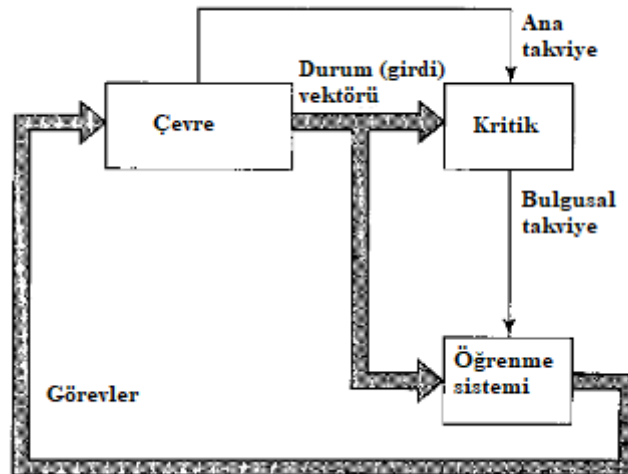


Şekil 11. Danışmansız öğrenme biçimini gösteren bir diyagram.

Gözetimsiz öğrenme gerçekleştirmek için rekabetçi bir öğrenme kuralı kullanabiliriz. Örneğin, bir giriş katmanı ve bir rekabetçi katman olmak üzere iki katmandan oluşan bir sinir ağı kullanabiliriz. Giriş katmanı mevcut verileri alır. Rekabetçi katman, girdi verilerinde yer alan özelliklere cevap verme fırsatı için bir öğrenme (kuralına göre) birbiriyle rekabet eden nöronlardan oluşur. En basit şekliyle, ağ "kazanan her şeyi alır" stratejisine uygun olarak çalışır. (Haykin, 1999, s.66).

2.5.2.1.3. Destekleyici Öğrenme

Destekleyici öğrenme, ağın çevresinden geri bildirimler alması sebebiyle denetimli öğrenmenin bir biçimi olarak ele alınır. Ancak geri bildirim öğretici olmaktan ziyade sadece değerlendircidir (Gumma, 2004, s.24). Bir ağ denetimli öğrenme ile eğitildiğinde her bir giriş için istenen çıkış değerleri kesin olarak bilinir. Ancak bazı durumlarda daha az ayrıntı mevcuttur. Destekleyici öğrenme sadece kısmi bilgi ağına tepki doğruluğu hakkında bilgi var olduğunda kullanılır. En uç durumda, ağa talimat sağlamak için yalnızca "doğru" veya "yanlış" bilgiler mevcuttur. Bu gibi durumlarda ortamdan alınan geri bildirim miktarı çok sınırlıdır, çünkü danışmanın geribildirimi öğretici olmaktan ziyade değerlendircidir (Zurada, 1992, s.73). Takviye öğrenmede, bir girdi-çıkı eşleşmesinin öğrenilmesi, skaler bir performans indeksini en aza indirmek için çevre ile sürekli etkileşim yoluyla gerçekleştirilir. Destekleyici öğrenme yaklaşımında ağı her iterasyonu sonucunda elde ettiği sonucun uygun olup olmadığına dair edinilen bilgiler doğrultusunda ağ kendini yeniden düzenler. Böylece ağ herhangi bir girdi dizisiyle hem öğrenerek hem de sonuç çıkararak işlemeye devam eder (Çayıroğlu, 2015, s.3). Şekil 12, destekleyici öğrenme biçimini gösteren bir diyagramı göstermektedir (Haykin, 1999, s.64).



Şekil 12. Destekleyici öğrenme biçimini gösteren bir diyagram.

Sistem, gecikmeli takviye altında öğrenmek için tasarlanmıştır, yani sistemin, çevreden alınan geçici bir uyarın dizisini (yani durum vektörleri) gözlemlediği, sonuçta sezgisel takviye sinyalinin üretilmesine neden olduğu gözlemlenir. Öğrenmenin amacı, sadece acil maliyet yerine bir dizi adımda atılan toplam işlem maliyetinin beklentisi olarak tanımlanan bir gidilecek maliyet işlevini en aza indirmektir. Bu zaman adımlarında daha önce yapılan bazı eylemlerin aslında genel sistem davranışının en iyi belirleyicileri olduğu ortaya çıkabilir. Sistemin ikinci bileşenini oluşturan öğrenme makinesinin işlevi, bu eylemleri keşfetmek ve onları çevreye geri beslemektir (Haykin, 1999, s.65).

2.5.2.2. Öğrenme Zamanına Göre Yapay Sinir Ağları

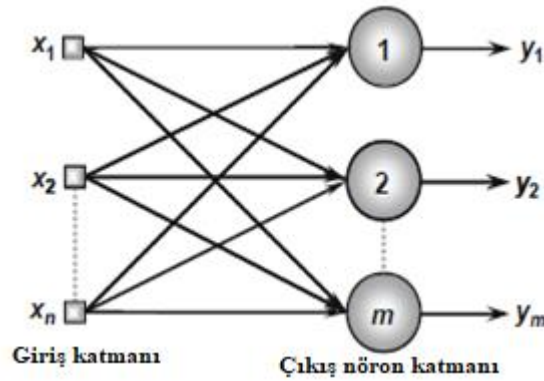
Öğrenme zamanına göre YSA'lar statik ve dinamik olmak üzere iki gruba ayrılır. Statik zamanlı ağ yapısında önceden eğitilen yapay sinir ağı, hesaplanan ağırlıklarında değişiklik olmadan mevcut haliyle daha sonraki durumlar için değişim göstermeden kullanılabilir. Dinamik zamanlı ağ yapısında ise ağı eğitimidenden elde edilen ağırlıklar, ağı tekrar kullanımında değişim gösterebilir (Çayiroğlu, 2015, s. 4). Bir ileri beslemeli ağı ağırlıkları eğitimden sonra sabittir, bu nedenle nöronun durumu, nöronun başlangıç ve geçmiş durumları tarafından değil, girdi/çıktı örüntüsü tarafından kesin olarak belirlenir, yani herhangi bir dinamik yoktur. Sonuç olarak, ileri beslemeli mimarinin türü statik olmayan sinir ağı olarak sınıflandırılır. Statik ileri beslemeli sinir ağının avantajı, ağı basit bir optimizasyon algoritması ile inşa edilebilmesidir ve günümüzde kullanılan en popüler sinir ağı mimarisidir. Bununla birlikte, statik-beslemeli sinirsel ağı da bazı uygulamalar için çeşitli dezavantajları vardır. İlk olarak, eğitim verilerinin boyutu yetersiz olduğu için tatmin edici bir çözüm üretilemeyebilir. İkinci olarak, statik-beslemeli sinir ağı hiç öğrenilmemiş büyük değişikliklerle iyi bir şekilde baş edemez eğitim aşaması. Son olarak, statik-beslemeli sinir ağı kolayca lokal bir en düşük değer olur ve veri setindeki girdi sayısı büyük olduğunda ağı yakınsama hızı çok yavaş olabilir (Chiang, L. C. Chang, & Chang, 2004). Tekrarlayan sinir ağı (Recurrent Neural Network; RNN), durağan sistemler için modelleme ve model genelleme kabiliyeti nedeniyle gerçek zamanlı kontrol, sistem tanımlama ve süreç koşulu

izleme alanlarında kullanılan bir tür dinamik sinir ağıdır (Yang & Ni, 2005). Statik sinir ağları akım girişleri ve akım çıkışları arasında doğrusal olmayan eşleme sağlar. Statik sinir ağlarında gecikmeli geri besleme veya gecikmeli giriş yoktur. Dinamik sinir ağları yapılarında gecikme geri bildirimi şeklinde veya giriş veya çıkışlarda belleğe sahiptir. Dinamik sinir ağı yapısı büyük ölçüde sistemin sırasına ve karmaşıklığına bağlıdır (Haykin, 1999, 36).

2.5.2.3. Katman Sayısına Göre Yapay Sinir Ağı Modelleri

2.5.2.3.1. Tek Katmanlı Ağlar

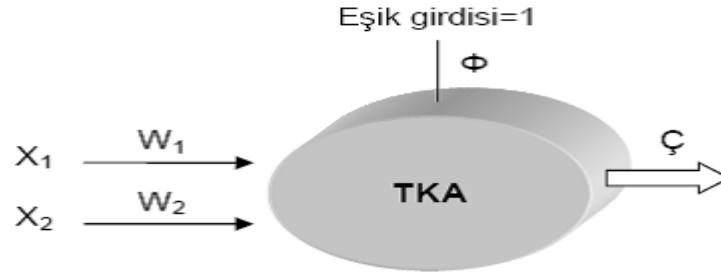
En basit haliyle tek katmanlı bir YSA modelinin her bir giriş düğümüne bağlı olan tek bir çıkış nöron tabakası vardır (Kantardzic, 2011, s.222). Başka bir ifadeyle tek katmanlı ağlar hem giriş hem de çıkış katmanı olan tek bir sinir katmanına sahiptir. Bu YSA modellerinde bilgiler her zaman, giriş katmanından çıkış katmanına olan tek bir yönde (tek yönlü) akar. Şekil 13'te, n giriş ve m çıkıştan oluşan basit bir ileri beslemeli ağı göstermektedir (Da Silva, vd. 2017, s.22).



Şekil 13. Basit ileri beslemeli ağ modeli.

Şekil 13 incelendiğinde bu mimariye ait ağlarda, ağ çıkışlarının sayısının her zaman nöron miktarıyla çakışacağını görmek mümkündür. Bu ağlar genellikle kalıp sınıflandırma ve doğrusal filtreleme problemlerinde kullanılır (Da Silva, vd. 2017, s.22). Yalnızca ikili sonuçlar üreten ve çıktı olarak 1 veya 0 değerini kullanan bu ağlar uyarıcı veya engelleyici

tipte yönlendirilmiş ağırlıksız kenarlardan oluşur. Sadece ikili bilgi üretilip iletilebildiğinden dolayı çok sınırlı görünen bu modeller aslında daha karmaşık modelleri uygulamak için gerekli tüm özelliklere sahiptir (Rojas, 1996, s.22). Tek katmanlı YSA modellerinin matematiksel formunu aşağıdaki biçimde ifade edebiliriz (Öztemel, 2016, s.59).



Şekil 14. Tek katmanlı YSA modellerinin matematiksel formu.

$X = (x_1, x_2, x_3, \dots, x_n)$ girdiler, $W = (w_1, w_2, w_3, \dots, w_n)$ girdilere karşılık gelen ağırlık değerleri ve ϕ eşik değer olmak üzere, çıktı katmanındaki değerler $\zeta = (\zeta_1, \zeta_2, \zeta_3, \dots, \zeta_n)$ aşağıdaki formülle hesaplanır ($i=1, 2, 3, \dots, n$).

$$\zeta = f\left(\sum_{i=1}^n w_i x_i + \phi\right)$$

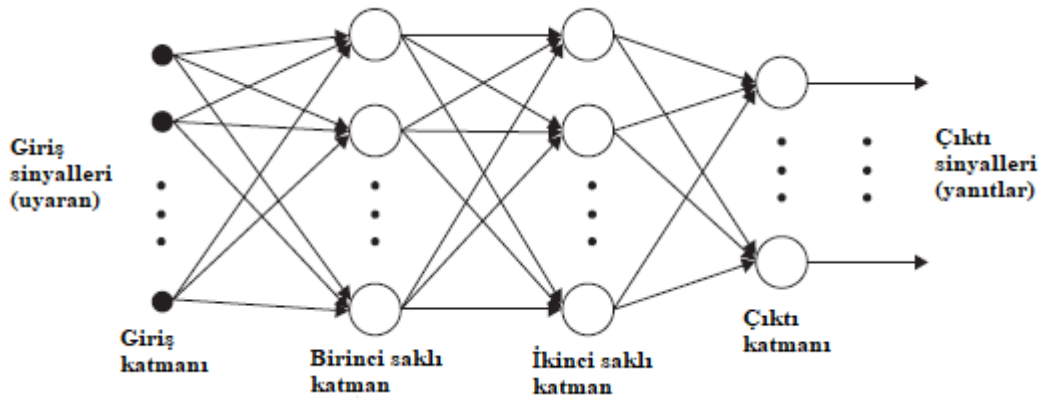
Formülden görüldüğü üzere çıktı değerleri, ağırlıklandırılmış girdi değerleri ile eşik değerinin toplamının aktivasyon fonksiyonundan geçirilmesi ile elde edilir. Formülde yer alan eşik değer, ağırlıklandırılmış girdi değerlerinin sıfıra eşit olmasını önlemek amacıyla daima 1 olarak alınır. Dolayısıyla ağırlıklandırılmış girdi değerlerinin 1 veya -1'dir. Bu durum ağırlıklandırılmış girdi değerlerinin 1 olması durumunda birinci sınıf, -1 olduğunda ise ikinci sınıf olacağı anlamını taşır.

$$f(g) = \begin{cases} 1, & \zeta > 0 \\ -1, & \zeta \leq 0 \end{cases}$$

2.5.2.3.2. Çok Katmanlı Ağlar

Marvin Minsky ve Seymour Papert'ın YSA ile ilgili yapılan çalışmaları durma noktasına getiren eleştirilerinin temeli, basit algılayıcı modellerin doğrusal olmayan ilişkilerde çözüm üretememesidir. Olaylar arasındaki doğrusallığın olamaması durumu, alanyazında XOR problemi olarak adlandırılır. Bu probleme göre çıktıların arasına çizilen doğrularla onları iki

veya daha fazla sınıfa ayırmak mümkün değildir. Rumhelt ve arkadaşları bu problemi çözmek amacıyla, girdi katmanı, ara katman ve çıktı katmanından meydana gelen çok katmanlı ağı geliştirmiştir (Öztemel, 2016, s.76). Bu ağlarda amaç ara katmanları ayrı ayrı eğiterek, çıkış katmanını elle hesaplamak zorunda kalmadan çok katmanlı bir ağ oluşturmaktır (Gurney, 2014, s.41). Şekil 15'te, n giriş ve m çıkıştan oluşan basit bir ileri beslemeli ağ göstermektedir (Kantardzic, 2011, s.214).



Şekil 15. Basit ileri beslemeli ağ modeli.

Şekil 15, YSA sürecindeki işlemler için iki ara katmana ve bir çıkış katmanına sahip çok katmanlı bir algılayıcının mimari grafiğini göstermektedir. Burada gösterilen ağ tamamen bağlı. Bu, ağın herhangi bir katmanındaki nöronun, önceki katmandaki tüm düğümlere (nöronlara) bağlandığı anlamına gelir. Ağ üzerinden veri akışı, soldan sağa ve katman katman bazında ileri yönde ilerler (Kantardzic, 2011, s.214). *Giriş katmanı*, gelen girdilerin hiçbir işleme tabii tutulmadan ara katmanlara gönderildiği yerdir. Bu katmanda bilgi işleme olmadığı için her bir girdinin sadece bir çıktısı olup, bu çıktılar bir sonraki katmanda yer alan bütün elemanlara iletilir. *Ara katman*, girdi katmanından gelen çıktıların işlenerek bir sonraki katmana veya doğrudan çıkışa gönderildiği yerdir. Her eleman bir sonraki aşamada yer alan bütün elemanlara bağlıdır. *Çıktı katmanı*, ara katmandan gelen bilgilerin işlenerek ağın ürettiği sonuçların yer aldığı katmandır. Her çıktı bir önceki katmanda yer alan bütün elemanlarla bağlantılıdır.

Çok katmanlı öğrenme makineleri çeyrek yüzyıldan fazla bir süredir bilinmesine rağmen, uygun eğitim algoritmalarının olmaması pratik görevler için başarılı uygulamalarını

engellemiştir. Denetimli eğitim algoritmalarındaki son gelişmeler, her tabakanın doğrusal tanımlayıcı işlevini kullanan çok katmanlı ağların yaygın ve başarılı mevcut kullanımlarına yol açmıştır. Çok katmanlı ağlara genellikle katmanlı ağlar denir. Bu ağlar örnekleri sınıflarını ayıran karmaşık girdi/çıktı eşlemelerine veya karar yüzeylerine uygulanabilirler (Zurada, 1992, s.164). Çok katmanlı YSA'ların üç ayırt edici özelliği vardır (Karakış, 2009, s.17):

- ✓ Ağ yapısında birden fazla sayıda katman vardır.
- ✓ Ağda doğrusal olmayan bir aktivasyon fonksiyonu kullanılır. Genellikle lojistik fonksiyonu kullanılır.
- ✓ Nöronlar arası bağlantılar yüksek derecelidir.

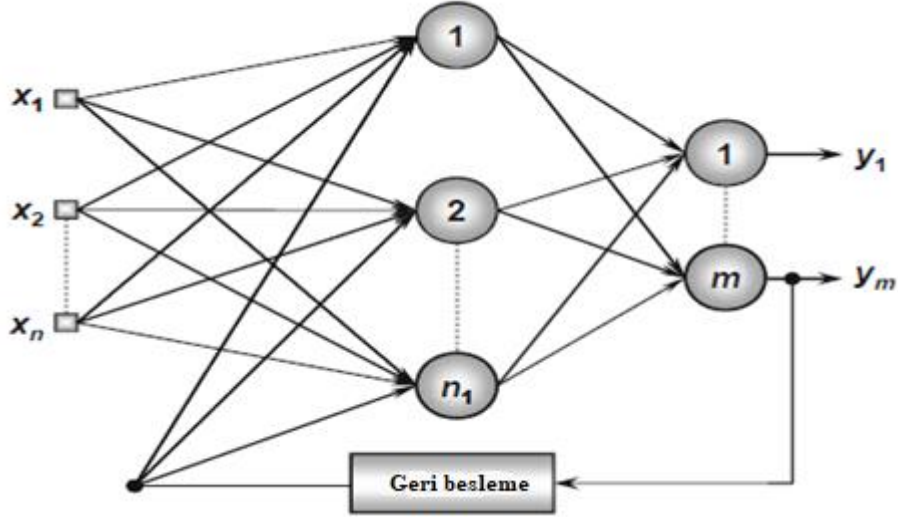
2.5.2.4. Topolojik Yapısına Göre Yapay Sinir Ağı modelleri

YSA, belirli bir topolojik yapıya göre bazı işlem nöronlarından ve genel olarak zamanla değişmeyen nöronlardan oluşan bir ağ modelidir. Geleneksel bir YSA bağlantı biçimine, nöronlar arasında bilgi aktarımının yönüne ve geri besleme durumuna göre ileri ve geri beslemeli olarak ikiye ayrılır (He & Xu, 2010, s.38).

2.5.2.4.1. Geri Beslemeli (Tekrarlayan) Ağlar

Geri beslemeli bir sinir ağında, giriş sinyali belirli bir başlangıç durumundan nöronlar arasında tekrar tekrar aktarılır ve birkaç kez dönüştürüldükten sonra, yavaş yavaş belirli bir sabit duruma geçiş gösterir. Geri besleme işlemi sinir ağı sadece bilgi geri bildirimli bir süreç sinir ağı modelidir ve tüm nöron düğümleri sistemin bilgi akış yönüne göre bağlanır. Bilgiler, belirli kurallarla önceki katmanlardaki düğümlere geri aktarılabilir ve çıktı bilgileri düğümlerin kendilerine geri beslenebilir (He & Xu, 2010, s.24). Bu ağlarda, nöronların çıkışları diğer nöronlar için geri besleme girdileri olarak kullanılır. Geri bildirim özelliği bu ağları dinamik bilgi işleme için nitelendirir, yani zaman serisi tahmini, sistem tanımlaması ve optimizasyonu, süreç kontrolü vb. gibi zaman değişkenli sistemlerde kullanılabilir. Geri

beslemeli ağların işleyişine yönelik bir örnek Şekil 16'da yer almaktadır (Da Silva, vd. 2017, s.25).



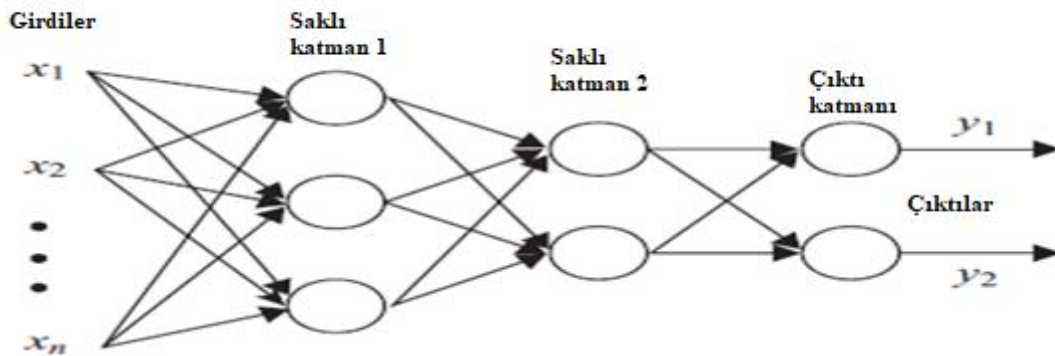
Şekil 16. Geri beslemeli ağların işleyişi modeli.

Tekrarlayan bağlantıları olan sinir ağı modelleri karmaşık zamana bağlı hesaplamalar yapabilir. Bu modellerin çoğu öğrenmeyi içermez, ancak belirli bir işlevi yerine getirmek için sabit ağırlıklar kullanır (Architecture Technology Corporation, 1991, s.31). Geri besleme işlemi sinir ağı bilgi aktardığında, ileri beslemeli sinir ağına olduğu gibi ileri akışlar ve ikinci katman düğümlerinden önceki katman düğümlerine olan zaman gecikmeli geri bildirim bilgileri vardır. Bu nedenle, geri besleme işlemi sinir ağı modeli, işlevsel bir tahminci olarak ve aynı zamanda ilişkilendirilebilir bir bellek makinesi olarak veya akıllı bir gerçek zamanlı denetleyici olarak kullanılabilir. Bu nedenle, geribildirim süreci sinir ağı modelinin geniş uygulama olanaklarına sahip olması beklenmektedir. (He & Xu, 2010, s.125). Tekrarlayan bir sinir ağının (RNN) geri bildirim yolları vardır, bu da onu 'birleşik' ağdan ziyade 'sıralı' bir ağ haline getirir ve geçici davranış sergilemesine izin verir. Nöronlar arasındaki tüm olası bağlantılara izin verilir ve döngüsel bağlantılar dâhil edildiğinde, bir sinir ağı doğrusal olmayan dinamik bir sistem haline gelir. Böyle bir sistem çok zengin zamansal ve mekânsal davranışa sahiptir. Çevrimsel bağlantıları olan RNN'leri analiz etmek ve tanımlamak, doğrusal olmayan dinamik sistemler için sınırlı matematiksel araçların zorluklarını yansıtan İleri Beslemeli Yapay Sinir Ağı'ndan (Feed Forward Network; FFN)

çok daha zordur. (Khare & Nagendra, 2006, s.32). RNN’de bazı teknik detayların gerekliliğini vurgular; geribildirimlerin veya harici girdilerin düğümlere gönderilmesine izin verecek bazı mekanizmalar olmalı ve geri bildirim sırasında yeni ağ çıktılarının oluşturulmasını sağlamak zorundayız (Gurney, 2014, s.68). İleri beslemeli sinir ağının aksine, RNN’ler zaman içinde dağıtılan gizli durumları içerir. Bu, geçmiş hakkında birçok bilgiyi verimli bir şekilde depolamalarını sağlar. RNN, daha önceki bir zaman adımında gizli birimler tarafından hesaplananlar hakkında bilgi yakalayan bir “belleğe” sahiptir (Lewis, 2017, s.64-67). RNN’ler karmaşık tahmin görevlerinde iyi değildir (Mehrotra, Mohan, & Ranka, 1997, s.139). Nedenselliğin devam edebilmesi için, bağlantı grafiğinin her döngüsünün sıfır olmayan gecikmeli en az bir bağlantısı olmalıdır (Dreyfus, 2005, s.8)

2.5.2.4.2. İleri Beslemeli Ağlar

İlk geliştirilen YSA modellerinden biri olan ileri beslemeli ağlar, regresyon ve sınıflandırma problemlerinin her ikisinde de çözüm üretebildiği için en yaygın kullanılan modeldir (Bataineh, 2012, s.25). Tek ve çok katmanlı YSA’nın her ikisinde de uygulanabilir (Özen, 2018, s.27). İleri beslemeli ağlarda işlem, herhangi bir döngü veya geri besleme olmadan giriş tarafından çıkış tarafında doğru ilerleme gösterir. Bu ağlarda aynı katmandaki düğümler (nodes) arasında bağlantı yoktur ve belirli bir katmandaki düğümlerin çıkışı, bir sonraki katmana girdi olarak bağlanır. İleri beslemeli ağların işleyişine yönelik bir örnek Şekil 17’de yer almaktadır. (Kantardzic, 2011, s.205).



Şekil 17. İleri beslemeli ağların işleyiş modeli.

İleri beslemeli ağlar herhangi bir girdi için tahmini genelleştirebilir. Başka bir ifadeyle, ağ eğitildikten sonra, eğitim verileri ve eğitim dışındaki veriler de dâhil olmak üzere herhangi bir yeni girdiyi tahmin edebilir. Birçok uygulamada özellikle zaman serisi verilerinin (farklı zaman ve değerlerde gelen veriler) eğri uyumunu sağlamada iyi performans gösterir. Ancak ileri beslemeli ağların gizli katmanında genellikle diğer YSA modellerinden daha fazla nöron bulunur. Bu nedenle yerel optimizasyon çözümünün ileri beslemeli ağlarda meydana gelme olasılığı daha yüksektir. Dolayısıyla optimizasyondan gelen yerel minimal çözüm nedeniyle yanlış sonuçlar üretilebilir. Ayrıca, optimizde edilecek daha fazla nöronun olması eğitim sürecinde bellek/hafıza sorunlarının oluşmasına ve eğitim süresinin uzamasına neden olur. Ancak ileri beslemeli ağların gizli katmanında genellikle diğer YSA modellerinden daha fazla nöron bulunur. Bu nedenle ağın daha fazla zamana gereksinim duyması ve yerel optimizasyon çözümünün ileri beslemeli ağlarda meydana gelme olasılığı daha yüksektir. Dolayısıyla optimizasyondan gelen yerel minimal çözüm nedeniyle yanlış sonuçlar üretilebilir (Bataineh, 2012, s.26).

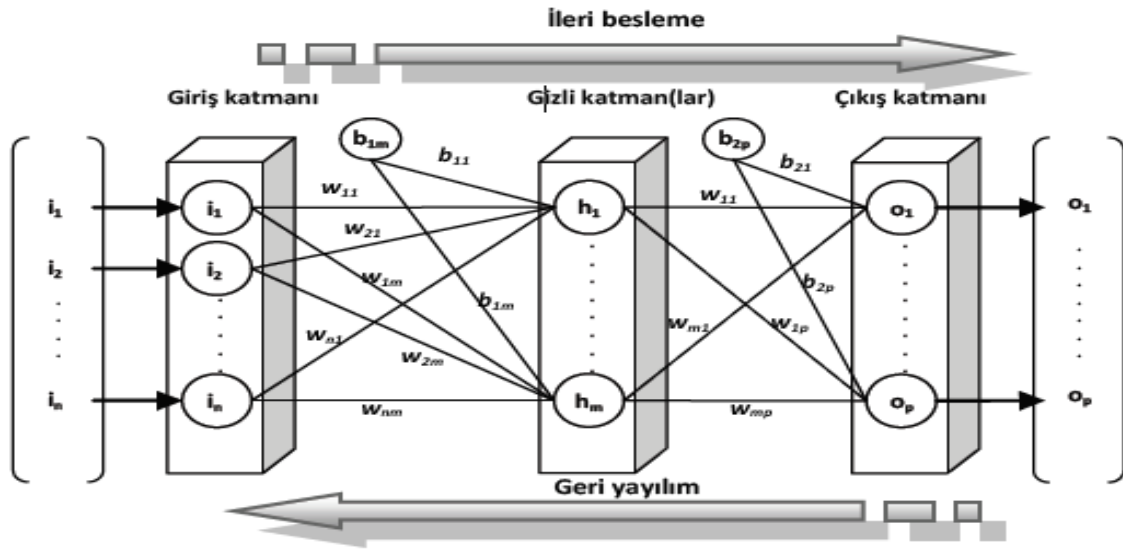
2.5.2.5. Çok Katmanlı Algılayıcı (Multilayer Perception; MLP)

MLP'ler, ağın hata geri yayılım algoritması olarak bilinen oldukça popüler bir algoritma ile denetimli bir şekilde eğitilmesiyle bazı farklı ve çeşitli sorunları çözmek için başarıyla uygulanmıştır. Bu algoritma geri yayılım ağı (Backpropagation network) kuralına dayanır ve genelleştirmesi olarak görülebilir (Kantardzic, 2011, s.214). Bu kural ilk olarak Werbos tarafından tanımlanmış daha sonra Parker, Rumelhart ve McClelland tarafından geliştirilmiştir. Sınıflandırmada kullanılan etkin YSA modellerinden biri olup çok katmanlı, ileri beslemeli ve denetimli bir ağ modelidir (Akpınar, 2017, s.323). Temel olarak, geri yayılım öğrenimi ağın farklı katmanları üzerinden gerçekleştirilen iki aşamadan oluşur: ileri geçiş ve geri geçiş. İleri geçişte, ağın giriş düğümlerine bir eğitim örneği (giriş verileri vektörü) uygulanır ve etkisi ağ katmanı boyunca katmana yayılır. Son olarak, ağın gerçek yanıtı olarak bir dizi çıkış üretilir. İleri aşama sırasında, ağın sinaptik ağırlıkları sabitlenir (Kantardzic, 2011, s.214). İleri doğru hesaplama, öncelikle eğitim setinde yer alan örnekler

($G_1, G_2, G_3, \dots, G_k$) girdi katmanına dâhil edilir. Girdi katmanında herhangi bir işlem olmadan burada yer alan girdiler doğrudan ara katmana gönderilir. Buradaki k. elemanın çıktısı $\zeta_k^i = G_k$ olmak üzere, girdi katmanındaki elemanlardan gelen bilgiler ($A_1, A_2, A_3, \dots, A_k$) bağlantı ağırlığı ile ara katmanda yer alan her elemana iletilir. Ara katmanda yer alan her eleman için gelen net girdi; A_{kj} girdi katmanında yer alan k. elemanını j. ara katmanına bağlayan bağlantı ağırlığı olmak üzere ve $NET_j^a = \sum_{k=1}^n A_{kj} \zeta_k^i$ formülü ile hesaplanır. Ardından elde edilen net girdi değerine aktivasyon fonksiyonu uygulanarak j. ara katman elemanının çıktısını hesaplanır. Aktivasyon fonksiyonu olarak türevi alınabilecek bir fonksiyonun kullanılması yeterlidir ancak genellikle sigmoid fonksiyonu kullanılır. Sigmoid fonksiyonunu kullanılması durumunda, β_j^a ara katmandaki j. elemana bağlanan eşik değerin ağırlığı olmak üzere j. ara katman elemanının çıktısı $\zeta_j^a = \frac{1}{1+e^{-(NET_j^a + \beta_j^a)}}$ olarak elde edilir. Eşik değerlerin çıktısı sabittir ve 1'e eşittir. Ağırlık değeri ise eğitim esnasında ağ tarafından belirlenir. Ara katmanlarda ve çıktılarında yer alan bütün işlem elemanları için NET girdilerin hesaplanması ve sigmoid fonksiyonunun uygulanmasından sonra ($\zeta_1, \zeta_2, \zeta_3, \dots$) çıktı katmanı değerleri elde edilir (Öztemel, 2016, s.78). İleri beslemeli çok katmanlı bir ağı matematiksel açılımını aşağıdaki biçimde verebiliriz (Krenker, Beşter, & Kos, 2011).

Geriye doğru aşamada ise ağırlıkların tümü bir hata düzeltme kuralına göre ayarlanır. Özellikle ağın gerçek yanıtı, bir hata sinyali üretmek için eğitim örneğinin bir parçası olan arzu edilen (hedef) bir yanıtta çıkarılır. Bu hata sinyali daha sonra sinaptik bağlantıların yönüne karşı ağ üzerinden geriye doğru yayılır (Kantardzic, 2011, s.214). Geriye doğru hesaplama, Sinaptik ağırlıklar, ağın gerçek tepkisini istenen yanıtta daha yakın yapacak şekilde ayarlanır. Ağa sunulan ($G_1, G_2, G_3, \dots, G_k$) girdi değerlerine uygulanan fonksiyonlar sonucunda elde edilen ağın ürettiği çıktı değerleri ($\zeta_1, \zeta_2, \zeta_3, \dots$) ile ağın beklenen çıktıları ($B_1, B_2, B_3, \dots$) ile karşılaştırılır. Beklenen çıktılar ile ağın ürettiği çıktılar arasındaki bu farka hata olarak ele alınır. Üretilen ağ ile bu hatanın mümkün olabildiğince küçük elde edilmesi amaçlandığından dolayı her bir sonraki iterasyonda bu hata ağın ağırlıklarına dağıtılarak azaltılır. Bu durumda m. eleman için elde edilen hata (E_m) olmak üzere, $E_m = B_m - \zeta_m$ 'dir. Sonda elde edilen çıktı katmanındaki toplam hata (TH) bütün hataların yer aldığı $TH =$

$\frac{1}{2} \sum_m E_m^2$ formülü ile elde edilir. S onda hesaplanan çıktı katmanındaki hataların bazıları negatif olabileceğinden dolayı toplam hatanın sıfır olmasının önüne geçmek için, toplam hata hesabı yapılırken hataların karesine dayalı bir işlem yapılarak son adımda karekök alınır. Ağıın eğitilmesinin temel amacı bu toplam hatayı sıfır yapmaktır. Bu amaçla işlemde yer alan elemanların ağırlıkları değiştirilerek hata işlemde yer alan elemanlara dağıtılır (Öztemel, 2016, s.79). Çok katmanlı ileri beslemeli geri yayımlı bir YSA modeli Şekil 18’de yer almaktadır (Aydoğan, Ayat, Öztürk, Çevik & Yüksel, 2010).



Şekil 18. Çok katmanlı ileri beslemeli geri yayımlı bir YSA modeli.

Bu YSA modelleri ileri beslemeli ağılar ile geri yayımlı öğrenme algoritmalarının birleşiminden meydana gelen, doğrusal olmayan bir aktivasyon fonksiyonuna sahip ve yeteri kadar nöron tanımlandığında süreksizlik içeren her türlü problem çözebilme yeteneğine sahiptir (Elge, 2018, s.30).

2.6. Destek Vektör Makinesi

Destek vektör makinesi, istatistiksel öğrenme teorisi bağlamında geliştirilen sınıflandırma ve regresyon amaçlı kullanılan bir yöntemdir (Vercellis, 2009, s.262). DVM, optimizasyon teorisinden faydalanılarak bir öğrenme algoritması ile eğitilmiş, yüksek boyutlu bir özellik uzayında lineer fonksiyonların hipotez alanını uygulayan, istatistiksel öğrenme teorisine

dayalı bir öğrenme sistemidir (Cristianini & Shawe-Taylor, 2000, s.7). Temelleri Vapnik-Chervonenkis Teorisine dayanan DVM, Boser, Guyon ve Vapnik tarafından 1992 tarafından geliştirilmiştir (Akpınar, 2017, s.291). DVM ortaya çıkışından sonra çeşitli uygulamalarda diğer sistemlerin çoğundan daha iyi performans gösteren ilkeli ve çok güçlü bir yöntemdir (Cristianini & Shawe-Taylor, 2000, s.7). Başlangıçta ikili bir sınıflandırıcı olarak geliştirilen DVM daha sonra çok kategorili sınıflandırma, regresyon, kümeleme ve değişken seçimine yönelik olarak kullanılmıştır. DVM'lerin görüntü ve metin sınıflandırması, karakter tanıma, protein sınıflandırması ve kanser tespiti için son derece yararlı olduğu kanıtlanmış ve son zamanlarda az da olsa sosyal bilimler alanında da kullanılmaya başlanmıştır (Attewell vd. 2015, s.147).

DVM'lerin sınıflandırma stratejisi tamamen istatistiksel bir ölçütten ziyade marj tabanlı geometrik bir ölçüt uygulamasına bağlıdır. Başka bir deyişle, DVM'leri sınıflandırma görevini yerine getirirken sınıfların istatistiksel dağılımlarının tahmin edilmesine ihtiyaç duymaz bunun yerine marj maksimizasyonu kavramını kullanarak sınıflandırma modelini tanımlarlar (Melgani & Bruzzone, 2004). DVM çoğunlukla sayısal verilerin ikili sınıflandırması için tasarlanmıştır. Ancak çeşitli yöntemler kullanılarak çok sınıflı duruma genelleştirilebileceği gibi bağımlı değişkenin kategorisinin binarizasyon yaklaşımıyla ikili verilere dönüştürülmesiyle de ele alınabilir (Aggarwal, 2015, s.313). DVM bir dizi etiketli eğitim verisinden girdi-çıktı haritalama (mapping) fonksiyonları üreten danışmanlı öğrenme yöntemleridir. Haritalama fonksiyonu bir sınıflandırma fonksiyonu veya regresyon fonksiyonu olabilir. Giriş verileri çoğunlukla doğası gereği karmaşık ve doğrusal olmayan ilişkilere sahiptir. Bu nedenle sınıflandırma işleminin yapılabilmesi amacıyla giriş verilerini daha ayrılabilir hale getirmek için yani doğrusallığı sağlamak için çoğunlukla Kernel fonksiyonlarını kullanır. Daha sonra, maksimum marjlı hiperdüzlemleri eğitim verilerindeki sınıfları en uygun şekilde ayırmak için yapılandırılır. İki paralel hiper düzlem arasında mesafeyi en üst düzeye çıkararak verileri ayıran hiper düzlemin her iki tarafında iki paralel hiperplan inşa edilmiştir. Bu paralel hiperplanlar arasındaki marj veya mesafe ne kadar büyük olursa, sınıflandırıcının genelleme hatasının o kadar iyi olacağı varsayılır. (Olson &

Delen, 2008, s.111). Tüm lineer modellerde olduğu gibi, SVM'ler iki sınıf arasındaki karar sınırı olarak ayırıcı hiperdüzlemleri kullanır. SVM'ler söz konusu olduğunda, bu hiperdüzlemleri belirleme optimizasyon problemi ve marj kavramı ile belirlenir. Öng örüsel olarak, bir maksimum marjin hiperdüzlem, iki sınıfı açık bir biçimde ayıran, her iki tarafa ait sınırdaki geniş bir bölge bulunan ve içinde eğitim veri noktası olmayan bir bölgedir. Doğrusal olarak ayrılabilir verilerde, iki sınıfa ait veri noktalarını açık bir biçimde ayıran doğrusal bir hiper düzlem oluşturmak mümkündür. Ancak bu durum uygulamada her zaman mümkün olmayabilir. Çünkü gerçek verilerde tamamen keskin bir ayrılmanın mümkün olmadığı gibi yanlış sınıflanmış, eksik ve uç değerlerin olmasından kaynaklı olarak doğrusal ayrılabilirliğin sağlanması zordur (Aggarwal, 2015, s.313).

DVM yapısal algoritmasında maksimal marj kullanılmasından dolayı öğrenme algoritmasının bir çıktısı olarak genelleme üzerindeki sınırı tanımlayabiliriz. Eğitim verilerinin küçük bir alt kümesi aracılığıyla çözüm üretebilmesinden dolayı çok büyük hesaplama kolaylığına sahiptir. DVM'daki marj maksimizasyonu hedef fonksiyon ve dağılım arasındaki iyi derecede ilişkinin sağladığı avantajdan faydalanarak yüksek boyutlu drumlarda çıkan zorlukların çözümünü kolaylaştırır (Cristianini & Shawe-Taylor, 2000, s.99-112-152). Doğrusal sınıflandırıcılar için tasarlanan yöntemleri kullanarak doğrusal olmayan karar sınırları oluşturma ve çekirdek fonksiyonlarını kullanarak belirli bir sabit boyutlu vektör alanı ile temsil edilemeyen verileri sınıflama becerilerine sahiptir (Ben-Hur & Weston, 2010). Çeşitli uygulama alanlarındaki diğer sınıflandırıcılara göre doğruluk açısından daha iyi performans elde ettikleri ve büyük problemler için verimli bir şekilde ölçeklendirilebilecekleri gösterilmiştir. Diğer bir önemli özelliği ise üretilen sınıflandırma kurallarının yorumlanması ile ilgilidir. Destek vektör makineleri, her hedef sınıf için en temsili gözlemler gibi görünen destek vektörleri adı verilen bir dizi örneği tanımlar. Bir bakıma, sınıflayıcı tarafından nitelik alanında üretilen ayırma yüzeyinin konumunu tanımladıkları için diğer örneklerden daha kritik bir rol oynarlar (Vercellis, 2009, s.262). Öte yandan yüksek güvenilirliğe, ezber öğrenme yeteneğine, lineer olmayan sınıflandırmada yüksek başarıya sahip bir yöntemdir. Ancak eğitim süresi uzundur (Akpınar, 2017, s.291).

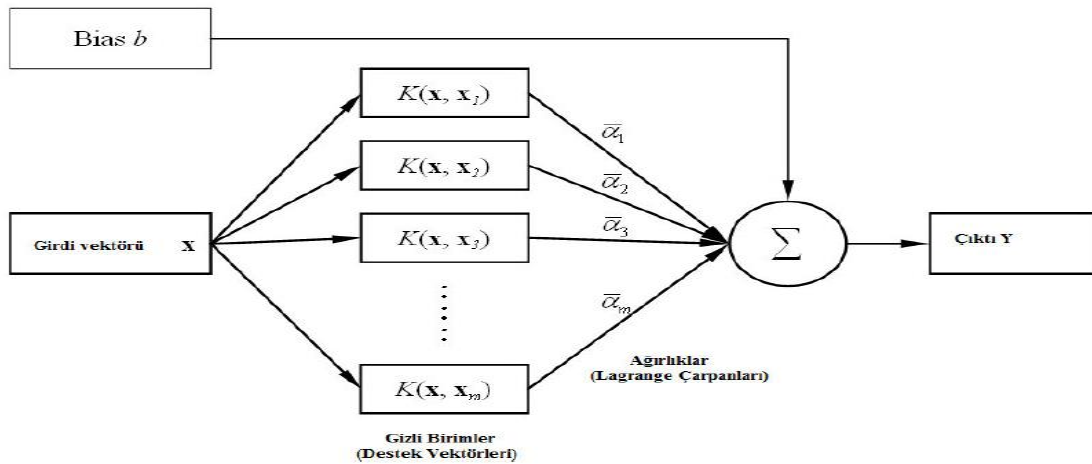
SVM algoritmasının, veri sınıflandırmada kullanılan diğer algoritmaların çoğundan daha iyi performans göstermesini sağlayan önemli faktörler den biri, sınıflara ait sınırın oluşturulması esnasında sınıra çok yakın veya sınırda bulunan veriler hariç eğitim verilerini gereksiz olarak kabul etmesi veya algoritmanın çalışmasını tamamlamak için bunlara güvenmemesidir (Wang, 2005, s.101). Çok boyutlu doğrusal olmayan gelişmiş gösterimler kernel fonksiyonları sayesinde yeni özellikleri tanımlanmadan ve hesaplamadan doğrusal modeller için kullanılabilir. Yeni özelliklerin tanımlanmaması alandan bağımsız geliştirilmiş temsillerin belirtilmesini ve bunları ürün tabanlı doğrusal modelleme algoritmalarına dâhil etmeyi kolaylaştırırken, hesaplanmaması dolaylı olarak kullanılan çok sayıda yeni özelliğin değerlerini hesaplama ihtiyacından kaçınarak, doğrusal model oluşturma ve tahmin için bu tür gösterimlerin kullanılmasını hesaplama açısından verimli hale getirir (Cichosz, 2014, s.484). DVM'de verilerin özellik alanına eşlenmesinden dolayı olarak overfitting artar. Çünkü çok boyutun olması durumunda daha fazla esnekliğin olmasından kaynaklı olarak verilerin uyum gösterebileceği daha fazla sayıda model vardır. Ancak DVM bir düzenleme terimi kullanarak overfitting sorununu çözer (Campbell & Ying, 2011, s.34). DVM'nin dezavantajı iki boyutlu sınıf problemleri ve sonuçları için tasarlanmış olması ve sonraki sınıf üyelik olasılıklarını doğal bir biçimde modelleyememesidir (Canty, 2014, s.228). DVM'ne ait bir sınırlama çekirdek fonksiyon seçimidir. İkinci bir sınırlama hem eğitim hem de testte hız ve boyuttur. Karmaşık ve zaman gerektiren hesaplamalar içerir. Pratik açıdan bakıldığında, SVM'lerle ilgili belki de en ciddi sorun, büyük ölçekli görevlerde yüksek algoritmik karmaşıklık ve geniş bellek gereksinimleridir. Diğer bir sınır ise kategorik verilerin işlenmesidir. Bu sınırlamalara rağmen, SVM'ler sağlam teorik temele dayandığından ve ürettiği çözümler sayesinde günümüzde veri madenciliğinde kullanılan en popüler yöntemlerden biridir (Olson & Delen, 2008, s.122).

DVM'nin algoritma adımlarını aşağıdaki biçimde belirtilmiştir (Hazım, 2018, s.43):

1. Veri setini okutun.
2. İşlem sınıfı, işlem zamanı ve ardışık işlemlerdeki farklılıklara göre verileri yeniden sıralayın.

3. Her işlem iki alanın vektörü olan bir veri formuna dönüştürülür.
4. Pozitif ve negatif işlem grupları olarak adlandırılan iki ayrı veri kümesi oluşturulur.
5. Çekirdek fonksiyonu seçilir.
6. DVM eğitilir.
7. Uygulama sonrası sınıflandırma performansı kaydedilir.
8. Mevcut işlem değerlendirilir.
9. Uygulama geçerli işlem verileri için 1'den 3'e kadar olan adımlarda yeniden başlatılır.
10. Kaydedilmiş sınıflandırıcı ve şu anda sınıflandırıcı tarafından üretilen vektör yer değiştirilir.
11. Üretilen karar kabul edilir.

DVM sınıflandırma işlemi yaparken, veri setini iki gruba ayıran ve bağımlı değişkene göre belirlenen karar doğrusu tarafından tanımlanmış destek noktaları arasındaki uzaklığı maksimize etmeyi amaçlar (Akpınar, 2017, s.291). Bu amaç doğrultusunda işlem sürecini gerçekleştiren DVM'nın işleyişi Şekil 19'da gösterilmektedir.



Şekil 19. DVM'nın işleyiş modeli.

Şekil 19'da (Ayhan, 2013, s.55) yer alan $K(\mathbf{x}_i, \mathbf{x}_j)$, girdilerin içi çarpımlarının hesaplanmasında kullanılan kernel fonksiyonlarını, $\bar{\alpha}_m$ 'ler langrange çarpanlarını yani ağırlıkları

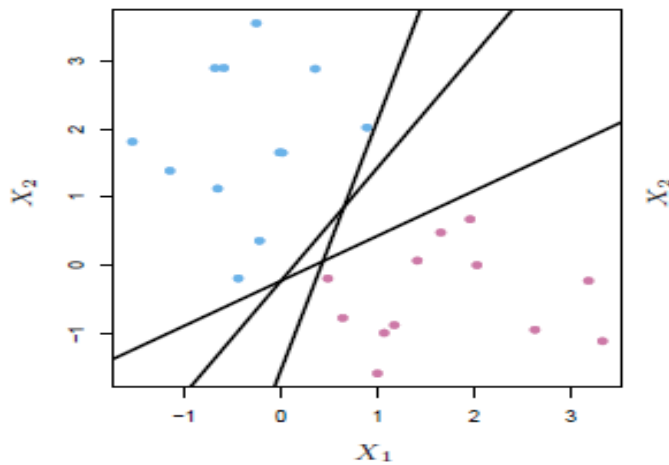
ve Y çıktı değerleri ise kernel fonksiyonu ile elde edilen iç çarpımlar ile ağırlık değerlerinin kombinasyonlarından elde edilen toplamaları ifade etmektedir (Ayhan, 2013, s.56).

2.6.1. Destek Vektör Makinelerinde Sınıflandırma

DVM sınıflandırma işlemi sert marjin lineer sınıflandırma (tamamıyla doğrusal ayrılabilen), soft marjin lineer sınıflandırma ve lineer olmayan sınıflandırma olmak üzere üç gruba ayrılır (Abe, 2005, s.28).

2.6.1.1. Sert Marjin Lineer Sınıflandırma (Hard Margin)

DVM'nin en basit ve ilk hali olan bu modelin gerçek yaşamdan elde edilen veriler için sağlanması zordur (Karakaynak, 2014, s.14). Bu modelin temel amacı, tanımlanmış iki sınıfa birbirinden ayıran fonksiyondan faydalanılarak farklı sınıflarda yer alan ancak birbirine en yakın olan noktalar arasındaki uzaklığı maksimum hale getirecek optimum hiperdüzlemi bulmaktır. Bulunan hiperdüzleme göre sınıflara ait sınıfları belirleyen noktalar üzerindeki doğrulara destek vektör adı verilir. (Akgün, 2016, s.22). Veri setini birbirinden ayıran üç olası hiperdüzleme ait gösterim Şekil 20'de yer almaktadır (James, Witten, Hastie & Tibshirani, 2013, s.345).



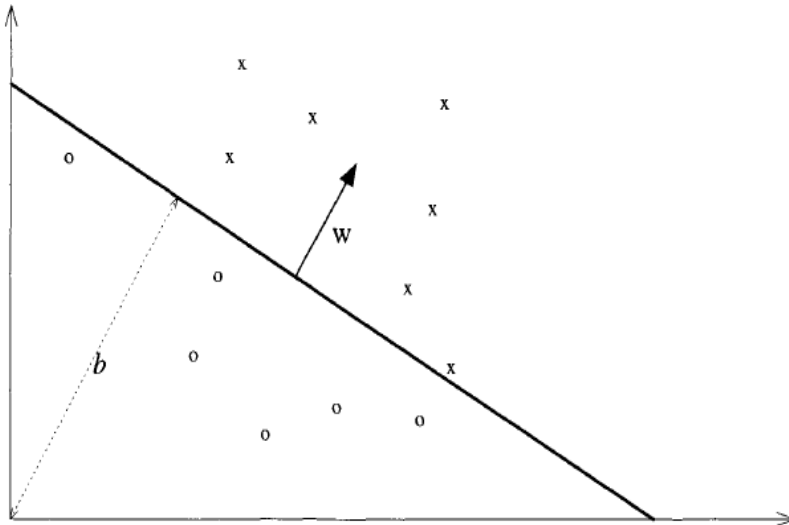
Şekil 20. Veri setini birbirinden ayıran üç olası hiperdüzleme ait gösterimi.

Şekil 20'ye göre, veri setini ayıran bir hiperdüzlem herhangi bir noktaya temas etmeden az bir miktar yukarı veya aşağı hareket ettirilebilir veya döndürülebilir. Bu nedenle verileri hatasız biçimde ayıran bir hiperdüzlem olduğu zaman verileri ayırabilecek sonsuz sayıda hiperdüzlem vardır (James vd., 2013, s.345). İkili sınıflandırma çoğunlukla $f: X \subseteq \mathbb{R}^n \rightarrow \mathbb{R}$ biçiminde gerçek değerli bir fonksiyonu kullanılarak gerçekleştirilir. Bu fonksiyona göre $X = (x_1, \dots, x_n)$ giriş değerleri $f(x) > 0$ ise pozitif sınıfa $f(x) < 0$ ise negatif sınıfa atanır. $f(x)$, $x \subseteq X$ olmak üzere lineer bir fonksiyon olarak ele alınırsa;

$$f(x) = \langle w, x \rangle + b$$

$$f(x) = \langle w, x \rangle + b = \sum_{i=1}^n w_i x_i + b$$

biçiminde elde edilebilir. İki boyutlu bir hiperdüzlem Şekil 21'de yer almaktadır.



Şekil 21. İki boyutlu bir hiperdüzlem modeli.

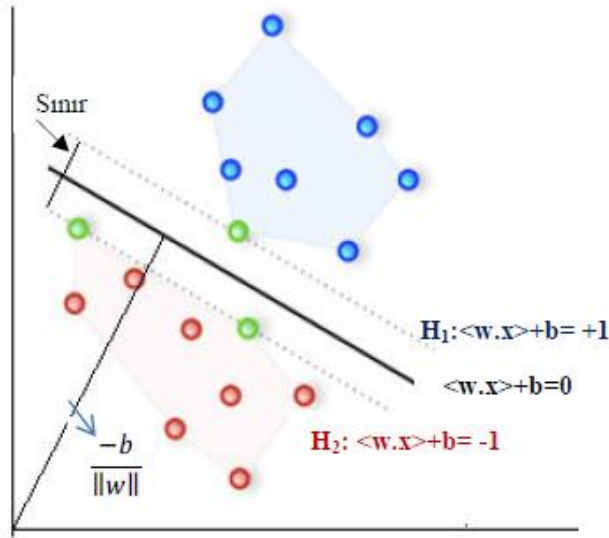
Burada $(w, b) \in \mathbb{R}^n * \mathbb{R}$ fonksiyonu kontrol eden parametrelerdir ve karar kuralı $sgn(f(x))$ tarafından elde edilir. Şekil X'te $sgn(0) = 1$ olacak biçimde bir kural kullanılmıştır ancak öğrenme yönteminde bu kuralın verilerden elde edilen parametrelerden öğrenilmesi gerekir. X'ler girişleri ve Y'ler çıktıları temsil etmek üzere, genellikle ikili sınıflama için $Y = \{-1, 1\}$,

m-sınıflı sınıflama için $Y=\{1,2,\dots,m\}$ ve $X \subseteq \mathbb{R}^n$ 'dir. Eğitim seti, eğitim verileri olarak da adlandırılan eğitim örneklerinin bir koleksiyonudur ve aşağıdaki biçimde ifade edilebilir:

$$S = ((x_1, y_1), \dots, (x_\ell, y_\ell)) \subseteq (X \times Y)^\ell$$

Burada ℓ örneklem sayısı, x örnek ve y sınıftır (Cristianini & Shawe-Taylor, 2000, s.9-11).

“Eğitim veri seti doğrusal olarak ayrılabilirse en büyük sınıra sahip ayırma hiperdüzlemini bulmaya çalışır” (Ayhan, 2013, s.58). Eğitim verilerini hatasız olarak birbirinden ayırabilen doğrusal fonksiyon bir ayırıcı hiperdüzlemdir. Bu hiperdüzlem bulunurken n örneklemden oluşan bir eğitim verisi için, $(x_1, y_1), \dots, (x_n, y_n)$, $x \in \mathbb{R}^d$, $y \in \{-1, +1\}$ olmak üzere, w ve b uygun katsayıları için hiperdüzlem karar fonksiyonu, $D(x)=(w.x)+b$ biçiminde tanımlanabilir (Cherkassky & Mulier, 2007, s.418). İki sınıf için Lineer sınıflandırma modeli Şekil 22’de yer almaktadır (Ayhan, 2013, s.60).



Şekil 22. Lineer sınıflandırma modeli.

Şekil 22’de yer alan H_0 , H_1 ve H_2 düzlemleri, $H_0 = \langle w.x \rangle + b = 0$, $H_1 = \langle w.x \rangle + b = +1$, $H_2 = \langle w.x \rangle + b = -1$ biçiminde tanımlanır. Hiperdüzlemlerden elde edilen denklemler verilen sınıflara göre, $H_1 = \langle w.x \rangle + b > 0$, $y_1 = +1$ ve $H_2 = \langle w.x \rangle + b < 0$, $y_2 = -1$ eşitsizlikleri ile ifade edilebilir. H_1 veya H_2 hiperdüzlemleri üzerinde bulunan bir destek vektör ile H_0 düzlemi üzerindeki bir P noktası arasındaki uzaklık d olmak üzere,

$$d = \frac{|wx_p \pm b|}{\|w\|} = \frac{|w_1x_{1p} + w_2x_{2p} + \dots + w_nx_{np} + b|}{\sqrt{w_1^2 + w_2^2 + \dots + w_n^2}}$$

Biçiminde hesaplanır. Buradan hareketle bir X_1 destek vektörü ile H_0 hiperdüzlemi arasındaki uzaklık $d = \frac{|wx_1+b|}{\|w\|} = \frac{1}{\|w\|}$ olup H_1 ve H_2 hiperdüzlemleri arasındaki uzaklık $m = 2d = \frac{2}{\|w\|}$ olarak elde edilir. H_1 ve H_2 hiperdüzlemleri arasındaki boşluğu maksimum hale getirerek optimal ayırıcı hiperdüzlemi bulmak için $\|w\|$ değeri minimal hale getirilerek m 'nin maximum değeri, w ve b değerlerine göre $y_i = [(w.x_i) + b] \geq 1$ kısıtlaması altında elde edilir (Özkan, 2016, s.172). Optimizasyon problemine ait bu durumun çözülebilmesi için sağladığı hesaplama kolaylığı ve doğrusal olmayan durumlar için de genelleştirilebileceğinden dolayı Lagrange fonksiyonu kullanılır. Fonksiyon ve hesaplanması aşağıdaki gibidir (Burges, 1998):

$$L(x_1, x_2, \dots, x_n, a_1, a_2, \dots, a_n) = f(x_1, x_2, \dots, x_n) - \sum_{i=1}^n a_i [g_i(x_1, x_2, \dots, x_n) - b_i]$$

$$L(w, b, a) = \frac{1}{2}(w.w) - \sum_{i=1}^n a_i \{y_i[(w.x_i) + b] - 1\}$$

$L(w, b, a)$ fonksiyonunu w ve b değerlerine göre minimize etmek için konveks bir problem haline dönüşen bu durum Karush-Kuhn-Tucker (KKT) koşulları ile çözülür. Aşağıda verilen bu eşitsizlik kısıtlamaları ve Lagrange çarpanları arasındaki ilişkilere KKT tamamlayıcılık koşulları denir (Abe, 2005, s.18).

$$\frac{\partial L(w_1, b, a)}{\partial w} = 0$$

$$\frac{\partial L(w_1, b, a)}{\partial b} =$$

$$\alpha_i \{y_i (w \cdot x_i + b) - 1\} = 0 \quad \text{for } i = 1, \dots, M,$$

$$\alpha_i \geq 0 \quad \text{for } i = 1, \dots, M.$$

Yukarıda yer alan lagrange fonksiyonunun türevleri alınıp sıfıra eşitlenirse;

$$\frac{\partial L(w_1, b, a)}{\partial w} = w - \sum_{i=1}^n y_i a_i x_i = 0, \quad \frac{\partial L(w_1, b, a)}{\partial b} = \sum_{i=1}^n y_i a_i = 0 \text{ olur.}$$

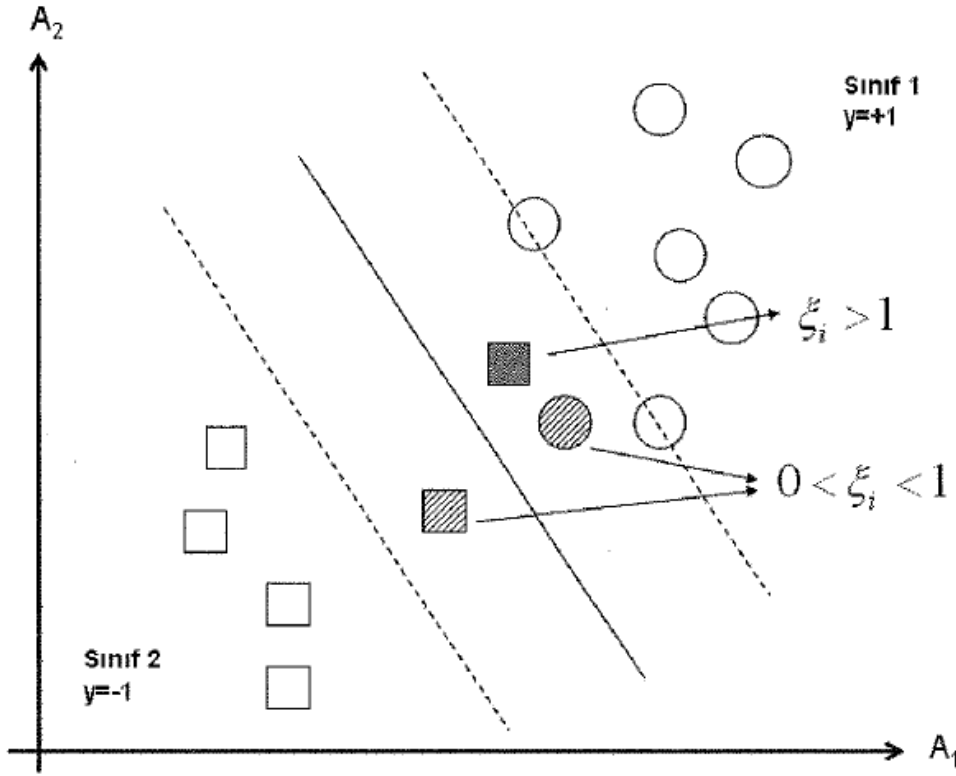
Buradan, $w = \sum_{i=1}^n y_i a_i a x_i$ ve $\sum_{i=1}^n y_i a_i = 0$ elde edilir. Son olarak gerekli düzenlemeler sonucunda,

$$L(w_1, b, a) = \frac{1}{2}(ww) - \sum_{i=1}^n a_i [y_i(wx_i + b) - 1] = \frac{1}{2} \sum_{i,j=1}^n y_i y_j a_i a_j (x_i x_j) - \sum_{i,j=1}^n y_i y_j a_i a_j (x_i x_j) + \sum_{i=1}^n a_i = \sum_{i=1}^n a_i - \frac{1}{2} \sum_{i,j=1}^n y_i y_j a_i a_j (x_i x_j)$$

elde edilir (Özkan, 2016, s.174).

2.6.1.2. Soft Marjin Lineer Sınıflandırma (Doğrusal Olarak Ayrılamama)

Maksimal marj sınıflandırıcısı DVM için önemli bir kavramdır. Ancak verilerin gürültülü olması durumunda değişkenler için doğrusal bir ayırım söz konusu olamayacağından birçok gerçek veride kullanılamaz (çok güçlü kernel fonksiyonu kullanmanın mümkün olmadığı durumlarda). Maksimal marj sınıflandırıcısındaki ana problem, her zaman eğitim hatası olmayan mükemmel bir tutarlı hipotez üretmesidir. Gürültünün her zaman mevcut olabileceği gerçek verilerde, bu durum hatalı bir tahminle sonuçlanabilir. Ayrıca, verilerin değişken kategorilerinde doğrusal olarak ayrılamadığı durumlarda, başlıca (primal) boş bir uygulanabilir bölge ve ikili sınırsız bir objektif fonksiyon mevcut olduğundan optimizasyon problemi çözülemez (Cristianini & Shawe-Taylor, 2000, s.103). Verinin doğrusal olarak ayrılmama durumu Şekil 23'te yer almaktadır (Özkan, 2016, s.179).



Şekil 23. Verinin doğrusal olarak ayrılmama durumuna ait model.

Optimizasyon probleminin çözümü için Şekil 23'te görüldüğü gibi negatif olmayan ve hata değerlerini temsil eden ξ_i gevşek değişkenleri modele eklenir. Bu gevşek değişkenler yardımıyla, $y_i(w^T x + b) \geq 1, \forall i$ yerine $y_i(w^T x + b) \geq 1 - \xi_i$ ($i=1,2,\dots,n$), $\xi_i \geq 0$ veya $y_i(w^T x + b) \geq 1 - \xi_i$, $y_i = +1$ için $y_i(w^T x + b) \leq -1 + \xi_i$, $y_i = -1$ için biçiminde yazılabilir (Özkan, 2016, s.179). Doğrusallığın sağlanamadığı durumlarda bazı eğitim örneklerinin yanlış sınıflandırılmasına izin verilmesi ile büyük bir marj ve küçük bir sınıflandırma hatası arasında belirli bir denge sağlanması koşuluyla bir ayırıcı hiperdüzlem seçilebilir. Böyle bir dengeyi sağlayarak optimizasyon probleminin çözümü aşağıdaki gibi gerçekleştirilir (Khemchandani & Chandra, 2017, s.4).

$$\text{Min}_{(w,b,\xi)} \quad \frac{1}{2} w^T w + C \sum_{i=1}^m \xi_i$$

$$\begin{aligned} y_i(w^T x^{(i)} + b) + \xi_i &\geq 1, \quad (i = 1, 2, \dots, m), \\ \xi_i &\geq 0, \quad (i = 1, 2, \dots, m), \end{aligned}$$

Burada $\sum_{i=1}^m \xi_i$ terimi yanlış sınıflandırma hatasının değeri, C hiperparametresi $\frac{1}{2} w^T w$ marj terimine göre $\sum_{i=1}^m \xi_i$ yanlış sınıflandırma teriminin önem derecesini gösterir. Bu hiper parametre kullanıcı tarafından seçilir ve genellikle bir ayarlama seti kullanılarak belirlenir, yani iyi bir C seçimi yapmak için eğitim setinin küçük bir örneği üzerinde çalışma yapılır (Khemchandani & Chandra, 2017, s.4). Büyük bir C değeri az sayıda yanlış sınıflandırmaya, daha büyük $w^T w$ ve sonuç olarak daha küçük marja neden olur. Yani $C=\infty$ alınması yanlış sınıflandırılan veri sayısının sıfır olması anlamına gelir. Bu nedenle $C<\infty$ olarak seçilmelidir (Huang, Kecman & Kopriva, 2006, s.32). Denklemin çözümüne devam edilirse,

$$\begin{aligned} \text{Max}_{\alpha} \quad & -\frac{1}{2} \sum_{i=1}^m \sum_{j=1}^m y_i y_j (x^{(i)})^T x^{(j)} \alpha_i \alpha_j + \sum_{j=1}^m \alpha_j \\ \text{subject to} \quad & \\ & \sum_{i=1}^m y_i \alpha_i = 0, \\ & 0 \leq \alpha_i \leq C, \quad (i = 1, 2, \dots, m). \end{aligned}$$

Sondaki ifadenin en uygun çözümü $\alpha^* = (\alpha_1^*, \alpha_2^*, \dots, \alpha_m^*)$ olmak üzere K.K.T koşulları altında yumuşak marj sınıflandırıcısı $x^T w^* + b^* = 0$ 'dır. Burada $w^* = \sum_{i=1}^m \alpha_i^* y_i x^{(i)}$ dir.

Bu durumda destek vektörlerine karşılık gelen $x^{(j)}$ çarpanı $0 < \alpha_j < C$ olup K.K.T koşullarına göre, $y_j (w^{*T} x^{(j)} + b) = 1$ 'dir. Hiperdüzlem üzerinde bulunan destek vektörler $w^{*T} x^{(j)} + b^* = \pm 1$ 'dir. Yumuşak marj DVM b^* herhangi bir destek vektörü yani $x^{(j)}$ çarpanı $0 < \alpha_j < C$ ve b^* tüm destek vektörlerinin ortalaması olmak üzere,

$b^* = y_j - \sum_{i=1}^m \alpha_i^* y_i (x^{(i)})^T x^{(j)}$ olarak elde edilir (Khemchandani & Chandra, 2017, s.4).

Elde edilen yumuşak marj sınıflamasına göre,

α_i ve ξ_i değerleri sıfır ise doğru sınıflandırma yapılmıştır.

$0 < \alpha_i < C$, $y_i (w^T x_i + b) - 1 + \xi_i = 0$ ve

$\xi_i = 0$ olduğunda $y_i (w^T x_i + b) = 1$ ve x_i

bir destek vektördür. $C > 0 > \alpha_i$ süper vektörüne sınırsız süper vektör denir.

$a_i=C, y_i (\mathbf{w}^T \mathbf{x}_i + b) - 1 + \xi_i = 0$; ve $\xi_i \geq 0$. $\mathbf{x}_i$ bir destek vektördür. $a_i=C$ ile elde edilen bu destek vektöre sınırlı vektör denir. Eğer $0 \leq \xi_i < 1$ ise $\mathbf{x}_i$ doğru sınıflandırılmıştır ve $\xi_i \geq 1$ ise yanlış sınıflandırılmıştır (Abe, 2005, s.24). Yumuşak marjlı sınıflama ile sert marjlı sınıflama arasındaki fark Lagrange çarpanlarına üst sınırının eklenmesidir (Mavroforakis & Theodoridis, 2006).

2.6.1.3. Lineer Olmayan Sınıflandırma

Doğrusallığın sağlanamadığı durumlarda veri setinden elde edilen gözlem vektörü kendisinden daha yüksek dereceli bir uzaya dönüştürülerek doğrusal sınıflandırıcılar elde etmemiz mümkündür. Matematiksel olarak bunu gerçekleştirmek için $\emptyset$ fonksiyonunu kullanırız. $\mathbf{x} \in R^n$ olmak üzere, bu vektörün dönüştürüleceği daha yüksek dereceli uzaydaki vektör $\mathbf{z} \in R^F$ olsun, $R^n \rightarrow R^F$ yani $\mathbf{z} = \emptyset(\mathbf{x})$ eşlemesini yapan $\emptyset$ fonksiyonu $\mathbf{x} \in R^n \rightarrow \mathbf{z}(\mathbf{x}) = [a_1, \emptyset_1(\mathbf{x}), \dots, \emptyset_n(\mathbf{x})]^T \in R^F$ biçiminde tanımlanır. Doğrusal olmayan sınıflandırma için kullanılan bu $\emptyset$ fonksiyonlarının tanımlanabilmesi için kernel fonksiyonlarından yararlanılır. (Özkan, 2016, s.182-184). Kernel fonksiyonlarında temel amaç verinin uygun fonksiyon aracılığıyla daha üst boyutlardaki bir uzaya dönüştürmektir. Bu amaç doğrultusunda, kernel (mapping) fonksiyonu ile doğrusal olarak ayrılmayan düşük boyuttaki verilerin daha üst boyuttaki uzaya ait yeni koordinatları belirlenir (Akpınar, 2017, s.300).. Bir fonksiyonun bir uzayda kernel fonksiyonu olabilmesi için Mercer teoreminin şartlarını taşıması gerekir (Emir, 2013, s.155). Bu şartlar, fonksiyonun sürekli, pozitif tanımlı ve simetrik olmasıdır (Özkan, 2016, s. 184). Literatürde yer alan bazı Kernel fonksiyonları aşağıda Tablo 4'te yer almaktadır (Souza, 2010).

Tablo 4

Literatürde Yer Alan Bazı Kernel Fonksiyonları

Kernelin Adı	Kernel fonksiyonu	Serbest Parametreler
Doğrusal kernel	$k(x, y) = x' \cdot y$	yok
Polinomial kernel	$k(x, y) = (\gamma x' \cdot y + r)^d$	$d \geq 2, \gamma > 0, r > 0$
Gaussian (Radyal Tabanlı) kernel	$k(x, y) = \exp(-\gamma x - y ^2)$	$\gamma > 0$
Sigmoid kernel	$k(x, y) = \tanh(\alpha x' y + c)$	$\alpha < 0, c > 0$

2.6.2. Çok Sınıflı Destek Vektör Makinesi

Destek vektör makinelerinin standart teorisi sadece ikili sınıflandırma problemlerini destekler. Ancak gerçek yaşama ait sınıflandırma probleminde genel olarak ikiden fazla sınıf söz konusudur (Hamel, 2011, s.185). İki sınıf problemleri için formüle edilen DVM'leri doğrudan karar işlevleri kullandığından, çok sınıflı problemlerin çözümüne yönelik genişletilmesi kolay değildir (Abe, 2005, s.83). DVM'yi çok sınıf problemlerine genişletmek için farklı bir dizi yöntemler sunulmuş olsa da en popüler ikisi bire-karşı-hepsi (one-against-all) ve ikili sınıflama/bire-karşı-bir (pairwise classification/one versus one) yaklaşımlardır (James vd., 2013, s.359).

2.6.2.1. Bire-Karşı-Hepsi (One-Against-All)

Bu yöntemde her bir sınıfı diğer sınıflardan ayırmak için bir grup ikili sınıflandırıcılar oluşturulur. Bu, çok sınıflı ayırmanın tamamlanabilmesi için modeldeki her sınıfın ayrı ayrı değerlendirileceği anlamına gelir (Duman, 2010, s.20). Bu yöntemde her defasında K sınıflarından biri diğer K-1 sınıflarıyla karşılaştırılır. Bir k. sınıfın (+1) diğer sınıflarla (-1) DVM tarafından karşılaştırılmasından elde edilen parametreler $\beta_{0k}, \beta_{1k}, \dots, \beta_{pk}$ ve bir test gözlemi x^* olsun. x^* gözlemini $\beta_{0k} + \beta_{1k}x_1^* + \beta_{2k}x_2^* + \dots + \beta_{pk}x_p^*$ nin en büyük olduğu sınıfa atarız. Çünkü bu durum test gözleminin diğer sınıflardan ziyade K. sınıfa ait olduğu noktasında daha yüksek bir güven düzeyi olduğunu göstermektedir (James, Witten, Hastie

& Tibshirani, 2013, s.360). Bu yöntem işleyişinde öncelikle K-sınıf sınıflandırıcılar elde etmek için, her biri bir sınıfı diğerlerinden ayırmak üzere eğitilmiş $f^1, \dots, f^K$ biçiminde bir dizi ikili sınıflandırıcı oluşturulur. Daha sonra sgn işlevini uygulamadan önce maksimum çıktıya göre çok sınıflı sınıflandırma yapılır. En sonunda ise elde edilen bu iki sonuç birleştirilir. $f^j(x) = \text{sgn}(g^j(x))$ olmak üzere $\arg\max_{j=1, \dots, K} g^j(x)$, burada $g^j(x) = \sum_{i=1}^k y_i \alpha_i^j k(x, x_i) + b^j$ biçiminde elde edilir. $g^j(x)$ değerleri reddetme kararları için de kullanılabilir. Bunun için, iki $g^j(x)$ arasındaki farkı, en büyük x'in sınıflandırılmasına duyulan güvenin bir ölçüsü olarak kabul ediyoruz. Bu hesaplama bir eşiğin altına düşerse bir sınıfa atama durumu olmaz (Scholkopf & Smola, 2001, s.211).

2.6.2.2. İkili Sınıflama/Bire-Karşı-Bir (Pairwise Classification/One Versus One)

Bu yöntemde sınıflandırıcılar, m sınıfın mümkün olabilecek tüm çiftleri için yapılandırılır ve bunun sonucunda $(m-1)m/2$ sınıf elde edilir (Duman, 2010, s.21). Bir Bire-karşı-bir veya bire-karşı-hepsi yaklaşımı, her biri bir çift sınıfı karşılaştıran kombinasyonu kadar DVM oluşturur. Her bir $\binom{M}{2}$ sınıflandırıcısını kullanarak bir test gözlemini sınıflandırırız ve test gözleminin K sınıflarının her birine kaç kez atandığını kaydederiz. En sınıflandırmada ise, test gözlemi bu $\binom{M}{2}$ ikili sınıflandırmalar arasında en sık atandığı sınıfa atanır (James vd., 2013, s.360). Bu yöntem fazla sayıda ikili karşılaştırma yaptığı için büyük eğitim süreleri gerektirse de üzerinde çalışılan problem önemli ölçüde daha kolaydır ve doğrusallık net bir biçimde sağlanıyorsa aslında zamandan tasarruf etmekte mümkündür. Bir veri setini sınıflandırmaya çalıştığımızda, sınıf sayısının ikili kombinasyonu kadar sınıflandırıcıyı da değerlendiririz ve sınıflardan en fazla oyu alan sınıflara göre sınıflandırırız. Bununla birlikte, diğer sınıflandırıcılar genellikle geri kalanlara karşı daha az destek vektöre sahiptir. Bunun iki nedeni vardır: Birincisi, eğitim setleri daha küçüktür ve ikincisi, sınıfların daha az çakışması olduğu için öğrenilecek problemler genellikle daha kolaydır (Scholkopf & Smola.

2001, s.212). İkili sınıflandırmalar için matematiksel yapı aşağıdaki gibidir (Hamel, 2011, s.189).

$$D = \{(\bar{x}_1, y_1), (\bar{x}_2, y_2), \dots, (\bar{x}_l, y_l)\} \subset \mathbb{R}^n \times \{1, 2, \dots, M\}$$

İkili sınıflandırmada, $M(M-1)/2$ karar yüzeyi oluşturmamız gerekmektedir: her bir olası sınıf çifti için bir karar yüzeyi. $g^{p,q}: \mathbb{R}^n \rightarrow \{p, q\}$ karar yüzeyini temsil etsin ve bu karar yüzeyi p ve q sınıf çiftlerini $p \neq q$ ve $\{p, q\} \subset \{1, 2, \dots, M\}$ ile ayırsın. Her bir karar yüzeyi eğitilirse,

$$g^{p,q}(\bar{x}) = \sum_{i=1}^{|D^{p,q}|} y_i \alpha_i^{p,q} k(\bar{x}_i, \bar{x}) - b^{p,q}$$

veri seti üzerinde

$$D^{p,q} = D^p \cup D^q$$

burada,

$$D^p = \{(\bar{x}, y) \mid (\bar{x}, y) \in D \wedge y = p\}$$

ve

$$D^q = \{(\bar{x}, y) \mid (\bar{x}, y) \in D \wedge y = q\}$$

D^p seti, p etiketiyle D 'deki gözlemlerin tümünden oluşur ve $D^{p,q}$ etiketiyle D 'deki gözlemlerin tümünden oluşur. p ve q sınıf çiftleri için $D^{p,q}$ eğitim seti bu iki setin basitçe birleşimidir.

2.7. Lojistik Regresyon

Lojistik regresyon analizi veri madenciliğinde kullanılan en yaygın yöntemlerden biri olan Lojistik regresyon bir veya daha fazla bağımsız değişken ile bir ikili veya çoklu sınıfa sahip bağımlı değişken arasındaki ilişkiyi inceler (Mohamed, 2015, s.21). 1944, 1953 ve 1955 yıllarında Berkson tarafından yapılan çalışmalar LR analizinin ilk aşamalarını

oluşturmaktadı (Çokluk vd., 2012, s.57). Bağımlı değişken ile bağımsız değişken arası ilişkileri matematiksel olarak modelleyen (Hosmer, Lemeshow & Sturdivant, 2013, s.1) denetimli öğrenme yöntemleri arasında yer alır (Kaur & Ginige, 2018). Bağımlı değişkenin kategorik olması nedeniyle ayırma analizine (discriminant analysis) alternatif olan LR'de temel amaç, bağımsız değişkenlerden yola çıkarak iki veya daha fazla gruba ilişkin üyelik tahmini yapmaktır. Dolayısıyla LR analizi ile aynı anda hem sınıflama yapılır hem de bağımlı değişken ile bağımsız değişkenler arasındaki ilişkiler incelenir (Mertler & Vannatta, 2017, s.304). LR, sosyal bilimlerde, sağlık bilimlerinde, ekonomi, pazarlama ve bankacılık gibi birçok alanda kullanılan bir yöntemdir (Burmaoğlu, 2009, s.50).

LR, ayırma ve çoklu regresyon analizleri ile ilişkilidir. Ancak bu analizlerden farklı olarak, bağımlı ve bağımsız değişkenlerin normal dağılım göstermesi, bağımlı değişken ile bağımsız değişken/ler arasında doğrusal ilişkinin olması ve grupların eşit varyansa sahip olması gibi varsayımlara ihtiyacı yoktur. Ayrıca LR modelleri bu yöntemlere göre daha esnek olup sürekli, kesikli veya her ikisinde bulunduğu karma değişken gruplarıyla çözüm üretebilecek yapıya sahiptir (Tabachnick vd., 2013, s.439). Bahsedilen avantajlarının yanında LR yönteminin model kurulurken dikkat edilmesi gereken bazı sınırlılıkları vardır. Nedenselliğe yönelik olarak incelenen özelliklerin kuramsal konular bakımından ele alınmasına, örneklem sayısının değişken sayısına oranına, analizin gücüne etkisinden dolayı beklenen frekanslara, logit dönüşümünün doğrusallığına, çoklu bağlantılılık sorununa, uç değerlere ve hataların bağımsız oluşuna duyarlı bir yöntemdir (Tabachnick vd., 2013, s.443-445; Mertler & Vannatta, 2017, s.304).

2.7.1. Lojistik Regresyon Modelleri ve Matematiksel Yapıları

Doğrusal regresyon modelinde tek bir bağımsız değişkenin yer aldığı basit doğrusal regresyon modeli; $Y = \beta_0 + \beta_1 X + \varepsilon$ birden fazla bağımsız değişkenin yer aldığı çoklu doğrusal regresyon modeli ise; $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \varepsilon$ biçiminde gösterilir. Doğrusal regresyona ait bazı varsayımların sağlanamadığı durumlarda, normal dağılım ve süreklilik koşullarını gerektirmeden kullanılan LR analizi kullanılır (Özdamar,

2013, s.468-479-523). LR analizi ise matematiksel olarak olasılıklara, odds oranına ve odds oranının logaritmasına (lojit fonksiyonu) dayanır (George & Mallery, 2020, s.326) ve doğrusal model denklemini logit olarak adlandırılan logaritmik terim dönüşümleri ile tanımlar (Field, 2013, s.1116).

Odds oranı bir olayın meydana gelme olasılığının meydana gelmeme olasılığına oranı olarak tanımlanır (Mertler & Vannatta, 2017, s.308). Odds oranı, LR analizinden elde edilen sonuçların yorumlanması için çok önemlidir. Odds oranı beta katsayılarının üstel değeridir ve bağımsız değişkendeki bir birimlik değişimin sonucu nasıl etkilendiğini ifade eder. Regresyon katsayısındaki beta değeri gibi ele alınabilir ancak anlaşılması daha kolaydır çünkü logaritmik bir dönüşüm gerektirmez ve bağımsız değişkenin kategorik olması durumunda odds oranını açıklamak daha kolaydır. Odds oranı aşağıdaki biçimde tanımlanabilir (Field, 2013, s.1120):

$$Odds = \frac{p(X)}{1 - p(X)}$$

Odds oranının doğal logaritması bağımsız değişkenlerin doğrusal bir fonksiyonudur (Kalaycı, 2006, s.280). Lojit fonksiyonu kullanılarak elde edilen değer eksi ve artı sonsuz aralığında değer alırken olasılık değerleri sıfır ile bir arasında değişir (Abacıoğlu, 2018, s.81). Olasılık, odds ve lojit fonksiyon değerlerini bir araya getirdiğimizde elde edilen LR denklemi aşağıdaki gibi tanımlanır (Mertler & Vannatta, 2017, s.309).

$$\ln\left(\frac{\hat{Y}}{1 - \hat{Y}}\right) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \varepsilon$$

Denklemden $\hat{Y}_i = \frac{e^u}{1 + e^u}$ bağımlı değişkenin bir kategorisinin meydana gelme olasılığı ve

$u = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \varepsilon$ doğrusal regresyon eşitliğidir.

Lojistik regresyon analizinde, ikili lojistik regresyon (Binary Logistic Regression), sıralı lojistik regresyon (Ordinal Logistic Regression) ve çok kategorili lojistik regresyon (Nominal Logistic Regression, Multinomial Logistic Regression) olmak üzere üç temel yöntem vardır (Özdamar, 2013, s.525).

İkili Lojistik regresyon, bağımlı değişkenin iki kategoriye sahip olduğu durumlarda kullanılan LR yöntemidir. Bir veya birden fazla bağımsız değişken ile ikili bağımlı değişken arasındaki ilişkilerin açıklanması amaçlanır (Özdamar, 2013, s.525). Matematiksel olarak aşağıdaki biçimde ifade edilir.

$$P(Y_j = 1) = \frac{e^{\beta_0 + \beta_1 X_{j1} + \beta_2 X_{j2} + \dots + \beta_k X_{jk}}}{1 + e^{\beta_0 + \beta_1 X_{j1} + \beta_2 X_{j2} + \dots + \beta_k X_{jk}}}$$

n: birim sayısı, $j:1,2,\dots,n$, $P(Y_j = 1)$ j. birimin incelenen sınıfa ait olma olasılığı, β_0 bağımsız değişkenlerin sıfır olması durumundaki sabit değer, $\beta_1, \beta_2, \dots, \beta_k$ bağımsız değişkenlere ait regresyon katsayıları, $X_{j1}, X_{j2}, \dots, X_{jk}$ j. birime ait bağımsız değişkenler ve k bağımsız değişken sayısıdır (Feinstein, 1996, s.27).

Sıralı Lojistik regresyon, bağımlı değişkenin sıralı ölçek türünde olduğu durumlarda kullanılan LR yöntemidir. Doğal sıralaması korunmak üzere en az üç kategoriye sahip bağımlı değişken ile bir veya birden fazla bağımsız değişken arasındaki ilişkilerin açıklanması amaçlanır. (Özdamar, 2013, s.525). Sıralı lojistik regresyon modeli için model oluşturma aşamaları, esasen ikili lojistik regresyon modeli ve çok uluslu lojistik regresyon modeline benzer olarak elde edilir (Hosmer vd., 2013, s.305).

Çok kategorili lojistik regresyon ise ikili bağımlı değişkene yönelik kullanılan LR modelinin geliştirilmesi sonucunda ortaya çıkmıştır. McFadden (1974) tarafından önerilen bu geliştirme sonucunda elde edilen bu model kesikli seçim (discrete choice) model olarak adlandırılmıştır. Alan yazında çok sınıflı (multinomial), çok kategorili (polychotomous /polytomous) lojistik regresyon model olarak da ifade edilmektedir (Hosmer vd., 2013, s.269). Bağımlı değişkenin sınıflamalı ölçek türünde olduğu durumlarda kullanılan LR yöntemidir. En az üç kategoriye sahip bağımlı değişken ile bir veya birden fazla bağımsız değişken arasındaki ilişkilerin açıklanması amaçlanır (Özdamar, 2013, s.526).

M kategorisi sayısına sahip bağımlı değişkenler için, bağımlı değişken ile bağımsız değişkenler arasındaki ilişkiyi tanımlamak için M-1 eşitlik hesaplanmasını gerektirir,

referans kategori ile ilgili her bir kategori için bir tane. Referans kategori hariç bağımlı değişkenin her bir kategorisi için aşağıdaki eşitlik yazılabilir:

$$g_k(X_1, X_2, \dots, X_k) = e^{a_h + b_{h1}X_1 + b_{h2}X_2 + \dots + b_{hk}X_k}, h=1,2,\dots,M-1 \text{ 'dir.}$$

Burada k alt indisi ilgili bağımsız değişkeni ve h alt indisi Y bağımlı değişkenine ait değerleri ifade eder. Referans kategori için, $g_0(X_1, X_2, \dots, X_k)$ dir. Y'nin hariç tutulan h_0 değerinden ziyade h'nin herhangi bir değerine eşit olma olasılığı aşağıdaki gibidir:

$$P(Y = h | X_1, X_2, \dots, X_k) = \frac{e^{a_h + b_{h1}X_1 + b_{h2}X_2 + \dots + b_{hk}X_k}}{1 + \sum_{h=1}^{M-1} e^{a_h + b_{h1}X_1 + b_{h2}X_2 + \dots + b_{hk}X_k}} \quad h=1,2,\dots,M-1 \text{ 'dir. Ve hariç tutulan kategori için } h_0=M \text{ ya da } 0,$$

$$P(Y = h_0 | X_1, X_2, \dots, X_k) = \frac{1}{1 + \sum_{h=1}^{M-1} e^{a_h + b_{h1}X_1 + b_{h2}X_2 + \dots + b_{hk}X_k}} \quad h=1,2,\dots,M-1 \text{ 'dir. } M=2$$

olduğunda, iki kategorili bağımlı değişken için lojistik regresyon modeline sahip olduğumuzu, referans kategorinin ilk kategori olduğu, $h_0=0$ olduğunu ve ilişkiyi tanımlamak için toplamda $M-1=1$ eşitliğe sahip olduğumuzu göz önünde bulundurmalıyız (Menard, 2002, s.92).

2.8. İlgili Araştırmalar

Kwon ve Sim (2013) veri setine ait özelliklerin sınıflama algoritmaları üzerindeki etkilerini incelemişlerdir. Çoklu regresyon metodu ile bağımsız değişken olarak alınan veri seti özellikleri ve bağımlı değişken olan performans değerlendirme kriterleri arasındaki nedensellik incelenmiştir. Çalışmada, örneklem büyüklüğü, sınıf sayısının ikili ya da çoklu olması, tanımlanmayan sınıf, kayıp veri, verinin dengesizliği, değişken sayısı ve konu alanı gibi veri seti özellikleri ele alınmıştır. 223 farklı veri setinin incelenmiştir. C4.5, RIPPER (repeated incremental pruning to produce error reduction), NRBNF (normalized Gaussian radial basis function network), DVM, Bayes ağı (Bayesian network), lojistik regresyon, k-enyakın komşuluk sınıflama algoritmaları seçilmiş ve WEKA programı ile analizler gerçekleştirilmiştir. Genel anlamda yöntemler arasında en iyi performansı Bayes ağı ve C4.5 göstermiştir. Çalışma sonucunda veri dengesizliği durumunda bütün sınıflama

algoritmalarında sınıflama doğruluğu azaldığı, örneklem büyüklüğü arttıkça genel olarak sınıflama doğruluğunun arttığı belirlenmiştir.

Kılıç-Depren, Aşkın ve Öz (2016) eğitimsel veri madenciliği yöntemlerinin performansını TIMSS 2011 Türkiye veri seti üzerinde değerlendirmişlerdir. TIMSS matematik başarı puanları genel ortalamaya göre başarılı/başarısız olarak yeniden kodlanarak çalışmanın bağımlı değişkeni ve matematik başarısı ile ilişkili olduğu düşünülen yine TIMSS veri setinde yer alan 11 değişken, bağımsız değişkenler olarak ele alınmıştır. Veri seti, RF, J48, Naive Bayes (NB), YSA ve LR yöntemleri ile WEKA 3.7 programı ile analiz edilmiştir. Her bir yöntem için hesaplanan performans değerlendirme kriterleri kullanılarak yöntemlerin sınıflama performansları arasındaki farkın anlamlılığı Friedman analizi ile incelenmiştir. Yapılan bu analiz sonucunda ilgili veri seti için sınıflama performansı en iyi olan yöntemin lojistik regresyon olduğu tespit edilmiştir. Çalışmada ayrıca, WEKA değişken seçim yöntemleri ile matematik başarısını en çok etkileyen değişkenlerin neler olduğu ve bunların önem sırası sunulmaktadır.

Kumar ve Chong (2018) çalışmalarında psikolojik veri üzerinde hastaların sağlık durumlarını, yüksek korelasyona sahip değişkenlerin modelde yer alması ile uygulanan makine öğrenmesi algoritmaları ile tanımlamayı amaçlamıştır. Uygulanan algoritmalar, DVM, RO ve LR'dir. Bağımsız değişkenler ile bağımlı değişkenler arasındaki korelasyonlar hesaplanmış ve bu sayede hangi bağımsız değişkenin bağımlı değişken üzerinde önemli etkisi olduğunun tespitinin yapılabileceği ifade edilmiştir. Bu sayede bağımlı değişkenin daha etkili bir şekilde yordandığından bahsedilmiştir. Test ettikleri modelde, bağımlı değişkenle yüksek korelasyona sahip değişken seti kullandıklarında ve zayıf korelasyona sahip değişkenler çıkarıldığında modelin doğrulunun arttığını tespit etmişlerdir. Yöntemler genel anlamda birbirine göre kıyaslandığında, RO yönteminin, çalışmada ele alınan depresyon türünün yordanmasında en uygun yöntem olduğu bulunmuştur.

Güneri ve Apaydın (2004), başarılarının sınıflandırılmasına yönelik yaptıkları çalışmada YSA ve LR yöntemlerini karşılaştırmışlardır. Çalışmalarında Gazi Üniversitesi Ticaret ve Turizm Eğitim Fakültesinde öğrenim gören 352 öğrenciden elde edilen gerçek verileri

kullanmışlardır. Sınıflandırma yöntemlerini karşılaştırırken bağımlı değişken olarak 3 yıla ait not ortalamalarını, bağımsız değişken olarak ise öğrenim görülen program, cinsiyet, lise puan ortalaması, ÖSS puanı, ailesinin yaşadığı şehir, mezun olunan lise türü ve yaş bağımsız değişkenlerini modellerde kullanmışlardır. Araştırma bulgularına göre yapay sinir ağları analizi ve lojistik regresyonlarından elde edilen sınıflandırma işlemleri sonucunda her iki yöntemin de sınıflandırma performanslarının benzer ve genel doğruluk değerlerinin %95,17 olduğu sonucuna ulaşmışlardır.

Tosun (2007) araştırmasında adlı araştırmasında yapay sinir ağları ve karar ağaçları yöntemlerini karşılaştırılmıştır. Demografik, eğitim, akademik ve kişisel bilgiler olarak dört ana başlık altında topladığı bağımsız değişkenlerden faydalanarak üniversite öğrencilerinin not ortalamalarına göre oluşturduğu kategorilere göre sınıflarını tahmin etmeye çalışmıştır. Karar ağaçları ile öğrenci başarılarına göre yapılan sınıflandırma sonrasında %86'lık bir doğruluk oranı elde ederken, aynı verilerle yapay sinir ağlarına göre yapılan sınıflandırma sonrasında %92'lik bir doğruluk oranı elde etmiştir. Bu sonuçlara göre öğrencilerin başarı puanlarına göre sınıflandırılmasında yapay sinir ağının karar ağaçlarına göre daha iyi performans gösterdiğini ifade etmiştir.

Büyükışıklar (2014) Karar ağaçları ile gerçekleştirdiği çalışmasında yöntemlerin sınıflandırma performanslarını bağımlı değişkenin 2, 3 ve 5 kategori olması durumuna göre karşılaştırmış ve tüm yöntemlerin en başarılı performansını bağımlı değişkenin 2 kategori olması halinde gerçekleştirdiği sonucuna ulaşmıştır. Yöntemler arasında yapılan karşılaştırmada ise RO yönteminin diğer karar ağaçlarına göre her koşulda daha iyi sınıflandırma performansı gösterdiğini ifade etmiştir.

Korkmaz (2018) çalışmasında SVRA ve RO yöntemlerini kullanarak Botnet tespiti yapmıştır. Kullanılan sınıflandırma modellerinde Devam Süresi, Protokol Tipi, Kaynak ve Hedef IP, Kaynak ve Hedef Port, Syn Bayrak Durumu, Reset Bayrak Durumu, Ack Bayrak Durumu, Toplam Paket, Reset Bağlantı Sayısı, Ortalama Bayt, Ortalama Paket Oranı, Paket Boyutu bağımsız değişkenler ve Botnet bağımlı değişken olarak ele alınmıştır. Yapılan uygulamalar sonucunda RO değişkeninin SVRA yöntemine göre daha iyi performans

gösterdiği sonucuna ulaşılmıştır. Ayrıca çalışmalarında bir önemli değişkenin bulunması durumunda diğer değişkenlerin çıkarılmasının sonuçları çok fazla etkilemediğini tespit etmişlerdir.

Toprak (2017) yapay sinir ağı, karar ağaçları ve ayırma analizi yöntemlerinin sınıflandırma performansını PISA 2012 matematik başarı puanları üzerinde karşılaştırmalı olarak değerlendirmişlerdir. Bağımlı değişken olarak matematik başarı puanı önce PISA tarafından tanımlanan kesme puanlarına göre 6 sınıfa ayrılmıştır. Ardından kategoriler birleştirilerek bağımlı değişken 3 ve 2 sınıfa ayrılmıştır. Matematik başarısı ile ilişkili olduğu düşünülen, eksik verisi olmayan, normal dağılıma sahip olan, uç değerler bulundurmeyen, çoklu bağlantısı olmayan ve PISA veri setinde yer alan 17 değişken, bağımsız değişkenler olarak ele alınmıştır. Veri seti, SPSS 23 programı ile analiz edilmiştir. Değerlendirmeler ve karşılaştırmalar, her bir yöntem ve koşul için gerçekleştirilen 50 analiz ve bu analizler sonucunda elde edilen doğru sınıflama oranlarının ortalamaları alınarak yapılmıştır. Yapılan karşılaştırmalar sonucunda ele alınan geniş, orta ve küçük örneklem büyüklüklerinde ve bağımlı değişkenin 6, 3 ve 2 kategori olması durumunda, sınıflandırmada en iyi performansı gösteren yöntemin yapay sinir ağı olduğu, bunu sırasıyla karar ağaçları ve ayırma analizinin izlediği tespit edilmiştir. Benzer biçimde varyans-kovaryans matrislerinin homojenliğinin sağlandığı koşullar altındaki tüm örneklem büyüklüklerinde de yapay sinir ağı yönteminin diğer yöntemlere göre daha iyi performans gösterdiği sonucuna ulaşılmıştır. Öte yandan araştırmada kullanılan yöntemlerin farklı örneklem büyüklükleri ve alt grup sayılarına sınıflama yapmada gösterdikleri performansların birbirleriyle karşılaştırılmalarının yanında, her bir yöntemin en iyi performansı hangi koşullarda gösterdiği incelenmiş ve yapay sinir ağı yönteminin varyans-kovaryans matrislerinde homojenliğin sağlanmadığı çok küçük örneklemde en iyi performansı sağladığı tespit edilmiştir. Karar ağaçları yöntemi en yüksek sınıflama performansı büyük örneklemde 2 kategorili sınıflandırma yapmada elde edilmiştir. Diğer bir yöntem olarak ele alınan ayırma analizi ise en yüksek sınıflandırma performansına varyans-kovaryans matrislerinde homojenliğin sağlandığı örneklemde 2 kategorili sınıflandırma yapmada ulaşmıştır.

Özmen (2018) çalışmasında PISA 2012 matematik okuryazarlığı performansları bakımından başarı sıralaması ilk sırada yer alan Çin-Şangay, ortada olan İspanya ve son sırada yer alan Peru'daki öğrencilerin duyuşsal özelliklerinin akademik sınıflandırmayı ne derecede etkilediğini belirleyerek, lojistik regresyon analizi ve ayırma analizlerini kullanılıp elde edilen sonuçları karşılaştırılmıştır. Bağımlı değişken sınıflara ayrılırken alt ve üst yeterlik düzeyleri olmak üzere 2 kategorili hale getirilmiştir. Varsayımlar açısından ayırma analizine ait varsayımların sağlanmasının yeterli olacağı ifade edilmiş ve bu bağlamda veri setindeki örneklem büyüklüğü, uç değerler, normal dağılım, varyans-kovaryans matrislerinin homojenliği ve çoklu doğrusal bağlantı varsayımlarının sağlandığı ve kayıp verilerin yer almadığı veri setleri ile analizler gerçekleştirilmiştir. Lojistik regresyon analizi ve ayırma analizi ile sınıflama doğruluğunun değerlendirilirken matematiğe yönelik duyuşsal özelliklerin, öğrencileri alt ve üst yeterlik düzeylerine göre ne düzeyde doğru sınıfladığı incelenmiş ve elde edilen sınıflama doğrulukları karşılaştırılmıştır. Yapılan analizler sonucunda her üç ülke için de ayırma analizi ile lojistik regresyon analizinden elde edilen sonuçların oldukça benzer olduğu görülmüştür. Matematiğe yönelik duyuşsal özellikler, bireyleri alt ve üst yeterlik düzeylerine göre Çin-Şangay ve İspanya'da anlamlı bir biçimde doğru sınıflarken, Peru'da sınıflandıramamıştır. Ayrıca bireyleri alt ve üst yeterlik düzeylerine göre ayırmada en etkili değişkenler Çin-Şangay'da matematik öz yeterliği ve matematik niyetleri olurken İspanya'da matematik öz yeterliği, matematik benlik kavramı ve matematiğe yönelik araçsal motivasyon düzeyi olmuştur. Peru'da ise matematik benlik kavramı ve matematik öz yeterliğin en etkili değişkenler olduğu tespit edilmiştir.

Koyuncu (2018) Naive Bayes, en yakın komşuluk, yapay sinir ağları ve lojistik regresyon yöntemlerinin sınıflandırma performansını PISA 2012 matematik başarı puanları üzerinde örneklem büyüklüğü ve test verisi oranı açısından karşılaştırmalı olarak incelemiştir. Araştırmada veri seti olarak, OECD ülkelerinden PISA (2012) uygulamasına katılan ve ilgili değişkenlere ait kayıp verisi olmayan 62728 öğrenciye ait veriler kullanılmıştır. Belirlenen bu hedef evrenden yerine koyma yöntemiyle 500 örneklem büyüklüğü için 100 veri seti, 1000 örneklem büyüklüğü için 50 veri seti ve 5000 örneklem büyüklüğü için 30 veri seti

olmak üzere toplamda 180 veri setine ait dosyalar oluşturulmuştur. Araştırma kapsamında ele alınan koşullardan biri olan test verisi oranları %11, %22, %33, %44 ve %55 olacak biçimde diğer bir koşul olan örneklem büyüklükleri ise 500,1000 ve 5000 olacak biçimde ele alınmıştır. Analizler her bir veri seti için test verisinin her defasında rastgele seçildiği 100 farklı sette gerçekleştirilmiştir. Yöntemlerin karşılaştırılması için değerlendirme ölçütleri olarak Kappa hata matrisi uyumu, ROC eğrisinin altında kalan alan ile doğruluk oranları ve standart sapma değerleri kullanılmış ve elde edilen farkların manidarlığı istatistiksel olarak test edilmiştir. Bağımlı değişken olarak birinci matematik başarı puanı kullanılmış ve bağımlı değişken OECD raporunda yer alan kesme puanlarına göre orta düzeyin altında ve üstünde yer alan öğrenciler şeklinde iki kategorili sınıflara ayrılmıştır. Bağımsız değişkenlerin belirlenmesinde filtreleme yöntemi kullanılmış ve tüm analizler WEKA Version 3.9.0 programında gerçekleştirilmiştir. Yapılan analizler sonucunda, örneklem büyüklüğü arttıkça yöntemlerin sınıflandırma performansında artış görülürken, test verisi oranının artması yöntemlerin performanslarında farklı etkiler yaratmıştır. Lojistik regresyon analizi büyük örneklemelerde en etkili yöntem iken küçük örneklemelerde düşük performans göstermiştir. Yapay sinir ağları benzer bir eğilim gösterirken, genel olarak Naive Bayes ve lojistik regresyon yöntemlerine göre daha düşük performans göstermiştir. Tüm koşullarda en düşük performanslar en yakın komşuluk yöntemi ile elde edilmiştir.

Aksu ve Doğan (2018) çalışmalarında veri madenciliği ve makine öğrenme yaklaşımının eğitim alanında kullanılmasını ve bu algoritmalara dayalı olarak elde edilen sonuçların güvenilirlik ve geçerlilik değerlerinin ne düzeyde olduğunu belirlemeyi amaçlamıştır. Araştırmada, Decision Stump, Hoeffding Tree, J.48, Lojistik Model, RepTree, Random Forest, Random Tree ve Ridge lojistik Regresyon yöntemleri ele alınmıştır. Bağımlı değişken olarak, PISA 2015 fen okuryazarlığı puanına göre Türkiye ortalamasının altında kalan öğrenciler başarısız, üstünde kalan öğrenciler ise başarılı olacak biçimde iki sınıfa ayrılmıştır. Alanyazın incelenerek fen okuryazarlığı ile ilişkili olduğu düşünülen 29 değişken ele alınmış ardından Best First –forward, Best First –backward ve Greedy Stepwise yöntemleri kullanılarak sınıflandırma analizlerinin yapılmasında kullanılacak en iyi 12

değişkene karar verilmiştir. Değerlendirmeler ve karşılaştırmalar yapılırken güvenilirlik ölçütleri olarak doğru sınıflama sayısı, doğru sınıflama oranı, Kappa istatistiği, ortalama mutlak hata, karekök hata, göreceli mutlak hata ve göreceli karekök hata değerleri kullanılırken ve geçerlilik ölçütleri olarak doğru pozitif oranı, yanlış pozitif oranı, duyarlılık, geri çağırma, F-ölçütü, Matthew korelasyon katsayısı, ROC eğrisi altında kalan alan ve Duyarlılık-Geri çağırma eğrisi altında kalan alan değerleri kullanılmıştır. Analizler WEKA programında gerçekleştirilmiştir. İlk olarak eğitim ve test verisi olarak ayırmadan doğrudan 10 katlı çapraz geçişleme yöntemi yardımıyla yapılan analizler sonucunda yöntemlere ait elde edilen güvenilirlik ve geçerlilik değerleri karşılaştırılmıştır. Sonuç olarak, sınıflama performanslarına göre sırasıyla Random forest, Ridge lojistik regresyon, lojistik model ve Hoeffding tree yöntemlerinin hem hataya dayalı güvenilirlik değerleri hem de geçerlilik ölçütleri bakımından diğer yöntemlere göre daha başarılı olduğu belirlenmiştir. İkinci olarak; farklı yöntemler yardımıyla test edilen verilerin %5, %10, %15, %20, %25, %30 ve %35 olması durumunda elde edilen sonuçların güvenilirlik ve geçerlilik değerleri karşılaştırılmış ve en iyi kestirimde bulunan Ridge regresyon yönteminin tüm test verilerinden oldukça kararlı kestirimlerde bulunduğu belirlenmiştir. Doğru pozitif oranları bakımından en iyi sonuçlar sırasıyla Ridge lojistik regresyon, Random forest, Lojistik model, Decision stump, Hoeffding tree, J4.8, Rep tree ve Random tree iken, ortalama hataların karekökleri incelendiğinde en iyi sonuçlar sırasıyla Ridge lojistik regresyon, Random forest, Lojistik model, Hoeffding tree, Decision stump, J4.8, Rep tree ve Random tree olacak biçimde elde edilmiştir. Ayrıca en iyi kestirimde bulunan Ridge regresyon yönteminin örneklem büyüklükleri bakımından tüm test verilerinde oldukça kararlı kestirimlerde bulunduğu tespit edilmiştir. Son olarak aynı veri seti üzerinde karışıklık matrisi yardımıyla elde edilen sonuçların doğruluk dereceleri karşılaştırılmış ve Random forest ve Random tree yöntemlerinin tüm örnekleri doğru olarak sınıflayarak en yüksek sınıflama oranına sahip olduğu belirlenmiştir. Ardından bu yöntemleri sınıflama oranlarına göre sırasıyla J4.8, Rep Tree, Lojistik model, Ridge regresyon, Hoeffding tree ve Decision stump yöntemleri takip etmiştir.

Agarwal, Pandey ve Tiwari (2012) eğitimde verimadenciliğinin kullanılmasındaki amacın sadece belirtilen soruna bir çözüm bulmak değil, aynı zamanda en iyi yaklaşımı tanımlamak için de bir takım çalışmalar yapmak olması gerektiğini ifade etmektedirler. Bu amaçla çalışmada, belirtilen sorun için en iyi sınıflandırma yöntemini seçmek amacıyla birkaç yöntem bir araya getirilmiş ve veri kümesine uygulanmıştır. Ayrıca optimum yaklaşımı oluşturmak için bazı karşılaştırmalı analizlerin yapılmıştır. Bir üniversitenin veri tabanından alınan öğrencilere ait veriler çalışma verisi olarak kullanılmıştır. Veri seti 2000 öğrenciye ait matematik puanları, sözel yetenek puanları, sayısal yetenek puanları ve yerleşme olasılıklarından oluşmaktadır. Analizler WEKA programı kullanılarak gerçekleştirilmiştir. Yöntemlerin performansının karşılaştırmalı olarak incelenmesi için araştırmada sekiz farklı yöntem ele alınmıştır: lojistik regresyon, multilayer perception, LIBSVM, RBFNetwork, lineer regresyon, Voted Perception, Winnow ve SVO. Bunlara ek olarak karar ağacı yöntemi de uygulanmıştır. Araştırmada ortaya konan en önemli bulgu, öğrencilerin yerleştirme işlemlerini değerlendirmek için veri madenciliği tekniklerinin kullanılmasının alana katkı sunacağına ortaya konmasıdır. Ayrıca sınıflama performansı açısından en iyi yöntemin LIBSVM olduğu belirtilmiştir.

Yükseltürk, Özekeş ve Türel (2014) çevrim içi eğitim programında veri madenciliği yöntemlerini uyguladıkları çalışmalarında öğrencilerin eğitimi bırakma durumlarının tahmin edilmesini incelemişlerdir. Çalışma grubu, 2007-2009 senelerinde çevrim içi Bilgi Teknolojileri Sertifika Programına kayıt yaptıran 189 öğrenciden oluşmaktadır. Çevrim içi anketler yoluyla toplanan veride cinsiyet, yaş, eğitim düzeyi, daha önceki çevrim içi deneyimi, iş, özyeterlik, hazırbulunuşluk, önsel bilgi, kontrol odağı ve sınıf düzeyi olarak da bırakma durumu (bıraktı/bırakmadı) bulunmaktadır. Öğrencileri sınıflamak için k-en yakın komşuluk, karar ağacı, Naive Bayes ve yapay sinir ağları yöntemleri kullanılmıştır. Yöntemlerden elde edilen duyarlılık değerleri birbirinden çok farklı olmamakla birlikte k-en yakın komşuluk ve karar ağacı yöntemlerinin daha duyarlı olduğu bulunmuştur. Çalışma sonucunda bu yöntemlerin eğitim araştırmalarında kullanımının uygunluğu ve faydaları

özetlenmiş ve çevrim içi eğitimin artması ve buradan öğrenciler hakkında daha fazla bilgi edinilmesinin kolaylaşması ile yöntemlerin kullanılabilirliğinin artacağı vurgulanmıştır.

Minaei-Bidgoli vd. (2003) çevrim içi eğitime katılan 227 öğrencinin birçok özelliğine ve ders notlarına ilişkin verileri veri madenciliği ile inceledikleri bir çalışma gerçekleştirmişlerdir. Ders notlarına göre öncelikle 9 sınıf oluşturmuşlar ve daha sonra da kategoriler birleştirilerek bağımlı değişkenin 2 ve 3 sınıflı halleri de elde edilmiştir. Ardından ilgili bağımsız değişkenler ve bağımlı değişkenin her 3 durumu için kurulan sınıflama modelleri ile yöntemlerin hem değişen sınıf sayısına göre hem de birbirlerine göre performansları karşılaştırılmıştır. Çalışmada kullanılan sınıflama yöntemleri, Karesel Bayes sınıflayıcısı, I-enyakın komşuluk (I-NN), k-en yakın komşuluk (k-NN), Parzen-window, çok katmanlı yapay sinir ağı ve karar ağaçları yöntemleridir. Yapılan analizler sonucunda bütün sınıflandırma yöntemlerinin sınıf sayısı arttıkça sınıflama doğruluk oranlarının düştüğü yani daha kötü performans gösterdiği tespit edilmiştir. Yöntemlerin birbirlerine göre karşılaştırmaları değerlendirildiğinde ise bağımlı değişken sınıf sayısının her üç durumu için de (2, 3 ve 9 sınıf) en iyi sonuçların Karar ağaçları ve yapay sinir ağları yöntemleri altında elde edildiği sonucuna ulaşılmıştır.

Herzog (2006) üniversite öğrencilerinin okul terki durumlarını incelemek için birçok makine öğrenmesi tekniğini karşılaştırmışlardır. Çalışmada ele alınan yöntemler, yapay sinir ağları (simple topology, multitopology, three-hidden-layer pruned), karar ağaçları (C&RT, CHAID-based, and C5.0) ve çok kategorili lojistik regresyondur. Sonuçlar, incelenen yöntemlerin hepsinin benzer doğruluk oranlarına sahip olduğunu göstermektedir.

Cortez ve Silva (2015) 29 adet yordayıcı değişken kullanarak ortaokul öğrencilerinin matematik ve Portekizce derslerindeki başarılarını tahmin etmeyi amaçlamıştır. Kurulacak sınıflama modellerinde gerçekleştirilecek analizler için karar ağaçları, rastgele orman, yapay sinir ağı ve destek vektör makinesi yöntemleri kullanılmıştır. Modelde bağımlı değişken olarak yer alan Matematik ve Portekizce başarı puanları 2 ve 5 kategorili değişkenler haline getirilmiş ardından 788 öğrenciye ait veri seti üzerinde analizler gerçekleştirilmiştir. Sonuç olarak hem 2 hem de 5 sınıflı durumda yapay sinir ağları ve karar ağacı yöntemlerine ait

sınıflama doğruluk oranları diğer yöntemlere göre daha yüksek bulunmuştur. Matematik ve Portekizce başarı puanlarına ait sınıf sayıları arasında yapılan karşılaştırmada ise 2 sınıflı durumda elde edilen doğruluğun 5 sınıflı duruma göre daha yüksek olduğu sonucuna ulaşılmıştır.

Guo, Zhang, Xu, Shi ve Yang (2015) çalışmalarında derin öğrenme yaklaşımını kullanarak öğrenci performansını tahmin etmek için bir model geliştirmişlerdir. Veri setinde 100 farklı ortaokuldan 120000 öğrenciye ait veri bulunmaktadır. Lise giriş sınav puanının tahmin edilmesi için öğrencilere ait farklı özellikler kullanılmıştır. Bu özellikler arasında demografik bilgiler, geçmiş akademik başarı puanları, okula ilişkin bilgiler ve duyuşsal özellikler bulunmaktadır. Öğrencilerin başarı puanlarını tahmin etmek amacı ile Naive Bayes, çok katmanlı yapay sinir ağı, destek vektör makinesi ve araştırmacılar tarafından derin öğrenmeye göre modellenen kendi geliştirdikleri yöntem kullanılmıştır. Yapılan analizler sonucunda, araştırmacıların geliştirdiği yöntem en iyi tahminleme doğruluk değerini üretirken en kötü tahminlemeyi Naive Bayes yöntemi göstermiştir. Çok katmanlı yapay sinir ağı ve destek vektör makinesi ise benzer tahminleme performansları sergilerken görece yapay sinir ağına dayalı model daha iyi bir tahminleme performansı gösterdiği tespit edilmiştir.

Asif, Merceron, Ali ve Haider (2017) çalışmalarında Bilgi Teknolojileri alanında 4 yıllık lisans programına devam eden öğrencilerin performansını iki bağlamda ele almışlardır. İlki öğrencilerin 4 yıllık eğitim sonrasındaki akademik başarıları ve ikincisi ise öğrencilerin gelişimi ve bu gelişimin tahmin sonuçları ile ilişkisidir. Bu doğrultuda lisans öğrencilerinin performanslarını incelemek için veri madenciliği yöntemleri kullanılmıştır. Çalışmada kullanılan veri madenciliği yöntemleri, karar ağaçları, yapay sinir ağları, 1-en yakın komşuluk, Naive Bayes ve rastgele orman yöntemleridir. Öğrenci performansı düşük ve yüksek başarı olmak üzere iki sınıfa ayrılmıştır. Çalışma grubu 220 üniversite öğrencisinden oluşmaktadır. Bağımlı ve bağımsız değişkenler kullanılarak oluşturulan sınıflama modellerine veri madenciliği yöntemleri uygulanarak doğruluk ve kappa değerleri elde edilmiştir. Bu değerlerin incelenmesi sonucunda en iyi performansı Naive Bayes yönteminin

gösterdiği sonucuna ulaşılmıştır. Karar ağaçları, yapay sinir ağı ve rastgele orman yöntemlerinden elde edilen doğruluk ve kappa değerleri incelendiğinde ise rastgele orman yönteminin diğerlerine göre daha iyi performans gösterdiği tespit edilmiştir. Ayrıca, iyi ya da kötü performansın göstergesi olan az sayıda derse odaklanarak, düşük başarılı öğrencilerin desteklenmesi ve yüksek başarılı öğrencilere başka olanaklar sağlanmasının mümkün olduğu belirtilmektedir.

İlgili araştırmalarda veri madenciliğinde kullanılan sınıflandırma yöntemlerinin farklı koşullar ve değerlendirme kriterlerine göre karşılaştırıldığı çalışmalara yer verilmiştir. Eğitim bilimleri ve farklı alanlarda yapılan araştırmalardan elde edilen sonuçlar incelendiğinde YSA, RO, DVM, SVRA ve LR yöntemlerinin her birinin eğitim alanında da başarılı bir şekilde kullanılabileceği söylenebilir. Ancak yöntemlerin performansları ele alınan koşullar bağlamında farklılık gösterebilmektedir. Bu bağlamda, yapılan bu çalışma ile YSA, RO, DVM, SVRA ve LR yöntemlerinin hangi koşullarda en iyi performansı gösterecekleri karşılaştırmalı olarak ele alınmıştır.

BÖLÜM III

YÖNTEM

Bu bölümde araştırmanın türü, çalışma grubu, veri toplama araçları ve verilerin analizine yönelik bilgiler, SPSS tarafından belirlenen Veri Madenciliği için Alanlararası Standart Prosedürler sürecinde yer alan aşamalar ile ilişkilendirilerek ele alınmıştır.

3.1. Araştırmanın Türü

Araştırmanın türü, SPSS tarafından tanımlanan CRISP-DM sürecinde İşin Anlaşılması aşamasında karşılık gelmektedir. Bu çalışmada, bağımsız değişkenler yardımıyla bireylerin sınıflandırılması amacı ile kullanılan YSA, RO, DVM, SVRA ve LR yöntemlerinin bağımlı değişkenin kategori sayısına, bağımsız değişken sayısına ve örneklem büyüklüğüne göre farklı koşullar altında performanslarının değerlendirilmesi amaçlandığından dolayı betimsel bir çalışmadır. Betimsel çalışmalarda, geniş bir kitlenin belirli bir konuya yönelik görüşlerinin ya da özelliklerinin betimlenmesi amaçlanır (Büyüköztürk, Kılıç-Çakmak, Akgün, Karadeniz ve Demirel, 2013, s.177).

3.2. Çalışma Grubu

Çalışma grubu, SPSS tarafından tanımlanan CRISP-DM sürecinde Veri Hazırlama aşamasında karşılık gelmektedir. Çalışma grubu veri setindeki değişkenler ve kayıp veriler doğrultusunda elde edilmiştir. Uygulanan işlem adımları aşağıdaki gibidir.

1. Adım: OECD'nin PISA ile ilgili dokümanlarının bulunduğu internet sitesinden öğrenci anketlerinin yer aldığı SPSS veri dosyaları indirilmiştir. Bu veri dosyasında yer almayan Arnavutluk, Arjantin, Kazakistan ve Malezya ülkelerine ait veriler de aynı internet sitesinden indirilerek diğer dosyalar ile birleştirilmiştir. Bu işlem sonucunda toplamda 542385 öğrenciye ait veriye ulaşılmıştır. Elde edilen veri setinde yer alan öğrencilerin ülkelere göre dağılımı Tablo 5'de yer almaktadır.

Tablo 5

PISA 2015 Öğrenci Verisinde Yer Alan Öğrencilerin Ülkelere Göre Dağılımı

Ülke	N	Ülke	N	Ülke	N
Albania	5215	Indonesia	6513	Qatar	12083
Algeria	5519	Ireland	5741	Romania	4876
Argentina	6349	Israel	6598	Rusya	6036
Australia	14530	Italy	11583	Singapore	6115
Austria	7007	Japan	6647	Çek Cumhuriyeti	6350
Belgium	9651	Kazakhstan	7841	Vietnam	5826
Brazil	23141	Jordan	7267	Slovenia	6406
Bulgaria	5928	Korea	5581	Spain	6736
Canada	20058	Kosovo	4826	Sweden	5458
Chile	7053	Lebanon	4546	Switzerland	5860
Chinese Taipei	7708	Latvia	4869	Thailand	8249
Colombia	11795	Lithuania	6525	Trinidad and Tobago	4692
Costa Rica	6866	Luxembourg	5299	United Arab Emirates	14167
Croatia	5809	Macao	4476	Tunisia	5375
Czech Rep.	6894	Malaysia	8861	Turkey	5895
Denmark	7161	Malta	3634	FYROM	5324
Dominik Cum.	4740	Mexico	7568	United Kingdom	14157
Estonia	5587	Moldova	5325	United States	5712
Finland	5882	Montenegro	5665	Uruguay	6062
France	6108	Netherlands	5385	B-S-J-G (China)	9841
Georgia	5316	New Zealand	4520	Spain (Regions)	32330
Germany	6504	Norway	5456	USA (Massachusetts)	1652
Greece	5532	Peru	6971	USA(North Carolina)	1887
Hong Kong	5359	Poland	4478	Argentina (Ciudad Autónoma de Buenos)	1657
Hungary	5658	Portugal	7325		
Iceland	3371	Puerto Rico	1398		

2. Adım: Bütün ülkelere uygulanmayan anketlere ilişkin değişkenler ile buna bağlı olarak ilgili öğrenciler veri setinden çıkarılmış ve 169326 öğrencinin yer aldığı içinde kayıp veri bulunmayan veri setine ulaşılmıştır.

3. Adım: Araştırma kapsamında ele alınan YSA, RO, DVM, SVRA ve LR yöntemlerinin sınıflandırma performansları bağımlı değişkenin kategori sayılarına göre incelenmesi amaçlandığı için bu amaç doğrultusunda bağımlı değişkenin sırasıyla 2, 3 ve 6 kategori olması durumuna göre çalışma grupları oluşturulmuştur. Çalışma gruplarının oluşturulması sürecindeki aşamalar aşağıdaki gibidir.

1. Aşama: PISA veri setinden yer alan 10 adet fen başarı puanının yerine bu puanlara ait ortalamaların kullanılmasının uygun olup olmadığını belirlemek amacıyla, 10 adet fen başarı puanının temel bileşenler analizi sonucundaki regresyon faktör puanları ile ortalamalarından elde edilen standart puanlar arasındaki ilişki korelasyon analizi ile incelenmiştir. Uygulanan korelasyon analizine ait sonuçlar Tablo 6’da yer almaktadır.

Tablo 6

Fen Başarı Puanlarına Ait Ortalama Standart Puanları ile Regresyon Faktör Puanları İçin Hesaplanan Korelasyon Analizi Sonuçları

Ortalama standart puanları		
Regresyon faktör puanları	r	1,000**
	p	,000
	N	169326

Tablo 6’da yer alan korelasyon analizine ait sonuçlar incelendiğinde, 10 adet fen başarı puanının temel bileşenler analizi sonucunda elde edilen regresyon faktör puanları ile 10 adet fen başarı puanının ortalamaları sonucunda elde edilen standart puanları arasında istatistiksel olarak anlamlı ve mükemmel düzeyde bir ilişki vardır ($p < 0,05$). Bu bulgu sonucunda 10 adet fen başarı puanını temsilen bu puanlardan elde edilen ortalama puanın kullanabileceği söylenebilir.

2. Aşama: Bağımlı değişkenin 2, 3 ve 6 düzeye göre kategorileştirilmesi için OECD tarafından belirlenen kesme puanları ele alınarak öncelikle 6 düzeyli grup belirlenmiştir. Ardından belirlenen 6 düzeyli gruptaki 1 ile 2, 3 ile 4 ve 5 ile 6’ncı düzeyler birleştirilerek sırasıyla 1. 2. ve 3. düzeye sahip olan 3 kategorili bağımlı değişken oluşturulmuştur. Daha sonra ise 6 düzeyli gruptaki 1, 2 ile 3 ve 4, 5 ile 6’ncı düzeyler birleştirilerek sırasıyla 1 ve

2. düzeye sahip olan 2 kategorili bağımlı değişken elde edilmiştir. Bağımlı değişkenin 6, 3 ve 2 kategori olduğu durumlara ait frekans ve yüzdelere Tablo 7’de yer almaktadır.

Tablo 7

Bağımlı Değişkenin 6, 3 ve 2 Kategori Olduğu Durumlara Ait Frekans ve Yüzdelere

Düzy	N	%	Düzy	N	%	Düzy	N	%
1. Düzy	24867	14,7	1. Düzy	65306	38,6	1. Düzy	116125	68,6
2. Düzy	40439	23,9	2. Düzy	89705	53,0			
3. Düzy	50819	30,0	3. Düzy	14315	8,5	2. Düzy	53201	31,4
4. Düzy	38886	23,0						
5. Düzy	12940	7,6						
6. Düzy	1375	0,8						
Toplam				169326				

Tablo 7’den görüldüğü üzere OECD tarafından belirlenen kesme puanlarına göre elde edilen 6 düzeyli gruptaki öğrencilerin 24867’si (%14,7) 1. düzeyde, 40439’u (%23,9) 2. düzeyde, 50819’u (%30,0) 3. düzeyde, 38886’sı (%23,0) 4. düzeyde, 12940’ı (%7,6) 5. düzeyde ve 1375’i (%0,8) 6. düzeyde yer almaktadır. Ayrıca Tablo 3 incelendiğinde 6 düzeyli grubun 1 ile 2. düzeylerinin birleştirilmesi ile elde edilen 1. düzeyde 65306 (%38,6) , 3 ile 4. düzeylerinin birleştirilmesi ile elde edilen 2. düzeyde 89705 (%53,0) , 5 ile 6. düzeylerinin birleştirilmesi ile elde edilen 3. düzeyde 14315 (%8,5) öğrenci olduğu görülmektedir. Benzer biçimde 6 düzeyli grubun 1, 2 ile 3. düzeylerinin birleştirilmesi ile elde edilen 1. düzeyde 116125 (%68,6) ve 3, 4 ile 5. düzeylerinin birleştirilmesi ile elde edilen 2. düzeyde 53201 (%34,1) öğrenci yer almaktadır

Bağımlı değişkenin kategori sayısının değişimine göre çalışma grupları belirlendikten sonra analizlerin yapılabilmesi için bağımlı değişkenin kategori sayılarındaki düzeylere ait frekanslar ve yüzdelikler sabit kalmak üzere, koşullara bağlı olarak rastgele seçim ile her bir koşul için 25’er adet farklı veri setleri oluşturulmuştur. Bu veri setlerinin oluşumuna ait detaylı bilgiler “Koşullar” başlığı altında verilmiştir.

3.3. Veri Analizi

Bu bölümde YSA, RO, DVM, SVRA ve LR yöntemlerinin performansları değerlendirilirken hangi değişkenlerin kullanılacağına, hangi koşullar altında değerlendirileceğine, değerlendirmede kullanılacak kriterlere ve bu kriterler doğrultusunda elde edilen değerlerin karşılaştırılmasında kullanılacak fark testlerine yer verilmiştir.

3.3.1. Değişken Seçimi

Bağımsız değişken seçiminde öncelikle, analiz için kullanılacak çalışma gruplarının mümkün olduğu kadar veri setini temsil etmesi ve kayıp verinin olmaması amacıyla PISA öğrenci anketi veri setindeki, veri kaybı çok olan, uygulandığı ülke sayısı az olan, birden fazla değişkenin birleştirilmesi ile elde edilen (örneğin anne baba eğitim durumunda hangisinin daha yüksek indekse sahip olduğu) ve bağımlı değişken olarak ele alınacak fen bilgisi başarı puanı ile ilişkisi olmadığı öngörülen (yabancı dil çalışma süresi) değişkenler çıkarılmıştır.

Öncelikle bütün ülkelere uygulanan değişkenler içerisinde sürekli olanların çalışmada ele alınmasına karar verilmiştir. Ardından yapılacak analizlerde yöntemlerden olabildiğince iyi performans elde etmek için, bağımlı değişkenle ilişkisi diğer bağımsız değişkenlere göre daha yüksek olan bağımsız değişkenlerin kullanılmasına karar verilmiştir. Değişkenlere ait betimsel istatistikler Ek 1’de yer almaktadır.

Yöntemlerin kıyaslanmasındaki koşullardan birinin bağımlı değişkendeki kategori sayısının farklılaşması olduğu için kategori sayılarından elde edilen sonuçların karşılaştırılmasında herhangi bir yanlılığa sebep olmamak amacıyla bağımsız değişken seçiminde bağımlı değişken olarak fen başarı puanlarının sürekli hali kullanılmıştır. Modellerde yer alacak bağımsız değişkenler, bağımlı değişken ve birbirleri ile olan ilişkileri incelenerek yapılan korelasyon analizleri sonucunda belirlenmiştir. EK 1’de yer alan betimsel istatistiklere göre normallik varsayımının sağlandığı durumlarda Pearson korelasyon katsayısı aksi durumda Spearman Brown korelasyon katsayısı hesaplanmıştır. Veri setinde yer alan değişkenler ile

fen başarı puanı arasında hesaplanan korelasyon analizi sonuçları Tablo 8’de yer almaktadır. Burada sadece korelasyon katsayısı 0,20’den büyük olanlar gösterilmektedir. Hesaplanan korelasyon katsayısı 0,20’den küçük olanlara ilişkin sonuçlar Ek 2’de sunulmaktadır.

Tablo 8

Bağımlı Değişken ile Bağımsız Değişkenler Arasında Hesaplanan Korelasyon Katsayıları

Bağımsız değişken	Fen ortalama	
	Pearson korelasyon	
	r	p
Matematik başarısı (MATHACH)*	0,93	0,00
Okuma Başarısı (READ)	0,92	0,00
Ekonomik, sosyal ve kültürel statü indeksi (ESCS)*	0,36	0,00
Epistomolojik inançlar (EPIST)*	0,34	0,00
Evde bulunan kültürel eşyalar (CULTPOSS)*	0,27	0,00
Feni sevmeye (JOYSCIE)*	0,25	0,00
Evdeki eğitim kaynakları (HEDRES)*	0,23	0,00
Bağımsız değişken	Spearman korelasyon	
	r	p
Evde sahip olunan eşyalar (HOMEPOS)*	0,34	0,00
Çevresel farkındalık (ENVAWARE)*	0,28	0,00
Bilgi iletişim teknoloji kaynakları (ICTRES)	0,24	0,00
Öğrenme süresi (haftalık dakika) (SMINS)*	0,24	0,00
Fen özyeterlik inancı (SCIEEFF)*	0,21	0,00
Aile geliri (WEALTH)	0,20	0,00
Genel kapsamlı fen konularına ilgi duyma (INTBRSCI)	0,20	0,00
*Modele seçilen değişkenler		

Tablo 8’de yer alan korelasyon analizi sonuçları incelendiğinde fen ortalama puanları ile matematik ve okuma becerisi ortalama puanlarının yüksek düzeyde, ESCS, EPIST ve HOMEPOS değişkenlerine ait puanların ise orta düzeyde, pozitif yönde ve istatistiksel olarak anlamlı bir ilişkiye sahip oldukları görülmektedir. Geriye kalan değişkenlerin ise fen ortalama puanları ile istatistiksel olarak anlamlı olacak biçimde pozitif yönlü ve düşük düzeyde ilişkiye sahip olduğu söylenebilir. Cohen (1988) sosyal bilimlerde korelasyon büyüklükleri hakkında yorum yaparken, $[0.10, 0.30)$ aralığının düşük, $[0.30, 0.50)$ aralığının orta ve $[0.50, 1.00]$ aralığının yüksek düzey olacak biçimde yorumlanmasını önermiştir. Korelasyon katsayısı 0.10 ve altında olan ilişkilerin ise özellikle büyük örneklemelerde istatistiksel olarak anlamlı olsa bile uygulamada pratik öneminin olmadığını ifade etmiştir. Matematik ortalama ve okuma becerisi ortalama puanları arasındaki ilişkinin düzeyi 0,80

den büyük olduğu için bu değişkenlerden bağımlı değişkenle diğerine görece daha yüksek ilişkiye sahip olan matematik ortalama puanının modellerde kullanılmasına karar verilmiştir. Ardından diğer değişkenler arasındaki korelasyonlar incelenmiş ve HOMEPOSS ile ICTRES, HOMEPOSS ile WEALTH ve WEALTH ile ICTRES arasındaki korelasyonların düzeylerinin 0.80'den büyük olduğu tespit edilmiştir. Ek olarak çoklubağlantılılık durumu için VIF değerleri incelenmiş, HOMEPOSS, ICTRES ve WEALTH değişkenlerinin VIF değerlerinin sırasıyla 24,28, 4,64 ve 14,70 olduğu tespit edilmiştir. Bağımlı değişkenle olan ilişkileri göz önüne alındığında, hem bağımlı değişkenle diğerlerine göre daha düşük düzeyde ilişkiye sahip olan hem de VIF değeri 10'nun üzerinde ve tolerans değeri 0,20'nin altında olan WEALTH değişkeninin sınıflama modellerinde kullanılıp kullanılamayacağı incelenmiştir. WEALTH değişkeni çıkarıldıktan sonra HOMEPOSS ve ICTRES değişkenlerinin VIF değerlerinin sırasıyla 7,53 ve 4,17 olacak biçimde ilk duruma göre düştüğü görülmüş ve WEALTH değişkeninin sınıflama modellerinde kullanılmamasına karar verilmiştir. Bu aşamada sonra ise HOMEPOSS ve ICTRESS değişkenleri arasındaki ilişkinin düzeyi 0,80'den büyük olduğu için sınıflandırma modellerinde bağımlı değişken hakkında diğerine göre daha fazla bilgi veren HOMEPOSS değişkeninin kullanılmasına karar verilmiştir. Hall (1999) ve Ooi, Chetty ve Teng (2007) de çalışmalarında diğer bağımsız değişkenlerle yüksek ilişkiye sahip değişkenin gereksiz değişken olduğunu belirtmektedir. Diğer değişkenlere ait VIF, tolerans ve birbirleri arasındaki ilişkileri gösteren korelasyon katsayıları, uygun aralıklarda olduğu için bu çalışmadaki sınıflama modellerinde kullanılmıştır. Bu değişkenler, Matematik başarısı, ESCS, EPIST, CULTPOSS, JOYSCIE, HEDRES, HOMEPOS, ENVAWARE, SMINS ve SCIEEFF'tir.

3.3.2. Koşullar

Araştırma kapsamında YSA, RO, DVM, SVRA ve LR yöntemlerinin sınıflandırma performansları bağımlı değişkenin kategori sayısının, modelde yer alan bağımsız değişken sayısının ve örneklem büyüklüğünün farklılaşmasının göre değişimi ele alınmıştır. Bu amaç

doğrultusunda belirlenen koşullara göre yöntemlerin hem kendi içindeki hem de diğer yöntemlere göre performansları incelenmiştir.

Yöntemlerin performansları ilk olarak bağımlı değişkenin kategori sayısının değişimine göre incelenmiştir. Yöntemlerin performansları bağımlı değişkenin kategori sayısına göre çalışma grubunun büyüklüğü 5000 ve bağımsız değişken sayısı 10 olacak biçimde değerlendirilmiştir. Bunun için yöntemlerin uygulanacağı sınıflandırma modellerinde, değişken seçimi aşamasında belirlenen 10 adet değişken kullanılmıştır (Matematik başarıları, ESCS, EPIST, CULTPOSS, JOYSCIE, HEDRES, HOMEPOS, ENVAWARE, SMINS ve SCIEEFF). Veri setleri oluştururken ise bağımlı değişkenin 6 sınıfa sahip olduğu durumdaki sınıflara ait oranlar göz önünde bulundurularak, random seçim ile 5000 öğrencinin bulunduğu 25 tane farklı veri seti oluşturulmuştur. Benzer işlemler bağımlı değişkenin 2 ve 3 kategoriye sahip olması durumlarında da gerçekleştirilmiştir. Bu işlemler sonucunda bağımlı değişkenin 2, 3 ve 6 olması durumuna göre her biri için 25 toplamda ise 75 tane farklı veri seti elde edilmiştir. Veri setleri Tablo 9’da gösterilmektedir.

Tablo 9

Bağımlı Değişkenin Kategori Sayısının 6, 3 ve 2 Olması Durumlarında Üretilen Veri Setleri

	Bağımlı değişken kategori sayısı		
	2	3	6
Çalışma grubu 5000 Bağımsız değişken sayısı 10	x25	x25	x25

İlk olarak yöntemlerin performansları bağımlı değişkenin kategori sayısına göre değerlendirilmiş ve bütün yöntemler en iyi performanslarını bağımlı değişkenin kategorisi sayısının 2 olması durumunda gerçekleştirmiştir. Bu nedenle yöntemlerin performansları bağımsız değişkenin sayısına göre çalışma grubunun büyüklüğü 5000 ve bağımlı değişkenin kategori sayısı 2 olacak biçimde değerlendirilmiştir. Modellerde yer alacak bağımsız değişken sayısı ise

- bağımlı değişkenle yüksek düzeyde ilişkiye sahip olan matematik ortalama puanının modelde kalması şartıyla, bağımsız değişken sayısının 4, 7 ve 10

- bağımlı değişkenle yüksek düzeyde ilişkiye sahip olan matematik ortalama puanının modelde kalmaması şartıyla, bağımsız değişken sayısının 3, 6 ve 9 olması

durumlarına göre ele alınmıştır. Veri madenciliğindeki aşamalar birbirleri ile bağlantılı olarak optimum düzeyde sonuçlara ulaşmak için süreç içinde değişim gösterebilir. Yapılan bu çalışmada matematik ortalama puanının bağımlı değişkenle olan ilişkisi modeldeki diğer bağımsız değişkenlere göre çok yüksek olduğu için araştırma sürecinde ortaya çıkan bu durumun da incelenmesinin önemli katkı sağlayacağı düşünülerek bağımsız değişken sayısındaki değişim iki farklı durum altında ele alınmıştır. Bu işlemler sonucunda bağımsız değişkenin 4, 7 ve 10 ile 3, 6 ve 9 olması durumlarının her biri aynı 25 veri seti üzerinde değerlendirilmiştir. Koşullara göre oluşan veri setleri Tablo 10’ da gösterilmektedir.

Tablo 10

Bağımsız Değişkenin 4, 7 ve 10 ile 3, 6 ve 9 Olması Durumlarında Üretilen Veri Setleri

Çalışma grubu 5000			25x veri seti		
Bağımlı değişken kategori sayısı 2					
Bağımsız değişken sayısı			Bağımsız değişken sayısı		
4	7	10	3	6	9
Bağımsız değişken kodları			Bağımsız değişken kodları		
Matematik başarısı EPIST JOYSCIE SMINS	Matematik başarısı HOMEPOS EPIST ENVAWARE JOYSCIE HEDRES SMINS	Matematik başarısı ESCS EPIST CULTPOSS JOYSCIE HEDRES HOMEPOS ENVAWARE SMINS SCIEEFF	EPIST JOYSCIE SMINS	HOMEPOS EPIST ENVAWARE JOYSCIE HEDRES SMINS	ESCS EPIST CULTPOSS JOYSCIE HEDRES HOMEPOS ENVAWARE SMINS SCIEEFF

Bağımlı değişkenin kategori sayısı ve modeldeki bağımsız değişken sayısının yöntemler üzerindeki etkileri incelendiğinde yöntemlerin en iyi performansı bağımlı değişkenin kategori sayısının 2 bağımsız değişken sayısının ise 10 olma durumunda gösterdiği sonucuna ulaşılmıştır. Bu nedenle yöntemlerin performansları örneklem büyüklüğüne göre incelenirken bağımlı değişken 2 kategori bağımsız değişken sayısı ise 10 olarak ele alınmış ve örneklem büyüklükleri 100, 250, 500, 1000, 2500 ve 5000 olacak biçimde değerlendirilmiştir. Bu durumun değerlendirilmesi amacıyla bağımlı değişkenin 2 olması

durumdaki kategori yüzdelikleri sabit kalmak üzere 100, 250, 500, 1000, 2500 ve 5000 örneklem büyüklüklerinin her biri için random seçimle 25'er adet farklı veri seti oluşturulmuştur. Koşullara göre oluşan veri setleri Tablo 11' de gösterilmektedir.

Tablo 11

Örneklem Büyüklüklerinin 100, 250, 500, 1000, 2500 ve 5000 Olması Durumlarında Üretilen Veri Setleri

	Örneklem büyüklüğü					
	100	250	500	1000	2500	5000
Bağımlı değişken kategori sayısı	x25	x25	x25	x25	x25	x25
Bağımsız değişken sayısı 10						

3.3.3. Sınıflama Yöntemlerinin Uygulanması

Araştırma kapsamında ele alınan her bir koşul için sınıflama analizleri SPSS Modeller demo sürümünde gerçekleştirilmiştir. Analizler bütün veri setleri için ayrı ayrı yapıp modellerden elde edilen hata matrisleri Excell ortamına aktarılmıştır. Daha sonra araştırma kapsamında kullanılan performans değerlendirme kriterlerine ait değerler ilgili formüller kullanılarak hesaplanmıştır.

3.3.4. Karşılaştırma testleri

Yöntemler uygulandıktan sonra elde edilen performans değerlendirme kriteri değerleri karşılaştırılırken uygulanacak analiz yöntemleri koşul ve yöntemlere göre ayrı ayrı ele alınmıştır.

Bağımlı değişkenin kategori sayısı 6, 3 ve 2 olması durumunda yapılan sınıflama işlemleri sonucunda YSA, RO, DVM, SVRA ve LR yöntemlerinden elde edilen doğruluk, duyarlılık, belirleyicilik, kesinlik, f ölçütü, negatif öngörü değeri ve kappa indeksi değerlerinin her bir yöntem için bağımlı değişkenin kategori sayısına göre istatistiksel olarak anlamlı farklılık gösterip göstermediğini incelemek amacıyla, betimsel istatistikler (Ek 3) incelenmiş ve varsayımlar sağlanmadığı için Kruskal Wallis testi kullanılmıştır. Elde edilen değerlendirme kriterlerine aitt değerlerin bağımlı değişkenin her bir kategori sayısı durumunda YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını

incelemek amacıyla betimsel istatistikler (Ek 3) doğrultusunda varsayımlar sağlanmadığı için Friedman Testi kullanılmıştır.

Bağımsız değişken sayısının etkisinin incelenmesinde iki durum ele alınmıştır. İlk durumda fen başarı puanı ile yüksek korelasyona sahip matematik başarısının dâhil edildiği 10, 7 ve 4 bağımsız değişken ikinci durumda ise matematik başarısı değişkeninin çıkarıldığı 9, 6 ve 3 bağımsız değişken olması durumu incelenmiştir. Yapılan sınıflama işlemleri sonucunda YSA, RO, DVM, SVRA ve LR yöntemlerinden elde edilen doğruluk, duyarlılık, belirleyicilik, kesinlik, f ölçütü, negatif öngörü değeri ve kappa indeksi değerlerinin her bir yöntem için iki durumdaki bağımsız değişken sayısına göre istatistiksel olarak anlamlı farklılık gösterip göstermediğini incelemek amacıyla, betimsel istatistikler (Ek 4-5) incelenmiş ve varsayımlar sağlanmadığı için Friedman Testi kullanılmıştır. Benzer biçimde, her iki durum için elde edilen değerlendirme kriterlerine aitt değerlerin her bir bağımsız değişken sayısı durumunda YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla betimsel istatistikler (Ek 4-5) doğrultusunda varsayımlar sağlanmadığı için Friedman Testi kullanılmıştır.

Örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması durumunda yapılan sınıflama işlemleri sonucunda YSA, RO, DVM, SVRA ve LR yöntemlerinden elde edilen doğruluk, duyarlılık, belirleyicilik, kesinlik, f ölçütü, negatif öngörü değeri ve kappa indeksi değerlerinin her bir yöntem için örneklem büyüklüğüne göre istatistiksel olarak anlamlı farklılık gösterip göstermediğini incelemek amacıyla, betimsel istatistikler (Ek 5) incelenmiş ve varsayımlar sağlanmadığı için Kruskal Wallis testi kullanılmıştır. Elde edilen değerlendirme kriterlerine ait değerlerin her bir örneklem büyüklüğünde YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla betimsel istatistikler (Ek 5) doğrultusunda varsayımlar sağlanmadığı için Friedman Testi kullanılmıştır.

3.3.5. Değerlendirme Kriterleri

Veri madenciliğinde kullanılan algoritmaların performansının değerlendirilmesi, sürecin önemli bir aşamasını oluşturmaktadır. Yüksek performanslı bir model oluşturmak ve model geliştirmenin sonraki aşamalarında da oluşabilecek beklenmedik sorunlardan kaçınmak için başlangıçtan itibaren modelin değerlendirilmesini ön planda tutmak gerekmektedir. Algoritmaların performanslarının değerlendirilmesinde birçok kriter kullanılmaktadır. Bu çalışmada yöntemlerin sınıflandırma performansının incelenmesinde ele alınacak değerlendirme kriterleri bu kısımda açıklanmaktadır. Bu kriterlere ait değerlerin çoğu confusion matris ile elde edilir. Sınıflama iki düzey olduğunda kullanılan confusion matris Tablo 12’de ve ikiden fazla sınıf olduğunda kullanılan matris Tablo 12’de gösterilmektedir.

Tablo 12

Sınıflama İki Düzey Olduğunda Kullanılan Confusion Matris

		Gerçek sınıf		
Tahmin edilen sınıf	Pozitif	Pozitif	Negatif	Toplam
	Pozitif	GP	YP	P
	Negatif	YN	GN	N
	Toplam	P	N	P+N

Bir modelin hata matrisi, modelin performansını görselleştiren bir tablodur. Matrisin her satırı tahmin edilen sınıfın durumunu temsil ederken, her sütun da ilgili sınıftaki gerçek durumları temsil eder. Başka bir ifadeyle modelin, ilgili bağımlı değişkenle ilgili tahmin edilen sınıfının gerçek sınıflamalara kıyasla ne durumda olduğunu gösterir. Gerçek pozitif (GP), gerçek durumu pozitif olan ve model tarafından pozitif olarak tahmin edilen veriyi ifade etmektedir. Gerçek negatif (GN), gerçek durumu negatif olan ve model tarafından negatif olarak tahmin edilen veriyi ifade etmektedir. Yanlış pozitif (YP) gerçek durumu negatif olan ve model tarafından pozitif olarak tahmin edilen veriyi göstermektedir. Yanlış negatif (YN) gerçek durumu pozitif olan ve model tarafından negatif olarak tahmin edilen veriyi göstermektedir.

Tablo 13

Çoklu Sınıflama Olduğunda Kullanılan Hata Matrisi

		Gerçek sınıf		
		Sınıf 1	Sınıf 2	Sınıf 3
Tahmin edilen sınıf	Sınıf 1	8	1	2
	Sınıf 2	0	10	1
	Sınıf 3	0	1	12

Tablo 13, üç farklı sınıf için hata matrisinin bir örneğini göstermektedir; burada köşegende yer alan sayılar, doğru olarak tanımlanmış durumları gösterirken, geri kalan hücreler yanlış tanımlanmış veri sayısını ve bunların hangi sınıflara yanlış atandığını gösterir.

Bu çalışmada ele alınan performans değerlendirme kriterleri doğruluk (accuracy), duyarlılık (sensitivity/recall), belirleyicilik (specificity), kesinlik/pozitif öngörü değeri (precision/positive predicted value), f ölçütü, negatif öngörü değeri (negative predicted value) ve Kappa indeksleri olup aşağıda kısaca özetlenmektedir.

Doğruluk: Doğruluk, tahmin edilen durumun, gerçek durumla ne kadar yakın olduğunu gösterir. Doğru tahminlerin, toplam sayıya bölünmesi ile elde edilir ve iki sınıf için hesaplama Eşitlik 1’de sunulmaktadır.

$$Doğruluk_M = \frac{GP+GN}{GP+GN+YP+YN} \quad \text{Eşitlik 1}$$

Burada M, test edilen modeli göstermektedir. GP doğru olarak tahmin edilen pozitif durumları, GN doğru olarak tahmin edilen negatif durumları, YP gerçek durumu negatifken yanlış olarak pozitif tahmin edilen durumları ve YN ise gerçek durumu negatif iken yanlış olarak pozitif tahmin edilen durumların sayılarını ifade etmektedir.

Çoklu-sınıf (multi-class) durumunda, N sayıda sınıfa sahip bir confusion matrisi için doğruluk, matrisin köşegeninde yer alan değerlerin toplamının toplam veri sayısına bölünmesi ile elde edilir:

$$Doğruluk_M(\text{çoklu} - \text{sınıf}) = \frac{\sum_{i=1}^N e_{ii}}{\sum_{i=1}^N \sum_{j=1}^N e_{ij}} \quad \text{Eşitlik 2}$$

Burada e matrisin bir hücresi ve ii matrisin köşegenini temsil etmektedir. Köşegen değerleri toplam veriye bölünür. Bu indeks, her bir sınıf için Eşitlik 4 ile hesaplanan iki sınıflı recall değeri kullanılarak da hesaplanabilir:

$$Doğruluk_M(\text{çoklu} - \text{sınıf}) = \frac{\sum_{i=1}^N \text{duyarlılık}_i * \text{ağırlık}_i}{\sum_{i=1}^N \sum_{j=1}^N e_{ij}} \quad \text{Eşitlik 3}$$

Burada duyarlılık_i i sınıfı için hesaplanan duyarlılık değerini ve ağırlık_i ise gerçekte i sınıfının kişi sayısını gösterir. Her sınıfın duyarlılık değeri o sınıftaki kişi sayısına göre ağırlıklandırılarak toplanır ve toplam sınıf sayısına (N) bölünür.

Duyarlılık: Gerçekte pozitif olan durumların model tarafından doğru tahmin edilme oranını ifade etmektedir. Modelin pozitif sonuçları tanımlayabilmesini yansıtmaktadır. İki sınıflı durum için duyarlılık Eşitlik 4 ile hesaplanmaktadır:

$$\text{Duyarlılık}_M = \frac{GP}{GP+YN} \quad \text{Eşitlik 4}$$

Çoklu-sınıf için duyarlılık, doğruluk değerine eşittir ve Eşitlik 3 ile hesaplanır.

Belirleyicilik: Gerçekte negatif olan durumların model tarafından doğru tahmin edilme oranını ifade etmektedir. Modelin negatif sonuçları tanımlayabilmesini yansıtmaktadır. İki sınıflı durum için belirleyicilik Eşitlik 5 ile hesaplanmaktadır:

$$\text{Belirleyicilik}_M = \frac{GN}{GN+YP} \quad \text{Eşitlik 5}$$

Çoklu-sınıf için belirleyicilik, her bir sınıf için hesaplanan belirleyicilik değerinin o sınıftaki kişi sayısı ile ağırlıklandırılarak toplanması ve toplam kişi sayısına bölünmesi ile elde edilir:

$$\text{Belirleyicilik}_M(\text{çoklu} - \text{sınıf}) = \frac{\sum_{i=1}^N \text{Belirleyicilik}_i * \text{ağırlık}_i}{\sum_{i=1}^N \sum_{j=1}^N e_{ij}} \quad \text{Eşitlik 6}$$

Kesinlik: Modelin pozitif olarak belirlediği durumlardaki gerçek pozitif durumların oranını ifade etmektedir. GP değerinin, test edilen modelin tahmin ettiği toplam pozitif değere bölünmesi ile hesaplanmaktadır. Eşitlik 7 ile gösterilmektedir.

$$\text{Kesinlik}_M = \frac{GP}{GP+YP} \quad \text{Eşitlik 7}$$

Çoklu-sınıf için kesinlik, her bir sınıf için hesaplanan kesinlik değerinin o sınıftaki kişi sayısı ile ağırlıklandırılarak toplanması ve toplam kişi sayısına bölünmesi ile elde edilir:

$$\text{Kesinlik}_M(\text{çoklu} - \text{sınıf}) = \frac{\sum_{i=1}^N \text{Kesinlik}_i * \text{ağırlık}_i}{\sum_{i=1}^N \sum_{j=1}^N e_{ij}} \quad \text{Eşitlik 8}$$

Negatif Öngörü Değeri: Modelin negatif olarak belirlediği durumlardaki gerçek negatif durumların oranını ifade etmektedir. GN değerinin, test edilen modelin tahmin ettiği toplam negatif değere bölünmesi ile hesaplanmaktadır. Eşitlik 9 ile gösterilmektedir.

$$\text{Negatif Öngörü Değeri}_M = \frac{GP}{GP + YP} \quad \text{Eşitlik 9}$$

Çoklu-sınıf için Negatif öngörü değeri, her bir sınıf için hesaplanan Negatif öngörü değeri değerinin o sınıftaki kişi sayısı ile ağırlıklandırılarak toplanması ve toplam kişi sayısına bölünmesi ile elde edilir:

$$\text{Negatif öngörü değeri}_M(\text{çoklu} - \text{sınıf}) = \frac{\sum_{i=1}^N \text{negatif öngörü değeri}_i * \text{ağırlık}_i}{\sum_{i=1}^N \sum_{j=1}^N e_{ij}} \quad \text{Eşitlik 10}$$

F ölçütü: Test edilen modelin doğruluğunu ölçmek için kullanılır. Kesinlik ve duyarlılık değerlerinin harmonik ortalamasıdır. Her ikisinden de katkı olan F ölçütü ne kadar yüksekse o kadar iyidir. Matematiksel gösterimi Eşitlik 11’de sunulmaktadır.

$$F \text{ ölçütü}_M = 2 * \frac{\text{kesinlik} * \text{duyarlılık}}{\text{kesinlik} + \text{duyarlılık}} \quad \text{Eşitlik 11}$$

Çoklu-sınıf için F ölçütü, her bir sınıf için hesaplanan F ölçütü değerinin o sınıftaki kişi sayısı ile ağırlıklandırılarak toplanması ve toplam kişi sayısına bölünmesi ile elde edilir:

$$F \text{ ölçütü}_M(\text{çoklu} - \text{sınıf}) = \frac{\sum_{i=1}^N F \text{ ölçütü}_i * \text{ağırlık}_i}{\sum_{i=1}^N \sum_{j=1}^N e_{ij}} \quad \text{Eşitlik 12}$$

Kappa: Şansla uyum olasılığını (chance-agreement probability; P_c) dikkate alarak doğruluk ölçütündeki hataları düzeltmeyi amaçlar. Eşitlik 13 ile gösterilmektedir:

$$K_M = \frac{\text{doğruluk} - P_c}{1 - P_c} \quad \text{Eşitlik 13}$$

Eşitlikte yer alan P_c matematiksel olarak şöyle ifade edilmektedir:

$$P_c = \sum_{i=1}^N \left(\frac{\sum_{i=1}^N e_i}{\sum_{i=1}^N \sum_{j=1}^N e_{ij}} * \frac{\sum_{i=1}^N e'_i}{\sum_{i=1}^N \sum_{j=1}^N e_{ij}} \right) \quad \text{Eşitlik 14}$$

Burada toplam $\sum_{i=1}^N e_i$ gerek durumda bir sınıftaki kiři sayısını gsterirken $\sum_{i=1}^N e'_i$ ise o sınıfın yordanan kiři sayısını gstermektedir. Bu deęerler toplam kiři sayısına blnerek arpılır. Bu iřlem sınıf sayısınca (iki ya da oklu-sınıf) gerekleřtirilir ve toplanır.

BÖLÜM IV

BULGULAR

Bu bölümde, araştırmanın amaçları doğrultusunda elde edilen bulgulara yer verilmiştir. Araştırma kapsamında bağımlı değişken sınıfları, bu sınıflara göre çalışma grubu dağılımı, modellerde kullanılacak bağımsız değişkenler belirlendikten sonra, YSA, RO, DVM, SVRA ve LR analizleri ile sınıflandırma işlemleri gerçekleştirilmiştir. Ardından belirlenen değerlendirme kriterleri ve karşılaştırma testlerine göre yöntemlerin hem koşullara göre genel performansları hem de birbirlerine göre karşılaştırmalı performansları incelenmiştir. Bu doğrultuda elde edilen bulgular koşullar doğrultusunda üç ana başlıkta ele alınmıştır.

4.1. Bağımlı Değişken Kategorisi Sayısının 6, 3 ve 2 Olması Durumuna Göre Elde Edilen Bulgular

Araştırma kapsamında bağımlı değişkenin kategori sayısının 6, 3 ve 2 olması durumunda YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonrasında elde edilen karşılaştırma kriterlerine ait betimsel istatistikler Ek 3'te yer almaktadır. Betimsel istatistiklere göre varsayımların incelenmesi doğrultusunda, sınıflandırma performanslarının her birinin bağımlı değişkenin kategori sayısının değişimine yönelik uygulanan Kruskal Wallis Testi sonuçları Tablo 14'ta yer almaktadır.

Tablo 14

Yöntemlerin Sınıflandırma Performanslarının Bağımlı Değişkenin Kategori Sayısının Değişimine Göre Değerlendirilmesine Yönelik Uygulanan Kruskal Wallis Testi

Kriter	Yöntem									
	YSA		RO		DVM		SVRA		LR	
	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark
Doğruluk	66,14*	I<II<III	64,78*	I<II<III	66,24*	I<II<III	66,35*	I<II<III	67,23*	I<II<III
Kesinlik	68,13*	I<II<III	65,85*	I<II<III	67,22*	I<II<III	68,28*	I<II<III	66,22*	I<II<III
F ölçütü	69,74*	I<II<III	66,73*	I<II<III	66,32*	I<II<III	65,48*	I<II<III	66,82*	I<II<III
Kappa	65,81*	I<II<III	57,73*	I<II<III	65,80*	I<II<III	6,53*	I<II	65,22*	I<II<III
Belirleyicilik	0,19	-	55,64*	I>II>III	0,68	-	67,25*	I>II>III	23,41*	I>II=III
Negatif öngörü değeri	13,40*	I=II>III	53,75*	I>II>III	11,44*	I>III=II	67,75*	I>II>III	44,30*	I>II=III

I: 6 kategori, II:3 kategori ve III: 2 kategori'den oluşan bağımlı değişkeni temsil etmektedir.
*p<0,05

Tablo 14'te YSA yöntemine ait değerlendirme kriterlerinin bağımlı değişkenin kategori sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Kruskal Wallis testine ait sonuçlar ele alındığında, YSA yönteminin 6, 3 ve 2 kategorili bağımlı değişkenin yer aldığı sınıflama modellerine uygulanması sonucunda belirleyicilik değeri hariç diğer tüm değerlendirme kriterlerine ait değerlerin bağımlı değişkenin kategori sayısının değişimine göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 66,14$, $\chi^2_{Kesinlik} = 66,13$, $\chi^2_{F-ölçütü} = 66,13$, $\chi^2_{Kappa} = 65,81$, $\chi^2_{Negatif\ öngörü} = 13,40$; $sd = 2$; $n = 75$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA yöntemine ait doğruluk, kesinlik, F-ölçütü ve Kappa değerlendirme kriterlerine ait değerlerin 2 kategorili bağımlı değişkenin yer aldığı modellerde 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu sonucuna ulaşılmıştır. 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellerdeki ikili karşılaştırmalar ele alındığında ise YSA yöntemine ait doğruluk, kesinlik, f-ölçütü ve Kappa değerlerinin 3 kategorili bağımlı değişkenin yer aldığı modellerde daha yüksek olduğu belirlenmiştir. Negatif öngörü değerine ait ikili karşılaştırmalar incelendiğinde ise 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellerde YSA yöntemine ait negatif öngörü değerinin benzer ve 2 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek oldukları sonucuna ulaşılmıştır.

RO yöntemine ait değerlendirme kriterlerinin bağımlı değişkenin kategori sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Kruskal Wallis testine ait Tablo 14'te yer alan sonuçlar ele alındığında, RO yönteminin 6, 3 ve 2 kategorili bağımlı değişkenin yer aldığı sınıflama modellerine uygulanması sonucunda tüm değerlendirme kriterlerine ait değerlerin bağımlı değişkenin kategori sayısının değişimine göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 64,78$, $\chi^2_{Kesinlik} = 65,85$, $\chi^2_{F-ölçütü} = 66,73$, $\chi^2_{Kappa} = 57,73$, $\chi^2_{Belirleyicilik} = 55,64$, $\chi^2_{Negatiföngörü} = 53,75$; $sd = 2$; $n = 75$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde RO yöntemine ait doğruluk, kesinlik, F-ölçütü ve Kappa değerlendirme kriterlerine ait değerlerin 2 kategorili bağımlı değişkenin yer aldığı modellerde 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu sonucuna ulaşılmıştır. 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellerdeki ikili karşılaştırmalar ele alındığında ise RO yöntemine ait doğruluk, kesinlik, f-ölçütü ve Kappa değerlerinin 3 kategorili bağımlı değişkenin yer aldığı modellerde daha yüksek olduğu belirlenmiştir. Belirleyicilik ve negatif öngörü değerine ait ikili karşılaştırmalar incelendiğinde ise bu değerlerin 6 kategorili bağımlı değişkenin yer aldığı modellerde 3 ve 2 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu sonucuna ulaşılmıştır. 3 ve 2 kategorili bağımlı değişkenin yer aldığı modellerdeki ikili karşılaştırmalar ele alındığında ise RO yöntemine ait belirleyicilik ve negatif öngörü değerlerinin 3 kategorili bağımlı değişkenin yer aldığı modellerde daha yüksek olduğu belirlenmiştir.

Tablo 14'te DVM yöntemine ait değerlendirme kriterlerinin bağımlı değişkenin kategori sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Kruskal Wallis testine ait sonuçlar ele alındığında, DVM yönteminin 6, 3 ve 2 kategorili bağımlı değişkenin yer aldığı sınıflama modellerine uygulanması sonucunda belirleyicilik değeri hariç diğer tüm değerlendirme kriterlerine ait değerlerin bağımlı değişkenin kategori sayısının değişimine göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 66,24$,

$\chi^2_{Kesinlik} = 67,22$, $\chi^2_{F-ölçütü} = 66,32$, $\chi^2_{Kappa} = 65,80$, $\chi^2_{Negatiföngörü} = 11,44$; $sd = 2$; $n = 75$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde DVM yöntemine ait doğruluk, kesinlik, F-ölçütü ve Kappa değerlendirme kriterlerine ait değerlerin 2 kategorili bağımlı değişkenin yer aldığı modellerde 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu sonucuna ulaşılmıştır. 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellerdeki ikili karşılaştırmalar ele alındığında ise DVM yöntemine ait doğruluk, kesinlik, f-ölçütü ve Kappa değerlerinin 3 kategorili bağımlı değişkenin yer aldığı modellerde daha yüksek olduğu belirlenmiştir. Negatif öngörü değerine ait ikili karşılaştırmalar incelendiğinde ise 6 kategorili bağımlı değişkenin yer aldığı modellerde DVM yöntemine ait negatif öngörü değerinin 3 ve 2 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu tespit edilmiştir. Ayrıca 3 kategorili bağımlı değişkenin yer aldığı modellerdeki negatif öngörü değerinin de 2 kategorili bağımlı değişkenin yer aldığı modeller ile benzer olduğu sonucuna ulaşılmıştır.

SVRA yöntemine ait değerlendirme kriterlerinin bağımlı değişkenin kategori sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Kruskal Wallis testine ait Tablo 14'te yer alan sonuçlar ele alındığında, SVRA yönteminin 6, 3 ve 2 kategorili bağımlı değişkenin yer aldığı sınıflama modellerine uygulanması sonucunda tüm değerlendirme kriterlerine ait değerlerin bağımlı değişkenin kategori sayısının değişimine göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 66,35$, $\chi^2_{Kesinlik} = 68,28$, $\chi^2_{F-ölçütü} = 65,48$, $\chi^2_{Kappa} = 6,53$, $\chi^2_{Belirlyicilik} = 67,25$, $\chi^2_{Negatiföngörü} = 67,75$; $sd = 2$; $n = 75$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde SVRA yöntemine ait doğruluk, kesinlik ve F-ölçütü değerlendirme kriterlerine ait değerlerin 2 kategorili bağımlı değişkenin yer aldığı modellerde 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu sonucuna ulaşılmıştır. 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellerdeki ikili karşılaştırmalar ele alındığında ise SVRA yöntemine ait doğruluk, kesinlik, f-ölçütü ve

Kappa değerlerinin 3 kategorili bağımlı değişkenin yer aldığı modellerde daha yüksek olduğu belirlenmiştir. Belirleyicilik ve negatif öngörü değerine ait ikili karşılaştırmalar incelendiğinde ise bu değerlerin 6 kategorili bağımlı değişkenin yer aldığı modellerde 3 ve 2 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu sonucuna ulaşılmıştır. 3 ve 2 kategorili bağımlı değişkenin yer aldığı modellerdeki ikili karşılaştırmalar ele alındığında ise SVRA yöntemine ait belirleyicilik ve negatif öngörü değerlerinin 3 kategorili bağımlı değişkenin yer aldığı modellerde daha yüksek olduğu belirlenmiştir.

Tablo 14'te LR yöntemine ait değerlendirme kriterlerinin bağımlı değişkenin kategori sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Kruskal Wallis testine ait sonuçlar ele alındığında, LR yönteminin 6, 3 ve 2 kategorili bağımlı değişkenin yer aldığı sınıflama modellerine uygulanması sonucunda tüm değerlendirme kriterlerine ait değerlerin bağımlı değişkenin kategori sayısının değişimine göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 67,23$, $\chi^2_{Kesinlik} = 66,22$, $\chi^2_{F-ölçütü} = 66,82$, $\chi^2_{Kappa} = 65,22$, $\chi^2_{Belirleyicilik} = 23,41$, $\chi^2_{Negatiföngörü} = 44,30$; $sd = 2$; $n = 75$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde LR yöntemine ait doğruluk, kesinlik, F-ölçütü ve Kappa değerlendirme kriterlerine ait değerlerin 2 kategorili bağımlı değişkenin yer aldığı modellerde 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu sonucuna ulaşılmıştır. 6 ve 3 kategorili bağımlı değişkenin yer aldığı modellerdeki ikili karşılaştırmalar ele alındığında ise LR yöntemine ait doğruluk, kesinlik, f-ölçütü ve Kappa değerlerinin 3 kategorili bağımlı değişkenin yer aldığı modellerde daha yüksek olduğu belirlenmiştir. Belirleyicilik ve negatif öngörü değerine ait ikili karşılaştırmalar incelendiğinde ise bu değerlerin 6 kategorili bağımlı değişkenin yer aldığı modellerde 3 ve 2 kategorili bağımlı değişkenin yer aldığı modellere göre daha yüksek olduğu sonucuna ulaşılmıştır. 3 ve 2 kategorili bağımlı değişkenin yer aldığı modellerdeki ikili karşılaştırmalar ele alındığında ise LR yöntemine ait belirleyicilik ve negatif öngörü değerlerinin benzer olduğu belirlenmiştir.

Bağımlı değişkenin kategori sayısının değişimine göre yöntemlerin sınıflama performanslarına ait değerlendirme kriterlerinden elde edilen bulgular genel olarak incelendiğinde doğruluk, kesinlik F ölçütü ve Kappa değerlerinin bütün yöntemler için bağımlı değişkenin kategori sayısının artmasına bağlı olarak düşüş gösterdiği söylenebilir. Belirleyicilik ve negatif öngörü değerleri ele alındığında ise tüm yöntemlerde bu değerler bağımlı değişkenin kategori sayısı arttıkça yükselme eğiliminde olduğu ifade edilebilir. Ancak bu artış RO ve SVRA yöntemlerinde belirgin bir anlamlı farklılığa neden olurken YSA, DVM ve LR yöntemlerinde belirgin bir farklılık oluşturmamıştır. Özetle, bağımlı değişkenin kategori sayısı arttıkça bütün yöntemlerin sınıflama performanslarının olumsuz yönde etkilendiği söylenebilir.

Bağımlı değişkenin kategori sayısının 6, 3 ve 2 olması durumunda YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonrasında elde edilen karşılaştırma kriterlerine ait Ek 3'te yer alan betimsel istatistiklere göre varsayımların incelenmiştir. Bu inceleme sonunda her bir örneklem büyüklüğünde yöntemlerin sınıflandırma performanslarının birbirlerine göre değerlendirilmesine yönelik uygulanan Friedman Testi sonuçları Tablo 15'te yer almaktadır.

Tablo 15

Bağımlı Değişkenin Kategori Sayısının 6,3 ve 2 Olduğu Durumlarda Yöntemlerin Sınıflandırma Performanslarının Birbirlerine Göre Değerlendirilmesine Yönelik Uygulanan Friedman Testi Sonuçları

Kriter	Bağımlı değişken kategori sayısı					
	6 kategori		3 kategori		2 kategori	
	İst.	Fark	İst.	Fark	İst.	Fark
Doğruluk	80,92	I, III, V>II>IV	83,12	I, III, V>II>IV	81,40	I, III, V>II>IV
Kesinlik	82,11	I, III, V>II>IV	83,36	I, III, V>II>IV	82,57	I, III, V>II>IV
F ölçütü	81,12	I, III, V>II>IV	82,66	I, III, V>II>IV	77,56	I, III, V>II>IV
Kappa	80,57	I, III, V>II>IV	82,02	I, III, V>II>IV	81,80	I, III, V>II>IV
Belirleyicilik	83,84	I, III, V>II>IV	85,47	I, III, V>II>IV	83,72	I, III, V>II>IV
Negatif öngörü	79,97	I, III, V>II>IV	81,63	I, III, V>II>IV	71,19	I, III, V>II>IV
I: YSA, II: RO, III: DVM, IV: SVRA, V:LR						

Tablo 15'te bağımlı değişkenin kategori sayısının 6 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine

göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, bağımlı değişkenin kategori sayısı 6 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 80,92$, $\chi^2_{Kesinlik} = 82,11$, $\chi^2_{F-ölçütü} = 81,12$, $\chi^2_{Kappa} = 80,57$, $\chi^2_{Belirleyicilik} = 83,84$, $\chi^2_{Negatiföngörü} = 79,97$; $sd = 4$; $n = 25$; $p < 0,05$).

Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise negatif öngörü değeri hariç diğer tüm değerlendirme kriterlerine ait değerlerin RO yönteminde daha yüksek olduğu belirlenmiştir.

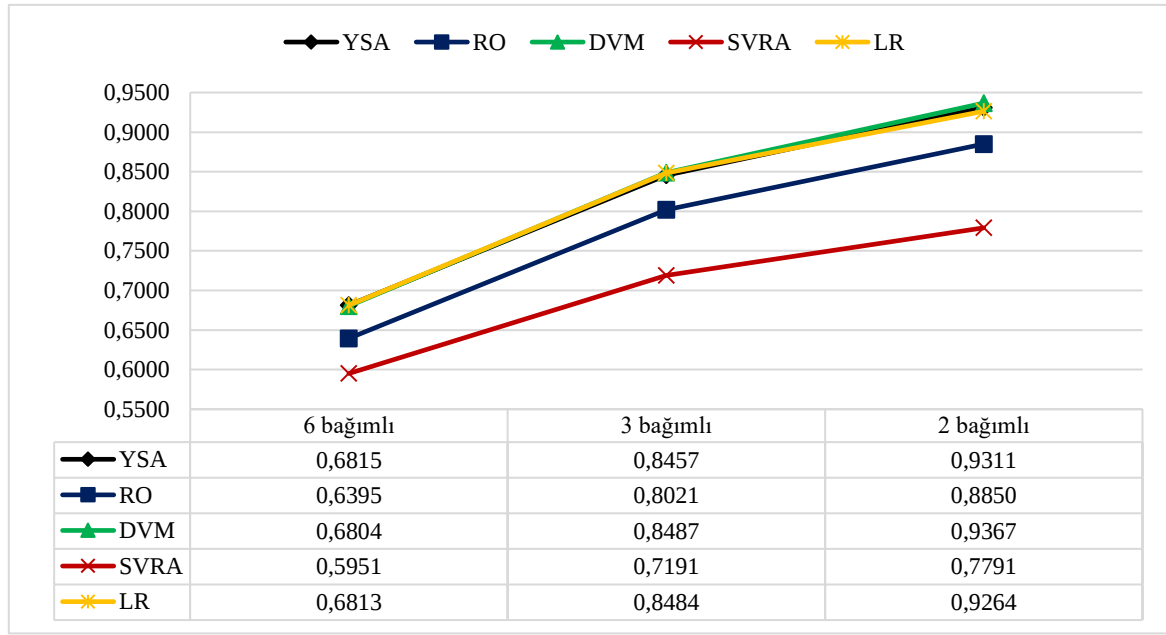
Bağımlı değişkenin kategori sayısının 3 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait Tablo 15'te yer alan sonuçlar ele alındığında, bağımlı değişkenin kategori sayısı 3 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 83,12$, $\chi^2_{Kesinlik} = 83,36$, $\chi^2_{F-ölçütü} = 82,66$, $\chi^2_{Kappa} = 82,02$, $\chi^2_{Belirleyicilik} = 85,47$, $\chi^2_{Negatiföngörü} = 81,63$; $sd = 4$; $n = 25$; $p < 0,05$).

Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerin RO yönteminde daha yüksek olduğu belirlenmiştir.

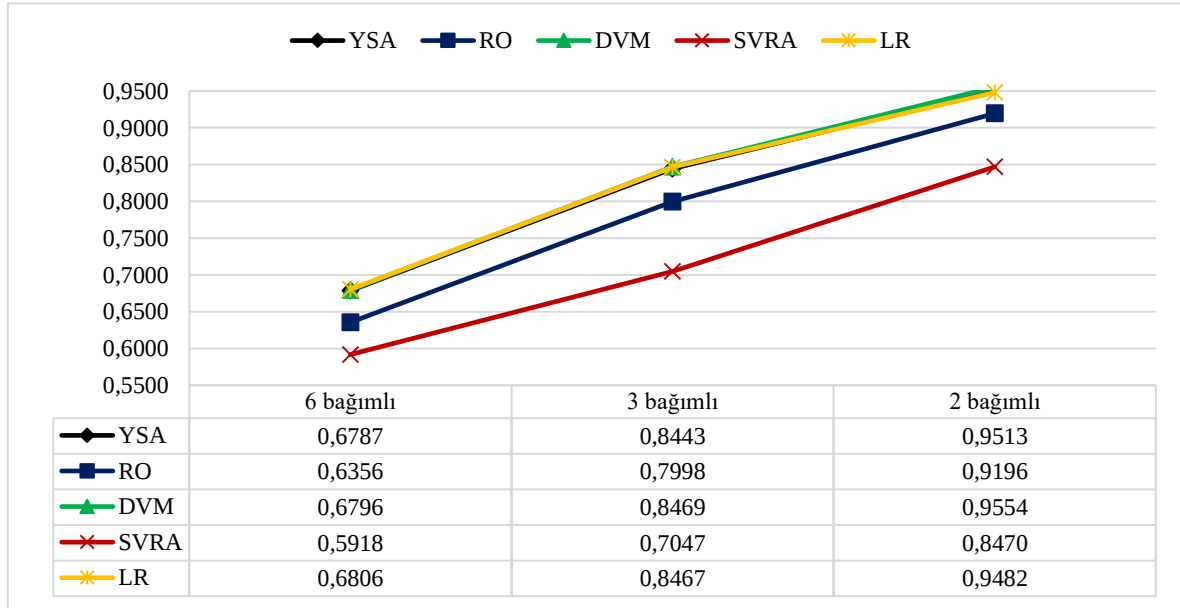
Tablo 15'te bağımlı değişkenin kategori sayısının 2 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, bağımlı değişkenin kategori sayısı 2 iken tüm değerlendirme

kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 81,40$, $\chi^2_{Kesinlik} = 82,57$, $\chi^2_{F-ölçütü} = 77,56$, $\chi^2_{Kappa} = 81,80$, $\chi^2_{Belirleyicilik} = 83,72$, $\chi^2_{Negatiföngörü} = 71,19$; $sd = 4$; $n = 25$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerin RO yönteminde daha yüksek olduğu belirlenmiştir.

Özetle, fark testlerinden elde edilen sonuçlara göre bağımlı değişkenin her bir kategori sayısı için bütün değerlendirme kriterlerine göre yöntemler arasında ortaya çıkan farkın kaynağını belirlemek amacıyla uygulanan ikili karşılaştırmalar sonucunda, değerlendirme kriterlerine ait değerlere göre YSA, DVM ve LR yöntemlerinin kendi aralarında benzer ancak RO ve SVRA yöntemine göre daha iyi sınıflama performansı gösterdiği sonucuna ulaşılmıştır. RO ve SVRA yöntemleri karşılaştırıldığında ise genel olarak RO yönteminin SVRA yöntemine göre daha iyi sınıflama performansına sahip olduğu tespit edilmiştir. Yöntemlerin bağımlı değişkenin her bir kategori sayısı için sahip oldukları doğruluk ve F ölçütü değerlerine ait grafikler sırasıyla Şekil 24 ve Şekil 25'te yer almaktadır. Diğer değerlendirme kriterlerine ait grafikler Ek 7'de yer almaktadır.



Şekil 24. Bağımlı değişkenin kategori sayısı koşulu için her bir yöntemle ait doğruluk değerleri.



Şekil 25. Bağımlı değişkenin kategori sayısı koşulu için her bir yöntemle ait F ölçütü değerleri.

4.2. Bağımsız Değişken Sayısının 10, 7 ve 4 Olması Durumuna Göre Elde Edilen Bulgular

Araştırma kapsamında bağımsız değişken sayısının 10, 7 ve 4 olması durumunda YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonrasında elde edilen değerlendirme

kriterlerine ait betimsel istatistikler Ek 4'te yer almaktadır. Betimsel istatistikler doğrultusunda varsayımların incelenmesi sonucunda, yöntemlerin sınıflandırma performanslarının kendi içinde bağımsız değişken sayısının değişimine göre değerlendirilmesine yönelik uygulanan Friedman Testi sonuçları Tablo 16'da yer almaktadır.

Tablo 16

Yöntemlerin Sınıflandırma Performanslarının Bağımsız Değişken Sayısının Değişimine Göre Değerlendirilmesine Yönelik Uygulanan Friedman Testi Sonuçları

Kriter	Yöntem									
	YSA		RO		DVM		SVRA		LR	
	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark
Doğruluk	2,16	-	2,88	-	1,02	-	6,86	-	0,43	-
Duyarlılık	3,51	-	3,29	-	0,16	-	5,96	-	0,41	-
Kesinlik	3,92	-	1,52	-	5,04	-	6,04	-	3,12	-
F ölçütü	2,16	-	3,44	-	0,56	-	6,36	-	0,32	-
Kappa	2,16	-	4,56	-	1,68	-	7,03	-	0,56	-
Belirleyicilik	3,90	-	1,37	-	5,04	-	6,05	-	2,41	-
Negatif öngörü	4,88	-	3,12	-	0,32	-	6,86	-	1,28	-

I: 10 bağımsız değişken II: 7 bağımsız değişken III: 4 bağımsız değişken

Tablo 16'da görüldüğü üzere her bir yönteme ait değerlendirme kriteri değerlerinin bağımsız değişken sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan fark testleri sonucunda, YSA, RO, DVM, SVRA ve LR yöntemlerinin bütün değerlendirme kriterlerine ait değerlerinin bağımsız değişkenin sayısına göre istatistiksel olarak anlamlı biçimde farklılık göstermediği tespit edilmiştir ($p > 0,05$).

Bu bulgulara göre sınıflandırma modelinde bağımlı değişkenle yüksek korelasyona sahip bir bağımsız değişken olması durumunda bağımsız değişken sayısındaki artış veya azalışın doğruluk, duyarlılık, kesinlik, F ölçütü, Kappa, belirleyicilik ve Negatif öngörü değerlerinde anlamlı bir değişime neden olmadığı söylenebilir. Özetle, bağımlı değişkenle yüksek korelasyona sahip bir bağımsız değişken olması durumunda bağımsız değişken sayısı yöntemlerin sınıflama performansında anlamlı bir farklılığa neden olmadığı ifade edilebilir.

Bağımsız değişken sayısının 10, 7 ve 4 olması durumunda YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonrasında elde edilen değerlendirme kriterlerine ait Ek 4'te yer

alan betimsel istatistiklere göre varsayımların incelenmiştir. Bu inceleme sonunda her bir bağımsız değişken sayısında yöntemlerin sınıflandırma performanslarının birbirlerine göre karşılaştırılmasına yönelik uygulanan Friedman Testi sonuçları Tablo 17’de yer almaktadır.

Tablo 17

Her Bir Bağımsız Değişken Sayısında Yöntemlerin Sınıflandırma Performanslarının Birbirlerine Göre Karşılaştırılmasına Yönelik Uygulanan Friedman Testi Sonuçları

Kriter	Bağımsız değişken sayısı					
	10 bağımsız		7 bağımsız		4 bağımsız	
	İst.	Fark	İst.	Fark	İst.	Fark
Doğruluk	81,40*	I, III, V>II>IV	81,71*	I, III, V>II>IV	84,20*	I, III, V>II>IV
Duyarlılık	62,72*	I, III, V>II>IV	67,25*	I, III, V>II>IV	78,96*	I, III, V>II>IV
Kesinlik	82,57*	I, III, V>II>IV	80,13*	I, III, V>II>IV	80,61*	I, III, V>II>IV
F ölçütü	77,56*	I, III, V>II>IV	81,95*	I, III, V>II>IV	84,99*	I, III, V>II>IV
Kappa	81,80*	I, III, V>II>IV	81,18*	I, III, V>II>IV	81,38*	I, III, V>II>IV
Belirleyicilik	83,72*	I, III, V>II>IV	80,35*	I, III, V>II>IV	81,32*	I, III, V>II>IV
Negatif öngörü	71,19*	I, III, V>II>IV	78,62*	I, III, V>II>IV	83,62*	I, III, V>II>IV

I: YSA, II: RO, III: DVM, IV: SVRA, V: LR yöntemlerini temsil etmektedir.
*p<0,05

Tablo 17’de bağımsız değişken sayısının 10 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, bağımsız değişken sayısı 10 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir

$(\chi^2_{Doğruluk} = 81,40, \chi^2_{Duyarlılık} = 62,72, \chi^2_{Kesinlik} = 82,57, \chi^2_{F-ölçütü} = 77,56, \chi^2_{Kappa} = 81,80, \chi^2_{Belirleyicilik} = 83,72, \chi^2_{Negatif\text{öngörü}} = 71,19; sd = 4; n = 25; p < 0,05)$. Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA’na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerin RO yönteminde daha yüksek olduğu belirlenmiştir.

Bağımsız değişken sayısının 7 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını

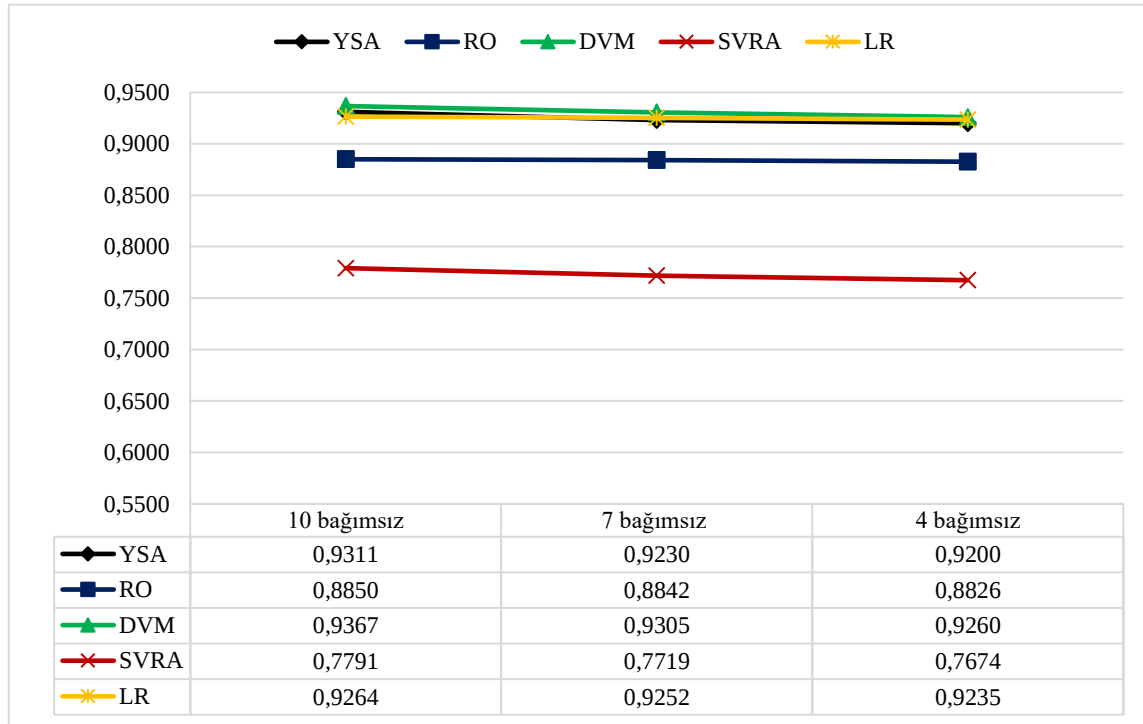
incelemek amacıyla uygulanan Friedman testine ait Tablo 17’de yer alan sonuçlar ele alındığında, bağımsız değişken sayısı 7 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 81,71$, $\chi^2_{Duyarluluk} = 67,25$, $\chi^2_{Kesinlik} = 80,13$, $\chi^2_{F-ölçütü} = 81,95$, $\chi^2_{Kappa} = 81,18$, $\chi^2_{Belirleyicilik} = 80,35$, $\chi^2_{Negatiföngörü} = 78,62$; $sd = 4$; $n = 25$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA’na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerin RO yönteminde daha yüksek olduğu belirlenmiştir.

Tablo 15’te bağımsız değişken sayısının 4 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, bağımsız değişken sayısı 4 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 84,20$, $\chi^2_{Duyarluluk} = 78,96$, $\chi^2_{Kesinlik} = 80,61$, $\chi^2_{F-ölçütü} = 84,99$, $\chi^2_{Kappa} = 81,38$, $\chi^2_{Belirleyicilik} = 81,32$, $\chi^2_{Negatiföngörü} = 83,62$; $sd = 4$; $n = 25$; $p < 0,05$). Farkın

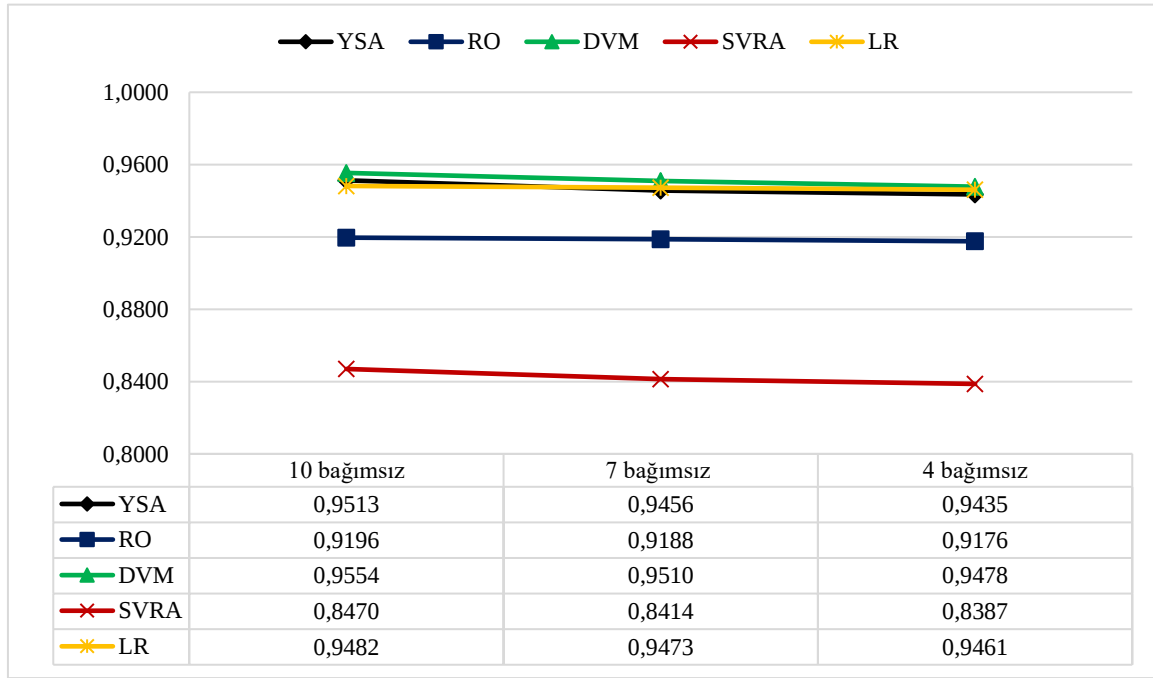
kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA’na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerin RO yönteminde daha yüksek olduğu belirlenmiştir.

Özetle, fark testlerinden elde edilen sonuçlara göre 10, 7 ve 4 bağımsız değişken sayısının olduğu durumlar için bütün değerlendirme kriterlerine göre yöntemler arasında ortaya çıkan farkın kaynağını belirlemek amacıyla uygulanan ikili karşılaştırmalar sonucunda, değerlendirme kriterlerine ait değerlere göre YSA, DVM ve LR yöntemlerinin kendi

aralarında benzer ancak RO ve SVRA yöntemine göre daha iyi sınıflama performansı gösterdiği sonucuna ulaşılmıştır. RO ve SVRA yöntemleri karşılaştırıldığında ise genel olarak RO yönteminin SVRA yöntemine göre daha iyi sınıflama performansına sahip olduğu tespit edilmiştir. Yöntemlerin 10, 7 ve 4 bağımsız değişken sayısı için sahip oldukları doğruluk ve F ölçütü değerlerine ait grafikler sırasıyla Şekil 26 ve Şekil 27’de yer almaktadır. Diğer değerlendirme kriterlerine ait grafikler Ek 8’de yer almaktadır.



Şekil 26. Bağımsız değişken sayısının 10, 7 ve 4 olması durumunda her bir yöntemle ait doğruluk değerleri.



Şekil 27. Bağımsız değişken sayısının 10, 7 ve 4 olması durumunda her bir yöntemle ait F ölçütü değerleri.

4.3. Bağımsız Değişken Sayısının 9, 6 ve 3 Olması Durumuna Göre Elde Edilen Bulgular

Araştırma kapsamında bağımsız değişken sayısının 9, 6 ve 3 olması durumunda YSA, RO, DVM, LR ve SVRA yöntemlerinin uygulanması sonrasında elde edilen değerlendirme kriterlerine ait betimsel istatistikler Ek 5'te yer almaktadır. Betimsel istatistikler doğrultusunda varsayımların incelenmesi doğrultusunda, yöntemlerin sınıflandırma performanslarının kendi içinde bağımsız değişken sayısının değişimine göre değerlendirilmesine yönelik uygulanan Friedman Testi sonuçları Tablo 18'de yer almaktadır.

Tablo 18

Yöntemlerin Sınıflandırma Performanslarının Kendi İçinde Bağımsız Değişken Sayısının Değişimine Göre Değerlendirilmesine Yönelik Uygulanan Friedman Testi Sonuçları

Kriter	Yöntem									
	YSA		RO		DVM		SVRA		LR	
	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark
Doğruluk	41,04*	I>II>II	40,88*	I>II=II	42,00*	I>II>III	34,75*	I>II=III	43,28*	I>II>II
Duyarlılık	36,16*	I>II>III	26,16*	I>II=III	12,24*	I>III	35,25*	I>II=III	22,64*	I=II>III
Kesinlik	44,72*	I>II>III	39,92*	I>II=III	34,16*	I>II=III	5,25	-	38,48*	I>II=III
F ölçütü	38,48*	I>II>III	39,94*	I>II=III	40,88*	I>II>III	36,75*	I>II=III	42,78*	I>II>III
Kappa	48,10*	I>II>III	40,18*	I>II=III	38,48*	I>II>III	36,01*	I>II=III	42,00*	I>II>III
Belirleyicilik	35,28*	I>II>III	45,72*	I>II>III	30,57*	I>II>III	6,75	-	34,08*	I>II>III
Negatif öngörü	41,56*	I>II>III	41,83*	I>II=III	42,27*	I>II>III	35,72*	I>II=III	43,78*	I>II>III
I: 9 bağımsız değişken II: 6 bağımsız değişken III: 3 bağımsız değişken										
*p<0,05										

Tablo 18’de YSA yöntemine ait değerlendirme kriterlerinin bağımsız değişkenin 9, 6 ve 3 olması durumuna göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, YSA yönteminin 9, 6 ve 3 bağımsız değişkenin yer aldığı sınıflama modellerine uygulanması sonucunda tüm değerlendirme kriterlerine ait değerlerin bağımsız değişken sayısına göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 40,04$, $\chi^2_{Duyarlılık} = 36,16$, $\chi^2_{Kesinlik} = 44,72$, $\chi^2_{F-ölçütü} = 38,48$, $\chi^2_{Kappa} = 48,10$, $\chi^2_{Belirleyicilik} = 35,28$, $\chi^2_{Negatif\text{öngörü}} = 41,56$; $sd = 2$; $n = 25$; $p < 0,05$). Farkın kaynağının belirlenmesi için gerçekleştirilen ikili karşılaştırmalar, ele alınan bütün değerlendirme kriterleri için 9, 6 ve 3 bağımsız değişken sayısı durumlarının hepsinin birbiri ile anlamlı olarak farklılaştığını göstermiştir. Buna göre en iyi performansı gösteren model 9 bağımsız değişkenin olduğu ve sonrasında sırasıyla 6 ve 3 bağımsız değişkenin yer aldığı modeller olduğu bulunmuştur.

RO yöntemi için de değerlendirme kriterlerinin hepsinde bağımsız değişken sayısına göre anlamlı fark olduğu görülmektedir ($\chi^2_{Doğruluk} = 40,88$, $\chi^2_{Duyarlılık} = 26,16$, $\chi^2_{Kesinlik} = 39,92$, $\chi^2_{F-ölçütü} = 39,94$, $\chi^2_{Kappa} = 40,18$, $\chi^2_{Belirleyicilik} = 45,72$, $\chi^2_{Negatif\text{öngörü}} = 41,83$; $sd = 2$; $n = 25$; $p < 0,05$). Yapılan ikili karşılaştırmalar sonucu elde edilen bulgular belirleyicilik hariç diğer bütün kriterler için aynı olup 9 bağımsız değişken olan

modelin diğer iki modelden daha yüksek değerlere sahip olduğunu, 6 ve 3 bağımsız değişken durumunda ise modellerin birbirinden anlamlı olarak farklılaşmadığını göstermektedir. Belirleyicilik için ise diğerleri ile benzer şekilde 9 bağımsız değişkenli modelin diğerlerinden yüksek olduğu fakat burada 6 ile 3 bağımsız değişken olan durumda da anlamlı fark olduğu ve 6 bağımsız değişkenli modelin de 3 bağımsız değişkenli modelden daha yüksek belirleyicilik değerine sahip olduğu tespit edilmiştir.

Tablo 18’de DVM yöntemine ait değerlendirme kriterlerinin bağımsız değişken sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, DVM yönteminin 9, 6 ve 3 bağımsız değişkenin yer aldığı sınıflama modellerine uygulanması sonucunda tüm değerlendirme kriterlerine ait değerlerin bağımsız değişkenin sayısının değişimine göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir

($\chi^2_{Doğruluk} = 42,00$, $\chi^2_{Duyarlılık} = 12,24$, $\chi^2_{Kesinlik} = 34,16$, $\chi^2_{F-ölçütü} = 40,88$, $\chi^2_{Kappa} = 38,48$, $\chi^2_{Belirleyicilik} = 30,57$, $\chi^2_{Negatiföngörü} = 42,27$; $sd = 2$; $n = 25$; $p < 0,05$). Farkın kaynağının belirlenmesi için yapılan ikili karşılaştırmalar, doğruluk, F ölçütü, Kappa, belirleyicilik ve negatif öngörü değerleri için farkın bütün karşılaştırmalarda anlamlı olduğunu ve en yüksek değer 9 bağımsız değişkenli durumda elde edildiğini, bunu sırasıyla 6 ve 3 bağımsız değişkenli modelin izlediğini göstermektedir. Kesinlik değeri için 9 bağımsız değişkenli model diğer modellerden daha yüksek değere sahipken 3 ve 6 bağımsız değişkenli modeller arasında kesinlik değeri açısından anlamlı fark bulunmamaktadır. Duyarlılık değeri için ise anlamlı fark sadece 9 ve 3 bağımsız değişkenli modeller arasında olup 9 bağımsız değişkenli modele ait duyarlılık değeri daha yüksektir.

SVRA yöntemine ait değerlendirme kriterlerinin bağımsız değişken sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, SVRA yönteminin 9, 6 ve 3 bağımsız değişkenin yer aldığı sınıflama modellerinde doğruluk, duyarlılık, F ölçütü, Kappa ve negatif öngörü değerleri açısından farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 34,75$, $\chi^2_{Duyarlılık} = 35,25$, $\chi^2_{F-ölçütü} = 36,75$, $\chi^2_{Kappa} = 36,01$, $\chi^2_{Negatiföngörü} = 35,72$; $sd = 2$; $n = 25$; $p < 0,05$). Yapılan

ikili karşılaştırmalar, 9 bağımsız değişkenli modelde elde edilen doğruluk, duyarlılık, F ölçütü, Kappa ve negatif öngörü değerlerinin 6 ve 3 bağımsız değişkenli modelde elde edilen değerlerden anlamlı olarak daha yüksektir. Bağımsız değişken sayısının 6 ve 3 olduğu modeller arasında ise ilgili değerler açısından anlamlı farklılık yoktur. SVRA yöntemi ile yapılan analizler sonucu elde edilen kesinlik ve belirleyicilik değerleri için ise bağımsız değişken sayısına göre anlamlı bir farklılık tespit edilememiştir ($p > 0,05$).

Benzer şekilde LR yöntemi için de bağımsız değişken sayısının değişimine göre doğruluk, duyarlılık, kesinlik, F ölçütü, Kappa, belirleyicilik ve negatif öngörü değerlerinin hepsi için yapılan fark analizlerinin anlamlı olduğu belirlenmiştir ($\chi^2_{Doğruluk} = 43,28$, $\chi^2_{Duyarlılık} = 22,64$, $\chi^2_{Kesinlik} = 38,48$, $\chi^2_{F-ölçütü} = 42,78$, $\chi^2_{Kappa} = 42,00$, $\chi^2_{Belirleyicilik} = 34,08$, $\chi^2_{Negatiföngörü} = 43,78$; $sd = 2$; $n = 25$; $p < 0,05$). Yapılan ikili karşılaştırmalar, doğruluk, F ölçütü, Kappa, belirleyicilik ve negatif öngörü değerleri için 9 bağımsız değişkenli modelin diğer modellerden daha yüksek değerlere sahip olduğunu ve buna ek olarak 6 bağımsız değişkenli modelin de 3 bağımsız değişkenli modelden daha yüksek olduğunu göstermektedir. Duyarlılık değeri için yapılan karşılaştırmalar ele alındığında 9 ve 6 bağımsız değişkenli modellere ait duyarlılık değerlerinin birbirine benzer ve 3 bağımsız değişkenli modele ait duyarlılık değerinden daha yüksek olduğu tespit edilmiştir. Kesinlik değeri için ise 3 ve 6 bağımsız değişkenli modellerden elde edilen kesinlik değerleri birbirine benzer ve 9 bağımsız değişkene sahip modele ait kesinlik değerinden daha düşük olduğu bulunmuştur.

Performans değerlendirme kriterlerine ait değerlerin bağımsız değişken sayısına göre değişimlerinden elde edilen bulgular genel olarak değerlendirildiğinde YSA, RO, DVM, LR ve SVRA yöntemlerinin bağımsız değişken sayısının artması durumunda anlamlı olarak daha yüksek değerler ürettiği dolayısıyla daha iyi bir sınıflama performansı sergilediğini söylenebilir. Ancak YSA, DVM ve LR yöntemlerinin performansları bağımsız değişken sayısı 3'ten büyük olduğunda artış gösterirken RO ve SVRA yöntemlerinde bağımsız değişken sayısı 6'dan büyük olduğunda anlamlı bir artış meydana getirmiştir.

Bağımsız değişken sayısının 9, 6 ve 3 olması durumunda YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonrasında elde edilen değerlendirme kriterlerine ait Ek 5’te yer alan betimsel istatistiklere göre varsayımların incelenmiştir. Bu inceleme sonunda her bir bağımsız değişken sayısında yöntemlerin sınıflandırma performanslarının birbirlerine göre karşılaştırılmasına yönelik uygulanan Friedman Testi sonuçları Tablo 19’da yer almaktadır.

Tablo 19

Her Bir Bağımsız Değişken Sayısında Yöntemlerin Sınıflandırma Performanslarının Birbirlerine Göre Karşılaştırılmasına Yönelik Uygulanan Friedman Testi Sonuçları

Kriter	Bağımsız değişken sayısı					
	9 bağımsız		6 bağımsız		3 bağımsız	
	İst.	Fark	İst.	Fark	İst.	Fark
Doğruluk	80,38*	I, III, V>II>IV	68,11*	I, III, V>II>IV	72,41*	I, III, V>II>IV
Duyarlılık	40,89*	I, III, V>II>IV	60,82*	I, III, V>II>IV	67,09*	I, III, V>II>IV
Kesinlik	61,71*	I, III, V>II>IV	51,57*	I, III, V>II=IV	58,44*	I, III, V>II=IV
F ölçütü	77,07*	I, III, V>II>IV	66,28*	I, III, V>II>IV	70,68*	I, III, V>II>IV
Kappa	68,94*	I, III, V>II>IV	73,38*	I, III, V>II>IV I>V	63,14*	I, III, V>II>IV
Belirleyicilik	39,54*	I, III, V>II=IV	35,53*	I, III, V>II=IV	35,45*	I, III, V>II=IV
Negatif öngörü	77,97*	I, III, V>II>IV	66,54*	I, III, V>II>IV	63,44*	I, III, V>II>IV

I: YSA, II: RO, III: DVM, IV: SVRA, V: LR yöntemlerini temsil etmektedir.
*p<0,05

Tablo 19’da bağımsız değişken sayısının 9 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, bağımsız değişken sayısı 9 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} =$

$$80,38, \chi^2_{Duyarlılık} = 40,89, \chi^2_{Kesinlik} = 61,71, \chi^2_{F-ölçütü} = 77,07, \chi^2_{Kappa} = 68,94,$$

$$\chi^2_{Belirleyicilik} = 39,54, \chi^2_{Negatiföngörü} = 77,97; sd = 2; n = 25; p < 0,05).$$

Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA yöntemlerine ait ikili karşılaştırmalar incelendiğinde ise belirleyicilik hariç tüm değerlendirme kriterlerine ait değerlerin RO yönteminde daha yüksek olduğu belirlenmiştir.

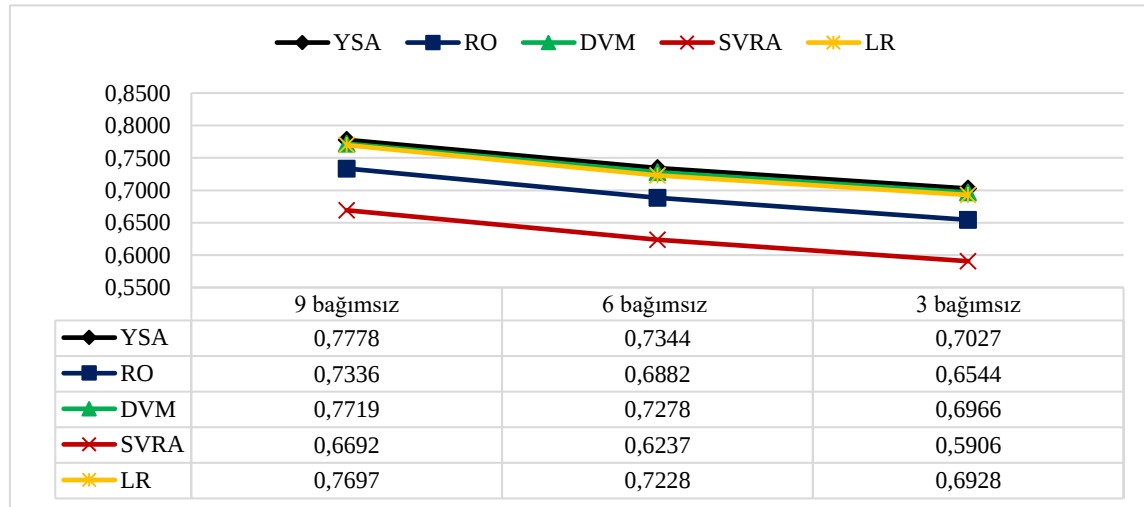
Belirleyicilik değeri için ise RO ve SVRA yöntemlerine ait belirleyicilik değerleri arasında anlamlı bir farklılık yoktur.

Değişken sayısının 6 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, bağımsız değişken sayısı 6 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 68,11$, $\chi^2_{Duyarluluk} = 60,82$, $\chi^2_{Kesinlik} = 51,57$, $\chi^2_{F-ölçütü} = 66,28$, $\chi^2_{Kappa} = 73,38$, $\chi^2_{Belirleyicilik} = 35,53$, $\chi^2_{Negatiföngörü} = 66,54$; $sd = 2$; $n = 25$; $p < 0,05$). Yapılan ikili karşılaştırmalar, 6 bağımsız değişken durumunda YSA, DVM ve LR yöntemlerine ait bütün değerlendirme kriterlerinin RO ve SVRA yöntemlerinden daha yüksek olduğunu göstermektedir. RO ve SVRA yöntemlerine ait değerler birbirine göre incelendiğinde ise RO yöntemine ait doğruluk, duyarlılık, F ölçütü, Kappa ve negatif öngörü değerlerinin SVRA yöntemi ile elde edilen değerlerden anlamlı olarak daha yüksek olduğu belirlenmiştir. RO ve SVRA yöntemleri için hesaplanan kesinlik ve belirleyicilik değerlerinde ise anlamlı bir farklılık yoktur. Ayrıca YSA yöntemi için elde edilen Kappa değerinin, LR yöntemi için hesaplanan Kappa değerinden anlamlı olarak daha yüksek olduğu belirlenmiştir.

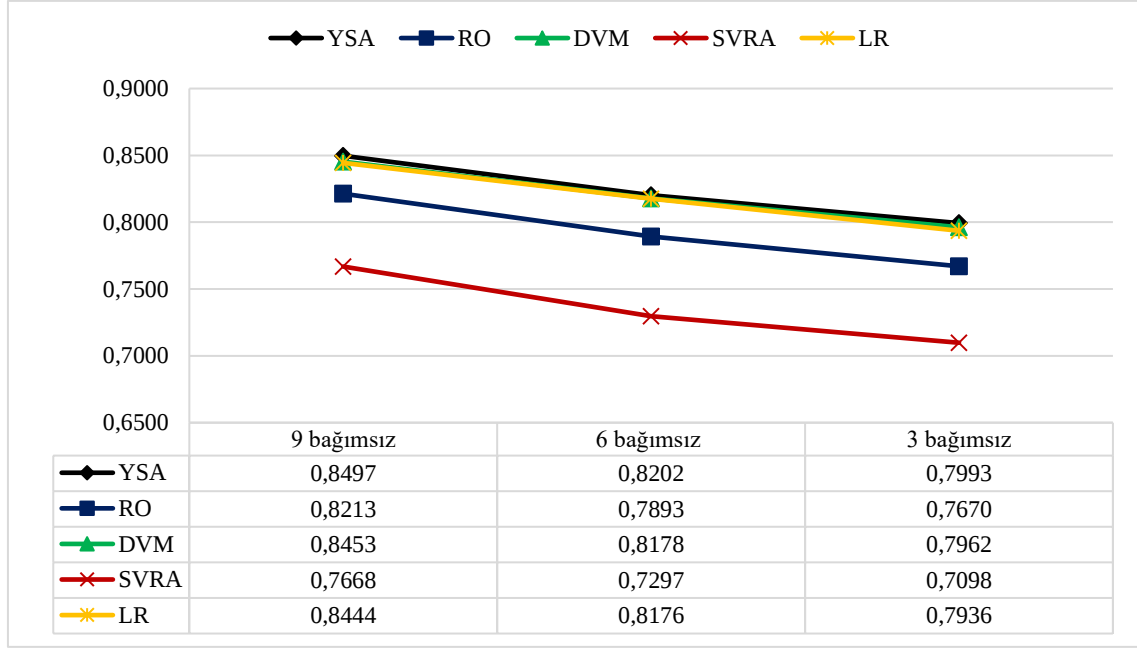
Bağımsız değişken sayısının 3 olması durumunda ise, diğerleri ile benzer şekilde ele alınan bütün kriterler için yöntemler arası yapılan fark analizi sonuçlarının anlamlı olduğu bulunmuştur ($\chi^2_{Doğruluk} = 72,41$, $\chi^2_{Duyarluluk} = 67,09$, $\chi^2_{Kesinlik} = 58,44$, $\chi^2_{F-ölçütü} = 70,68$, $\chi^2_{Kappa} = 63,14$, $\chi^2_{Belirleyicilik} = 35,45$, $\chi^2_{Negatiföngörü} = 63,44$; $sd = 2$; $n = 25$; $p < 0,05$). Farkın hangi gruplar arasında olduğunun belirlenmesi için yapılan ikili karşılaştırmalar, ele alınan tüm kriterler için YSA, DVM ve LR yöntemleri için hesaplanan değerlerin birbiri ile benzer, RO ve SVRA yöntemleri için elde edilen değerlerden daha yüksek olduğunu göstermektedir. RO ve SVRA yöntemleri birbirine göre değerlendirildiğinde ise RO yöntemi için hesaplanan doğruluk, duyarlılık, F ölçütü, Kappa

ve negatif öngörü değerlerinin SVRA yöntemi ile elde edilenden daha yüksektir. RO için elde edilen kesinlik ve belirleyicilik değerleri ise SVRA ile elde edilen değerlerden anlamlı olarak farklılaşmamaktadır.

Özetle, fark testlerinden elde edilen sonuçlara göre her bir bağımsız değişken sayısı için bütün değerlendirme kriterlerine göre yöntemler arasında ortaya çıkan farkın kaynağını belirlemek amacıyla uygulanan ikili karşılaştırmalar sonucunda, değerlendirme kriterlerine ait değerlere göre YSA, DVM ve LR yöntemlerinin kendi aralarında benzer ancak RO ve SVRA yöntemine göre daha iyi değerler ürettiği sonucuna ulaşılmıştır. RO ve SVRA yöntemleri karşılaştırıldığında ise RO yönteminin SVRA'a göre daha iyi sonuçlar ürettiği tespit edilmiştir. Yöntemlerin 9, 6 ve 3 bağımsız değişken sayısı için sahip oldukları doğruluk ve F ölçütü değerlerine ait grafikler sırasıyla Şekil 28 ve Şekil 29'da yer almaktadır. Diğer değerlendirme kriterlerine ait grafikler Ek 9'da yer almaktadır.



Şekil 28. Bağımsız değişken sayısının 9, 6 ve 3 olması durumunda her bir yönteme ait doğruluk değerleri.



Şekil 29. Bağımsız değişken sayısının 9, 6 ve 3 olması durumunda her bir yöntemle ait F ölçütü değerleri.

4.4. Örneklem Büyüklüğünün Değişimine Göre Elde Bulgular

Araştırma kapsamında örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması durumunda YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonrasında elde edilen karşılaştırma kriterlerine ait betimsel istatistikler Ek 6'da yer almaktadır. Betimsel istatistiklere göre varsayımların incelenmesi doğrultusunda, her bir yöntemin sınıflama performanslarının örneklem büyüklüğünün değişimine göre değerlendirilmesine yönelik uygulanan Kruskal Wallis Testi sonuçları Tablo 20'de yer almaktadır.

Tablo 20

Yöntemlerin Sınıflandırma Performanslarının Kendi İçinde Örneklem Büyüklüğünün Değişimine Göre Değerlendirilmesine Yönelik Uygulanan Kruskal Wallis Testi Sonuçları

Kriter	Yöntem									
	YSA		RO		DVM		SVRA		LR	
	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark
Doğruluk	87,83*	VI, V, IV, III > II, I	88,71*	VI, V, IV > III, II, I	69,66*	VI, V, IV, III > II, I	35,00*	VI, V, IV> III, II, I	64,77*	VI, V, IV, III > II, I
Duyarlılık	58,30*	VI, V, IV, III > II, I	63,87	VI, V, IV> III, II, I	68,99*	VI, V, IV, III > II, I	67,05*	VI, V> IV, III, II, I	28,84*	VI, V, IV, III > II, I
Kesinlik	75,76*	VI, V, IV, III > II, I	85,61*	VI, V, IV> III, II, I	73,93*	VI, V, III > I, II IV > II	33,10*	IV, VI > I, II, V	45,64*	VI, V, IV, III > II, I
F ölçütü	94,36*	VI, V, IV, III > II, I	93,23*	VI, V, IV > III, II, I	75,24*	VI, V, IV, III > II, I	47,31*	VI, V, IV> III, II, I	64,02*	VI, V, IV, III > II, I
Kappa	83,10*	VI, V, IV, III > II, I	84,15*	VI, V, IV> III, II, I	68,99*	VI, V, IV, III > II, I	27,90*	V, IV> VI, III, II, I	59,32*	VI, V, IV, III > II, I
Belirleyicilik	57,75*	VI, V, IV, III > II, I	64,57*	VI, V, IV>I> III, II	65,05*	VI, V, IV, III > I>II	14,03*	IV > V, VI	20,83*	VI, IV > I, II III, V>II
Negatif öngörü	55,29*	VI, V, IV, III > II, I	74,09*	VI, V, IV> III, II, I	66,75*	VI, V, IV, III > II, I	46,87*	VI, V, IV > III VI, V>II, I	27,32*	VI, V, IV, III > II, I

I: 100, II: 250, III: 500, IV: 1000, V: 2500 ve VI: 5000 örneklem büyüklüğünü temsil etmektedir.
*p<,05

Tablo 20’de YSA yöntemine ait değerlendirme kriterlerinin örneklem büyüklüğüne göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Kruskal Wallis testine ait sonuçlar ele alındığında, YSA yönteminin farklı örneklem büyüklüğünün yer aldığı sınıflama modellerine uygulanması sonucunda tüm değerlendirme kriterlerine ait değerlerin örneklem büyüklüğünün değişimine göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 87,83$, $\chi^2_{Duyarlılık} = 58,30$, $\chi^2_{Kesinlik} = 75,76$, $\chi^2_{F-ölçütü} =$

94,36, $\chi^2_{Kappa} = 83,10$, $\chi^2_{Belirleyicilik} = 57,75$, $\chi^2_{Negatif\ddot{ö}ng\ddot{ö}r\ddot{u}} = 55,29$; $sd = 2$; $n = 150$; $p < 0,05$). Farkın kaynağının belirlenmesi için gerçekleştirilen ikili karşılaştırmalar, 500, 1000, 2500 ve 5000 örneklem büyüklüğünde elde edilen değerlerin birbiri ile benzerken 100 ve 250 örneklem için elde edilen değerlerden anlamlı olarak yüksek olduğunu göstermektedir. 100 ve 250 örneklem için hesaplanan bütün performans değerlendirme kriterleri de birbirleri ile benzerdir.

RO yöntemi için de değerlendirme kriterlerinin hepsinde örneklem büyüklüğüne göre anlamlı fark olduğu görülmektedir ($\chi^2_{Doğruluk} = 88,71$, $\chi^2_{Duyarluluk} = 63,87$, $\chi^2_{Kesinlik} = 85,61$, $\chi^2_{F-\ddot{ö}lçütü} = 93,23$, $\chi^2_{Kappa} = 84,15$, $\chi^2_{Belirleyicilik} = 64,57$, $\chi^2_{Negatif\ddot{ö}ng\ddot{ö}r\ddot{u}} = 74,09$; $sd = 2$; $n = 150$; $p < 0,05$). Yapılan ikili karşılaştırmalar sonucu elde edilen bulgular 1000, 2500 ve 5000 örneklem büyüklüğünde elde edilen değerlerin birbiri ile benzerken 100, 250 ve 500 örneklem için elde edilen değerlerden anlamlı olarak yüksek olduğunu göstermektedir. 100, 250 ve 500 örneklem için hesaplanan belirleyicilik hariç bütün performans değerlendirme kriterleri de birbirleri ile benzerdir. Belirleyicilik değeri için ise 100 ve 250 örneklem birbirine benzerken 500 örneklem ile hesaplanan belirleyicilik değerinden daha düşük belirleyicilik değerine sahiptir.

Tablo 20’de DVM yöntemine ait değerlendirme kriterlerinin örneklem büyüklüğüne göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Kruskal Wallis testine ait sonuçlar ele alındığında, DVM yönteminin farklı örneklem sayılarının yer aldığı sınıflama modellerine uygulanması sonucunda tüm değerlendirme kriterlerine ait değerlerin örneklem sayısının değişimine göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir

($\chi^2_{Doğruluk} = 69,66$, $\chi^2_{Duyarluluk} = 68,99$, $\chi^2_{Kesinlik} = 73,93$, $\chi^2_{F-\ddot{ö}lçütü} = 75,24$, $\chi^2_{Kappa} = 68,99$, $\chi^2_{Belirleyicilik} = 65,05$, $\chi^2_{Negatif\ddot{ö}ng\ddot{ö}r\ddot{u}} = 66,75$; $sd = 2$; $n = 150$; $p < 0,05$). Farkın kaynağının belirlenmesi için yapılan ikili karşılaştırmalar, 500, 1000, 2500 ve 5000 örneklem büyüklükleri için elde edilen doğruluk, duyarlılık, F ölçütü, Kappa, belirleyicilik ve negatif öngörü değerlerinin birbirine benzerken 100 ve 250 örneklem

büyüklikleri için hesaplanan değerlerden anlamlı olarak daha yüksek olduğunu göstermektedir. 250 örnekleme ait belirleyicilik değeri 100 örneklem için hesaplanan değerden daha yüksekken diğer kriterler için bu iki örneklem büyüklüğüne ait değerler arasında anlamlı fark yoktur. Kesinlik değeri için yapılan ikili karşılaştırmalar incelendiğinde 500, 2500 ve 5000 örneklem büyüklüklerine ait kesinlik değerlerinin birbirine benzerken 100 ve 250 örneklem büyüklükleri için hesaplananlardan anlamlı olarak daha yüksek olduğu tespit edilmiştir. Ortaya çıkan bir diğer fark da 1000 örneklemden elde edilen kesinlik değerinin 250 örneklemden elde edilende daha yüksek olmasıdır.

SVRA yöntemine ait değerlendirme kriterlerinin örneklem sayısına göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Kruskal Wallis testine ait sonuçlar ele alındığında, SVRA yönteminin farklı örneklem büyüklüklerinin yer aldığı sınıflama modellerinde bütün performans değerlendirme kriterleri açısından farklılık gösterdiği tespit edilmiştir

($\chi^2_{Doğruluk} = 35,00$, $\chi^2_{Duyarlılık} = 67,05$, $\chi^2_{Kesinlik} = 33,10$, $\chi^2_{F-ölçütü} = 47,31$, $\chi^2_{Kappa} = 27,90$, $\chi^2_{Belirleyicilik} = 14,03$, $\chi^2_{Negatiföngörü} = 46,87$; $sd = 2$; $n = 150$; $p < 0,05$). Yapılan ikili karşılaştırmalar, 1000, 2500 ve 5000 örneklem büyüklüğü için hesaplanan doğruluk, belirleyicilik ve F ölçütü değerlerinin birbiri ile benzerken 100, 250 ve 500 örneklemden elde edilen değerlerden daha yüksektir. Kesinlik kriteri için 1000 ve 5000 örneklem büyüklüğüne ait değerler 100, 250 ve 2500 örneklem ile elde edilenden daha yüksektir. Kappa kriteri için farklar 1000 ve 2500 örneklem büyüklüğüne ait değerlerin diğerlerinden daha yüksek olduğu şeklindedir. Belirleyicilik için tek fark 1000 ile 2500 ve 5000 arasındadır ve 1000 örneklem büyüklüğü lehinedir. 1000, 2500 ve 5000 örneklem büyüklüğü için elde edilen negatif öngörü değeri 500 örneklem için hesaplanan değerden daha yüksek ve 2500 ve 5000 için hesaplanan değerler de 100 ve 250 için hesaplanan negatif öngörü değerlerinden anlamlı olarak daha yüksektir.

Benzer şekilde LR yöntemi için de örneklem büyüklüğünün değişimine göre doğruluk, duyarlılık, kesinlik, F ölçütü, Kappa, belirleyicilik ve negatif öngörü değerlerinin hepsi için yapılan fark analizlerinin anlamlı olduğu belirlenmiştir ($\chi^2_{Doğruluk} = 64,77$, $\chi^2_{Duyarlılık} =$

28,84, $\chi^2_{Kesinlik} = 45,64$, $\chi^2_{F-ölçütü} = 64,02$, $\chi^2_{Kappa} = 59,32$, $\chi^2_{Belirleyicilik} = 20,83$, $\chi^2_{Negatiföngörü} = 27,32$; $sd = 2$; $n = 150$; $p < 0,05$). Yapılan ikili karşılaştırmalar, 500, 1000, 2500 ve 5000 örneklem için hesaplanan doğruluk, duyarlılık, kesinlik, F ölçütü, Kappa ve negatif öngörü değerleri birbiri ile benzerken 100 ve 250 örneklem için hesaplanan değerlerden anlamlı olarak daha yüksektir. Belirleyicilik kriteri için 1000 ve 5000 örnekleme ait değer 100 ve 250 için hesaplanan değerden daha yüksektir. Ayrıca 500 ve 2500 için hesaplanan belirleyicilik değeri, 250 örneklem büyüklüğü için hesaplanan değerden anlamlı olarak daha yüksektir.

Özetle, performans değerlendirme kriterlerine ait değerlerin örneklem büyüklüğüne göre değişimlerinden elde edilen bulgular genel olarak değerlendirildiğinde YSA, DVM ve LR yöntemlerinin 500 ve daha fazla örneklem büyüklüğünde 100 ve 250 örneklem büyüklüğüne göre daha yüksek ve birbirine benzer değerler ürettiği dolayısıyla daha iyi bir sınıflama performansı sergilediğini söylenebilir. RO ve SVRA yöntemlerine ait değerler incelendiğinde ise örneklem büyüklüğün arttığı durumlarda çoğu performans kriteri olumlu yönde artış göstermiştir. Ancak bu artış YSA, DVM ve LR'dan farklı olarak 1000 ve daha büyük örneklemelerde gerçekleşmeye başlamıştır. Öte yandan örneklem büyüklüğü 500, 1000, 2500 ve 5000 olduğu durumlarda performans değerlendirme kriterlerine ait değerlerin anlamlı olarak bir farklılık göstermemesi, belirli bir örneklem büyüklüğüne ulaşıldıktan sonra örneklem büyüklüğünde meydana gelecek artışların yöntemlerin performansında bir farklılaşmaya sebep olmadığı şeklinde ifade edilebilir.

Örneklem büyüklüklerinin 100, 250, 500, 1000, 2500 ve 5000 olması durumunda YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonrasında elde edilen karşılaştırma kriterlerine ait Ek 6'da yer alan betimsel istatistiklere göre varsayımların incelenmiştir. Bu inceleme sonunda her bir örneklem büyüklüğünde yöntemlerin sınıflandırma performanslarının birbirlerine göre karşılaştırılmasına yönelik uygulanan Friedman Testi sonuçları Tablo 21'de yer almaktadır.

Tablo 21

Her Bir Örneklem Büyüklüğünde Yöntemlerin Sınıflandırma Performanslarının Birbirlerine Göre Karşılaştırılmasına Yönelik Uygulanan Friedman Testi Sonuçları

Kriter	Örneklem büyüklüğü											
	100		250		500		1000		2500		5000	
	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark	İst.	Fark
Doğruluk	39,13*	I, III, V >II, IV	42,25*	I, III, V >II, IV	81,89*	I, III, V >II, IV	79,10*	I, III, V >II>IV	81,28*	I, III, V >II>IV	81,40*	I, III, V >II>IV
Duyarlılık	32,99*	I, III, V >II, IV	38,38*	I, III, V >II, IV	81,24*	I, III, V >II>IV	65,10*	I, III, V >II>IV	65,36*	I, III, V >II>IV	62,72*	I, III, V >II>IV
Kesinlik	19,34*	III, V >IV	38,98*	I, III, V >II, IV	77,54*	I, III, V >II, IV	71,97*	I, III, V >II>IV	69,48*	I, III, V >II>IV	82,57*	I, III, V >II>IV
F ölçütü	30,00*	III, V >II, IV I>IV	40,87*	I, III, V >II, IV	82,02*	I, III, V >II, IV	80,80*	I, III, V >II>IV	80,87*	I, III, V >II>IV	77,56*	I, III, V >II>IV
Kappa	29,94*	I, III, V >II, IV	41,10*	I, III, V >II, IV	79,78*	I, III, V >II, IV	79,13*	I, III, V >II>IV	80,34*	I, III, V >II>IV	81,40*	I, III, V >II>IV
Belirleyicilik	17,42*	I, III, V >II, IV	34,71*	I, III, V >II, IV	76,55*	I, III, V >II, IV	69,85*	I, III, V >II>IV	67,22*	I, III, V >II>IV	83,72*	I, III, V >II>IV
Negatif öngörü	43,52*	I, III, V >II, IV	39,39*	I, III, V >II, IV	81,95*	I, III, V >II>IV	71,98*	I, III, V >II>IV	67,68*	I, III, V >II>IV	71,19*	I, III, V >II>IV

I: YSA, II: RO, III: DVM, IV: SVRA, V: LR yöntemlerini temsil etmektedir.
*p<0,05

Tablo 21’de örneklem büyüklüğünün 100 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, örneklem büyüklüğü 100 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 39,13$, $\chi^2_{Duyarlılık} = 32,99$, $\chi^2_{Kesinlik} = 19,34$, $\chi^2_{F-ölçütü} = 30,00$, $\chi^2_{Kappa} = 29,94$, $\chi^2_{Belirleyicilik} = 17,42$, $\chi^2_{Negatif\ öngörü} = 43,52$; $sd = 4$; $n = 25$; $p < 0,05$). Doğruluk, duyarlılık, Kappa, belirleyicilik ve negatif öngörü değerlerine ait elde edilen farkın

kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerine ait bu değerlerin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise doğruluk, duyarlılık, Kappa, belirleyicilik ve negatif öngörü kriterlerine ait değerlerin benzer olduğu belirlenmiştir. Kesinlik değerine ait farkın kaynağı incelendiğinde DVM ve LR yöntemlerinde bu değerlerin SVRA yöntemine göre daha yüksek olduğu sonucuna ulaşılmıştır. F-ölçütü değeri bakımından ise farkın kaynağı YSA, DVM ve LR yöntemlerinde bu değerlerin SVRA yöntemlerine göre ayrıca DVM ile LR yöntemlerinde RO yöntemine göre yüksek olmasından kaynaklandığı tespit edilmiştir. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerin 100 örneklem büyüklüğünde benzer olduğu belirlenmiştir.

Örneklem büyüklüğünün 250 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait Tablo 21'de yer alan sonuçlar ele alındığında, örneklem büyüklüğü 250 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} =$

$$42,25, \chi^2_{Duyarluluk} = 38,38, \chi^2_{Kesinlik} = 38,98, \chi^2_{F-ölçütü} = 40,87, \chi^2_{Kappa} = 41,10,$$

$$\chi^2_{Belirleyicilik} = 34,71, \chi^2_{Negatif\text{öngörü}} = 39,39; sd = 4; n = 25; p < 0,05). \quad \text{Farkın}$$

kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerin 250 örneklem büyüklüğünde benzer olduğu belirlenmiştir.

Tablo 21'de örneklem büyüklüğünün 500 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele alındığında, örneklem büyüklüğü 500 iken tüm değerlendirme kriterlerine ait değerlerin

yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 81,99$, $\chi^2_{Duyarluluk} = 81,24$, $\chi^2_{Kesinlik} = 77,54$, $\chi^2_{F-ölçütü} = 82,02$, $\chi^2_{Kappa} = 79,78$, $\chi^2_{Belirleyicilik} = 76,55$, $\chi^2_{Negatiföngörü} = 81,95$; $sd = 4$; $n = 25$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise duyarlılık ve negatif öngörü değerlerinde RO yönteminin daha yüksek değerlere sahip olduğu ancak geri kalan değerlendirme kriterlerinde RO ve SVRA yöntemlerine ait değerlerin benzer olduğu belirlenmiştir.

Örneklem büyüklüğünün 1000 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait Tablo 21'de yer alan sonuçlar ele alındığında, örneklem büyüklüğü 1000 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 79,10$, $\chi^2_{Duyarluluk} = 65,10$, $\chi^2_{Kesinlik} = 71,97$, $\chi^2_{F-ölçütü} = 80,80$, $\chi^2_{Kappa} = 79,13$, $\chi^2_{Belirleyicilik} = 69,85$, $\chi^2_{Negatiföngörü} = 71,98$; $sd = 4$; $n = 25$; $p < 0,05$).

Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerde RO yönteminin SVRA yöntemine göre daha yüksek değerlere sahip olduğu belirlenmiştir.

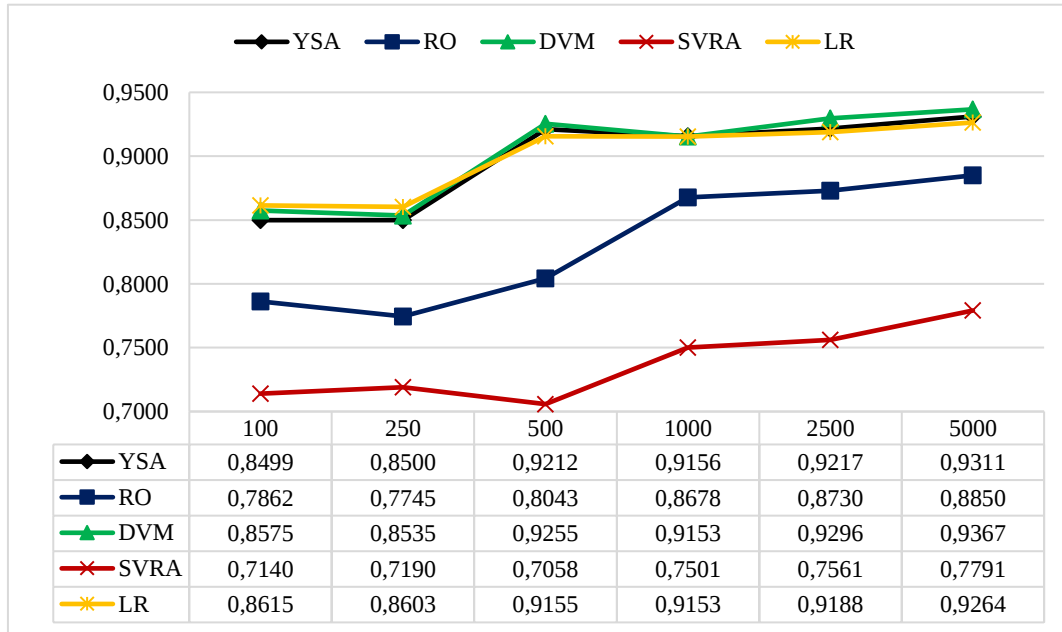
Tablo 21'de örneklem büyüklüğünün 2500 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait sonuçlar ele

alındığında, örneklem büyüklüğü 2500 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 81,28$, $\chi^2_{Duyarluluk} = 65,36$, $\chi^2_{Kesinlik} = 69,48$, $\chi^2_{F-ölçütü} = 80,87$, $\chi^2_{Kappa} = 80,34$, $\chi^2_{Belirleyicilik} = 67,22$, $\chi^2_{Negatiföngörü} = 67,68$; $sd = 4$; $n = 25$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerde RO yönteminin SVRA yöntemine göre daha yüksek değerlere sahip olduğu belirlenmiştir.

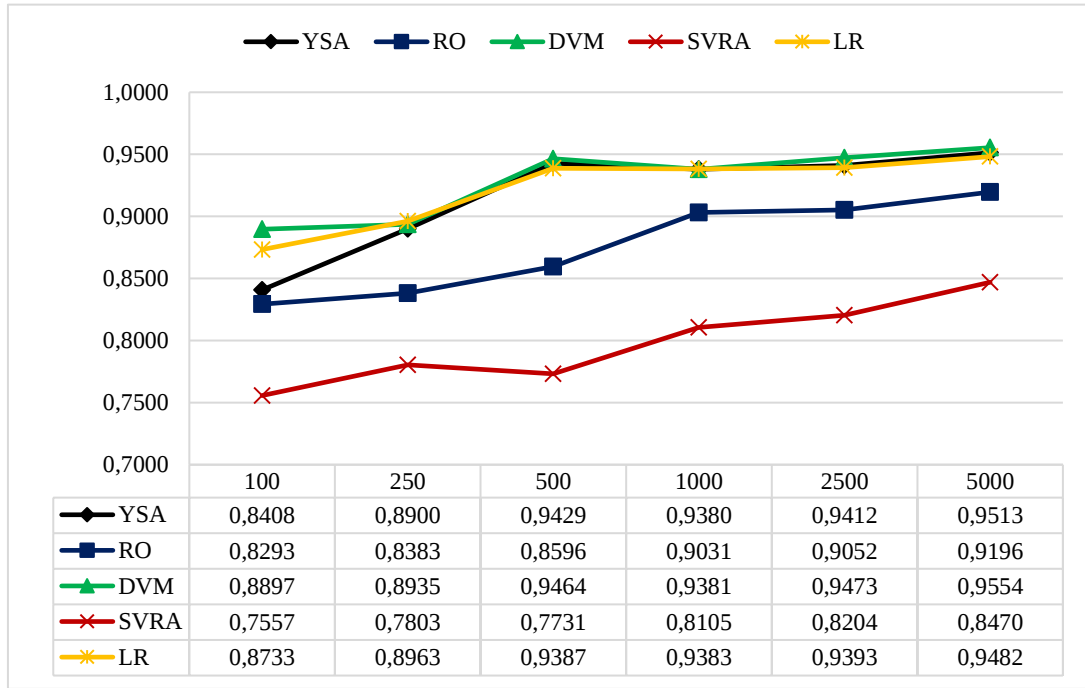
Örneklem büyüklüğünün 5000 olması durumunda sınıflama modellerine ait değerlendirme kriterlerinin YSA, RO, DVM, SVRA ve LR yöntemlerine göre farklılaşıp farklılaşmadığını incelemek amacıyla uygulanan Friedman testine ait Tablo 21'de yer alan sonuçlar ele alındığında, örneklem büyüklüğü 5000 iken tüm değerlendirme kriterlerine ait değerlerin yöntemlere göre istatistiksel olarak anlamlı farklılık gösterdiği tespit edilmiştir ($\chi^2_{Doğruluk} = 81,40$, $\chi^2_{Duyarluluk} = 62,72$, $\chi^2_{Kesinlik} = 82,57$, $\chi^2_{F-ölçütü} = 77,56$, $\chi^2_{Kappa} = 81,40$, $\chi^2_{Belirleyicilik} = 83,72$, $\chi^2_{Negatiföngörü} = 71,19$; $sd = 4$; $n = 25$; $p < 0,05$). Farkın kaynağını belirlemek amacıyla yapılan ikili karşılaştırmalar incelendiğinde YSA, DVM ve LR yöntemlerinin tüm değerlendirme kriterlerine ait değerlerinin birbirine benzer ancak RO ve SVRA yöntemlerine ait değerlere göre daha yüksek olduğu sonucuna ulaşılmıştır. RO ve SVRA'na ait ikili karşılaştırmalar incelendiğinde ise tüm değerlendirme kriterlerine ait değerlerde RO yönteminin SVRA yöntemine göre daha yüksek değerlere sahip olduğu belirlenmiştir.

Özetle, fark testlerinden elde edilen sonuçlara göre her bir örneklem büyüklüğü için bütün değerlendirme kriterlerine göre yöntemler arasında ortaya çıkan farkın kaynağını belirlemek amacıyla uygulanan ikili karşılaştırmalar sonucunda, değerlendirme kriterlerine ait değerlere

göre YSA, DVM ve LR yöntemlerinin kendi aralarında benzer ancak RO ve SVRA yöntemine göre daha iyi değerler ürettiği sonucuna ulaşılmıştır. RO ve SVRA yöntemleri karşılaştırıldığında ise bu iki yöntemin bütün değerlendirme kriterlerine göre genel olarak örneklem büyüklüğünün 100, 250 ve 500 olduğu durumlarda benzer ancak diğer örneklem büyüklüklerinde RO yönteminin SVRA yöntemine göre daha iyi sonuçlar ürettiği tespit edilmiştir. Yöntemlerin farklı örneklem büyüklükleri için sahip oldukları doğruluk ve F ölçütü değerlerine ait grafikler sırasıyla Şekil 30 ve Şekil 31’de yer almaktadır. Diğer değerlendirme kriterlerine ait grafikler Ek 10’da yer almaktadır.



Şekil 30. Farklı örneklem büyüklükleri koşulunda her bir yönteme ait doğruluk değerleri.



Şekil 31. Farklı örneklem büyüklükleri koşulunda her bir yöntem için F ölçütü değerleri.

BÖLÜM 5

TARTIŞMA, SONUÇ VE ÖNERİLER

5.1. Tartışma

Bu bölümde çalışma bulguları ışığında elde edilen bulgular yorumlanmış ve ilgili literatür ile birlikte tartışılmıştır. Çalışmada ele alınan koşullar, PISA 2015 öğrenci veri setinden elde edilen bağımlı değişkenin kategori sayısı, örneklem büyüklüğü ve bağımsız değişken sayısıdır. Çalışma kapsamını oluşturan bütün durumlar için ele alınan yöntemlerin (YSA, RO, LR, SVRA ve DVM) hepsi ile analizler gerçekleştirilmiş ve hem koşulların kendi içindeki durumları hem de yöntemlerin birbirine göre performansları arasındaki farklar istatistiksel olarak incelenmiştir.

Bu çalışmada ele alınan bağımlı değişken, PISA 2015 veri setinde yer alan 10 adet fen başarı puanının ortalaması alındıktan sonra PISA raporundaki düzeylere göre belirlenmiş 2, 3 ve 6 sınıfa sahip kategorik fen başarıları değişkenidir. Bağımsız değişkenler ise yapılan korelasyon analizi sonucunda ortalama fen başarıları ile ilişkisi tespit edilen 10 bağımsız değişkendir. Bağımsız değişkenlerin belirlenmesinden sonra, 10 bağımsız değişken ve 5000 öğrencinin yer aldığı çalışma grubu koşulunda, modelde yer alan bağımlı değişkenin kategori sayısının (2, 3 ve 6) farklılaşması durumuna göre YSA, RO, SVRA, DVM ve LR yöntemleri ile analizler yapılmıştır. Araştırma kapsamında elde edilen bulgular incelendiğinde bağımlı değişkenin kategori sayısı azaldıkça ele alınan bütün yöntemlerin sınıflama performanslarının çoğu değerlendirme kriterine göre artış gösterdiği tespit edilmiştir. Bu durumda 10 bağımsız değişken ve 5000 öğrencinin yer aldığı çalışma grubu koşulları altında en iyi sınıflama performansı, bağımlı değişkenin 2 kategorili olması durumunda elde

edilmiştir. Performansların değerlendirilmesi amacıyla araştırma kapsamında ele alınan değerlendirme kriterlerinden Kappa istatistiği ve doğru sınıflama oranları güvenilirlik ölçütü, geri kalan değerlendirmeye kriterleri ise geçerlilik ölçütü olarak tanımlanmaktadır (Aksu & Doğan, 2018). Bu bağlamda bütün yöntemlerin en iyi sınıflandırma performansını bağımlı değişkenin 2 kategorili olması durumunda göstermesi, ele alınan değerlendirme kriterlerine bağlı olarak yöntemlerin en başarılı güvenilir ve geçerli sonuçları da bu koşulda ürettiği anlamına gelmektedir. Bu sonuca göre, bağımlı değişkenin ikili veya çoklu olmasının sınıflama algoritmalarının performansını dolayısıyla güvenilirlik ve geçerliliği anlamlı olarak etkilediği söylenebilir. Çelikten ve Çakan'a (2019) göre sınıflama performansları ele alınan sınıf sayısından doğrudan etkilenmektedir ve sınıf sayısı arttıkça sınıflama doğruluğu ve tutarlılığı azalmaktadır. Minaei-Bidgoli vd (2003) de web-tabanlı eğitim sistemi için veri madenciliği yöntemlerini uyguladıkları çalışmalarında 2, 3 ve 9 kategorili bağımlı değişken olduğu durumda karşılaştırma yapmış ve kategori sayısı arttıkça sınıflama performansının azaldığını tespit etmişlerdir. Benzer şekilde, Cortez ve Silva (2008) matematik başarısını 2 ve 5 sınıflı olarak ele almış ve 2 sınıflı durumda veri madenciliği yöntemlerinin doğru sınıflandırma yüzdelerinin daha yüksek olduğunu bulmuşlardır. Bir başka çalışmada, karar ağaçları ile yapılan uygulamalar sonucunda sınıf sayısı arttıkça yöntemlerin sınıflandırma performanslarına ait başarılarının düştüğü sonucuna ulaşılmıştır. Büyükkışık (2014, s.57) da yöntemlerin sınıflandırma performanslarını bağımlı değişkenin 2, 3 ve 5 kategori olması durumuna göre karşılaştırmış ve tüm yöntemlerin en başarılı performansını bağımlı değişkenin 2 kategori olması halinde gerçekleştirdiği sonucuna ulaşmıştır.

Öte yandan araştırma kapsamında bağımlı değişkenin kategori sayısının artması durumunda, örnekleme yer alan bireylerin bu kategorilere olan dağılımlarında oransal olarak dengesizlikler olduğuna rastlanmıştır. Bu durumun bağımlı değişkenin kategori sayısının artması durumunda yöntemlerin sınıflama performanslarının düşmesine neden olduğu düşünülmektedir. Verinin dengesiz olması bireylerin sınıflardaki olan dağılım oranlarının aşırı derecede farklılaşmasıdır ve bütün sınıflama algoritmalarında verinin dengesizliği performansı etkiler. Bağımlı değişkenin kategori sayısının ikiden fazla olması durumunda

da veri daha dengesiz hale gelir. (Kwon & Sim, 2013). Veri dengesizliđi durumunda sınıflama dođruluđu dűűűk olma eđilimindedir (He & Garcia, 2009; Kuzey, 2012). Li ve Hung (2009) alıűmalarında verinin bađımlı deđiűkenin kategorilerine iliűkin dađılımlarının dengeli olup olmaması durumuna gűre yűntemlerin sınıflama performanslarını incelemiű ve dađılım dengesizleűtiđi durumlarda sınıflama performansının daha dűűűk olduđunu tespit etmiűtir.

Bađımlı deđiűkenin kategori sayısına gűre yűntemlerin performanslarının en iyi olduđu durumun bađımlı deđiűkene ait sınıf sayısının 2 olduđu durum olduđu tespit edildikten sonra bađımsız deđiűken sayısı ile ilgili analizlere geilmiűtir. Bu analizlerde, bađımlı deđiűken kategori sayısı 2 ve alıűma grubu 5000 olarak alınmıű ve farklı bađımsız deđiűken sayısı durumlarında YSA, RO, SVRA, DVM ve LR yűntemleri ile analizler yapılmıűtır.

Baűlangıta belirlenen 10 bađımsız deđiűkenden biri olan matematik baűarısı, fen baűarısı ile yűksek dűzeyde korelasyona ve diđer deđiűkenler ise fen baűarısı ile orta veya dűűűk dűzeyde iliűkiye sahiptir. alıűma kapsamında ele alınan koűullardan biri olan farklı bađımsız deđiűken sayısı koűulu iin űncelikle bađımlı deđiűkenle yűksek dűzeyde korelasyona sahip matematik baűarısı deđiűkeninin modelde bulunmadıđı ve sadece bađımlı deđiűkenle dűűűk veya orta dűzeyde korelasyona sahip sırasıyla 3, 6 ve 9 tane bađımsız deđiűkenin modelde olması durumu ele alınmıűtır. alıűmada ele alınan YSA, RO, SVRA, DVM ve LR yűntemleri iin bađımsız deđiűken sayısı arttıka sınıflama performansının genel olarak olumlu yűnde deđiűim gűsterdiđi sonucuna ulaűılmıűtır. İncelenen performans deđerlendirme kriterlerinin her biri iin yapılan fark testleri bu artıűın istatistiksel olarak anlamlı olduđunu gűstermiűtir. Yani bűtűn yűntemler iin bađımsız deđiűken sayısı 9 olduđu durumda performansın en yűksek olduđu sonucuna ulaűılmıűtır. Bu bađlamda bűtűn yűntemlerin deđiűken sayısı arttıka sınıflandırma performansının artması bađımsız deđiűken sayısı arttıka ele alınan deđerlendirme kriterlerine bađlı olarak yűntemlerin daha gűvenilir ve geerli sonular űrettiđi anlamına gelmektedir.

Bađımsız deđiűken sayısının sınıflama algoritmaları űzerinde űnemli bir etkiye sahip olduđu belirtilmektedir (Mannila, 2000; Kwon & Sim, 2013). Kumar ve Chong'a (2018) gűre

değişken sayısının artması değişken seti ve bağımlı değişken arasındaki korelasyonun arttığını gösterir. Bu da sınıflama doğruluğunun artmasına sebep olur. Diğer bir çalışmada Dolgun (2014) farklı koşullar altında genel olarak sınıflama yöntemlerinin performansının olumlu yönde değişim gösterdiğini tespit etmiştir. Fakat teorik olarak bağımsız değişken sayısındaki artışın sınıflama gücünü arttıracaklarını ancak makine öğrenme algoritmalarına ait pratik uygulamalarında bu durumun her zaman geçerli olmayacağı belirtilmektedir (Hall, 1999, s.25; Kotsiantis, 2007).

Bağımsız değişken sayısı koşulu için yapılan bir diğer analizde, bağımlı değişken ile yüksek düzeyde korelasyona sahip matematik başarıları değişkeni daha önce belirlenen bağımsız değişken gruplarının her birine eklenerek analizler yapılmıştır. Başka bir ifadeyle, bağımlı değişken kategori sayısı 2, çalışma grubu 5000 ve bağımlı değişkenle yüksek düzeyde korelasyona sahip matematik başarıları puanlarının ve düşük düzeyde korelasyona sahip diğer değişkenlerin bulunduğu 4, 7 ve 10 tane bağımsız değişkenin olması durumu incelenmiştir. Bu durumda gerçekleştirilen analiz sonucu, bağımsız değişken sayısının 4, 7 ve 10 olmasının bütün performans değerlendirme kriterlerinde anlamlı bir farklılık ortaya çıkarmadığını göstermektedir. Yani, test edilen bütün modellerde bağımlı değişken ile yüksek korelasyona sahip bir bağımsız değişken olduğu durumda, ilgili bütün yöntemler için sınıflama performansının değişken sayısındaki artıştan etkilenmediği söylenebilir. Bu sonuç ile benzer şekilde, Dağ vd. (2012) çalışmalarında korelasyona dayalı olarak bağımlı değişkenle ilişkili olan belirli sayıdaki (5) bağımsız değişkenle analize başlamış, değişken sayısını arttırdıkça sınıflama doğruluğunun değişmediği sonucuna ulaşmıştır. Diğer bir çalışmada Kumar ve Chong (2018) kullandıkları algoritmalarda, bağımlı değişkenle ilişkisi az olan bağımsız değişkenleri modelden çıkardıkları durumlarda sınıflama performanslarının değişime uğramadığı sonucuna ulaşmışlardır. Kwon ve Sim (2013) de çalışmasında bağımsız değişken sayısı artmasının sınıflama doğruluğunu değiştirmedığı sonucuna ulaşmıştır.

Bağımsız değişken sayısının 4, 7 ve 10 olması durumunda bağımsız değişken sayısındaki artış sınıflama performansında anlamlı bir etkiye neden olmazken, bağımsız değişken sayısının 3, 6 ve 9 olması durumunda bağımsız değişken sayısındaki artış sınıflama

performansının anlamlı bir biçimde artış göstermesine neden olmuştur. Bu durumun, 4, 7 ve 10 şeklinde belirlenen değişken setlerinin her birinde bağımlı değişkenle yüksek düzeyde korelasyona sahip olan matematik başarı puanı değişkenin yer alması olduğu düşünülmektedir. Buna ek olarak, bağımlı değişkenle yüksek korelasyona sahip matematik başarıları değişkeninin modelde yer aldığı durumlarda (4, 7 ve 10 bağımsız değişken) sınıflama performansının, matematik başarıları değişkeninin yer almadığı (3, 6 ve 9 bağımsız değişken) durumlardaki sınıflama performansına göre tüm sınıflama yöntemlerinde daha iyi olduğu sonucuna ulaşılmıştır. Bunun nedeninin matematik başarıları değişkeninin bağımlı değişkenle yüksek korelasyona sahipken diğer bağımsız değişkenlerin düşük korelasyona sahip olması olduğu düşünülebilir. Bağımlı değişkenle yüksek düzeyde ilişkiye sahip olan değişkenler, sınıflama için iyi değişkenlerdir (Liu & Motoda, 1998, s.28). Bu tür değişkenler bağımlı değişkenin kategorileri ile ilgili olup geri kalan değişkenler bağımlı değişken ile ilgisizdir (Blum & Langley 1997; Gennari, Langley & Fisher, 1989,). Bağımlı ve bağımsız değişken arasındaki ilişki ne kadar yüksekse, bağımsız değişkenin tahminleme gücü de o kadar yüksektir (Lutu, 2010, s.66).

Yu ve Liu (2003) çalışmalarında eğer bir bağımsız değişken ile bağımlı değişken arasındaki ilişki yüksek ve o değişkenin diğer bağımsız değişkenlerle arasındaki ilişki düşük ise bu değişkenin sınıflama için iyi bir bağımsız değişken olduğunu ifade etmiş ve korelasyona dayalı değişken seçimi yöntemi kullanılarak modele alınan değişkenlerle yapılan sınıflamanın, sınıflama performansını arttırdığı sonucuna ulaşmıştır. Diğer bir çalışmada Kumar ve Chong (2018) bağımlı değişkenle ilişkisi yüksek olan bağımsız değişkenleri modelde tutarak %40-50 oranında sınıflama doğruluğu artışı elde etmiş ve bağımlı değişkenle yüksek korelasyona sahip değişken seti kullanıldığında sınıflama doğruluğunun artacağını belirtmiştir. Benzer biçimde Khalili ve Moradi (2009) de çalışmalarında korelasyon temelli değişken seçimi yapmış ve sınıflandırma doğruluğunda artış olduğunu tespit etmiştir. Makine öğrenmesinde değişken seçimi veriye ulaşmadaki kolaylık ve sınıflandırma doğruluğu üzerinde önemli bir etkiye sahiptir.

Değişken seçimi ortaya konulan modelin doğruluğunu ve ulaşılan bilginin anlaşılabilirliğini arttırmaktadır (Kohavi & John, 1996). Bağımlı ve bağımsız değişkenler arasındaki ilişki bağımsız değişkenlerin bağımlı değişkenler üzerinde güçlü etkiye sahip olup olmadığının tanımlanmasına yardımcı olur. Bu durumda bağımsız değişkenlerin daha etkili bir biçimde yordamayı sağladığı belirtilmektedir. Güçlü bir ilişki olduğunda bağımsız değişken bağımlı değişkenin güçlü bir yordayıcısı olarak ele alınabilir (Kumar & Chong, 2018). Bu çalışmanın bulguları ve ilgili literatür incelendiğinde, yöntemlerin sınıflama performansları açısından modele alınan bağımsız değişkenin bağımlı değişken ile olan ilişkisinin, modele alınan bağımsız değişken sayısından daha etkili olduğu söylenebilir. Nitekim Emir (2013, s.216-219) Sınıflandırıcıların performansını arttırmada yapılabilecek işlemlerden birinin bağlamla ilişkisi olan değişken sayısının artırılması olduğunu ve daha az değişken kullanımıyla olası tüm değişkenlerin kullanıldığı duruma benzer performansların elde edilebileceğini ifade etmiştir.

Araştırma kapsamında ele alınan bir diğer koşul örneklem büyüklüğüdür. Kategori sayısı 2 olan fen başarısı bağımlı değişkeni ve ilgili 10 bağımsız değişken ile örneklem büyüklüğünün sırasıyla 100, 250, 500, 1000, 2500 ve 5000 olması durumu incelenmiştir. Örneklem büyüklüğünün her bir koşulu için YSA, RO, DVM, SVRA ve LR yöntemleri ile analizler gerçekleştirilmiştir. Sınıflama performansının belirlenen örneklem büyüklüklerinde farklılaşıp farklılaşmadığını incelemek amacıyla her bir yöntem için ayrı ayrı fark testleri yapılmıştır. Analiz sonucunda, tüm yöntemlere ait performans değerlendirme kriterleri değerlerinin, örneklem büyüklüklerine göre istatistiksel olarak anlamlı biçimde farklılık gösterdiği tespit edilmiştir. YSA, DVM ve LR için elde edilen sınıflama performansının, örneklem büyüklüğü 500 ve daha fazla olduğu durumda, 100 ve 250 örneklem büyüklüğüne göre anlamlı olarak daha iyi olduğu sonucuna ulaşılmıştır. Ayrıca RO ve SVRA yöntemleri için ise örneklem büyüklüğü 1000 ve üzeri olduğu durumda daha iyi sınıflama performansı elde edilmiştir. Bu bağlamda bütün yöntemlerin örneklem büyüklüğü arttıkça sınıflandırma performansının artması örneklem büyüklüğünün artışına bağlı olarak daha güvenilir ve geçerli sonuçlar ürettiği anlamına gelmektedir. Kwon ve Sim

(2013) de örneklem büyüklüğü ve modelin sınıflama doğruluğu arasında pozitif bir nedensellik olduğundan bahsetmektedirler. Örneklem büyüklüğündeki yetersizlik, sınıflama doğruluğunu azaltır ve dolayısıyla sınıflama hatasını arttırmış olur (Okamoto, 1963). Örneklem büyüklüğü belli bir sayıya ulaştıktan sonra performans değerlendirme kriterlerine ait değerler anlamlı olarak bir farklılık göstermemektedir. Bu durum, çalışmada ele alınan koşullar ve veri seti göz önüne alınarak YSA, RO, DVM, SVRA ve LR yöntemleri ile analiz yapıldığında belirli bir örneklem büyüklüğüne ulaşıldıktan sonra örneklem büyüklüğünde meydana gelecek artışların yöntemlerin performansında bir farklılaşmaya sebep olmadığı şeklinde ifade edilebilir. Aksu, Güzeller ve Eser (2019) de çalışmalarında YSA yöntemleri ile gerçekleştirilen sınıflama ve yordama analizlerinde örneklem büyüklüğünün anlamlı bir etkiye sahip olduğunu göstermektedirler. Çalışma sonucunda, eğitim alanında YSA ile yapılan çalışmalar için 1000'in üzerinde bir örneklem büyüklüğü kullanılması ile daha tutarlı sonuçlar elde edilebileceği bulunmuştur. Koyuncu (2018) PISA 2012 veri setinde veri madenciliği yöntemlerini uyguladığı çalışmasında örneklem büyüklüğündeki artışın (500'den 5000'e) daha güvenilir ve genellenebilir sonuçlar elde edilmesine olanak sağladığını göstermektedir.

Bu çalışmada incelenen bağımlı değişken kategori sayısı, bağımsız değişken sayısı ve örneklem büyüklüğü koşullarının her biri altında PISA 2015 verisi kullanılarak oluşturulan veri setlerinin her biri için YSA, RO, DVM, LR ve SVRA yöntemleri ile ayrı ayrı analizler gerçekleştirilmiştir. Her bir koşul içinde yer alan durumlar incelenirken aynı zamanda yöntemlerin de bu koşullardaki durumları ortaya konmuştur. Ele alınan bütün koşullar dâhilinde yapılan yöntem kıyaslamaları, YSA, RO, DVM, LR ve SVRA yöntemlerinden elde edilen performans değerlendirme kriterlerinin her biri için anlamlı fark olduğunu ortaya koymaktadır. Farkın kaynağını incelemek amacıyla yapılan ikili karşılaştırmalar, YSA, DVM ve LR yöntemlerinin sınıflama performanslarının birbirine benzer olduğunu göstermektedir. Ayrıca bu yöntemlerin her birinin RO ve SVRA yöntemine göre daha iyi değerler ürettiği sonucuna ulaşılmıştır. RO ve SVRA yöntemlerinin birbirlerine göre sınıflama performansı değerlendirildiğinde, bunlar arasındaki farkın da RO yöntemi lehine

anlamli olduđu g r lm şt r. Genel olarak, bu arařtırmada elde edilen bulgular ıřıřında, en iyi sınıflama performansına sahip y ntemlerin YSA, DVM ve LR olduđu, sonra gelen y ntemin RO ve diđerlerine g re en k t  performansı sergileyen y ntemin ise SVRA olduđu s ylenebilir. Ayrıca, g venirlik ve ge erlilik baęlamında ele alındıęında da sınıflandırma performansına baęlı olarak YSA, DVM ve LR y ntemlerinin RO ve SVRA y ntemlerine g re  ęrencileri sınıflandırmada daha g venilir ve ge erli sonu lar  rettięi ifade edilebilir.

Y ntemlerin birbirlerine g re karřılařtırıldıęı ilgili literat r incelendięinde, bu arařtırma ile benzer sonu lara rastlanmaktadır. Cortez ve Silva (2008)  alıřmalarında ele aldıkları durumlar i in genel anlamda YSA ve DVM y ntemlerinin RO ve karar aęa ları y ntemlerine g re daha iyi  alıřtıęını ifade etmektedir. Ayrıca bir bařka  alıřmada, DVM ve YSA y ntemlerinin baęımsız deęiřkenler s rekli olduęunda daha iyi performans sergiledięi ifade edilmektedir (Kotsiantis, 2007). Diđer bir  alıřmada ise M.Kayri, Kayri ve Gencoęlu (2017) YSA y nteminin RO y ntemine g re daha iyi bir performans sergiledięini ifade etmiřtir. Kılı -Depren vd. (2016) TIMSS matematik bařarısının sınıflandırılması i in kullandıkları y ntemlerden (LR, RO, Naive Bayes, multilayer perception ve J48) LR y nteminin ele alınan performans kriterleri a ısından RO y nteminden anlamli olarak daha iyi sonu lar  rettięini g stermiřtir.  ięřar ve  nal (2019) banka m řterileri ile ilgili ger ek verileri ele aldıkları  alıřmalarında altı sınıflama algoritmasını kullanmıř (Bayes network, Naive Bayes, J48, RO, multilayer perceptron ve LR) ve en iyi sınıflama performansı sergileyen y ntemin lojistik regresyon olduęunu ifade etmiřlerdir. Elde edilen bulgulara g re LR analizinin YSA ve DVM y ntemlerine benzer ve iyi bir sınıflandırma performansına sahip olmasının nedeni  oklubaęlantılılık varsayımının saęlanması olduđu d ř n lebilir.   nk  Tabachnick ve Fidell'e (2013, s.443) g re varsayımların saęlanması LR y nteminin performansını arttırmaktadır. Dolayısıyla varsayımların saęlanmadıęı durumlarda, bu durumlara karřı daha dayanıklı olan diđer y ntemlerin LR analizine g re daha iyi bir performans sergileyeceęi s ylenebilir.

N u (2018) sim lasyon verisi  zerinden LR ve DVM y ntemlerini karřılařtırmıř ve bu y ntemlerin birbiri ile benzer yordama g c ne sahip olduęunu belirtmiřtir. Diđer bazı

çalışmalarda da YSA ve LR yöntemlerinin benzer sınıflama performansına sahip olduğu bulunmuştur (Bartfay vd., 2006; Bayru, 2007; Chaveesuk vd., 1999; Efe, 2012; Güneri ve Apaydın, 2004). Bu araştırmada elde edilen bulgular, SVRA yönteminin diğerlerine göre daha kötü sınıflama performansına sahip olduğunu göstermektedir. Benzer şekilde bu yöntemlerin kıyaslanmasını ele alan bazı çalışmalarda YSA ve LR'nin SVRA'den daha iyi olduğu (Kuyucu, 2012) ve YSA modelinden elde edilen sonuçların karar ağaçlarından daha iyi olduğu (Köktürk, 2012) belirtilmiştir. Ayrıca, B. L. Zhang ve Zhang (2008) ve Acharya vd. (2012) yaptıkları çalışmada YSA ve DVM yöntemlerinin birbirine benzer ancak SVRA yöntemine göre daha iyi sınıflama performansı gösterdiği sonucuna ulaşmışlardır. Çelik'in (2009) çalışmasında ise YSA ve LR yöntemlerinin sınıflama performanslarının benzer olup SVRA yönteminden daha iyi olduğu belirtilmiştir. Bu araştırma ile ulaşılan sonuçlara çok benzer olarak Baumgartner vd (2004) tarafından yapılan çalışmada YSA, DVM ve LR yöntemlerinin sınıflandırma performanslarından elde edilen doğruluk değerlerinin birbirine çok benzer olduğu tespit edilmiştir. Tosun (2007) de öğrenci başarılarının sınıflandırılmasına yönelik yaptıkları çalışmada YSA yönteminin karar ağaçlarına göre daha iyi performans sergilediğini ifade etmiştir. Bu araştırma sonucunda yöntemlere ilişkin bir diğer sonuç da RO yönteminin, bir diğer karar ağacı algoritması olan SVRA'den daha iyi performans göstermesidir. Korkmaz (2018) de çalışmasında RF yönteminin SVRA yönteminden daha iyi performans gösterdiği sonucuna ulaşmıştır. Watts ve Lawrence (2008) ise, RO yönteminin gürültü faktöründen etkilenmemesi ve sabit bir modeli olmamasından ötürü diğer karar ağaçları yöntemlerinden üstün olduğunu ifade etmektedir.

Makine öğrenmesi yöntemlerinin sınıflama performanslarının birbirine göre ele alındığı birçok çalışma ve farklı sonuçlar bulunmaktadır (Çokluk ve Çırak, 2012; Hekimoğlu, 2017; Moroco vd, 2011; Ocakoğlu, 2006; Prancevicius & Marcinkevicius, 2017; Yakut & Gemici, 2017; Yılmaz, 2009). Bu çalışmalarda sınıflama algoritmalarının birçok farklı türe sahip veri yapısı, içerik ve bağlamda farklı performanslar sergilediği görülmektedir. Bu durum, sınıflama algoritmalarının seçiminde bağlamın dikkate alınmasının, en iyi algoritmayı belirlemede önemli olabileceğini göstermektedir. Bir başka ifadeyle, sınıflama

algoritmaları farklı bağlamlarda farklı performanslar sergileyebilirler. Kotsiantis (2007), sınıflamada esas problemin, bir algoritmanın diğerine göre üstün olup olmaması değil, hangi koşullar altında belirli bir yöntemin anlamlı olarak diğerlerinden daha iyi performans sergilediği olduğunu belirtmektedir. Böylece bir uygulama durumunda elde olan koşullar listelenir ve onlara uygun yöntem seçimi yapılabilir.

5.2. Sonuç

Bu bölümde çalışma bulguları ışığında elde edilen sonuçlar özetlenmiştir. Çalışmada ele alınan koşullar, PISA 2015 öğrenci veri setinden elde edilen bağımlı değişkenin kategori sayısı, örneklem büyüklüğü ve bağımsız değişken sayısıdır. Çalışma kapsamını oluşturan bütün durumlar için ele alınan yöntemlerin (YSA, RO, LR, SVRA ve DVM) hepsi ile analizler gerçekleştirilmiş ve hem koşulların kendi içindeki durumları hem de yöntemlerin birbirine göre performansları arasındaki farklar istatistiksel olarak incelenmiştir.

Bağımlı değişkenin kategori sayısının değişimine göre performans değerlendirme kriterlerinin her biri için istatistiksel olarak anlamlı bir farkın olup olmadığı incelenmiş ve elde edilen bulgular doğrultusunda bu farkların anlamlı olduğu tespit edilmiştir. Yani YSA, RO, DVM, LR ve SVRA yöntemlerinin hepsi için bağımlı değişkenin kategori sayısının değişmesi ile yöntemlerin performanslarının farklılaştığı tespit edilmiştir. Yapılan ikili karşılaştırmalar sonucunda genel olarak, bağımlı değişkenin kategori sayısı azaldıkça tüm yöntemlerde doğruluk, kesinlik F ölçütü ve Kappa değerlerinin artış gösterdiği söylenebilir. Yani bağımlı değişkenin kategori sayısı daha az olduğunda yöntemlerin sınıflama performansı daha iyi olduğu sonucuna ulaşılmıştır. Belirleyicilik ve negatif öngörü değerleri ele alındığında ise diğer değerlendirme kriterlerinin aksine bağımlı değişkenin kategorisi arttıkça bu değerlerin artış gösterdiği görülmüştür. Ancak bu artış RO ve SVRA yöntemlerinde belirgin bir anlamlı farklılığa neden olurken YSA, DVM ve LR yöntemlerinde belirgin bir farklılık oluşturmamıştır. Bu bağlamda, değerlendirme kriterlerinin çoğuna göre bağımlı değişkenin kategori sayısı azaldıkça bütün yöntemlerin sınıflama performanslarının olumlu yönde değişim gösterdiği söylenebilir.

Bağımsız değişken sayısı koşulu için iki farklı durum ele alınmıştır. Bunlardan ilki, bağımlı değişken ile yüksek korelasyona sahip matematik değişkeninin bütün modellerde yer aldığı, bağımsız değişken sayısı 4, 7 ve 10 olduğundaki karşılaştırmadır. Burada elde edilen bulgular ile bağımsız değişken sayısının artmasının ya da azalmasının yöntemlerin performansında anlamlı bir farklılığa sebep olmadığı sonucuna ulaşılmıştır. Bir diğer durumda ise matematik başarıları değişkeni çıkarılmış ve bağımsız değişken sayısının 3, 6 ve 9 olduğu durumlar birbirine göre kıyaslanmıştır. Performans değerlendirme kriterlerine ait değerler karşılaştırıldığında genel olarak bütün yöntemlerin performansının bağımsız değişken sayısına göre farklılaştığı ve bağımsız değişken sayısı arttıkça yöntemlerin daha iyi performans gösterdiği sonucuna ulaşılmıştır. Buradaki artışın aksine ilk durumda bir farklılığın olmamasının nedeninin, modelde yer alan matematik başarıları değişkeni olduğu düşünülmektedir. Ele alınan 10 değişkenden sadece matematik başarıları, bağımlı değişken olan fen başarıları ile yüksek düzeyde korelasyona sahipken, diğer 9 değişkenin fen başarıları ile orta ve yüksek düzeyde ilişkisi vardır. Bu sebeple, modelde bağımlı değişken ile yüksek düzeyde ilişkili bir değişkenin bulunması durumunda bağımsız değişken sayısının test edilen modelin performansı üzerinde bir değişikliğe sebep olmadığı sonucuna ulaşılmıştır. Bu özelliklerde bir değişken olmadığında ise bağımsız değişken sayısı arttıkça yöntemlerin daha iyi performans sergilediği sonucuna ulaşılmıştır. YSA, DVM ve LR yöntemlerinde sınıflama performansı bağımsız değişken sayısının ele alınan her bir artışında anlamlı bir iyileşme göstermiştir. Ancak RO ve SVRA yöntemlerinde 3 ve 6 bağımsız değişkenin olduğu durumlarda benzer 9 bağımsız değişkenin olduğu durumda ise diğer bağımsız değişken koşullarına göre daha iyi bir sınıflama performansı elde edilmiştir. Bu bağlamda, sınıflama modellerinde, bağımlı değişkenle orta ve düşük düzeyde korelasyona sahip bağımsız değişkenler olduğunda RO ve SVRA yöntemlerinin sınıflama performanslarının daha iyi hale gelebilmesi için YSA, DVM ve LR yöntemlerine göre daha fazla sayıda bağımsız değişkene ihtiyaç duyduğu sonucuna ulaşılmıştır.

Çalışmada ele alınan bir diğer koşul örneklem büyüklüğüdür. Çalışmada elde edilen bulgular, örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olarak farklılaşması

durumunda incelenen yöntemlerin hepsinde sınıflama performanslarının farklılaştığını göstermektedir. YSA, DVM ve LR yöntemleri için örneklem büyüklüğü 500, 1000, 2500 ve 5000 olduğunda sınıflama performanslarının benzer ve örneklem büyüklüğünün 100 ve 250 olduğu durumdan daha iyi olduğu sonucuna ulaşılmıştır. RO ve SVRA yöntemlerinde ise sınıflama performansının örneklem büyüklüğünün 1000 ve üzeri olduğu durumlarda benzer ve örneklem büyüklüğünün 500 ve altı olduğu durumlardan daha iyi olduğu sonucuna ulaşılmıştır. Elde edilen bu sonuçlara göre YSA, DVM ve LR yöntemlerinin daha iyi bir performansa sahip olabilmesi için en az 500 örnekleme, RO ve SVRA yöntemlerinin ise en az 1000 örnekleme ihtiyaç duyduğu yeterli olduğu söylenebilir. Öte yandan, YSA, DVM ve LR yöntemlerinin 500 örneklem büyüklüğünden, RO ve SVRA yöntemlerinin 1000 örneklem büyüklüğünden sonra sınıflama performanslarının değişmemesi, belli bir örneklem büyüklüğüne ulaşıldıktan sonra örnekleme artışı modellerin performansını değiştirmeyeceği şeklinde ele alınabilir. Özetle, genel olarak literatürdeki bulgulara paralel olarak örneklem büyüklüğü arttıkça bütün yöntemlerde sınıflama performansının arttığı sonucuna ulaşılmıştır.

Bağımlı değişkenin kategori sayısı, bağımsız değişken sayısı ve örneklem büyüklüğü koşullarının her biri altında yöntemler birbirine göre de kıyaslanmıştır. Bütün koşullar için elde edilen yöntem farklılıkları genel olarak birbirine benzerdir. Yöntemlerin kıyaslanması için gerçekleştirilen fark analizleri, yöntemlerin birbirine göre performanslarının farklılaştığını göstermektedir. Bu farka ilişkin yapılan ikili karşılaştırmalar ile YSA, DVM ve LR yöntemlerinin performanslarının farklılaşmadığı fakat bu yöntemlerin RO ve SVRA yöntemlerinden daha iyi performans gösterdiği sonucuna ulaşılmıştır. Ayrıca genel olarak RO yöntemi de SVRA yönteminden daha iyi performans göstermiştir. Bu durumda çalışma verisi ve uygulanan koşulların her biri için sınıflama performansı en iyi olan yöntemler YSA, DVM ve LR iken SVRA ise en kötü performans sergileyen yöntem olmuştur.

5.3. Öneriler

5.3.1. Uygulayıcılara Yönelik Öneriler

Bağımlı değişkenin kategori sayısının azalması durumunda, araştırma kapsamında ele alınan bütün yöntemlerin (YSA, RO, DVM, LR ve SVRA) sınıflandırma performanslarının artış gösterdiği sonucuna ulaşılmıştır. Bunun sebeplerinden birinin örneklemin bağımlı değişkenin kategorilerindeki dağılımının dengesiz olması diğerinin ise kategori sayısının artması durumunda karmaşıklığın da artacak olmasıdır. Yöntemlerin sınıflandırma performansını etkilediği düşünülen bu sebepler tartışma bölümünde alanyazında yer alan diğer çalışmalar ışığında desteklenmiştir. Elde edilen bu sonuçlar doğrultusunda, sınıflama işlemi yaparken örneklemin bağımlı değişkene ait kategorilere göre dağılımının daha dengeli olabilmesi ve daha iyi sınıflama performansına ulaşmak için mümkünse kategori birleştirme yani sınıf sayısı azaltma uygulayıcıların başvuracağı bir yöntem olabilir. Ancak kategori birleştirme mümkün değilse bu çalışma ile elde edilen sonuçlara göre araştırmanın koşullarına benzer veri setlerinin bulunduğu durumlarda daha iyi bir sınıflama performansı elde edebilmek için YSA, DVM ve LR yöntemleri, RO ve SVRA yöntemlerine tercih edilmesi önerilmektedir.

Bağımlı değişkenle düşük veya orta düzeyde korelasyona sahip bağımsız değişken sayısının artması durumunda, araştırma kapsamında ele alınan bütün yöntemlerin (YSA, RO, DVM, LR ve SVRA) sınıflama performanslarının artış gösterdiği sonucuna ulaşılmıştır. Elde edilen bu sonuç doğrultusunda, sınıflama performansını arttırmak dolayısıyla öğrencilerin sınıflarının belirlenmesinde daha doğru sonuçların elde edilmesi için sınıflama işlemi yaparken veri setinden bağımlı değişkenle düşük ve orta düzeyde ilişkiye sahip değişkenler bulunduğu zaman, sınıflama modelinde mümkün olabilecek tüm değişkenlerin dâhil edilmesi önerilmektedir.

Sınıflama modelinde bağımlı değişkenle yüksek düzeyde korelasyona sahip bağımsız değişken olması durumunda, bağımlı değişkenle düşük veya orta düzeyde ilişkiye sahip bağımsız değişken sayısının artması durumunda, araştırma kapsamında ele alınan bütün

yöntemlerin (YSA, RO, DVM, LR ve SVRA) sınıflandırma performanslarının değişim göstermediği sonucuna ulaşılmıştır. Elde edilen bu sonuç doğrultusunda, sınıflama işlemi yaparken bağımlı değişkenle yüksek düzeyde ilişkiye sahip değişkenlerin sınıflama modellerinde yer alması durumunda bağımlı değişkenle düşük ve orta düzeyde ilişkiye sahip değişkenler modellerden çıkarılabilir. Bağımlı değişkeni açıklama oranı yüksek olan az sayıda değişkenle sınıflama işleminin yapılması, daha iyi performans ve zaman tasarrufu sağlayacaktır. Öte yandan az sayıda değişkenle daha iyi performansta sınıflama sonuçlarının elde edilecek olması, veri toplarken daha az sayıda bağımsız değişken hakkında bilgi elde edilmesine dolayısıyla veri toplama aşamasında araştırmacıya daha kolaylık sağlayacaktır. Bu nedenle uygulayıcılara, sınıflama modellerinde bağımlı değişkenle yüksek düzeyde ilişkiye sahip bağımsız değişkenler kullanması önerilmektedir.

Örneklem büyüklüğünün artması durumunda araştırma kapsamında ele alınan bütün yöntemlerin sınıflama performanslarının anlamlı olarak artış gösterdiği sonucuna ulaşılmıştır. Ancak YSA, DVM ve LR yöntemleri 500 ve üzerindeki örneklem büyüklüklerinde iyi düzeyde bir sınıflama performansı sergilerken RO ve SVRA yöntemlerinin 1000 ve üzeri örneklem büyüklüklerinde daha iyi sınıflama performansı sergilediği tespit edilmiştir. Elde edilen bu sonuçlar doğrultusunda araştırma kapsamında kullanılan veri setindeki benzer koşullar altında sınıflama işlemi yapılırken YSA, DVM ve LR yöntemleri kullanıldığı zaman en az 500 örneklemle, RO ve SVRA yöntemi kullanıldığı zaman ise en az 1000 örneklemle çalışılması önerilmektedir.

Araştırma kapsamında ele alınan tüm koşullarda, SVRA ve RO yöntemi YSA, DVM ve LR yöntemlerine göre daha düşük sınıflandırma performansı göstermiştir. Bu sonuca göre, benzer koşullar için yapılacak uygulamalarda daha yüksek sınıflandırma performansı sağlayarak öğrencilerin ait olduğu sınıfları daha doğru belirleyebilmek için YSA, DVM ve LR yöntemleri, SVRA ve RO yöntemlerine tercih edilebilir. Öte yandan SVRA ve RO yöntemlerinin birbirlerine göre olan sınıflandırma performanslarının ele alındığında ise RO yönteminin SVRA yöntemine göre daha yüksek sınıflandırma performansı gösterdiği

sonucuna ulařılmıştır. Elde edilen bu sonu doęrultusunda arařtırma kapsamında alınan kořullara benzerlik gsteren veri setlerinde RO yntemi SVRA yntemine tercih edilmelidir.

Ayrıca tm kořullar altında YSA, LR, ve DVM yntemlerinin sınıflandırma performanslarının benzer olduęu sonucuna ulařılmıştır. Elde edilen bu sonu doęrultusunda, yapısı ve inřaası bakımından LR'nun YSA ve DVM'ne gre daha kolay olduęu gznne alındıęında, arařtırma kapsamındaki benzer kořulların yer aldıęı veri setlerinde LR yntemi dięer yntemlere tercih edilebilir. Ayrıca, yapılacak sınıflamaya dnk dięer alıřmalar iin bu arařtırmada olduęu gibi birden fazla sınıflama ynteminin karřılařtırmalı olarak ele alınması ve yntemlerin performansından kaynaklı hatalarınlemek iin sonuların bu doęrultuda deęerlendirilmesinlerilmektedir.

5.3.2. Arařtırmacılara Yneliknleriler

Bu arařtırma kapsamında elde edilen sonuların kararlılıęını incelemek amacıyla benzer veri setleri zerinde aynı arařtırma tekrar gerekleřtirilebilir.

Bařka bir alıřma iin arařtırma kapsamında sınıflandırma performansı incelenen YSA, RO, DVM, LR ve SVRA yntemlerine veri madencilięinde yer alan dięer sınıflama yntemleri de eklenebilir.

Arařtırma kapsamında baęımsız deęiřken seimi yapılırken korelasyon analizi kullanılmıştır. Bařka arařtırmalar iin dięer deęiřken seim yntemlerinin sınıflandırma zerindeki etkisi arařtırılabilir.

Gerek verilerle alıřılan bu arařtırmada baęımsız deęiřkenlerin baęımlı deęiřkenle gsterdięi korelasyon deęerlerini dřk, orta ve yksek olacak biimde sınıflandırılabilen az sayıda baęımsız deęiřken vardı. Baęımlı deęiřken ile baęımsız deęiřken arasındaki korelasyonun sınıflandırma zerindeki etkisinin daha detaylı arařtırılabileceęi farklı gerek veya simlasyon veriler incelenebilir.

Kořullara ve veri yapısına grezellikle gerek verilerde modellerin gsterdięi sınıflandırma performanslarının birbirine gre stnlkleri deęiřim gsterebileceęinden

eğer mümkünse araştırmacıların birden fazla modelleme ile sınıflama analizlerini gerçekleştirmesi daha doğru ve güvenilir sonuçlara ulaşmasına yardımcı olacaktır.

Açıklanan 2023 eğitim vizyonu paketinde öğrencilerin eğitim süreci hakkında yıllara ait tüm bilgilerin yer alacağı yani veri yığının fazla miktarda olacağı ve karmaşık hale geleceği e-portfolio veya EBA gibi çevrim içi sistemlerden elde edilecek boylamsal verilerin analizlerinde eğitimde veri madenciliğine ait yöntemlerin kullanılması zaman ve performans açısından kazanç, elde edilen kararlar açısından ise daha güvenilir ve geçerli değerlendirmeler sağlayacaktır.

Dijital eğitime geçişin başladığı bu süreçte veri setlerinin daha fazlalaşacağı ve karmaşık hale geleceği gözönüne alındığında, veri madenciliği uygulamalarına ülkemizdeki eğitim bilimleri alanında daha fazla ve detaylı olarak yer verilmelidir.

Bu çalışmada EVM yöntemleri eğitimde uluslararası bir uygulama olan PISA 2015 verileri üzerinde uygulanmış ve veri setine ait özelliklerin belli koşullarına göre yüksek doğrulukta sınıflandırma performansı elde edilmiştir. Ülkemizde de eğitim alanında ulusal düzeyde geniş ölçekli uygulamalar yer almaktadır (ABİDE, LGS, Özel yetenek vb.). Ayrıca, rehberlik birimleri tarafından duyuşsal özelliklere yönelik uygulamalar da bulunmaktadır. Dolayısıyla öğrenciler hakkında sınıflandırmaya yönelik karar verilirken sınıflandırma performansı iyi olan EVM yöntemlerinin tercih edilmesi geçerli ve güvenilir kararlar alınmasına yardımcı olacaktır.

KAYNAKLAR

- Abacıoğlu, S. (2018). *Finansal başarısızlığı belirlemede istatistiksel yöntemlerin sınıflandırma performanslarının karşılaştırılması: borsa İstanbul'da bir uygulama*. Yüksek Lisans Tezi, 19 Mayıs Üniversitesi, Sosyal Bilimler Enstitüsü, Samsun.
- Abe, S. (2005). *Support vector machines for pattern classification* (Vol. 2). London: Springer.
- Acharya, U. R., Sree, S. V., Krishnan, M. M. R., Molinari, F., Garberoglio, R., & Suri, J. S. (2012). Non-invasive automated 3D thyroid lesion classification in ultrasound: a class of ThyroScan™ systems. *Ultrasonics. Advanced Technology*, 52(4), 508-520.
- Aggarwal, C. C. (2015). *Data mining: the textbook*. New York: Springer.
- Agrawal, R., Imieliński, T., & Swami, A. (1993). Mining association rules between sets of items in large databases. In *Acm sigmod record* (Vol. 22, No. 2, pp. 207-216). ACM.
- Agarwal, S., Pandey, G. N., & Tiwari, M. D. (2012). Data mining in education: data classification and decision tree approach. *International Journal of e-Education, e-Business, e-Management and e-Learning*, 2(2), 140.
- Ahn, H. (1996): Log-Gamma regression modeling through regression trees. *Communications in Statistics – Theory and Methods*, 25, 295–311.
- Akgün, B. (2016). *Apache Spark tabanlı destek vektör makineleri ile akan büyük veri sınıflandırma*. Yüksek Lisans Tezi, İstanbul Teknik Üniversitesi, Fen Bilimleri Enstitüsü

- Akman, M. (2010). *Veri madenciliğine genel bakış ve random forests yönteminin incelenmesi: sağlık alanında bir uygulama*. Yüksek Lisans Tezi, Ankara Üniversitesi, Ankara
- Akpınar, H. (2000). Veri tabanlarında bilgi keşfi ve veri madenciliği. *İÜ İşletme Fakültesi Dergisi*, 29(1), 1-22.
- Akpınar, H. (2017). *Data: Veri madenciliği veri analizi*. İstanbul: Papatya.
- Aksu, G., & Doğan, N. (2018). Veri madenciliğinde kullanılan öğrenme yöntemlerinin farklı koşullar altında karşılaştırılması, *Ankara Üniversitesi Eğitim Bilimleri Fakültesi Dergisi*, 51(3), 71-100.
- Aksu, G. & Doğan, N. (2019). Comparison of decision trees used in data mining. *Pegem Eğitim ve Öğretim Dergisi*, 9(4), 1183-1208.
- Aksu, G, Güzeller, C. & Eser, M. (2019). The effect of the normalization method used in different sample sizes on the success of artificial neural network model. *International Journal of Assessment Tools in Education*, 6(2) , 170-192. DOI: 10.21449/ijate.479404
- Albayrak, A. S. & Koltan Yılmaz, Ş. (2009). Veri madenciliği: karar ağacı algoritmaları ve İMKB verileri üzerine bir uygulama. *Süleyman Demirel Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi*, 14(1).
- Anwar, A.S., Ghany, K. K. A., & Elmahdy, H. (2015). Human ear recognition using geometrical features extraction. *Procedia Computer Science*, 65, 529–537.
- Architecture Technology Corporation, (1991). *Introduction to Neural Networks*. Elsevier.
- Asif, R., Merceron, A., Ali, S. A., & Haider, N. G. (2017). Analyzing undergraduate students' performance using educational data mining. *Computers & Education*, 113, 177-194.

- Atasever, Ü. H. (2011). *Uydu görüntülerinin sınıflandırılmasında hızlandırma (boosting), destek vektör makineleri, rastgele orman (random forest) ve regresyon ağaçları yöntemlerinin kullanılması*. Yüksek Lisans Tezi, Erciyes Üniversitesi Fen Bilimleri Enstitüsü, Kayseri.
- Ateş, Y. (2008). *Hibrit yakıt hücresi/ultra-kapasitörlü taşıt güç sisteminin yapay sinir ağları ile kontrolü*. Yüksek Lisans Tezi, Yıldız Teknik Üniversitesi Fen Bilimleri Enstitüsü, İstanbul.
- Atılğan, E. (2011). *Karayollarında meydana gelen trafik kazalarının karar ağaçları ve birliktelik analizi ile incelenmesi*. Yüksek lisans tezi, Hacettepe Üniversitesi Fen Bilimleri Enstitüsü, Ankara.
- Attewell, P., Monaghan, D., & Kwong, D. (2015). *Data mining for the social sciences: An introduction*. California: University of California.
- Aydoğan, B., Ayat, B., Öztürk, M. N., Çevik, E. Ö., & Yüksel, Y. (2010). Current velocity forecasting in straits with artificial neural networks, a case study: Strait of Istanbul. *Ocean engineering*, 37(5-6), 443-453.
- Ayhan, S (2013). *Kaba küme ve destek vektör makineleri kullanılarak nitelik indirgeme ve sınıflandırma problemlerinin çözümü için bütünleşik bir yaklaşım*. Doktora Tezi, Eskişehir Osmangazi Üniversitesi Fen Bilimleri Enstitüsü, Eskişehir.
- Baker, R.S.J.D. (2010). Data Mining for Education. McGaw, B., Peterson, P., Baker, E. (Ed), *International Encyclopedia of Education* (3rd edition) içinde. Oxford, UK: Elsevier.
- Baker, R., & Yacef, K. (2009). The state of educational data mining in 2009: A review and future visions. *Journal of Educational Data Mining*, 1(1), 3–17.
- Balaban, M. E., & Kartal, E. (2018). *Veri Madenciliği ve Makine Öğrenmesi*. İstanbul: Çağlayan.

- Barnett, D. W., & Macmann, G. M. (1992). Decision reliability and validity: contribution and limitations of alternative assessment systems. *The Journal of Special Education*, 25(4), 431-452.
- Bartfay et al. (2006) Comparing the predictive value of neural network models to logistic regression models on the risk of death for small-cell lung cancer patients. *European Journal of Cancer Care*, 15, 115–124.
- Başokçu, T, O. (2011). *Bağıl ve mutlak değerlendirme ile DINA modeline göre yapılan sınıflamaların geçerliğinin karşılaştırılması*. Doktora Tezi, Hacettepe Üniversitesi Eğitim Bilimleri Enstitüsü, Ankara.
- Bataineh, M. H. (2012). *Artificial neural network for studying human performance*. Yüksek Lisans Tezi, University of Iowa, Iowa.
- Baumgartner, C., Böhm, C., Baumgartner, D., Marini, G., Weinberger, K., Olgemöller, B., ... Roscher, A. A. (2004). Supervised machine learning techniques for the classification of metabolic disorders in newborns. *Bioinformatics*, 20(17), 2985-2996.
- Bayru, P. (2007). *Elektronik basında tüketici tercihleri analizi: yapay sinir ağları ile lojit modelin performans değerlendirilmesi*. Doktora Tezi, İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, İstanbul.
- Ben-Hur, A., & Weston, J. (2010). A user's guide to support vector machines, data mining techniques for the: life sciences methods in molecular biology. *Humana Pressure*, 223-239.
- Berkhin, P. (2006). A survey of clustering data mining techniques. Kogan J., Nicholas C., Teboulle M. (Ed.) *Grouping multidimensional data* içinde. Berlin: Springer.
- Bhalla, D. (2014). Random Forest in R: Step by Step Tutorial. <http://www.listendata.com/2014/11/random-forest-with-r.html> sayfasından erişilmiştir.

- Blum, A. L. & Langley, P. (1997) Selection of relevant features and examples in machine learning. *Artificial Intelligence*, 97(1-2), 245-271.
- Bounsaythip, C., & Rinta-Runsala, E. (2001). Overview of data mining for customer behavior modeling. *VTT Information Technology Research Report, Version, 1*, 1-59.
- Boznar, M., Lesjak, M. and Malaker, P., 1993. A neural network based method for short- term predictions of ambient SO₂ concentrations in highly polluted industrial areas of complex terrain. *Atmospheric Environment*, 27B(2), 221-230.
- Breiman, L. & Cutler, t.y. A Manual-Setting Up, Using, And Understanding Random Forests. University of California, Berkeley. https://www.stat.berkeley.edu/~breiman/RandomForests/cc_home adresinden erişilmiştir.
- Breiman, L. (2001). Random forests. *Machine learning*, 45(1), 5-32.
- Breiman, L. (2004). *Consistency for a simple model of random forests*. Technical Report 670, Statistics Department, University Of California At Berkeley.
- Breiman, L., Friedman, J., Stone, C. J., & Olshen, R. A. (1984). *Classification and regression trees*. CRC.
- Burges, C. J. (1998). A tutorial on support vector machines for pattern recognition. *Data mining and knowledge discovery*, 2(2), 121-167.
- Burmaoğlu, S. (2009). *Birleşmiş-milletler kalkınma programı beşeri kalkınma endeksi verilerini kullanarak diskriminant analizi, lojistik regresyon analizi ve yapay sinir ağlarının sınıflandırma başarılarının değerlendirilmesi*. Doktora Tezi, Atatürk Üniversitesi Sosyal Bilimler Enstitüsü, Erzurum.
- Burmaoğlu, S., Oktay, E., Özen, Ü. (2009). Birleşmiş Milletler Kalkınma Programı Beşeri Kalkınma Endeksi Verilerini Kullanarak Diskriminant Analizi ve Lojistik Regresyon Analizinin Sınıflandırma Performanslarının Karşılaştırılması. *Savunma Bilimleri Dergisi*, 8(2), 23-49.

- Büyükışıklar, A. (2014). *Karar ağaçları sınıflandırma algoritması ile toprak özgül direnci tespitinde jeolojik veri kullanımı*. Yüksek Lisans Tezi, Bilecik Şeyh Edebali Üniversitesi, Fen Bilimleri Enstitüsü, Bilecik.
- Büyüköztürk, Ş., Çakmak, E. K., Akgün, Ö. E., Karadeniz, Ş., & Demirel, F. (2017). *Bilimsel araştırma yöntemleri*. Ankara: Pegem.
- Cabena, P., Hadjinian, P., Stadler, R., Verhees, J., & Zanasi, A. (1998). *Discovering data mining: from concept to implementation*. Upper Saddle River, NJ: Prentice-Hall, Inc.
- Campbell, C., & Ying, Y. (2011). Learning with support vector machines. *Synthesis lectures on artificial intelligence and machine learning*, 5(1), 1-95.
- Canan, S. (2006). *Yapay sinir ağları ile GPS destekli navigasyon sistemi*. Doktora Tezi, Selçuk Üniversitesi Fen Bilimleri Enstitüsü, Konya.
- Canty, M. J. (2014). *Image analysis, classification and change detection in remote sensing: with algorithms for ENVI/IDL and Python*. Crc.
- Chandrashekar, G., & Sahin, F. (2014). A survey on feature selection methods. *Computers & Electrical Engineering*, 40(1), 16-28.
- Chaveesuk et al. (1999). Alternative neural network approaches to corporate bond rating. *Engineering Valuation and Computational Intelligence of Journal of Engineering Valuation and Cost Analysis*, 2(2). 117-131.
- Chiang, Y. M., Chang, L. C., & Chang, F. J. (2004). Comparison of static-feedforward and dynamic-feedback neural networks for rainfall–runoff modeling. *Journal of hydrology*, 290(3-4), 297-311.
- Chizi, B., Rokach, L., & Maimon, O. (2009). A Survey of Feature Selection Techniques. J. Wang (Eds.) *Encyclopedia of Data Warehousing and Mining* içinde, s. 1888-1895. IGI Global.

- Chudzinska, M. & Baralkiewicz, D. (2011). Application of icp-ms method of determination of 15 elements in honey with chemometric approach for the verification of their authenticity. *Food and Chemical Toxicology*, 49(11), 2741-2749.
- Cichosz, P. (2014). *Data mining algorithms: explained using R*. John Wiley & Sons.
- Cohen, J. (1988) *Statistical power analysis for the behavioural sciences*. 2nd Edition. Hillsdale: New Jersey Lawrence Erlbaum Associates.
- Cooper, S. J. (1998): Multiple regimes in us output fluctuations. *Journal of Business and Economic Statistics*, 16(1), 92–100.
- Cortez, P., & Silva, A. (2008). Using data mining to predict secondary school student performance. *In the Proceedings of 5th Annual Future Business Technology Conference*, Porto, Portugal, 5-12.
- Cristianini, N., & Shawe-Taylor, J. (2000). *An introduction to support vector machines and other kernel-based learning methods*. Cambridge University.
- Crows. T. (1999). *Introduction to data mining and knowledge discovery*. A.B.D. Two Crows Corporation.
- Cutler, A., Cutler, D. R., & Stevens, J. R. (2012). Random Forests. C. Zhang and Y. Q. Ma (Ed), *Ensemble machine learning* içinde (s. 157-175). New York: Springer.
- Çayıroğlu, İ. (2015). İleri algoritma analizi-5 yapay sinir ağları. Karabük Üniversitesi Mühendislik Fakültesi. <http://www.ibrahimcayiroglu.com/Dokumanlar/IleriAlgoritmaAnalizi/IleriAlgoritmaAnalizi-5.Hafta-YapaySinirAglari.pdf> sayfasından erişilmiştir.
- Çelik, M.(2009). *Veri madenciliğinde kullanılan sınıflandırma yöntemleri ve bir uygulama*. Yüksek Lisans Tezi, İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, İstanbul.
- Çelik, Ö. (2018). A Research on Machine Learning Methods and Its Applications. *Journal of Educational Technology and Online Learning*, 1(3), 25-40.

- Çelikten, S., & Çakan, M. (2019). Bayesian ve nonbayesian kestirim yöntemlerine dayalı olarak sınıflama indekslerinin TIMSS-2015 matematik testi üzerinde incelenmesi. *Necatibey Eğitim Fakültesi Elektronik Fen ve Matematik Eğitimi Dergisi*, 13(1), 105-124.
- Çiğşar, B. & D. Ünal (2019). Comparison of data mining classification algorithms determining the default risk. *Scientific Programming*, 2019. <https://doi.org/10.1155/2019/8706505>
- Çokluk, Ö. T. D., & Çırak, G. Y. (2012). *Yükseköğretimde öğrenci başarılarının sınıflandırılmasında yapay sinir ağları ve lojistik regresyon yöntemlerinin kullanılması*. Doktora Tezi, Ankara Üniversitesi Eğitim Bilimleri Enstitüsü, Ankara
- Çokluk, Ö., Şekercioğlu, G., & Büyüköztürk, Ş. (2012). *Sosyal bilimler için çok değişkenli istatistik: SPSS ve LISREL uygulamaları*. Ankara: Pegem Akademi.
- Da Rosa, J. C., Veiga, Á., & Medeiros, M. C. (2003). Three-structured smooth transition regression models based on CART algorithm (No. 469). Texto para discussão. Department of Economics PUC-Rio (Brazil).
- Da Silva, I. N., Spatti, D. H., Flauzino, R. A., Liboni, L. H. B., & dos Reis Alves, S. F. (2017). *Artificial neural networks*. Cham: Springer International.
- Dağ, H., Sayın, K. E., Yenidoğan, I., Albayrak, S., & Acar, C. (2012, July). *Comparison of feature selection algorithms for medical data*. In 2012 International Symposium on Innovations in Intelligent Systems and Applications (pp. 1-5). IEEE.
- Dash, M., & Liu, H. (1997). Feature selection for classification. *Intelligent data analysis*, 1(3), 131-156.
- De'ath, G., & Fabricius, K. E. (2000). Classification and regression trees: a powerful yet simple technique for ecological data analysis. *Ecology*, 81(11), 3178-3192

- Dhurandhar, A. & Dobra, A. (2009). *Evaluating evaluation measures*. In Proceedings of evaluation methods in machine learning workshop in international conference on machine learning (ICML) 2009.
- Doğan, N., & Sevindik, H. (2011). İlköğretim 6. sınıflar için uygulanan seviye belirleme sınavı'nın uygunluk geçerliği. *Eğitim ve Bilim*, 36(160).
- Dolgun, M. Ö. (2014). Veri madenciliği sınıflama yöntemlerinin başarılarının; bağımlı değişken prevelansı, örneklem büyüklüğü ve bağımsız değişkenler arası ilişki yapısına göre karşılaştırılması. Doktora Tezi, Hacettepe Üniversitesi Sağlık Bilimleri Enstitüsü, Ankara.
- Dondurmacı, G. (2011). *Veri madenciliği'nde regresyon ağaçları ile sınıflandırma ve bir uygulama*. Doktora Tezi, Mimar Sinan Güzel Sanatlar Üniversitesi Fen Bilimleri Enstitüsü, İstanbul.
- Dreyfus, G. (2005). *Neural networks: methodology and applications*. Springer Science & Business Media.
- Du, K. L., & Swamy, M. N. (2013). *Neural networks and statistical learning*. Springer Science & Business Media.
- Duda, R. O., Hart, P. E., & Stork, D. G. (2012). *Pattern classification*. John Wiley & Sons.
- Duman, A. (2013). *Multiple criteria sorting methods based on support vector machines*, Yüksek Lisans Tezi, Orta Doğu Teknik Üniversitesi Fen Bilimleri Enstitüsü, Ankara.
- Dunham, M. H. (2003). *Data mining introductory and advanced topics*. Upper Saddle River, NJ: Pearson Education, Inc.
- EFE, A. (2012). Evaluation of obesity risk factors using logistic regression and artificial neural networks. Doktora Tezi, DEÜ Fen Bilimleri Enstitüsü, İzmir.

- Elge, M. (2018). *Çanakkale Boğazı Ege Denizi Ve Marmara Denizi Çıkışlarında Oluşan Akıntıların Yapay Sinir Ağları Metodu Kullanılarak Tahmin Edilmesine Yönelik Sistem Geliştirilmesi*. Doktora Tezi, İstanbul Üniversitesi Deniz Bilimleri ve İşletmeciliği Enstitüsü, İstanbul.
- Elmas, Ç. (2007). *Yapay zeka uygulamaları: yapay sinir ağı, bulanık mantık, genetik algoritma*. Ankara: Seçkin.
- Emel, G. & Taşkın, Ç. (2005). Veri madenciliğinde karar ağaçları ve bir satış analizi uygulaması. *Eskişehir Osmangazi Üniversitesi Sosyal Bilimler Dergisi*, 6(2).
- Emir, Ş. (2013). *Yapay Sinir Ağları ve Destek Vektör Makineleri Yöntemlerinin Sınıflandırma Performanslarının Karşılaştırılması: Borsa Endeks Yönünün Tahmini Üzerine Bir Uygulama*. Doktora Tezi, İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, İstanbul.
- Ercikan, K., & Julian, M. (2002). Classification accuracy of assigning student performance to proficiency levels: Guidelines for assessment design. *Applied Measurement in Education*, 15, 269-294.
- Ercire, M. (2019). *Kısa süreli güç kalitesi bozulmalarının dalgacık analizi ve rastgele orman yöntemi ile sınıflandırılması*. Yüksek Lisans Tezi, Kütahya Dumlupınar Üniversitesi Fen Bilimleri Enstitüsü, Kütahya.
- Erkuş, A. (2003). *Psikometri üzerine yazılar*. Ankara: Türk Psikologlar Derneği.
- Ersöz, F. (2019). *Veri Madenciliği Teknikleri ve Uygulamaları*. Ankara: Seçkin.
- Fakı, B. M.. & Öztayşı, B. (2015, Şubat). *Veri madenciliği yöntemlerini kullanarak anemi sınıflandırılmasına yönelik bir uygulama*. XVII. Akademik Bilişim Konferansı'nda sunulmuş bildiri, Anadolu Üniversitesi, Eskişehir.
- Fayyad, U., Piatetsky-Shapiro, G., & Smyth, P. (1996). From data mining to knowledge discovery in databases. *AI magazine*, 17(3), 37-37.
- Feinstein, A.R. (1996) *Multivariate Analysis: An Introduction*. New Haver: Yale University.

- Field, A. (2013). *Discovering statistics using IBM SPSS statistics*. Sage.
- Gennari, J. H. Langley, P. & Fisher, D. (1989). Models of incremental concept formation. *Artificial Intelligence*, 40(1), 11-61.
- George, D., & Mallery, P. (2020). *IBM SPSS statistics 23 step by step: A simple guide and reference*. Routledge.
- Gislason, P. O., Benediktsson, J. A., & Sveinsson, J. R. (2006). Random forests for land cover classification. *Pattern Recognition Letters*, 27(4), 294-300
- Gore, P.A. (2000). Cluster analysis. H.E.A. Tinsley & S.D. Brown (Ed), *Applied Multivariate Statistics and Mathematical Modeling* içinde (s. 297–321). San Diego: Academic.
- Gumma, S., 2004. A Radial Basis Function Neuro Controller for Permenent Magnet Stepper Motor (yüksek lisans tezi basılmamış) Master of Science in Electrical Engineering. Cleveland State University, India.
- Guo, B., Zhang, R., Xu, G., Shi, C., & Yang, L. (2015, July). Predicting students performance in educational data mining. In *2015 International Symposium on Educational Technology (ISET)* (pp. 125-128). IEEE.
- Gurney, K. (2014). *An introduction to neural networks*. CRC.
- Güler, Ç. (2017). *Ortaöğretim başarısını etkileyen faktörlerin karar ağacı ile sınıflandırılması*. Yüksek Lisans Tezi, Kahramanmaraş Sütçü İmam Üniversitesi Fen Bilimleri Enstitüsü, Kahramanmaraş.
- Güneri, N., & Apaydın, A. (2004). Öğrenci başarılarının sınıflandırılmasında lojistik regresyon analizi ve sinir ağı yaklaşımı. *Ticaret ve Turizm Eğitim Fakültesi Dergisi*, 1, 170-188.
- Hall, M. A. (1999) *Correlation-based feature selection for machine learning*. Doktora Tezi, Department of Computer Science Hamilton, New Zealand, University of Waikato.

- Hallett, M. J., Fan, J. J., Su, X. G., Levine, R. A., & Nunn, M. E. (2014). Random forest and variable importance rankings for correlated survival data, with applications to tooth loss. *Statistical Modelling*, 14(6), 523-547.
- Hamel, L. H. (2011). *Knowledge discovery with support vector machines (Vol. 3)*. John Wiley & Sons.
- Han, J., Kamber, M., 2001. *Data Mining Concepts and Techniques*. Morgan Kaufman.
- Han, J., Pei, J., & Kamber, M. (2012). *Data mining: concepts and techniques*. Elsevier.
- Hand, D. J., Mannila, H., & Smyth, P. (2001). *Principles of data mining (adaptive computation and machine learning)*. MIT.
- Haykin, S. (1999). *Neural networks: a comprehensive foundation*. Prentice Hall PTR.
- Haykin, S. S. (2009). *Neural networks and learning machines*. Prentice Hall
- Hazım, L.R. (2018). *Four classification methods naïve bayesian, support vector machine, k-nearest neighbors and random forest are tested for credit card fraud detection*. Altınbas University Graduate School Of Sciences Engineering
- He, H., & Garcia, E. A. (2009). Learning from imbalanced data. *IEEE Transactions on Knowledge and Data Engineering*, 21(9), 1263–1284.
- He, X. & Xu, S. (2010). *Process neural networks: Theory and applications*. Springer Science & Business Media.
- Hekimoğlu, C, K. (2017). Demand prediction in clothing industry with using neural networks. Yüksek Lisans Tezi, Yaşar Üniversitesi Fen Bilimleri Enstitüsü, İstanbul.
- Herzog, S. (2006). Estimating student retention and degree-completion time: Decision trees and neural networks vis-à-vis regression. In *New Directions for Institutional Research*, 2006(131), (pp. 17–33).

- Holden, J.E. & Kelley, K. (2010). The effects of initially misclassified data on the effectiveness of ayırma function analysis and finite mixture modeling. *Educational and Psychological Measurement*, 70(1), 36-55.
- Hosmer Jr, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). *Applied logistic regression* (Vol. 398). John Wiley & Sons.
- Huang, T. M., Kecman, V., & Kopriva, I. (2006). *Kernel based algorithms for mining huge data sets* (Vol. 1). Heidelberg: Springer.
- Ihantola, P., Vihavainen, A., Ahadi, A., Butler, M., Börstler, J., Edwards, S. H., ... Toll, D. (2015). *Educational Data Mining and Learning Analytics in Programming: Literature Review and Case Studies*. In Proceedings of the 2015 ITiCSE on Working Group Reports (pp. 41–63). New York, NY, USA: ACM. <https://doi.org/10.1145/2858796.2858798>
- Jain, N., & Srivastava, V. (2013). Data mining techniques: a survey paper. *IJRET: International Journal of Research in Engineering and Technology*, 2(11), 2319-1163.
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2013). *An introduction to statistical learning*. New York: Springer.
- Kalaycı, Ş. (2005). *SPSS uygulamalı çok değişkenli istatistik teknikleri*. Ankara: Asil.
- Kane, M. T. (2013). Validating the interpretations and uses of test scores. *Journal of Educational Measurement*, 50(1).
- Kantardzic, M. (2011). *Data mining: concepts, models, methods, and algorithms*. Şehir: John Wiley & Sons.
- Karakaynak, Z. (2014). *Destek vektör makineleri ile sınıflandırma ve görüntü tanıma üzerine bir uygulama*. Yüksek Lisans Tezi, Mimar Sinan Güzel Sanatlar Üniversitesi Fen Bilimleri Enstitüsü, İstanbul.

- Karakış, R. (2009). *Yapay sinir ağları ve lojistik regresyon yöntemleri ile meme kanseri koltuk altı lenf nodu durumunun belirlenmesi*. Yüksek Lisans Tezi, Gazi Üniversitesi Bilişim Enstitüsü, Ankara.
- Karunanithi, N., Whitley, D., & Malaiya, Y. K. (1992). Using neural networks in reliability prediction. *IEEE Software*, 9(4), 53-59.
- Kayri, M., Kayri, I., & Gencoglu, M. T. (2017). The performance comparison of multiple linear regression, random forest and artificial neural network by using photovoltaic and atmospheric data. In *2017 14th International Conference on Engineering of Modern Electric Systems (EMES)*, 1-4, IEEE.
- Kaur, R., & Ginige, J. A. (2018). Comparative evaluation of accuracy of selected machine learning classification techniques for diagnosis of cancer: a data mining approach. world academy of science, engineering and technology. *International Journal of Biomedical and Biological Engineering*, 12(2), 19-25.
- Kelecioğlu, H., ve Şahin, S. G. (2014). Geçmişten günümüze geçerlik. *Eğitimde ve Psikolojide Ölçme ve Değerlendirme Dergisi*, 5(2), 1-11.
- Khalili, Z. & Moradi, M.H. (2009). Emotion recognition system using brain and peripheral signals: Using correlation dimension to improve the results of EEG. In *Proceedings of the 2009 International Joint Conference on Neural Networks*, Atlanta, GA, USA, 14–19 June 2009; pp. 1571–1575.
- Khare, M., & Nagendra, S. S. (2006). *Artificial neural networks in vehicular pollution modelling (Vol. 41)*. Springer.
- Khemchandani, R., & Chandra, S. (2017). *Twin Support Vector Machines*. Berlin: Springer.
- Kılıç-Depren, S., Aşkın, Ö. E., & Öz, E. (2017). Identifying the classification performances of educational data mining methods: A case study for TIMSS. *Educational Sciences: Theory & Practice*, 17, 1605–1623. <http://dx.doi.org/10.12738/estp.2017.5.0634>

- Kıran, Z. B. (2010). *Lojistik regresyon ve CART analizi teknikleriyle sosyal güvenlik kurumu ilaç provizyon sistemi verileri üzerinde bir uygulama*. Yüksek lisans tezi, Gazi Üniversitesi Fen Bilimleri Enstitüsü, Ankara.
- Kohavi, R. and G. John. (1996). Wrappers for feature subset selection. *Artificial Intelligence, special issue on relevance*, 97(1–2), 273–324.
- Korkem, E. (2013). *Mikroarray gen ekspresyon veri setlerinde random forest ve naive bayes sınıflama yöntemleri yaklaşım*. Yüksek Lisans Tezi, Hacettepe Üniversitesi Sağlık Bilimleri Enstitüsü, Ankara.
- Korkmaz, D. (2018). *Sınıflandırma ve regresyon ağaçları ile rastgele orman algoritması kullanarak botnet tespiti: van yüzüncü yıl üniversitesi örneği*. Yüksek Lisans Tezi Yüzüncü Yıl Üniversitesi Fen Bilimleri Enstitüsü, Van.
- Kotsiantis S. B. 2007 Supervised machine learning: a review of classification techniques. *Informatica*, 31.
- Koyuncu, İ. (2018). *Öğrencilerin PISA matematik başarılarının yordanmasında veri madenciliği yöntemlerinin karşılaştırılması*. Doktora tezi, Hacettepe Üniversitesi, Eğitim Bilimleri Enstitüsü, Ankara.
- Köktürk, F. (2012). K-En yakın komşuluk, yapay sinir ağları ve karar ağaçları yöntemlerinden sınıflandırma başarılarının karşılaştırılması. Doktora Tezi, Bülent Ecevit Üniversitesi Sağlık Bilimleri Enstitüsü, Zonguldak.
- Köse, İ. (2018). *Veri madenciliği teori, uygulama ve felsefesi*. Ankara: Papatya.
- Krenker, A., Bešter, J., & Kos, A. (2011). Introduction to the artificial neural networks. *Artificial Neural Networks: Methodological Advances and Biomedical Applications*. InTech, 1-18.
- Kumar, S., & Chong, I. (2018). Correlation analysis to identify the effective data in machine learning: Prediction of depressive disorder and emotion states. *International journal of environmental research and public health*, 15(12), 2907.

- Kuyucu, Y. E. (2012). *Lojistik regresyon analizi (LRA), yapay sinir ağıları (YSA) ve sınıflandırma ve regresyon ağaçları (C&RT) yöntemlerinin karşılaştırılması ve tıp alanında bir uygulama*. Yüksek Lisans Tezi, Gaziosmanpaşa Üniversitesi Sağlık Bilimleri Enstitüsü, Tokat.
- Kuzey, C. (2012). *Veri madenciliğinde destek vektör makinaları ve karar ağaçları yöntemlerini kullanarak bilgi çalışanlarının kurum performansı üzerine etkisinin ölçülmesi ve bir uygulama*. Doktora Tezi, İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, İstanbul.
- Kwon, O., & Sim, J. M. (2013). Effects of data set features on the performances of classification algorithms. *Expert Systems with Applications*, 40(5), 1847-1857.
- Larose, D. T. (2005). *An introduction to data mining*. Traduction et adaptation de Thierry Vallaud.
- Lathrop, Q. N., & Cheng, Y. (2014). A nonparametric approach to estimate classification accuracy and consistency. *Journal of Educational Measurement*, 51, 318-334
- Lee, W., Hanson, B. A., & Brennan, R. L. (2000). *Procedures for computing classification consistency and accuracy indices with multiple categories*. ACT: Inc, Research Report.
- Lewis, C. 1986. Test theory and psychometrika: The past twenty-five years. *Psychometrika*, 51, 11-22.
- Lewis, N. D. (2017). *Neural networks for time series forecasting with r: an intuitive step by step blueprint for beginners*. AusCov.
- Li, Y., & Hung, E. (2009, November). Building a decision cluster forest model to classify high dimensional data with multi-classes. In *Asian Conference on Machine Learning* (pp. 263-277). Berlin, Heidelberg: Springer.
- Liaw, A., & Wiener, M. (2002). Classification and regression by randomForest. *R news*, 2(3), 18-22.

- Liu, H., & Motoda, H. (Eds.). (1998). *Feature extraction, construction and selection: A data mining perspective* (Vol. 453). Springer Science & Business Media.
- Luan, J. (2002). Data mining and its applications in higher education. *New directions for institutional research*, 113, 17-36.
- Lutu, P. E. N. (2010). *Dataset selection for aggregate model implementation in predictive data mining*. Doktora Tezi, University of Pretoria.
- Mannila, H. (2000). Theoretical frameworks for data mining. *Explorations*, 1(2), 30–32.
- Maroco, J., Silva, D., Rodrigues, A., Guerreiro, M., Santana, I. & de Mendoça, A. (2011). Data mining methods in the prediction of Dementia: A real-data comparison of the accuracy, sensitivity and specificity of linear discriminant analysis, logistic regression, neural networks, support vector machines, classification trees and random forests. *BMC Research Notes* 4, 299. <https://doi.org/10.1186/1756-0500-4-299>
- Mavroforakis, M. E., & Theodoridis, S. (2006). A geometric approach to support vector machine (SVM) classification. *IEEE transactions on neural networks*, 17(3), 671-682.
- Mehrotra, K., Mohan, C. K., & Ranka, S. (1997). *Elements of artificial neural networks*. MIT.
- Melgani, F., & Bruzzone, L. (2004). Classification of hyperspectral remote sensing images with support vector machines. *IEEE Transactions on geoscience and remote sensing*, 42(8), 1778-1790.
- Menard, S. (2002). *Applied logistic regression analysis* (Vol. 106). Sage.
- Mertler, C. , & Vannatta, R. A. (2017). *Advanced and multivariate statistical methods: Practical application and interpretation*. New York: Taylor & Francis.

- Minaei-Bidgoli, B., Kashy, D. A., Kortemeyer, G., & Punch, W. F. (2003). *Predicting student performance: an application of data mining methods with an educational web-based system*. In 33rd Annual Frontiers in Education, 2003. FIE, Westminster, CO (Vol. 1, pp. T2A-13). IEEE. <https://ieeexplore.ieee.org/document/1263284>
- Mislevy, R. J. (1994). Evidence and inference in educational assessment. *Psychometrika*, 59, 439-483.
- Mislevy, R. J., Behrens, J. T., Dicerbo, K. E., & Levy, R. (2012). Design and discovery in educational assessment: Evidence-centered design, psychometrics, and educational data mining. *Journal of educational data mining*, 4(1), 11-48.
- Mohamed, M. Y. (2015). *Comparison of classification algorithms for predicting diabetes. Turkish republic erciyes university graduate school of natural and applied science department of computer engineering*. Yüksek Lisans Tezi, Erciyes Üniversitesi Fen Bilimleri Enstitüsü, Kayseri.
- Morgan, J., & J. Sonquist (1963): Problems in the analysis of survey data and a proposal. *Journal of the American Statistical Association*, 58, 415–434.
- Morgan, J., Dougherty, R., Hilchie, A., & Carey, B. (2003). Sample size and modeling accuracy with decision tree based data mining tools. *Academy of Information and Management Sciences Journal*, 6(2), 77–91.
- Murphy, R. K. & Davidshofer, O. C. (2005). *Psychological testing: principles and application*. New Jersey: Prentice-Hall Book Company.
- Nöu, A. (2018). *Logistic regression versus support vector machines*. Lisans tezi, Stockholm University Matematik Enstitüsü, Stokholm.
- Ocañoğlu, G. (2006). *Lojistik regresyon analizi ve yapay sinir ağları tekniklerinin sınıflama özelliklerinin karşılaştırılması ve bir uygulama*. Yüksek Lisans Tezi, Uludağ Üniversitesi Sağlık Bilimleri Enstitüsü, Bursa.

- Ocuppaugh, J., Baker, R., Gowda, S., Heffernan, N., & Heffernan, C. (2014). Population validity for Educational Data Mining models: A case study in affect detection. *British Journal of Educational Technology*, 45(3), 487-501.
- Okamoto, M. (1963). An asymptotic expansion for the distribution of the linear discriminant function. *The Annals of Mathematical Statistics*, 34(4), 1286–1301.
- Olson, D. L., & Delen, D. (2008). *Advanced data mining techniques*. Springer Science & Business Media.
- Ooi, C. H., Chetty, M. & Teng, S. W. (2007). Differential prioritization in feature selection and classifier aggregation for multiclass microarray datasets. *Data Mining and Knowledge Discovery*, 14, 329-366.
- Orekeci-Temel, G. (2004). *Sınıflama ve regresyon ağaçları*. Yüksek Lisans Tezi, Mersin Üniversitesi Sağlık Bilimleri Enstitüsü, Mersin.
- Özçalıcı, M. & Ayriçay, Y. (2016). Bilgi işlemsel zekâ yöntemleri ile hisse senedi fiyat tahmini: bist uygulaması. *Pamukkale University Journal of Social Sciences Institute/Pamukkale Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, 25(1), 274-298.
- Özdamar, K. (2013). *Paket programlar ile istatistiksel veri analizi: SPSS-MINITAB*. Eskişehir: Kaan.
- Özen, S, B. (2018). *KASKİ Atık Su Arıtma Verilerinin Yapay Sinir Ağları Ve Bulanık Mantık Yöntemleri İle Tahmin Edilmesi*. Yüksek Lisans Tezi, Erciyes Üniversitesi Fen Bilimleri Enstitüsü, Kayseri.
- Özkan, Y. (2016). *Veri madenciliği yöntemleri*. Ankara: Papatya.
- Özmen, E. (2018). *PISA 2012'de yer alan duyuşsal özelliklerin matematik başarısını sınıflama doğruluğunun incelenmesi: Şangay, İspanya ve Peru örneği*. Yüksek lisans tezi, Hacettepe Üniversitesi, Eğitim Bilimleri Enstitüsü. Ankara.
- Öztemel, E. (2016). *Yapay sinir ağları*. İstanbul: Papatya.

- Pal, M. (2005). Random forest classifier for remote sensing classification. *International Journal of Remote Sensing*, 26(1), 217-222.
- Pal, M., & Mather, P. M. (2003). An assessment of the effectiveness of decision tree methods for land cover classification. *Remote sensing of environment*, 86(4), 554-565.
- Pektaş, O. A. (2013). SPSS 20 ile veri madenciliği. İstanbul: Dikeyksen.
- Pena-Ayala, A. (2014). Educational data mining: A survey and a data mining-based analysis of recent works. *Expert System with Applications*, 41(4), 1432–1462.
- Pranckevicius, T., & Marcinkevicius, V. (2017). Comparison of naive bayes, random forest, decision tree, support vector machines, and logistic regression classifiers for text reviews classification. *Balt. J. Mod. Comput.*, 5. DOI:10.22364/bjmc.2017.5.2.05
- Prasad, A. M., Iverson, L. R., & Liaw, A. (2006). Newer classification and regression tree techniques: bagging and random forests for ecological prediction. *Ecosystems*, 9(2), 181-199.
- Rao, V. & Rao, H., (1998). *C++ Neural network and fuzzy logic*. New Delhi: BPB.
- Raudys, S. Pikelis, V. (1980). On dimensionality, sample size, classification error, and complexity of classification algorithm in pattern recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-2(3), 242–252.
- Rojas, R. (1996). *Neural networks: a systematic introduction*. Springer.
- Romero, C., & Ventura, S. (2007). Educational data mining: A survey from 1995 to 2005. *Expert Systems with Applications*, 33(1), 135–146.
- Romero, C., & Ventura, S. (2010). Educational data mining: a review of the state of the art. *IEEE Transactions on Systems, Man, and Cybernetics, Part C (Applications and Reviews)*, 40(6), 601-618.
- S. P. S. S. (2000). *Step-by-step Data Mining Guide*.

- Sarle, W. S. (1994). *Neural networks and statistical models*. Proceedings of the Nineteenth Annual SAS Users Group International Conference, Cary, NC: SAS Institute, USA, pp. 1538-1550.
- Sayıcı, G. (2013). *Karar ağaçları, bayes ağları ve etki diyagramları aracılığı ile bilgi keşfi ve karar verme*. Yüksek Lisans Tezi, İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, İstanbul.
- Scholkopf, B., & Smola, A. J. (2001). *Learning with kernels: support vector machines, regularization, optimization, and beyond*. MIT.
- Segal, M. R. (1992): Tree-Structured methods for longitudinal data. *Journal of the American Statistical Association*, 87, 407–418.
- Shim. J., Warkentin. M., Courtney. J., Power. D., Sharda. R., & Carlson. C. (2002). Past, present, and future of decision support technology. *Decision Support Systems*, 33(2), 111-126.
- Silahtaroglu, G. (2016). *Veri madenciliği: Kavram ve algoritmaları*. Ankara: Papatya.
- Sinharay, S. (2016). An NCME instructional module on data mining methods for classification and regression. *Educational Measurement: Issues and Practice*, 35(3), 38–54.
- Souza, C. R. (2010). Kernel functions for machine learning applications. *Creative Commons Attribution-Noncommercial-Share Alike*, 3, 29.
- Souza, J., Matwin, S., and Japkowicz, N. (2002). *Evaluating data mining models: a pattern language*. In: 9th Conference on Pattern Language of Programs (PLOP'02), Monticello, Illinois, 8–12 September 2002.
- Strobl, C., & Zeileis, A. (2008). *Danger: High power! – exploring the statistical properties of a test for random forest variable importance*. COMPSTAT 2008-Proceedings in Computational Statistics.

- Şimşek-Gürsoy, U., T. (2009). *Veri madenciliği ve bilgi keşfi (1. Baskı)*. Ankara: Pegem Akademi.
- Şimşek-Gürsoy, U., T. (2012). *Uygulamalı veri madenciliği (3. Baskı)*. Ankara: Pegem Akademi.
- Tabachnick, B. G., Fidell, L. S., & Ullman, J. B. (2013). *Using multivariate statistics (Vol. 6)*. Boston, MA: Pearson.
- Taylor, G. (2002). Neural Networks. Taylor, G. Shadbolt, J (Ed.), *Neural Networks and the Financial Markets: Bpredicting, Combining, and Portfolio Optimisation* içinde (s. 87-95). Springer Science & Business Media.
- Thorndike R. L.ve Hagen, E. (1961). *Measurement and evaluation in psychology and education*. Newyork: John Wiley and sons.
- Toprak, E. (2017). *Yapay sinir ağı, karar ağaçları ve ayırma analizi yöntemleri ile PISA 2012 matematik başarılarının sınıflandırılma performanslarının karşılaştırılması*. Doktora tezi, Hacettepe Üniversitesi, Eğitim Bilimleri Enstitüsü, Ankara.
- Tosun, S. (2007). *Sınıflandırmada yapay sinir ağları ve karar ağaçları karşılaştırması: öğrenci başarıları üzerine bir uygulama*. Doktora Tezi, İTÜ Fen Bilimleri Enstitüsü, İstanbul.
- Turgut, M. F. & Baykul, Y. (1992). *Ölçekleme teknikleri*. Ankara: OSYM.
- Turgut, M. F. ve Baykul, Y. (2013). *Eğitimde Ölçme ve Değerlendirme*. Ankara: Pegem
- Vercellis, C. (2009). *Business intelligence: data mining and optimization for decision making*. New York: Wiley.
- Wang, L. ed., 2005. *Support vector machines: theory and applications (Vol. 177)*. Springer Science & Business Media.

- Watts, J. D. ve Lawrence, R. L., 2008. Merging random forest classification with an object oriented approach for analysis of agricultural lands, *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, XXXVII, B7.
- Wegman, E., J. (1988). Computational statistics: A new agenda for statistical theory and practice. *Journal of the Washington Academy of Science*, 78, 310-322.
- Wegman, E., J. (2000). Visions: New techniques and technologies in statistics. *Computer Statistics*, 15(1), 133-144.
- Weng, C. G. & Poon, J. (2006). *A data complexity analysis on imbalanced datasets and an alternative imbalance recovering strategy*. In *Proceedings of the 2006 IEEE/WIC/ACM international conference on web, intelligence* (pp. 270–276).
- Westphal, C., & Blaxton, T. (1998). *Data mining solutions: methods and tools for solving real-world problems*. New york: Wiley.
- Whitley, L. A. (2018). *Educational data mining and its uses to predict the most prosperous learning environment*. M.Sc. Thesis, East Carolina University in USA
- Wu, C. H., & McLarty, J. W. (Ed). (2012). *Neural networks and genome informatics*. Şehir: Elsevier.
- Xu, Y. (2003) *Using data mining in educational research: A comparison of bayesian network with multiple regression in prediction*. Doctoral thesis, Department of Educational Psychology, The University of Arizona, Arizona.
- Yakut, E. & Gemici, E. (2017). LR, C5.0, CART, DVM yöntemlerini kullanarak hisse senedi getiri sınıflandırma tahmini yapılması ve kullanılan yöntemlerin karşılaştırılması: Türkiye’de BIST’de bir uygulama. *Ege Akademik Bakış*, 17(4).
- Yan, J. (2017). *Multi-class ROC random forest for imbalanced classification*. Doctoral dissertation, The Graduate School, Stony Brook University: Stony Brook, NY.

- Yang, H., & Ni, J. (2005). Dynamic neural network modeling for nonlinear, nonstationary machine tool thermally induced error. *International Journal of Machine Tools and Manufacture*, 45(4-5), 455-465.
- Yeşilkanat, C. M. (2016). *Jeoistatistik analiz, yapay sinir ağları ve bulanık mantık yaklaşımları* kullanarak çevresel radyoaktivitenin aradeğerleme modellemesi ve haritalandırılması. Doktora Tezi, Karadeniz Teknik Üniversitesi Fen Bilimleri Enstitüsü, Trabzon.
- Yılmaz, I. (2009). Landslide susceptibility mapping using frequency ratio, logistic regression, artificial neural networks and their comparison: A case study from Kat landslides (Tokat—Turkey). *Computers & Geosciences*, 35(6), 1125–1138. doi: 10.1016/j.cageo.2008.08.007
- Yohannes, Y., & Webb, P. (1999). *Classification and regression trees, CART: a user manual for identifying indicators of vulnerability to famine and chronic food insecurity* (Vol. 3). Intl Food Policy Res Inst.
- Yu, L., & Liu, H. (2003). *Feature selection for high-dimensional data: A fast correlation-based filter solution*. In Proceedings of the 20th international conference on machine learning (ICML-03) (pp. 856-863).
- Yuhas, B., & Ansari, N. (Eds.). (2012). *Neural networks in telecommunications*. Springer Science & Business Media.
- Yükseltürk, E., Ozekeş, S., & Türel, Y. K. (2014). Predicting dropout student: an application of data mining methods in an online education program. *European Journal of Open, Distance and e-learning*, 17(1), 118-133.
- Zhang, B. L., & Zhang, Y. (2008). Classification of cerebral palsy gait by kernel fisher discriminant analysis. *International journal of hybrid intelligent systems*, 5(4), 209-218.
- Zhao, C. M., & Luan, J. (2006). Data Mining: Going beyond Traditional Statistics. *New Directions for Institutional Research*, 131, 7-16.

Zhi, D., Yan, H., & Yan, L. (2016). Analysis on a new data mining algorithm of the statistics work. *Telkomnika*, 14(3A), s. 42-46.

Zurada, J. M. (1992). *Introduction to artificial neural systems (Vol. 8)*. St. Paul: West.

EKLER

EK 1. Değişkenlere ait betimsel istatistikler

	Mean	Variance	Std. Deviation	Minimum	Maximum	Skewness	Kurtosis
Fenortalama	510,34	8083,70	89,91	212,99	817,87	-0,08	-0,46
matortalama	505,43	7690,37	87,69	163,61	826,34	-0,12	-0,30
okumaortalama	509,61	7404,54	86,05	145,89	785,62	-0,29	-0,20
Disciplinary climate in science classes (WLE)	0,06	0,90	0,95	-2,42	1,88	-0,15	0,02
Teacher support in a science classes of students choice (WLE)	0,08	0,94	0,97	-2,72	1,45	-0,31	-0,26
Inquiry-based science teaching an learning practices (WLE)	0,00	0,90	0,95	-3,34	3,18	0,05	2,46
Teacher-directed science instruction (WLE)	0,06	0,95	0,97	-2,45	2,08	-0,06	0,31
Environmental Awareness (WLE)	0,23	1,23	1,11	-3,40	3,29	0,63	1,52
Environmental optimism (WLE)	-0,02	1,47	1,21	-1,79	3,01	0,35	0,02
Enjoyment of science (WLE)	0,23	1,12	1,06	-2,12	2,16	-0,11	-0,12
Interest in broad science topics (WLE)	0,18	0,80	0,90	-2,58	2,73	-0,50	1,61
Instrumental motivation (WLE)	0,32	0,93	0,97	-1,93	1,74	-0,25	-0,43
Science self-efficacy (WLE)	0,15	1,32	1,15	-3,76	3,28	0,17	1,86
Epistemological beliefs (WLE)	0,09	0,91	0,95	-2,79	2,16	0,09	0,72
Index science activities (WLE)	0,16	1,19	1,09	-1,77	3,36	-0,15	-0,22
Out-of-School Study Time per week (Sum)	18,58	188,57	13,73	0,00	70,00	1,22	1,37
Learning time (minutes per week) - <science>	235,55	16652,67	129,05	0,00	1950,00	1,51	6,21
Learning time (minutes per week) - in total	1650,58	144056,35	379,55	0,00	3000,00	0,40	2,39
Subjective well-being: Sense of Belonging to School (WLE)	0,05	1,04	1,02	-3,15	2,64	0,40	1,11
Personality: Test Anxiety (WLE)	0,14	0,88	0,94	-2,51	2,55	0,08	0,78
Student Attitudes, Preferences and Self-related beliefs: Achieving motivation (WLE)	0,12	0,92	0,96	-3,09	1,85	0,15	-0,33
Collaboration and teamwork dispositions: Enjoy cooperation (WLE)	0,12	0,95	0,98	-3,33	2,29	0,42	-0,11
Collaboration and teamwork dispositions: Value cooperation (WLE)	0,11	0,98	0,99	-2,83	2,14	0,11	-0,02
Parents emotional support (WLE)	0,03	0,93	0,97	-3,08	1,10	-0,59	-0,32
Perceived Feedback (WLE)	0,08	0,98	0,99	-1,53	2,50	0,25	-0,25
Adaption of instruction (WLE)	0,08	0,95	0,98	-1,97	2,05	0,00	-0,12
Teacher Fairness (Sum)	9,53	12,26	3,50	1,00	24,00	1,35	1,90
Cultural possessions at home (WLE)	0,08	0,93	0,96	-1,84	2,63	0,45	0,11
Home educational resources (WLE)	0,02	0,94	0,97	-4,39	1,18	-0,34	-0,39
Home possessions (WLE)	-0,06	1,04	1,02	-8,94	5,63	-0,29	1,80
ICT Resources (WLE)	-0,07	1,02	1,01	-3,38	3,51	0,25	1,39
Family wealth (WLE)	-0,09	1,12	1,06	-6,98	4,44	-0,05	1,73
Index of economic, social and cultural status (WLE)	-0,04	1,05	1,02	-5,48	3,27	-0,48	-0,23

EK 2. Korelasyon katsayısı 0,20'den küçük olan değişkenler

Değişken	Fen ortalama Pearson korelasyon	
	r	p
Personality: Test Anxiety (WLE)	-0,17	0,00
Environmental optimism (WLE)	-0,14	0,00
Perceived Feedback (WLE)	-0,14	0,00
Disciplinary climate in science classes (WLE)	0,14	0,00
Collaboration and teamwork dispositions: Value cooperation (WLE)	-0,13	0,00
Collaboration and teamwork dispositions: Enjoy cooperation (WLE)	0,11	0,00
Teacher-directed science instruction (WLE)	0,11	0,00
Student Attitudes, Preferences and Self-related beliefs: Achieving motivation (WLE)	0,10	0,00
Instrumental motivation (WLE)	0,07	0,00
Parents emotional support (WLE)	0,06	0,00
Adaption of instruction (WLE)	0,06	0,00
Teacher support in a science classes of students choice (WLE)	-0,03	0,00
Index science activities (WLE)	-0,02	0,00

Değişken	Spearman korelasyon	
	r	p
Learning time (minutes per week) - <Mathematics>	0,09	0,00
Subjective well-being: Sense of Belonging to School (WLE)	0,07	0,00
Learning time (minutes per week) - in total	0,04	0,00
Learning time (minutes per week) - <test language>	0,03	0,00
Out-of-School Study Time per week (Sum)	-0,11	0,00
Inquiry-based science teaching an learning practices (WLE)	-0,12	0,00
Teacher Fairness (Sum)	-0,14	0,00

EK 3. Bağımlı değişken koşuluna göre değerlendirme kriterlerine ait betimsel istatistikler

Bağımlı değişkenin 6, 3 ve 2 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen doğruluk, kesinlik ve F ölçütü değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Doğruluk	6	YSA	0,6815	0,0001	0,0114	0,6593	0,6954	-0,2734	-0,9897
		RO	0,6395	0,0001	0,0104	0,6166	0,6600	-0,3557	-0,2329
		DVM	0,6804	0,0001	0,0095	0,6575	0,6984	-0,5892	0,4493
		SVRA	0,5951	0,0003	0,0167	0,5769	0,6166	0,1045	-1,6118
		LR	0,6813	0,0001	0,0094	0,6508	0,6941	-1,7251	3,8539
	3	YSA	0,8457	0,0001	0,0092	0,8211	0,8608	-0,6837	1,0985
		RO	0,8021	0,0013	0,0362	0,7448	0,8388	-0,6343	-1,3194
		DVM	0,8487	0,0001	0,0087	0,8309	0,8639	-0,3198	-0,1551
		SVRA	0,7191	0,0005	0,0216	0,6947	0,7460	0,3078	-1,6522
		LR	0,8484	0,0001	0,0077	0,8327	0,8614	-0,2193	-0,3705
	2	YSA	0,9311	0,0007	0,0265	0,8926	0,9579	0,0165	-2,0623
		RO	0,8850	0,0003	0,0175	0,8529	0,9072	-1,1330	-0,0370
		DVM	0,9367	0,0008	0,0282	0,8919	0,9609	-0,3352	-1,9323
		SVRA	0,7791	0,0000	0,0042	0,7766	0,7857	1,0437	-0,9976
		LR	0,9264	0,0002	0,0154	0,9029	0,9396	-0,3831	-1,8390
Kesinlik	6	YSA	0,6785	0,0001	0,0109	0,6551	0,6922	-0,2480	-0,6514
		RO	0,6393	0,0001	0,0113	0,6136	0,6637	-0,2978	0,2129
		DVM	0,6825	0,0001	0,0094	0,6590	0,6995	-0,6448	0,6199
		SVRA	0,5916	0,0003	0,0175	0,5722	0,6136	0,0296	-1,6631
		LR	0,6836	0,0001	0,0095	0,6525	0,6995	-1,6542	4,1050
	3	YSA	0,8448	0,0001	0,0095	0,8186	0,8600	-0,8159	1,4020
		RO	0,8009	0,0013	0,0357	0,7448	0,8375	-0,6102	-1,3658
		DVM	0,8481	0,0001	0,0090	0,8296	0,8643	-0,2566	-0,0715
		SVRA	0,7095	0,0021	0,0459	0,6385	0,7463	-0,9396	-1,0544
		LR	0,8477	0,0001	0,0079	0,8311	0,8610	-0,2078	-0,3208
	2	YSA	0,9584	0,0003	0,0183	0,9347	0,9766	-0,0779	-1,9660
		RO	0,9164	0,0002	0,0150	0,8884	0,9344	-1,2340	0,1511
		DVM	0,9595	0,0004	0,0200	0,9325	0,9767	-0,3063	-2,0112
		SVRA	0,8363	0,0003	0,0170	0,8096	0,8468	-1,0437	-0,9976
		LR	0,9506	0,0001	0,0108	0,9325	0,9596	-0,5284	-1,5845
F Ölçütü	6	YSA	0,6787	0,0001	0,0117	0,6564	0,6932	-0,2066	-1,0244
		RO	0,6356	0,0001	0,0102	0,6138	0,6553	-0,3018	-0,4272
		DVM	0,6796	0,0001	0,0097	0,6565	0,6984	-0,5633	0,4740
		SVRA	0,5918	0,0003	0,0173	0,5727	0,6138	0,0541	-1,6463
		LR	0,6806	0,0001	0,0095	0,6500	0,6930	-1,7478	3,7952
	3	YSA	0,8443	0,0001	0,0092	0,8191	0,8585	-0,7557	1,4343
		RO	0,7998	0,0012	0,0348	0,7446	0,8357	-0,6327	-1,3030
		DVM	0,8469	0,0001	0,0088	0,8290	0,8625	-0,3533	-0,1506
		SVRA	0,7047	0,0020	0,0442	0,6382	0,7453	-0,7531	-1,1540
		LR	0,8467	0,0001	0,0078	0,8303	0,8596	-0,2616	-0,3590
	2	YSA	0,9513	0,0004	0,0188	0,9235	0,9702	0,0156	-2,0566
		RO	0,9196	0,0001	0,0120	0,8978	0,9359	-1,0651	-0,1240
		DVM	0,9554	0,0004	0,0199	0,9233	0,9725	-0,3435	-1,9117
		SVRA	0,8470	0,0001	0,0075	0,8424	0,8588	1,0437	-0,9976
		LR	0,9482	0,0001	0,0109	0,9315	0,9575	-0,3830	-1,8386

Bağımlı değişkenin 6, 3 ve 2 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen Kappa, belirleyicilik ve negatif öngörü değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Kappa	6	YSA	0,5847	0,0002	0,0143	0,5561	0,6023	-0,2911	-0,8674
		RO	0,5287	0,0002	0,0139	0,4976	0,5574	-0,3189	0,0179
		DVM	0,5820	0,0002	0,0126	0,5523	0,6056	-0,5805	0,3580
		SVRA	0,3877	0,0004	0,0209	0,4477	0,4976	0,1603	-1,5737
		LR	0,5830	0,0002	0,0124	0,5432	0,5999	-1,6698	3,6334
	3	YSA	0,7230	0,0003	0,0161	0,6795	0,7484	-0,7049	1,3203
		RO	0,6437	0,0039	0,0626	0,5442	0,7083	-0,6367	-1,2938
		DVM	0,7271	0,0002	0,0157	0,6953	0,7541	-0,3850	-0,1966
		SVRA	0,4701	0,0277	0,1664	0,1750	0,5447	-0,4832	-1,7363
		LR	0,7268	0,0002	0,0139	0,6978	0,7499	-0,2695	-0,3625
	2	YSA	0,8337	0,0041	0,0639	0,7433	0,8982	0,0161	-2,0675
		RO	0,7176	0,0020	0,0443	0,6351	0,7682	-1,2141	0,0727
		DVM	0,8465	0,0047	0,0683	0,7407	0,9052	-0,3244	-1,9589
		SVRA	0,4484	0,0003	0,0170	0,4217	0,4588	-1,0437	-0,9976
		LR	0,8211	0,0014	0,0374	0,7650	0,8530	-0,3853	-1,8361
Belirleyicilik	6	YSA	0,9002	0,0000	0,0032	0,8925	0,9038	-0,8527	0,1540
		RO	0,8851	0,0000	0,0041	0,8761	0,8948	0,0020	0,6054
		DVM	0,8976	0,0000	0,0038	0,8889	0,9036	-0,5893	-0,1354
		SVRA	0,8701	0,0000	0,0042	0,8660	0,8761	0,4963	-1,3460
		LR	0,8976	0,0000	0,0036	0,8865	0,9022	-1,3584	2,5577
	3	YSA	0,8902	0,0003	0,0167	0,8639	0,9068	-0,0838	-1,9430
		RO	0,8333	0,0006	0,0254	0,7928	0,8643	-0,5986	-1,2056
		DVM	0,8901	0,0001	0,0078	0,8549	0,8809	-0,5023	-0,4751
		SVRA	0,7757	0,0003	0,0175	0,7486	0,7896	-0,9491	-1,0492
		LR	0,8698	0,0000	0,0071	0,8563	0,8818	-0,2570	-0,3424
	2	YSA	0,8984	0,0020	0,0448	0,8386	0,9427	-0,1018	-1,9328
		RO	0,7912	0,0016	0,0395	0,7176	0,8386	-1,2009	0,0955
		DVM	0,9004	0,0024	0,0492	0,8323	0,9427	-0,3146	-1,9922
		SVRA	0,5802	0,0052	0,0720	0,4671	0,6242	-1,0437	-0,9976
		LR	0,8783	0,0007	0,0267	0,8323	0,9002	-0,5656	-1,5101
Negatif öngörü değeri	6	YSA	0,9001	0,0000	0,0041	0,8922	0,9047	-0,5554	-0,8779
		RO	0,8874	0,0000	0,0034	0,8808	0,8926	-0,4422	-0,8098
		DVM	0,9006	0,0000	0,0034	0,8912	0,9067	-0,8904	1,3764
		SVRA	0,8818	0,0000	0,0056	0,8750	0,8880	-0,0426	-1,7210
		LR	0,9009	0,0000	0,0034	0,8895	0,9052	-1,8444	4,3632
	3	YSA	0,8845	0,0001	0,0092	0,8640	0,9028	-0,1240	-0,3198
		RO	0,8504	0,0011	0,0330	0,7992	0,8871	-0,5799	-1,4113
		DVM	0,8893	0,0001	0,0082	0,8726	0,9076	0,0355	0,1453
		SVRA	0,7758	0,0004	0,0194	0,7477	0,7955	-0,5340	-1,2672
		LR	0,8886	0,0001	0,0073	0,8759	0,9044	0,0483	-0,3938
	2	YSA	0,8670	0,0021	0,0463	0,7921	0,9136	-0,0127	-1,9881
		RO	0,8054	0,0009	0,0296	0,7578	0,8667	-0,3619	-0,4510
		DVM	0,8821	0,0024	0,0486	0,7952	0,9231	-0,4407	-1,6916
		SVRA	0,6308	0,0013	0,0361	0,6087	0,6875	1,0437	-0,9976
		LR	0,8676	0,0008	0,0276	0,8223	0,8908	-0,4400	-1,7524

EK 4. Bağımsız değişken sayısının 10, 7 ve 4 olması koşuluna göre değerlendirme kriterlerine ait betimsel istatistikler

Bağımsız değişken sayısının 10, 7 ve 4 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen doğruluk, duyarlılık ve kesinlik ve değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Doğruluk	10	YSA	0,9311	0,0007	0,0265	0,8926	0,9579	0,0165	-2,0623
		RO	0,8850	0,0003	0,0175	0,8529	0,9072	-1,1330	-0,0370
		DVM	0,9367	0,0008	0,0282	0,8919	0,9609	-0,3352	-1,9323
		SVRA	0,7791	0,0000	0,0042	0,7766	0,7857	1,0437	-0,9976
		LR	0,9264	0,0002	0,0154	0,9029	0,9396	-0,3831	-1,8390
	7	YSA	0,9230	0,0004	0,0208	0,9017	0,9499	0,5765	-1,7516
		RO	0,8842	0,0001	0,0084	0,8748	0,8974	0,1576	-1,5090
		DVM	0,9305	0,0005	0,0230	0,9005	0,9554	0,1363	-1,9733
		SVRA	0,7719	0,0002	0,0123	0,7552	0,7857	-0,5039	-1,4151
		LR	0,9252	0,0007	0,0264	0,8987	0,9591	0,5247	-1,7519
	4	YSA	0,9200	0,0004	0,0188	0,8938	0,9414	0,1259	-1,7828
		RO	0,8826	0,0000	0,0045	0,8773	0,8907	-0,0719	-1,5549
		DVM	0,9260	0,0006	0,0254	0,9005	0,9615	0,6769	-1,4653
		SVRA	0,7674	0,0004	0,0196	0,7399	0,7857	-0,6923	-1,4531
		LR	0,9235	0,0007	0,0273	0,8907	0,9618	0,6918	-1,4588
Duyarlılık	10	YSA	0,9443	0,0004	0,0198	0,9100	0,9640	-0,0654	-1,8888
		RO	0,9228	0,0001	0,0116	0,9075	0,9503	0,2100	-0,5001
		DVM	0,9514	0,0004	0,0203	0,9126	0,9683	-0,5167	-1,5067
		SVRA	0,8594	0,0012	0,0349	0,8380	0,9143	1,0437	-0,9976
		LR	0,9458	0,0001	0,0116	0,9263	0,9554	-0,4979	-1,6586
	7	YSA	0,9399	0,0004	0,0192	0,9169	0,9640	0,3935	-1,7320
		RO	0,9188	0,0001	0,0072	0,9135	0,9332	0,8001	-1,1285
		DVM	0,9488	0,0003	0,0181	0,9246	0,9683	0,1114	-1,9931
		SVRA	0,8509	0,0014	0,0372	0,8209	0,9143	1,1423	-0,4966
		LR	0,9453	0,0005	0,0215	0,9203	0,9726	0,4579	-1,7524
	4	YSA	0,9394	0,0004	0,0209	0,9040	0,9632	0,1034	-1,6376
		RO	0,9174	0,0001	0,0083	0,9032	0,9306	0,0052	-1,5992
		DVM	0,9444	0,0004	0,0201	0,9195	0,9726	0,6904	-1,4357
		SVRA	0,8512	0,0017	0,0417	0,8123	0,9143	0,7948	-1,1142
		LR	0,9451	0,0005	0,0231	0,9126	0,9773	0,6506	-1,4137
Kesinlik	10	YSA	0,9584	0,0003	0,0183	0,9347	0,9766	-0,0779	-1,9660
		RO	0,9164	0,0002	0,0150	0,8884	0,9344	-1,2340	0,1511
		DVM	0,9595	0,0004	0,0200	0,9325	0,9767	-0,3063	-2,0112
		SVRA	0,8363	0,0003	0,0170	0,8096	0,8468	-1,0437	-0,9976
		LR	0,9506	0,0001	0,0108	0,9325	0,9596	-0,5284	-1,5845
	7	YSA	0,9514	0,0001	0,0115	0,9379	0,9657	0,3349	-1,7092
		RO	0,9188	0,0000	0,0065	0,9111	0,9321	0,0397	-1,0586
		DVM	0,9534	0,0002	0,0149	0,9318	0,9691	-0,0210	-1,8287
		SVRA	0,8334	0,0002	0,0150	0,8096	0,8468	-0,7538	-0,9619
		LR	0,9494	0,0003	0,0165	0,9314	0,9701	0,4190	-1,7768
	4	YSA	0,9478	0,0001	0,0084	0,9261	0,9550	-0,9348	-0,0013
		RO	0,9177	0,0000	0,0036	0,9111	0,9255	0,4794	0,2386
		DVM	0,9514	0,0003	0,0166	0,9301	0,9734	0,3836	-1,5251
		SVRA	0,8280	0,0003	0,0162	0,8096	0,8468	0,1966	-1,8175
		LR	0,9471	0,0003	0,0161	0,9307	0,9689	0,5427	-1,5972

Bağımsız değişken sayısının 10, 7 ve 4 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen F ölçütü, kapa, belirleyicilik ve negatif öngörü değerlerine ait betimsel istatistikler

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Baskılık
F ölçütü	10	YSA	0,9513	0,0004	0,0188	0,9235	0,9702	0,0156	-2,0566
		RO	0,9196	0,0001	0,0120	0,8978	0,9359	-1,0651	-0,1240
		DVM	0,9554	0,0004	0,0199	0,9233	0,9725	-0,3435	-1,9117
		SVRA	0,8470	0,0001	0,0075	0,8424	0,8588	1,0437	-0,9976
		LR	0,9482	0,0001	0,0109	0,9315	0,9575	-0,3830	-1,8386
	7	YSA	0,9456	0,0002	0,0149	0,9305	0,9648	0,5707	-1,7511
		RO	0,9188	0,0000	0,0059	0,9123	0,9284	0,2203	-1,5071
		DVM	0,9510	0,0003	0,0163	0,9298	0,9687	0,1381	-1,9740
		SVRA	0,8414	0,0001	0,0121	0,8269	0,8588	0,1935	-1,1504
		LR	0,9473	0,0003	0,0187	0,9283	0,9713	0,5235	-1,7505
	4	YSA	0,9435	0,0002	0,0137	0,9238	0,9590	0,1281	-1,7667
		RO	0,9176	0,0000	0,0035	0,9135	0,9235	-0,1329	-1,7013
		DVM	0,9478	0,0003	0,0180	0,9294	0,9730	0,6845	-1,4611
		SVRA	0,8387	0,0003	0,0169	0,8165	0,8588	-0,2801	-1,4197
		LR	0,9461	0,0004	0,0193	0,9225	0,9731	0,6894	-1,4525
Kappa	10	YSA	0,8337	0,0041	0,0639	0,7433	0,8982	0,0161	-2,0675
		RO	0,7176	0,0020	0,0443	0,6351	0,7682	-1,2141	0,0727
		DVM	0,8465	0,0047	0,0683	0,7407	0,9052	-0,3244	-1,9589
		SVRA	0,4484	0,0003	0,0170	0,4217	0,4588	-1,0437	-0,9976
		LR	0,8211	0,0014	0,0374	0,7650	0,8530	-0,3853	-1,8361
	7	YSA	0,8140	0,0024	0,0492	0,7630	0,8780	0,5814	-1,7517
		RO	0,7175	0,0004	0,0205	0,6940	0,7487	0,0746	-1,4988
		DVM	0,8310	0,0031	0,0556	0,7591	0,8913	0,1324	-1,9717
		SVRA	0,4339	0,0005	0,0230	0,4089	0,4588	0,1209	-1,9867
		LR	0,8182	0,0040	0,0636	0,7557	0,9000	0,5252	-1,7541
	4	YSA	0,8063	0,0019	0,0439	0,7462	0,8560	0,1135	-1,7954
		RO	0,7135	0,0001	0,0095	0,7026	0,7321	0,1458	-1,1678
		DVM	0,8205	0,0038	0,0614	0,7597	0,9062	0,6634	-1,4715
		SVRA	0,4245	0,0015	0,0382	0,3700	0,4588	-0,3373	-1,5940
		LR	0,8143	0,0044	0,0664	0,7374	0,9073	0,6937	-1,4676
Belirleyicilik	10	YSA	0,8984	0,0020	0,0448	0,8386	0,9427	-0,1018	-1,9328
		RO	0,7912	0,0016	0,0395	0,7176	0,8386	-1,2009	0,0955
		DVM	0,9004	0,0024	0,0492	0,8323	0,9427	-0,3146	-1,9922
		SVRA	0,5802	0,0052	0,0720	0,4671	0,6242	-1,0437	-0,9976
		LR	0,8783	0,0007	0,0267	0,8323	0,9002	-0,5656	-1,5101
	7	YSA	0,8811	0,0008	0,0277	0,8471	0,9151	0,2558	-1,6922
		RO	0,7986	0,0003	0,0166	0,7792	0,8344	0,0985	-0,9824
		DVM	0,8851	0,0013	0,0367	0,8280	0,9236	-0,0595	-1,7770
		SVRA	0,5763	0,5763	0,5763	0,5763	0,5763	0,5763	0,5763
		LR	0,8752	0,0016	0,0403	0,8301	0,9257	0,3900	-1,7728
	4	YSA	0,8718	0,0004	0,0202	0,8153	0,8875	-1,2058	0,8951
		RO	0,7962	0,0001	0,0109	0,7749	0,8195	0,1883	0,0753
		DVM	0,8805	0,0017	0,0409	0,8259	0,9342	0,3310	-1,5205

Negatif öngörü değeri		SVRA	0,5654	0,0042	0,0651	0,4671	0,6242	-0,4791	-1,3576
		LR	0,8702	0,0016	0,0402	0,8280	0,9244	0,5227	-1,6028
	10	YSA	0,8670	0,0021	0,0463	0,7921	0,9136	-0,0127	-1,9881
		RO	0,8054	0,0009	0,0296	0,7578	0,8667	-0,3619	-0,4510
		DVM	0,8821	0,0024	0,0486	0,7952	0,9231	-0,4407	-1,6916
		SVRA	0,6308	0,0013	0,0361	0,6087	0,6875	1,0437	-0,9976
		LR	0,8676	0,0008	0,0276	0,8223	0,8908	-0,4400	-1,7524
	7	YSA	0,8564	0,0018	0,0429	0,8102	0,9112	0,4791	-1,7433
		RO	0,7989	0,0003	0,0162	0,7842	0,8301	0,6329	-1,2885
		DVM	0,8750	0,0019	0,0431	0,8182	0,9216	0,1318	-1,9893
		SVRA	0,6266	0,0020	0,0442	0,5717	0,6875	0,7988	-0,7718
		LR	0,8667	0,0026	0,0507	0,8106	0,9316	0,4959	-1,7512
	4	YSA	0,8547	0,0020	0,0451	0,7850	0,9067	0,1819	-1,7395
		RO	0,7960	0,0002	0,0151	0,7735	0,8184	-0,0338	-1,7454
		DVM	0,8652	0,0023	0,0476	0,8105	0,9322	0,7198	-1,4387
		SVRA	0,6199	0,0031	0,0556	0,5466	0,6875	0,2689	-1,2994
		LR	0,8664	0,0031	0,0555	0,7944	0,9442	0,6953	-1,4306

EK 5. Bağımsız değişken sayısının 9, 6 ve 3 olması koşuluna göre değerlendirme kriterlerine ait betimsel istatistikler

Bağımsız değişken sayısının 9, 6 ve 3 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen doğruluk, duyarlılık ve kesinlik ve değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Doğruluk	9	YSA	0,7778	0,0005	0,0222	0,7411	0,8028	-0,1353	-1,5078
		RO	0,7336	0,0000	0,0061	0,7192	0,7460	-0,5838	1,1339
		DVM	0,7719	0,0006	0,0240	0,7381	0,8028	-0,1030	-1,8233
		SVRA	0,6692	0,0009	0,0297	0,6422	0,7283	0,7980	-0,5709
		LR	0,7697	0,0007	0,0268	0,7387	0,8107	0,5050	-1,4295
	6	YSA	0,7344	0,0002	0,0146	0,7149	0,7613	0,7399	-0,3658
		RO	0,6882	0,0010	0,0321	0,6386	0,7167	-0,9290	-1,0632
		DVM	0,7278	0,0002	0,0149	0,6996	0,7430	-1,3300	0,1993
		SVRA	0,6237	0,0003	0,0187	0,6020	0,6422	-0,2605	-1,9073
		LR	0,7228	0,0002	0,0156	0,6996	0,7424	-0,6741	-1,1854
	3	YSA	0,7027	0,0004	0,0201	0,6807	0,7283	-0,1777	-2,0179
		RO	0,6544	0,0003	0,0181	0,6331	0,6862	0,0760	-1,1358
		DVM	0,6966	0,0007	0,0264	0,6386	0,7167	-1,7850	1,6476
		SVRA	0,5906	0,0011	0,0331	0,5504	0,6422	-0,1597	-1,5003
		LR	0,6928	0,0009	0,0298	0,6300	0,7186	-0,9529	0,1406
Duyarlılık	9	YSA	0,8801	0,0002	0,0142	0,8449	0,8895	-1,2221	0,2946
		RO	0,8592	0,0001	0,0113	0,8458	0,8817	0,0627	-1,2923
		DVM	0,8697	0,0014	0,0372	0,8216	0,9117	-0,5483	-1,7328
		SVRA	0,7882	0,0051	0,0716	0,7104	0,8689	-0,1698	-2,0636
		LR	0,8739	0,0008	0,0275	0,8216	0,8972	-1,3804	0,2370
	6	YSA	0,8524	0,0015	0,0388	0,7661	0,9186	0,1508	0,1210
		RO	0,8217	0,0016	0,0400	0,7609	0,8629	-0,8207	-1,1241
		DVM	0,8592	0,0020	0,0448	0,7969	0,9177	-0,4294	-1,5067
		SVRA	0,7124	0,0000	0,0051	0,7095	0,7224	1,5739	0,5566
		LR	0,8726	0,0003	0,0177	0,8466	0,8980	-0,5919	-1,1268
	3	YSA	0,8326	0,0012	0,0347	0,7961	0,8792	-0,0252	-1,9733
		RO	0,7983	0,0001	0,0085	0,7943	0,8303	2,8309	8,4772
		DVM	0,8334	0,0011	0,0337	0,7609	0,8629	-1,6319	1,2769
		SVRA	0,7055	0,0001	0,0121	0,6910	0,7224	-0,0559	-1,3874
		LR	0,8302	0,0014	0,0371	0,7926	0,8680	0,0826	-2,1702
Kesinlik	9	YSA	0,8214	0,0003	0,0175	0,7978	0,8430	0,2743	-1,6449
		RO	0,7866	0,0000	0,0062	0,7751	0,7934	-0,3662	-1,0918
		DVM	0,8257	0,0021	0,0456	0,7741	0,8775	0,0622	-1,9418
		SVRA	0,7517	0,0010	0,0308	0,7125	0,8022	-0,4274	-1,5661
		LR	0,8189	0,0016	0,0396	0,7787	0,8775	0,4980	-1,6113
	6	YSA	0,7917	0,0002	0,0136	0,7746	0,8270	1,1733	0,9230
		RO	0,7598	0,0002	0,0135	0,7394	0,7738	-0,7933	-1,1788
		DVM	0,7828	0,0008	0,0279	0,7584	0,8259	0,9298	-1,0578
		SVRA	0,7480	0,0004	0,0201	0,7257	0,7697	-0,0527	-1,8749
		LR	0,7692	0,0001	0,0078	0,7594	0,7819	-0,0257	-1,4055
	3	YSA	0,7691	0,0000	0,0065	0,7616	0,7846	1,1767	0,0302
		RO	0,7383	0,0002	0,0152	0,7194	0,7650	-0,2357	-1,4131
		DVM	0,7625	0,0001	0,0111	0,7394	0,7738	-1,4182	0,8555
		SVRA	0,7144	0,0010	0,0312	0,6764	0,7697	-0,1466	-1,4064
		LR	0,7606	0,0003	0,0162	0,7176	0,7674	-2,4706	4,5018

Bağımsız değişken sayısının 9, 6 ve 3 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen F ölçütü, kapa, belirleyicilik ve negative öngörü değerlerine ait betimsel istatistikler

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
F ölçütü	9	YSA	0,8497	0,0002	0,0148	0,8230	0,8656	-0,2972	-1,3868
		RO	0,8213	0,0000	0,0046	0,8124	0,8305	0,1260	-0,2035
		DVM	0,8453	0,0001	0,0113	0,8293	0,8656	0,5736	-0,4848
		SVRA	0,7668	0,0008	0,0283	0,7389	0,8201	0,4271	-1,0957
		LR	0,8444	0,0002	0,0146	0,8264	0,8700	0,8836	-0,5211
	6	YSA	0,8202	0,0002	0,0139	0,7954	0,8458	0,6320	-0,0849
		RO	0,7893	0,0006	0,0255	0,7500	0,8113	-0,9250	-1,0709
		DVM	0,8178	0,0001	0,0117	0,8006	0,8338	-0,2359	-1,4299
		SVRA	0,7297	0,0001	0,0102	0,7175	0,7389	-0,3906	-1,9534
		LR	0,8176	0,0001	0,0113	0,8006	0,8324	-0,7357	-1,1161
	3	YSA	0,7993	0,0003	0,0173	0,7803	0,8199	-0,1861	-2,0529
		RO	0,7670	0,0001	0,0106	0,7554	0,7904	0,4720	-0,5032
		DVM	0,7962	0,0004	0,0210	0,7500	0,8113	-1,7897	1,6394
		SVRA	0,7098	0,0005	0,0213	0,6836	0,7389	-0,2129	-1,5614
		LR	0,7936	0,0005	0,0223	0,7533	0,8146	-0,4358	-1,1297
Kappa	9	YSA	0,4255	0,0035	0,0595	0,3424	0,4961	0,0870	-1,6047
		RO	0,3020	0,0003	0,0184	0,2593	0,3292	-1,0051	0,2876
		DVM	0,4091	0,0100	0,0999	0,2791	0,5038	-0,1724	-2,0599
		SVRA	0,1999	0,0016	0,0403	0,1739	0,3104	1,6580	1,4451
		LR	0,4011	0,0092	0,0958	0,2979	0,5230	0,3219	-1,8859
	6	YSA	0,3120	0,0009	0,0295	0,2556	0,3753	-0,1026	-0,2882
		RO	0,1920	0,0037	0,0611	0,0985	0,2518	-0,8659	-1,1026
		DVM	0,2787	0,0042	0,0648	0,1990	0,3711	0,4374	-1,1910
		SVRA	0,1114	0,0035	0,0594	0,0442	0,1739	-0,1232	-1,8809
		LR	0,2490	0,0013	0,0361	0,1990	0,3008	-0,3734	-1,3909
	3	YSA	0,2277	0,0010	0,0319	0,1968	0,2810	0,4029	-1,4322
		RO	0,1016	0,0035	0,0588	0,0289	0,2008	-0,2371	-1,4504
		DVM	0,2060	0,0025	0,0501	0,0985	0,2518	-1,6125	1,2975
		SVRA	0,0552	0,0081	0,0900	-0,0931	0,1739	-0,0852	-1,3842
		LR	0,1962	0,0048	0,0689	0,0215	0,2394	-2,0809	3,3202
Belirleyicilik	9	YSA	0,5231	0,0025	0,0499	0,4544	0,5860	0,3918	-1,6505
		RO	0,4223	0,0007	0,0258	0,3822	0,4544	0,0897	-1,4629
		DVM	0,5286	0,0246	0,1567	0,3503	0,7072	0,0641	-1,9297
		SVRA	0,4065	0,0057	0,0753	0,3164	0,5117	-0,2233	-1,8500
		LR	0,5109	0,0173	0,1317	0,3694	0,7072	0,4876	-1,5644
	6	YSA	0,4421	0,0041	0,0637	0,3715	0,6030	0,9031	0,2819
		RO	0,3575	0,0005	0,0220	0,3079	0,3907	-0,3267	-0,8723
		DVM	0,4021	0,0136	0,1168	0,2824	0,5839	0,9687	-1,0178
		SVRA	0,3514	0,0040	0,0630	0,3355	0,4735	0,0329	-1,8740
		LR	0,4038	0,0004	0,0211	0,3227	0,3949	0,6001	-0,6567
	3	YSA	0,3806	0,0009	0,0294	0,3227	0,4310	-0,6434	-0,3135
		RO	0,2979	0,0028	0,0528	0,2314	0,3864	-0,3482	-1,4983
		DVM	0,3576	0,0004	0,0208	0,3079	0,3907	-0,3020	-0,4426

Negatíföngörú dögéri		SVRA	0,3116	0,0063	0,0796	0,2173	0,4735	0,0388	-1,1177
		LR	0,3524	0,0027	0,0522	0,2272	0,3949	-1,6316	2,1445
	9	YSA	0,6359	0,0018	0,0423	0,5575	0,6798	-0,4178	-1,3260
		RO	0,5479	0,0002	0,0149	0,5156	0,5783	-0,1408	0,4793
		DVM	0,6158	0,0010	0,0316	0,5719	0,6798	1,2222	0,6371
		SVRA	0,4665	0,0040	0,0629	0,3975	0,5392	-0,2192	-2,0764
		LR	0,6157	0,0017	0,0410	0,5634	0,6912	1,1257	0,0053
		YSA	0,5536	0,0021	0,0461	0,5054	0,6481	1,3002	0,6314
	6	RO	0,4532	0,0036	0,0598	0,3616	0,5100	-0,8766	-1,0895
		DVM	0,5382	0,0015	0,0393	0,4688	0,5919	-0,8290	-0,2561
		SVRA	0,3598	0,0013	0,0366	0,3179	0,3975	-0,1855	-1,8904
		LR	0,5285	0,0016	0,0405	0,4688	0,5854	-0,6314	-1,1582
		YSA	0,4832	0,0017	0,0410	0,4387	0,5389	-0,1511	-1,9839
	3	RO	0,3712	0,0022	0,0471	0,3132	0,4461	-0,2065	-1,4484
		DVM	0,4684	0,0024	0,0493	0,3616	0,5100	-1,6921	1,4611
		SVRA	0,2992	0,0034	0,0583	0,2290	0,3975	-0,1002	-1,4276
		LR	0,4598	0,0047	0,0684	0,3066	0,5157	-1,2416	0,9101

EK 6. Örneklem büyüklüğü koşuluna göre değerlendirme kriterlerine ait betimsel istatistikler

Örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen doğruluk değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Doğruluk	100	YSA	0,8499	0,0017	0,0407	0,7813	0,9375	0,3011	-0,3799
		RO	0,7862	0,0010	0,0312	0,7576	0,8438	1,1385	-0,0222
		DVM	0,8575	0,0046	0,0682	0,7188	0,9688	-0,2262	-0,8775
		SVRA	0,7140	0,0083	0,0913	0,5061	0,8750	0,3989	0,0249
		LR	0,8615	0,0057	0,0754	0,7500	0,9688	-0,0602	-0,9463
	250	YSA	0,8500	0,0002	0,0153	0,8375	0,8875	1,1833	0,2960
		RO	0,7745	0,0035	0,0595	0,7250	0,8625	0,4967	-1,6348
		DVM	0,8535	0,0004	0,0189	0,8125	0,9000	0,5766	1,2809
		SVRA	0,7190	0,0090	0,0951	0,6125	0,8875	0,5582	-0,9936
		LR	0,8603	0,0006	0,0239	0,8250	0,9063	0,2621	-0,6649
	500	YSA	0,9212	0,0003	0,0164	0,8882	0,9441	0,1387	-0,9005
		RO	0,8043	0,0014	0,0379	0,7640	0,8944	0,7847	-0,0690
		DVM	0,9255	0,0001	0,0072	0,9130	0,9317	-0,8825	-0,6451
		SVRA	0,7058	0,0039	0,0625	0,6584	0,8758	1,2330	1,1426
		LR	0,9155	0,0001	0,0120	0,8944	0,9317	-0,3967	-1,4483
	1000	YSA	0,9156	0,0001	0,0090	0,9012	0,9228	-0,5983	-1,5340
		RO	0,8678	0,0013	0,0358	0,7747	0,8981	-2,2827	2,0006
		DVM	0,9153	0,0000	0,0046	0,9074	0,9228	0,8479	-0,3663
		SVRA	0,7501	0,0006	0,0250	0,7284	0,7778	0,2575	-2,1097
		LR	0,9153	0,0001	0,0120	0,9006	0,9317	-0,2574	-1,7238
	2500	YSA	0,9217	0,0002	0,0148	0,9008	0,9335	-0,4568	-1,9195
		RO	0,8730	0,0021	0,0457	0,7526	0,9067	-2,4240	1,3743
		DVM	0,9296	0,0006	0,0250	0,8938	0,9545	-0,0695	-1,9599
		SVRA	0,7561	0,0014	0,0377	0,7398	0,8821	3,1416	1,9741
		LR	0,9188	0,0003	0,0159	0,8961	0,9358	0,0451	-1,8965
	5000	YSA	0,9311	0,0007	0,0265	0,8926	0,9579	0,0165	-2,0623
		RO	0,8850	0,0003	0,0175	0,8529	0,9072	-1,1330	-0,0370
		DVM	0,9367	0,0008	0,0282	0,8919	0,9609	-0,3352	-1,9323
		SVRA	0,7791	0,0000	0,0042	0,7766	0,7857	1,0437	-0,9976
		LR	0,9264	0,0002	0,0154	0,9029	0,9396	-0,3831	-1,8390

Örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen duyarlılık değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Duyarlılık	100	YSA	0,8724	0,0023	0,0482	0,8000	0,9524	-0,3097	-0,9739
		RO	0,8076	0,0054	0,0734	0,7143	0,9048	-0,1507	-1,4147
		DVM	0,8800	0,0049	0,0703	0,7619	1,0000	-0,2537	-0,4891
		SVRA	0,7324	0,0243	0,1560	0,0714	0,9048	-3,1862	1,4132
		LR	0,8678	0,0066	0,0811	0,7143	0,9524	-0,2248	-1,5001
	250	YSA	0,9007	0,0019	0,0431	0,8333	0,9630	0,2605	-1,1449
		RO	0,8470	0,0029	0,0534	0,8036	0,9444	0,6425	-1,2351
		DVM	0,9096	0,0003	0,0180	0,8704	0,9630	0,9296	3,0817
		SVRA	0,7600	0,0166	0,1290	0,5741	0,9630	-0,4396	-1,0167
		LR	0,8999	0,0029	0,0542	0,8148	0,9630	-0,4501	-1,0875
	500	YSA	0,9431	0,0003	0,0162	0,9099	0,9730	-0,0847	-0,4215
		RO	0,8627	0,0019	0,0435	0,8198	0,9459	0,8891	-0,4540
		DVM	0,9535	0,0001	0,0096	0,9279	0,9820	-0,1029	1,4107
		SVRA	0,7305	0,0039	0,0622	0,6847	0,9009	1,4210	1,8055
		LR	0,9384	0,0004	0,0196	0,9099	0,9730	-0,2758	-1,0145
	1000	YSA	0,9396	0,0001	0,0079	0,9182	0,9591	-0,5553	2,9163
		RO	0,9071	0,0009	0,0306	0,8500	0,9591	-0,2450	-0,4020
		DVM	0,9447	0,0000	0,0052	0,9273	0,9545	-1,1172	1,9735
		SVRA	0,7858	0,0000	0,0046	0,7818	0,7909	0,2575	-2,1097
		LR	0,9351	0,0004	0,0196	0,9099	0,9640	-0,2467	-1,4923
	2500	YSA	0,9453	0,0001	0,0120	0,9190	0,9542	-0,8449	-0,8639
		RO	0,9154	0,0015	0,0381	0,8187	0,9525	-2,0841	2,3875
		DVM	0,9545	0,0003	0,0178	0,9278	0,9718	-0,2010	-1,7671
		SVRA	0,8392	0,0006	0,0253	0,8187	0,9384	2,5740	2,6259
		LR	0,9477	0,0001	0,0120	0,9243	0,9595	-0,3363	-1,4357
	5000	YSA	0,9443	0,0004	0,0198	0,9100	0,9640	-0,0654	-1,8888
		RO	0,9228	0,0001	0,0116	0,9075	0,9503	0,2100	-0,5001
		DVM	0,9514	0,0004	0,0203	0,9126	0,9683	-0,5167	-1,5067
		SVRA	0,8594	0,0012	0,0349	0,8380	0,9143	1,0437	-0,9976
		LR	0,9458	0,0001	0,0116	0,9263	0,9554	-0,4979	-1,6586

Örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen kesinlik değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Kesinlik	100	YSA	0,8178	0,0174	0,1319	0,5714	0,9524	-1,2228	0,1217
		RO	0,8587	0,0015	0,0387	0,7917	0,9000	-0,8194	-0,6084
		DVM	0,9023	0,0034	0,0581	0,8000	1,0000	0,3022	-0,4457
		SVRA	0,8003	0,0048	0,0691	0,6563	0,9444	0,5815	-0,0114
		LR	0,8857	0,0128	0,1130	0,5714	1,0000	-1,6338	2,3977
	250	YSA	0,8821	0,0009	0,0307	0,8387	0,9231	-0,2709	-1,3864
		RO	0,8304	0,0014	0,0370	0,7969	0,9057	0,9076	-0,6703
		DVM	0,8783	0,0004	0,0209	0,8421	0,9259	0,7313	0,8085
		SVRA	0,8124	0,0037	0,0610	0,7679	0,9762	1,3938	0,9251
		LR	0,8966	0,0013	0,0365	0,8448	0,9500	-0,0740	-1,5480
	500	YSA	0,9431	0,0003	0,0183	0,9043	0,9720	-0,3891	-0,5577
		RO	0,8571	0,0003	0,0177	0,8349	0,9052	0,3129	0,6770
		DVM	0,9394	0,0001	0,0099	0,9008	0,9464	-2,6169	2,1645
		SVRA	0,8218	0,0015	0,0384	0,7917	0,9333	1,1052	1,0595
		LR	0,9393	0,0001	0,0118	0,9060	0,9537	-1,9887	2,4919
	1000	YSA	0,9365	0,0002	0,0138	0,9056	0,9486	-1,1064	-0,4268
		RO	0,8998	0,0010	0,0318	0,8238	0,9206	-1,7591	2,0155
		DVM	0,9316	0,0001	0,0085	0,9170	0,9452	0,6425	-0,5816
		SVRA	0,8371	0,0009	0,0297	0,8113	0,8700	0,2575	-2,1097
		LR	0,9417	0,0001	0,0087	0,9060	0,9537	-3,1011	1,2194
	2500	YSA	0,9372	0,0001	0,0115	0,9158	0,9459	-0,7626	-1,1648
		RO	0,8953	0,0011	0,0333	0,8101	0,9192	-2,1687	3,5232
		DVM	0,9403	0,0004	0,0203	0,9119	0,9600	-0,1777	-1,8461
		SVRA	0,8032	0,0015	0,0383	0,7805	0,9499	2,8432	3,2222
		LR	0,9312	0,0002	0,0137	0,9119	0,9445	-0,2104	-1,8623
	5000	YSA	0,9584	0,0003	0,0183	0,9347	0,9766	-0,0779	-1,9660
		RO	0,9164	0,0002	0,0150	0,8884	0,9344	-1,2340	0,1511
		DVM	0,9595	0,0004	0,0200	0,9325	0,9767	-0,3063	-2,0112
		SVRA	0,8363	0,0003	0,0170	0,8096	0,8468	-1,0437	-0,9976
		LR	0,9506	0,0001	0,0108	0,9325	0,9596	-0,5284	-1,5845

Örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen F ölçütü değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
F ölçütü	100	YSA	0,8408	0,0088	0,0938	0,6667	0,9524	-1,1896	0,0551
		RO	0,8293	0,0011	0,0327	0,7895	0,8780	0,1153	-1,1488
		DVM	0,8897	0,0029	0,0541	0,7805	0,9767	-0,2928	-0,9301
		SVRA	0,7557	0,0212	0,1456	0,1288	0,9048	-3,4072	1,5219
		LR	0,8733	0,0076	0,0871	0,6667	0,9756	-0,9266	0,6706
	250	YSA	0,8900	0,0002	0,0125	0,8738	0,9174	0,3773	-0,9000
		RO	0,8383	0,0017	0,0416	0,8036	0,9027	0,4885	-1,6460
		DVM	0,8935	0,0002	0,0134	0,8649	0,9259	0,4296	1,1123
		SVRA	0,7803	0,0075	0,0867	0,6667	0,9159	0,0229	-1,0891
		LR	0,8963	0,0004	0,0196	0,8679	0,9286	-0,0506	-1,1620
	500	YSA	0,9429	0,0001	0,0118	0,9204	0,9593	0,1654	-1,0363
		RO	0,8596	0,0009	0,0292	0,8273	0,9251	0,6125	-0,3660
		DVM	0,9464	0,0000	0,0051	0,9364	0,9507	-0,9764	-0,3968
		SVRA	0,7731	0,0026	0,0509	0,7343	0,9074	1,1937	0,9602
		LR	0,9387	0,0001	0,0091	0,9238	0,9511	-0,4103	-1,4435
	1000	YSA	0,9380	0,0000	0,0062	0,9279	0,9431	-0,5641	-1,5608
		RO	0,9031	0,0007	0,0258	0,8367	0,9265	-2,2051	3,7978
		DVM	0,9381	0,0000	0,0031	0,9315	0,9431	0,5953	-0,0135
		SVRA	0,8105	0,0003	0,0164	0,7963	0,8286	0,2575	-2,1097
		LR	0,9383	0,0001	0,0093	0,9266	0,9511	-0,2663	-1,6810
	2500	YSA	0,9412	0,0001	0,0111	0,9258	0,9500	-0,4579	-1,9179
		RO	0,9052	0,0012	0,0345	0,8144	0,9304	-2,4228	2,3703
		DVM	0,9473	0,0003	0,0186	0,9205	0,9659	-0,0668	-1,9630
		SVRA	0,8204	0,0007	0,0266	0,8115	0,9103	3,2833	3,5896
		LR	0,9393	0,0001	0,0118	0,9224	0,9520	0,0522	-1,8957
	5000	YSA	0,9513	0,0004	0,0188	0,9235	0,9702	0,0156	-2,0566
		RO	0,9196	0,0001	0,0120	0,8978	0,9359	-1,0651	-0,1240
		DVM	0,9554	0,0004	0,0199	0,9233	0,9725	-0,3435	-1,9117
		SVRA	0,8470	0,0001	0,0075	0,8424	0,8588	1,0437	-0,9976
		LR	0,9482	0,0001	0,0109	0,9315	0,9575	-0,3830	-1,8386

Örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen Kappa değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Kappa	100	YSA	0,6425	0,0094	0,0969	0,4815	0,8615	0,8689	-0,0332
		RO	0,5392	0,0041	0,0641	0,4815	0,6610	1,3590	0,2796
		DVM	0,6879	0,0216	0,1469	0,3898	0,9292	-0,1810	-0,8045
		SVRA	0,3870	0,0347	0,1863	0,0316	0,7229	0,6817	-0,4045
		LR	0,6903	0,0265	0,1627	0,4459	0,9322	0,1885	-1,0426
	250	YSA	0,6531	0,0014	0,0373	0,6182	0,7410	1,4242	0,5087
		RO	0,4658	0,0212	0,1456	0,3452	0,6897	0,5160	-1,6140
		DVM	0,6591	0,0021	0,0459	0,5595	0,7721	0,6868	1,3427
		SVRA	0,3860	0,0346	0,1860	0,2317	0,7461	1,0096	-0,8120
		LR	0,6803	0,0029	0,0534	0,5846	0,7966	0,3608	0,0608
	500	YSA	0,8159	0,0015	0,0391	0,7331	0,8702	0,0802	-0,7122
		RO	0,5361	0,0067	0,0820	0,4548	0,7465	1,0381	0,4495
		DVM	0,8243	0,0003	0,0175	0,7851	0,8396	-0,9012	-0,4104
		SVRA	0,3594	0,0163	0,1276	0,2630	0,7192	1,2791	1,3874
		LR	0,8029	0,0007	0,0269	0,7521	0,8378	-0,4326	-1,3411
	1000	YSA	0,8057	0,0005	0,0224	0,7699	0,8234	-0,6598	-1,4529
		RO	0,6948	0,0072	0,0848	0,4738	0,7621	-2,3187	1,0682
		DVM	0,8041	0,0001	0,0115	0,7887	0,8234	0,9707	-0,4758
		SVRA	0,4445	0,0040	0,0636	0,3893	0,5148	0,2575	-2,1097
		LR	0,8032	0,0007	0,0261	0,7600	0,8378	-0,2982	-1,6516
	2500	YSA	0,8242	0,0011	0,0331	0,7764	0,8506	-0,4607	-1,9088
		RO	0,7130	0,0104	0,1022	0,4438	0,7887	-2,4190	4,3583
		DVM	0,8413	0,0032	0,0567	0,7606	0,8976	-0,0735	-1,9554
		SVRA	0,4408	0,0083	0,0911	0,3937	0,7469	2,9050	2,0087
		LR	0,8168	0,0013	0,0362	0,7655	0,8553	0,0349	-1,8974
	5000	YSA	0,8337	0,0041	0,0639	0,7433	0,8982	0,0161	-2,0675
		RO	0,7176	0,0020	0,0443	0,6351	0,7682	-1,2141	0,0727
		DVM	0,8465	0,0047	0,0683	0,7407	0,9052	-0,3244	-1,9589
		SVRA	0,4484	0,0003	0,0170	0,4217	0,4588	-1,0437	-0,9976
		LR	0,8211	0,0014	0,0374	0,7650	0,8530	-0,3853	-1,8361

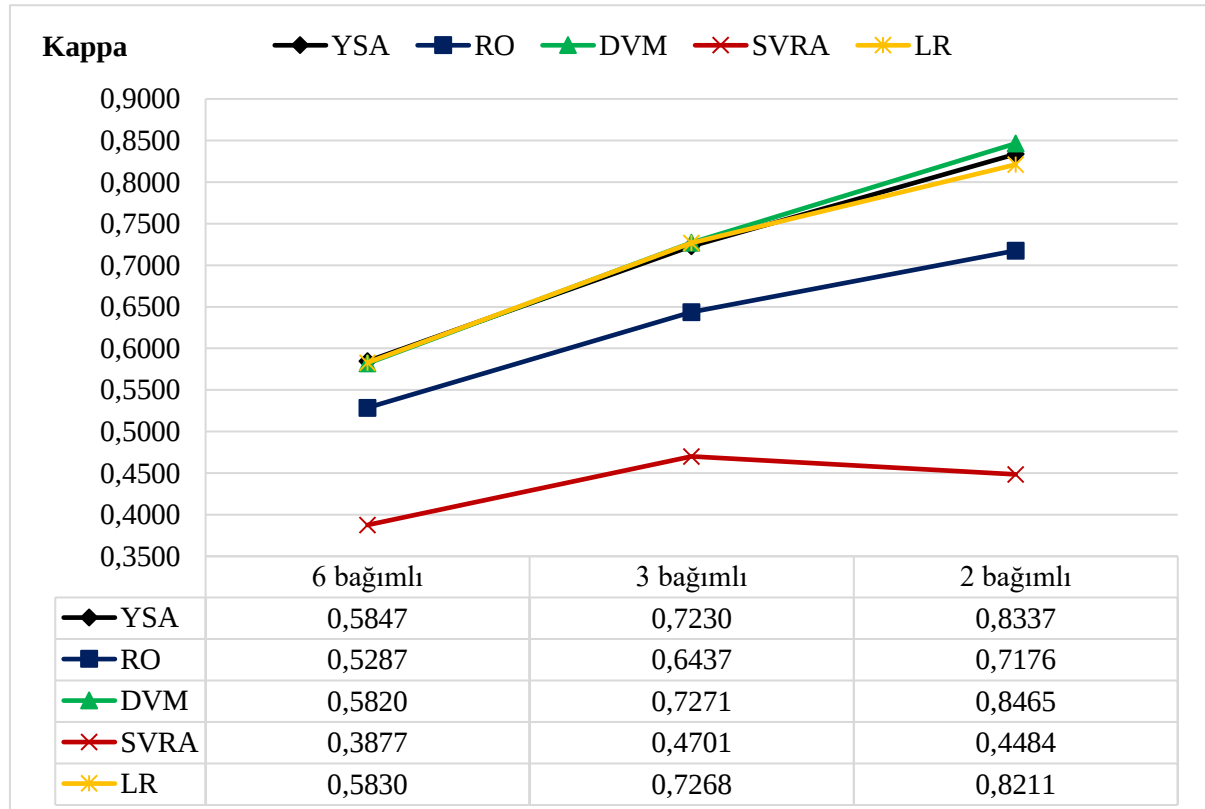
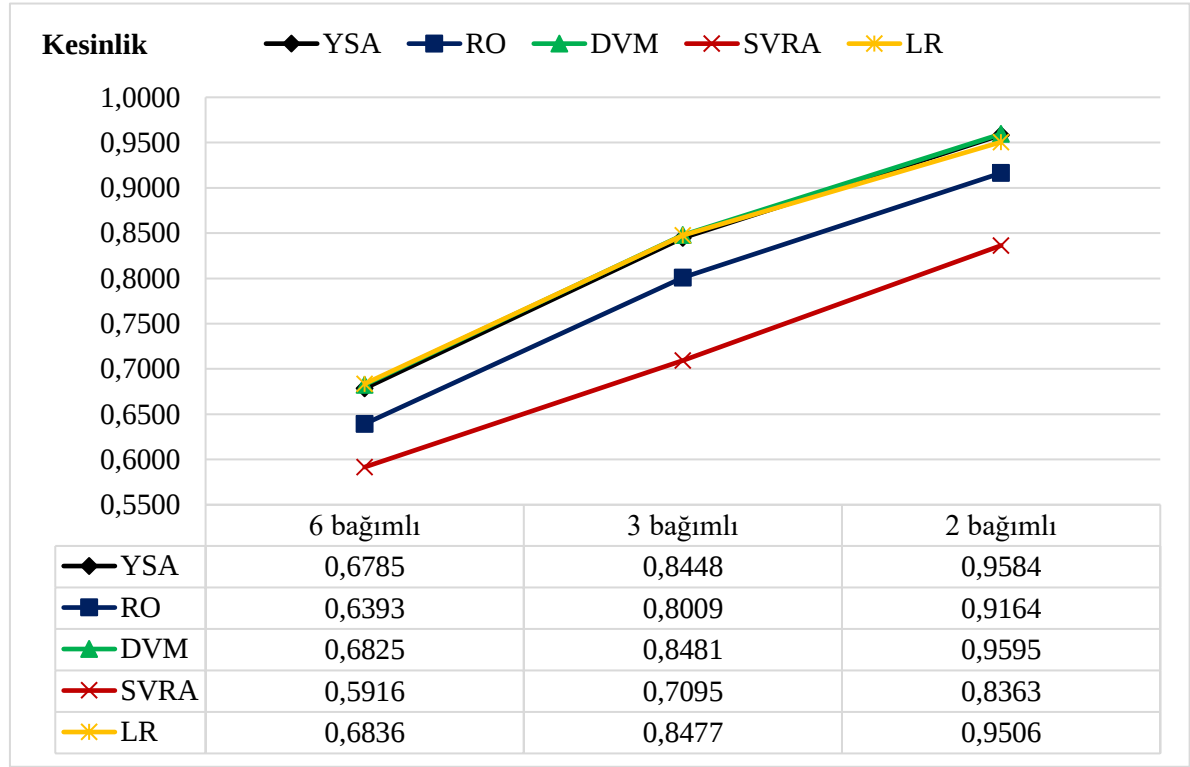
Örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen belirleyicilik değerlerine ait betimsel istatistikler.

	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Belirleyicilik	100	YSA	0,7859	0,0092	0,0961	0,5455	0,9091	-0,1963	-0,2892
		RO	0,7430	0,0121	0,1100	0,5455	0,8333	-1,0405	-0,3766
		DVM	0,8145	0,0127	0,1128	0,6364	1,0000	0,2236	-0,6581
		SVRA	0,6639	0,0167	0,1292	0,5455	0,9609	0,9281	-0,3379
		LR	0,8409	0,0138	0,1177	0,6364	1,0000	0,2283	-1,2957
	250	YSA	0,7446	0,0075	0,0866	0,6154	0,8462	-0,4901	-1,2856
		RO	0,6126	0,0098	0,0988	0,5000	0,8077	0,8108	-0,8660
		DVM	0,7369	0,0027	0,0517	0,6538	0,8462	0,3559	0,5387
		SVRA	0,6338	0,0188	0,1369	0,5000	0,9615	0,5811	-0,4253
		LR	0,7779	0,0091	0,0955	0,6538	0,9091	-0,1310	-1,5480
	500	YSA	0,8728	0,0019	0,0439	0,7800	0,9400	-0,4994	-0,5614
		RO	0,6701	0,0009	0,0300	0,6400	0,7800	1,9581	3,7896
		DVM	0,8632	0,0007	0,0256	0,7600	0,8800	-2,9111	1,9200
		SVRA	0,6304	0,0223	0,1492	0,0000	0,8600	-3,1461	1,4041
		LR	0,8648	0,0009	0,0302	0,7800	0,9000	-2,0299	3,7161
	1000	YSA	0,8646	0,0011	0,0325	0,7885	0,8942	-1,1660	-0,1811
		RO	0,7846	0,0055	0,0742	0,6154	0,8365	-1,5487	1,2865
		DVM	0,8531	0,0004	0,0201	0,8173	0,8846	0,5529	-0,5911
		SVRA	0,6746	0,0047	0,0682	0,6154	0,7500	0,2575	-2,1097
		LR	0,8712	0,0005	0,0224	0,7800	0,9000	-3,0561	1,7624
	2500	YSA	0,8754	0,0005	0,0234	0,8304	0,8927	-0,8512	-0,9386
		RO	0,7896	0,0044	0,0665	0,6228	0,8374	-2,0245	3,0335
		DVM	0,8807	0,0017	0,0411	0,8201	0,9204	-0,2087	-1,8095
		SVRA	0,5929	0,0076	0,0875	0,5329	0,9100	2,2893	3,6725
		LR	0,8621	0,0008	0,0283	0,8201	0,8893	-0,2588	-1,8123
	5000	YSA	0,8984	0,0020	0,0448	0,8386	0,9427	-0,1018	-1,9328
		RO	0,7912	0,0016	0,0395	0,7176	0,8386	-1,2009	0,0955
		DVM	0,9004	0,0024	0,0492	0,8323	0,9427	-0,3146	-1,9922
		SVRA	0,5802	0,0052	0,0720	0,4671	0,6242	-1,0437	-0,9976
		LR	0,8783	0,0007	0,0267	0,8323	0,9002	-0,5656	-1,5101

Örneklem büyüklüğünün 100, 250, 500, 1000, 2500 ve 5000 olması koşulları altında YSA, RO, DVM, SVRA ve LR yöntemlerinin uygulanması sonucu elde edilen Negativeöngörü değerlerine ait betimsel istatistikler.

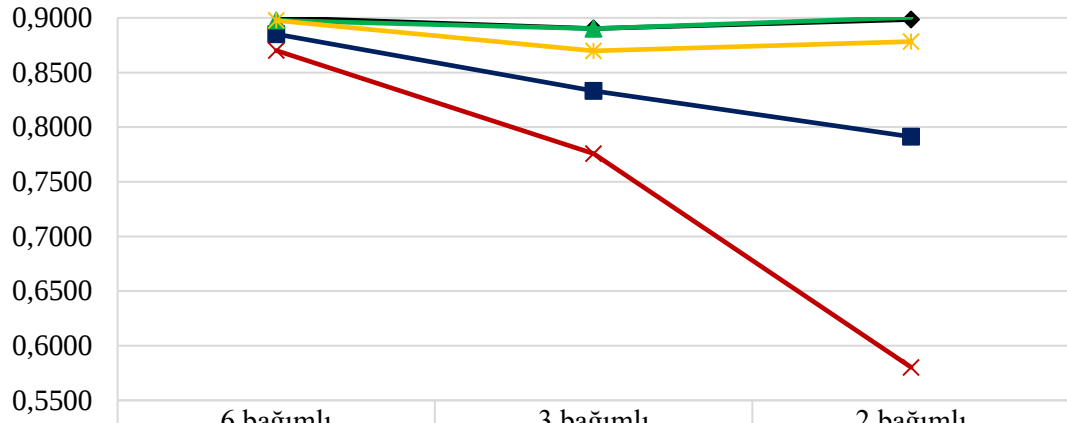
	Kategori sayısı	Yöntem	Ortalama	Varyans	Std. Sapma	Min.	Max.	Çarpıklık	Basıklık
Negative öngörü değeri	100	YSA	0,7977	0,0036	0,0601	0,6667	0,9091	-0,3255	-0,5301
		RO	0,6867	0,0031	0,0553	0,6250	0,7500	0,1631	-1,8230
		DVM	0,7890	0,0126	0,1124	0,5833	1,0000	0,0536	-0,6045
		SVRA	0,5878	0,0133	0,1153	0,4972	0,8182	1,0734	-0,6124
		LR	0,7917	0,0136	0,1167	0,6000	0,9167	-0,4455	-1,4390
	250	YSA	0,7934	0,0041	0,0638	0,7097	0,8889	0,4605	-1,1693
		RO	0,6461	0,0154	0,1241	0,5417	0,8571	0,4456	-1,7115
		DVM	0,7981	0,0012	0,0342	0,7308	0,8947	0,6911	2,0771
		SVRA	0,5817	0,0205	0,1433	0,4390	0,8947	0,9484	-0,3110
		LR	0,8022	0,0057	0,0752	0,6970	0,9091	-0,1003	-1,3596
	500	YSA	0,8745	0,0010	0,0312	0,8148	0,9333	0,1001	-0,6986
		RO	0,6850	0,0070	0,0839	0,6154	0,8667	1,2293	0,0309
		DVM	0,8938	0,0004	0,0190	0,8462	0,9500	-0,0516	3,9934
		SVRA	0,5075	0,0189	0,1375	0,0000	0,7679	-1,6836	1,8129
		LR	0,8653	0,0013	0,0355	0,8148	0,9286	-0,2041	-1,0367
	1000	YSA	0,8716	0,0002	0,0131	0,8364	0,9011	-0,8210	1,9532
		RO	0,8013	0,0040	0,0632	0,6598	0,8875	-1,1190	1,1677
		DVM	0,8796	0,0001	0,0082	0,8491	0,8958	-1,6701	1,4767
		SVRA	0,5968	0,0009	0,0292	0,5714	0,6290	0,2575	-2,1097
		LR	0,8599	0,0012	0,0353	0,8148	0,9149	-0,2056	-1,4226
	2500	YSA	0,8907	0,0005	0,0231	0,8456	0,9085	-0,6757	-1,3554
		RO	0,8274	0,0055	0,0741	0,6360	0,8889	-2,2262	2,7889
		DVM	0,9077	0,0013	0,0362	0,8536	0,9433	-0,1373	-1,8662
		SVRA	0,6509	0,0027	0,0524	0,6360	0,8622	3,6017	2,6474
		LR	0,8935	0,0006	0,0238	0,8562	0,9179	-0,1518	-1,7262
	5000	YSA	0,8670	0,0021	0,0463	0,7921	0,9136	-0,0127	-1,9881
		RO	0,8054	0,0009	0,0296	0,7578	0,8667	-0,3619	-0,4510
		DVM	0,8821	0,0024	0,0486	0,7952	0,9231	-0,4407	-1,6916
		SVRA	0,6308	0,0013	0,0361	0,6087	0,6875	1,0437	-0,9976
		LR	0,8676	0,0008	0,0276	0,8223	0,8908	-0,4400	-1,7524

EK 7. Bağımlı değişkenin kategori sayısının 6, 3 ve 2 olması durumunda yöntemlerin değerlendirme kriterlerinin değerlerine ait grafikler.



Belirleyicilik

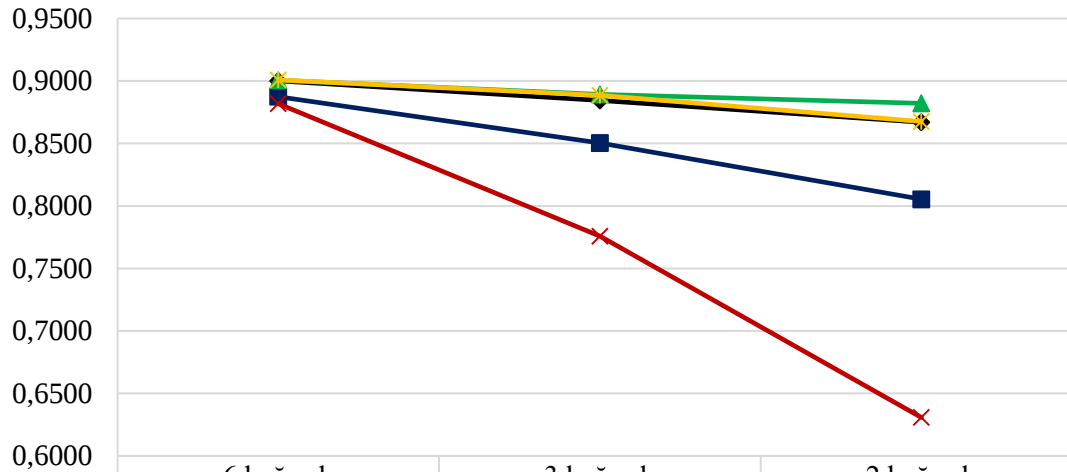
YSA RO DVM SVRA LR



	6 bağımlı	3 bağımlı	2 bağımlı
YSA	0,9002	0,8902	0,8984
RO	0,8851	0,8333	0,7912
DVM	0,8976	0,8901	0,9004
SVRA	0,8701	0,7757	0,5802
LR	0,8976	0,8698	0,8783

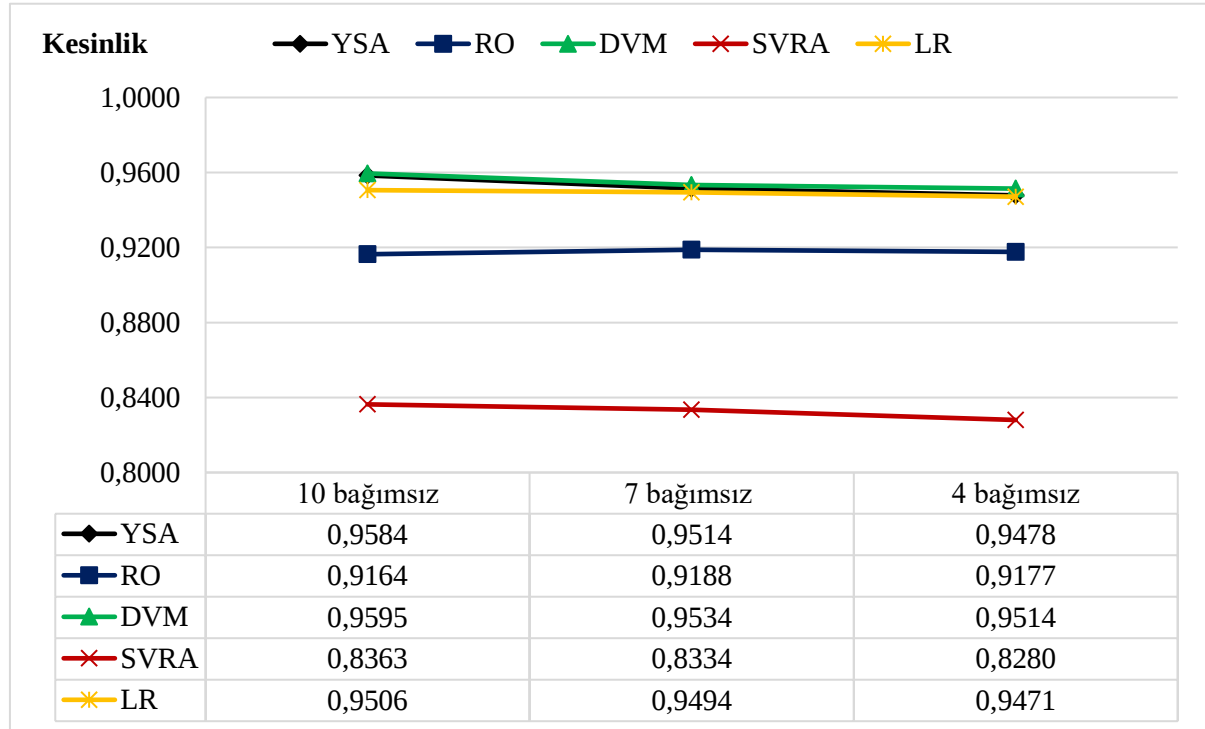
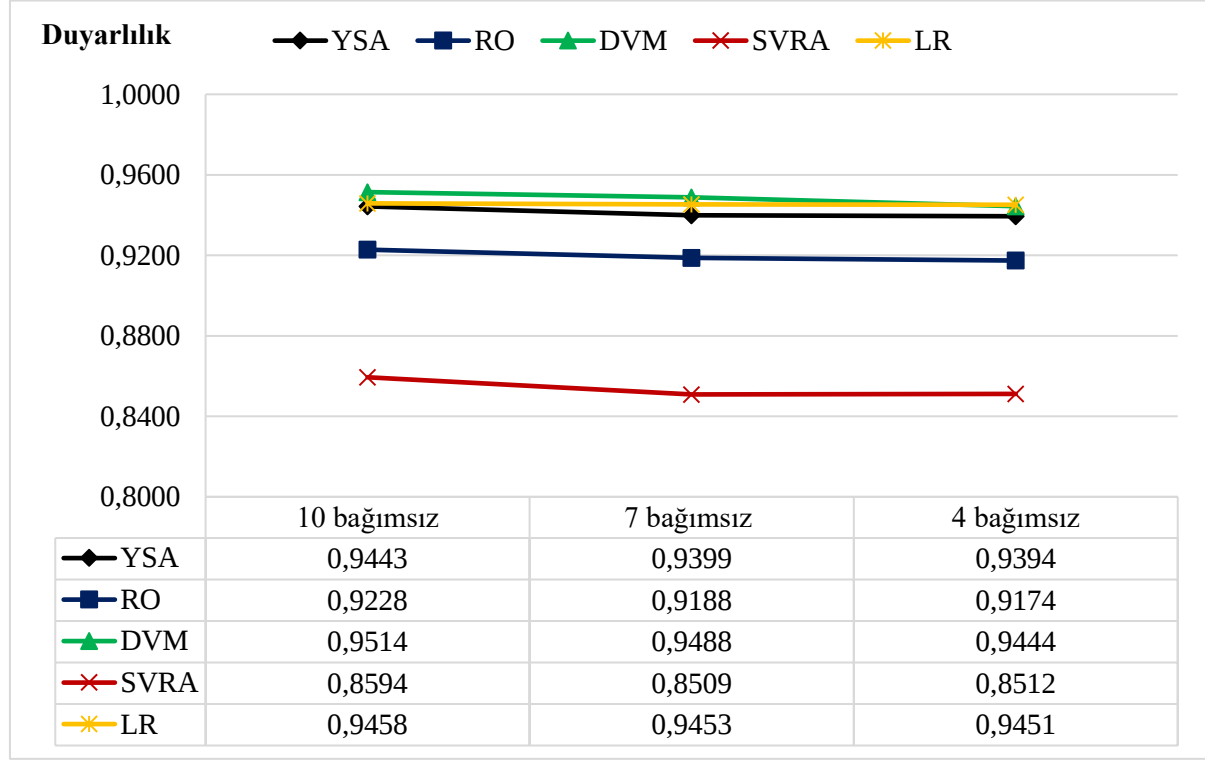
Negatif Öngörü değeri

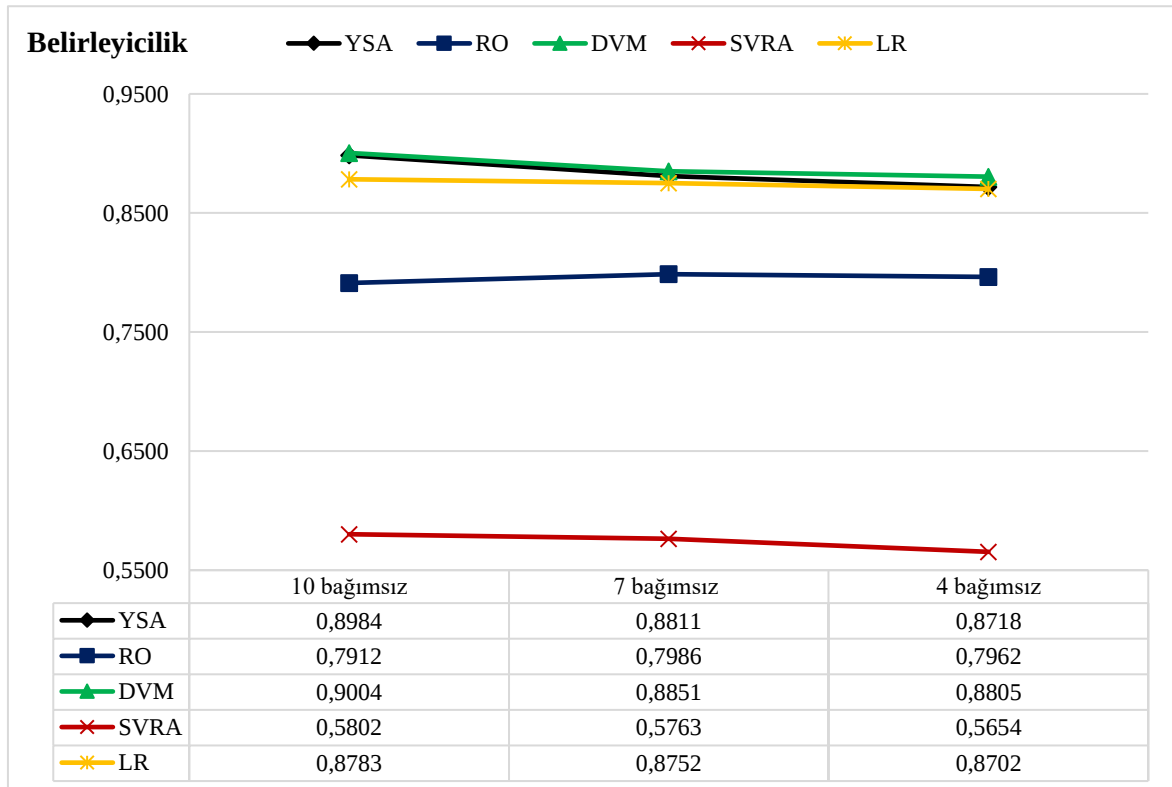
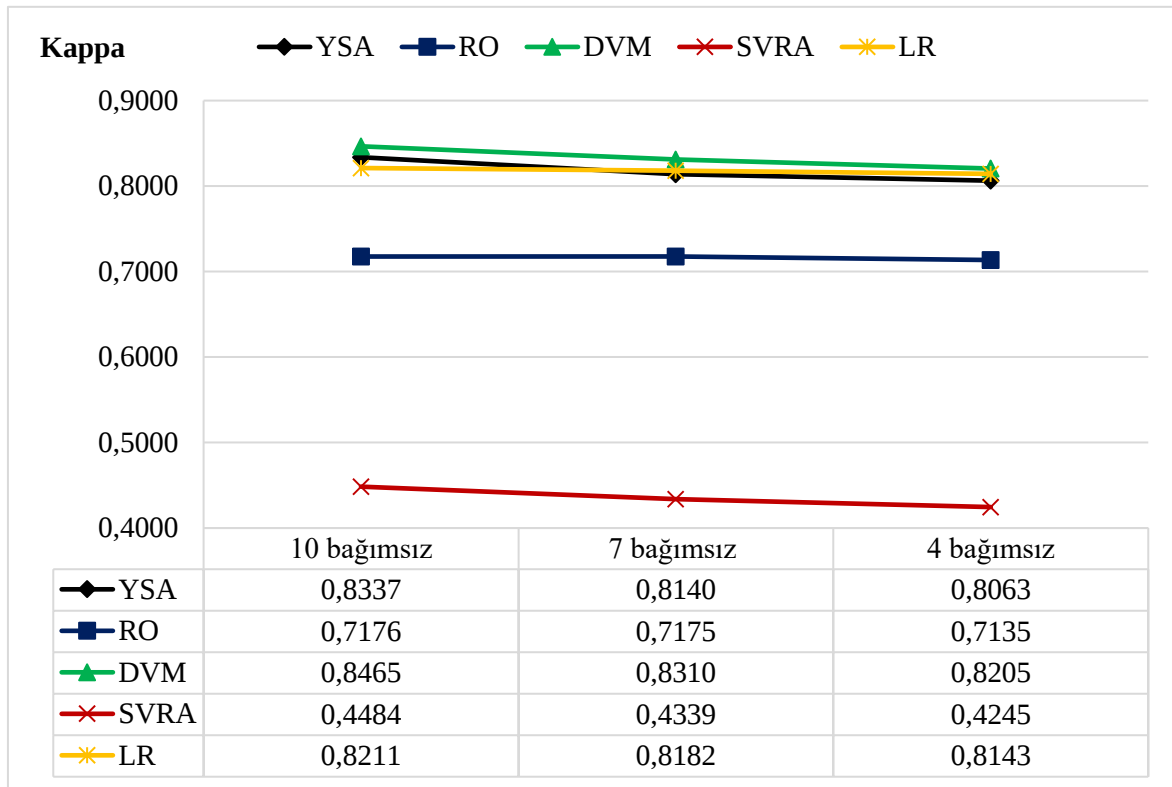
YSA RO DVM SVRA LR



	6 bağımlı	3 bağımlı	2 bağımlı
YSA	0,9001	0,8845	0,8670
RO	0,8874	0,8504	0,8054
DVM	0,9006	0,8893	0,8821
SVRA	0,8818	0,7758	0,6308
LR	0,9009	0,8886	0,8676

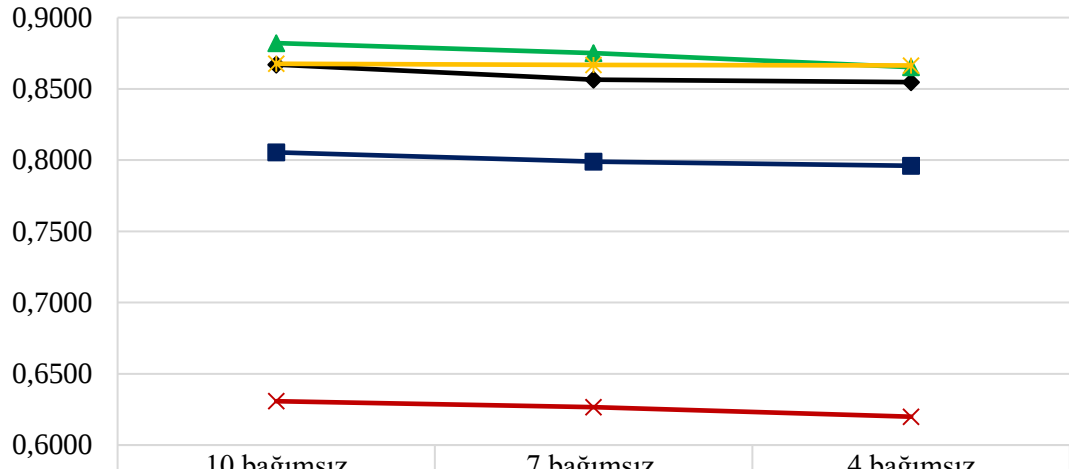
EK 8. Bağımsız değişken sayısının 10, 7 ve 4 olması durumunda yöntemlerin değerlendirme kriterlerinin değerlerine ait grafikler.





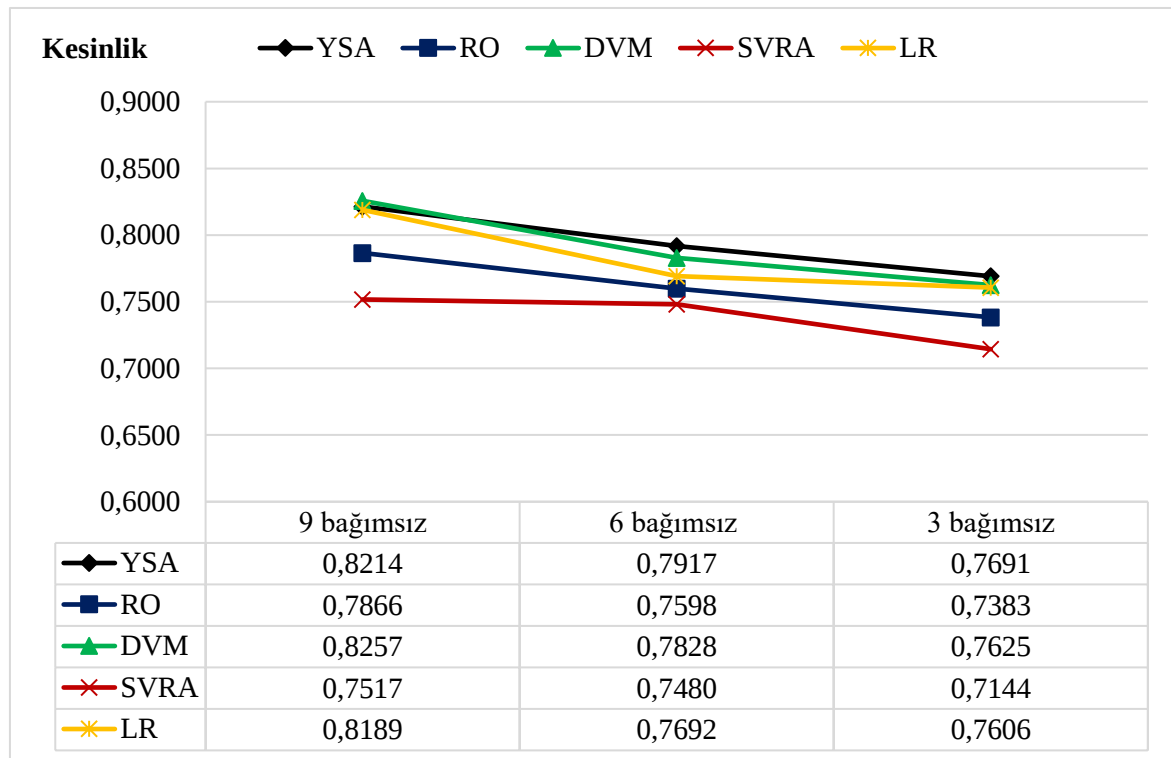
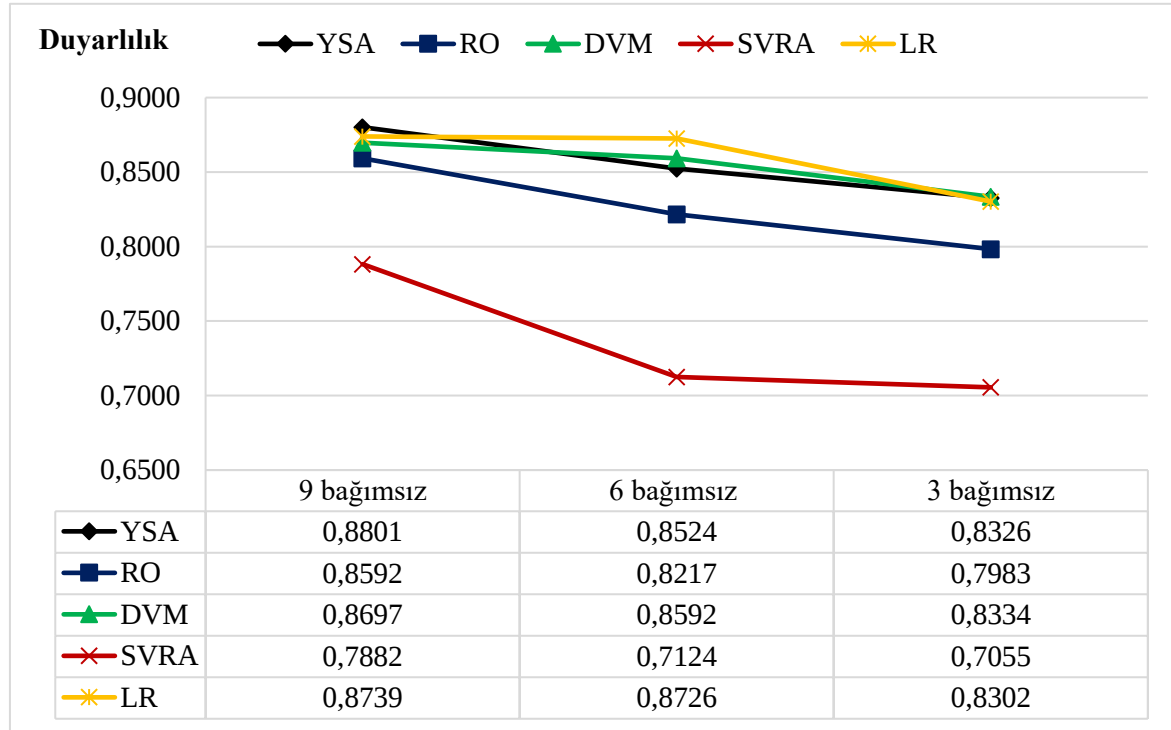
Negatif öngörü değeri

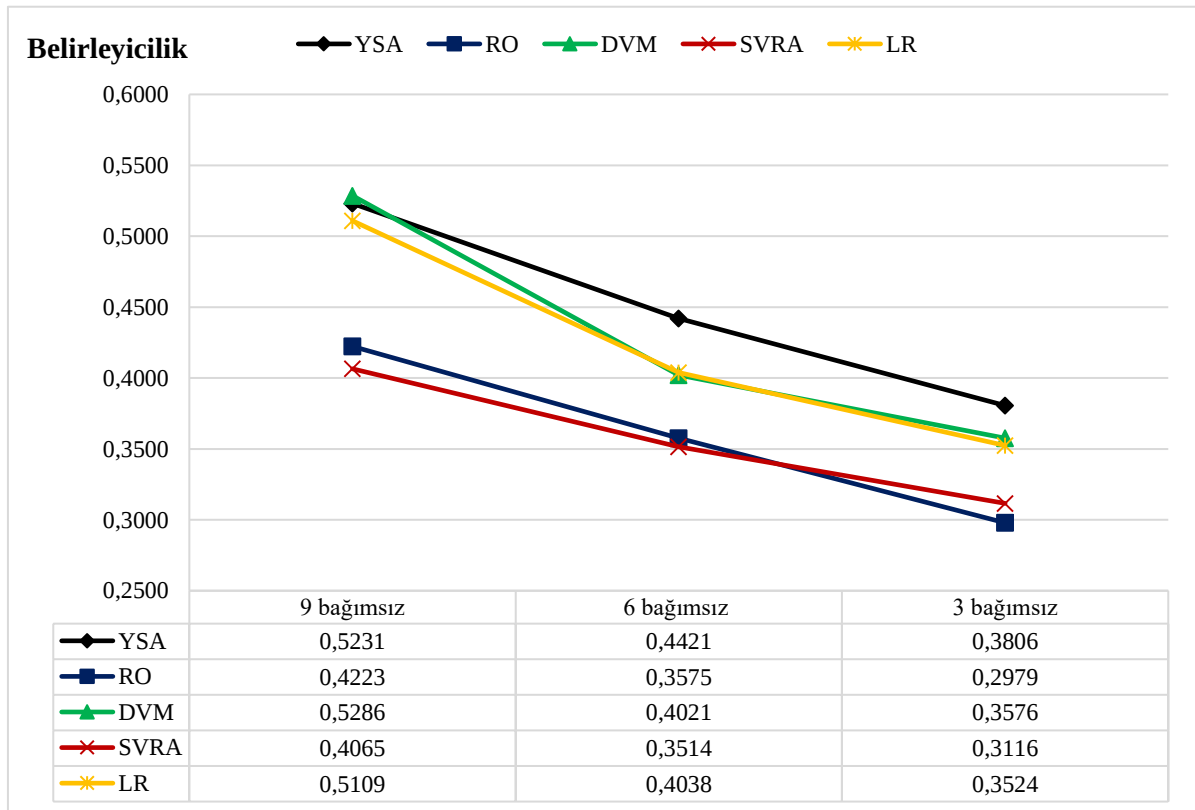
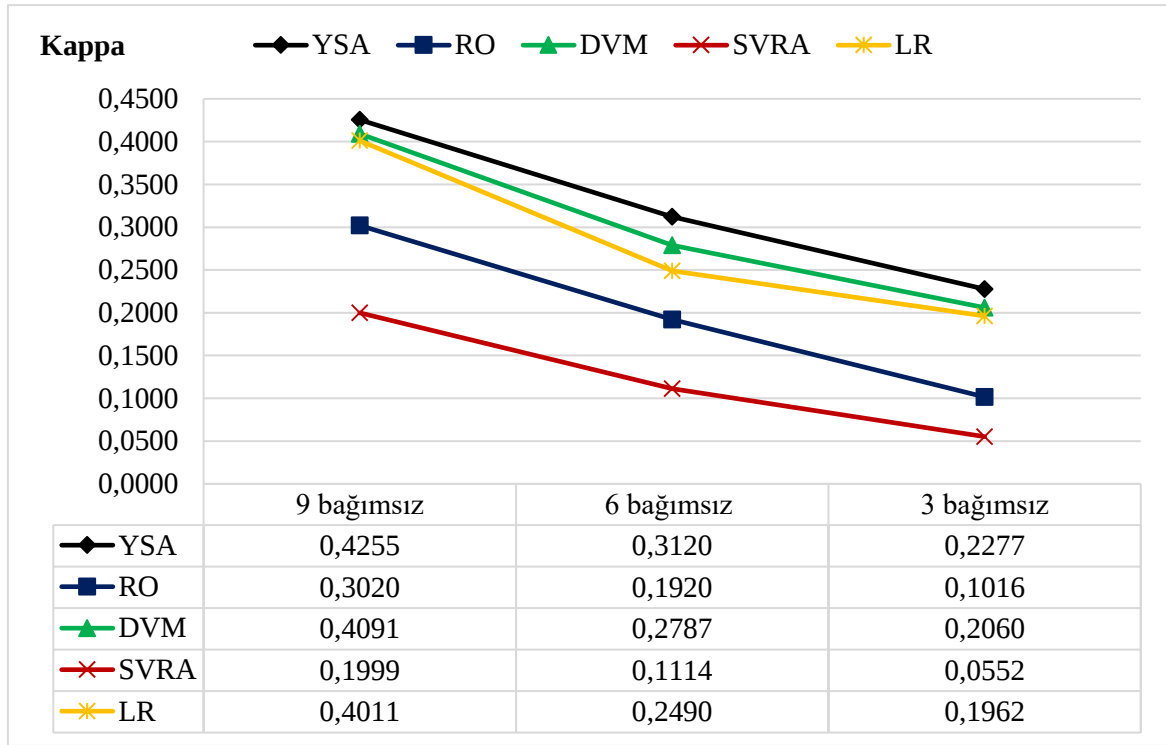
YSA RO DVM SVRA LR



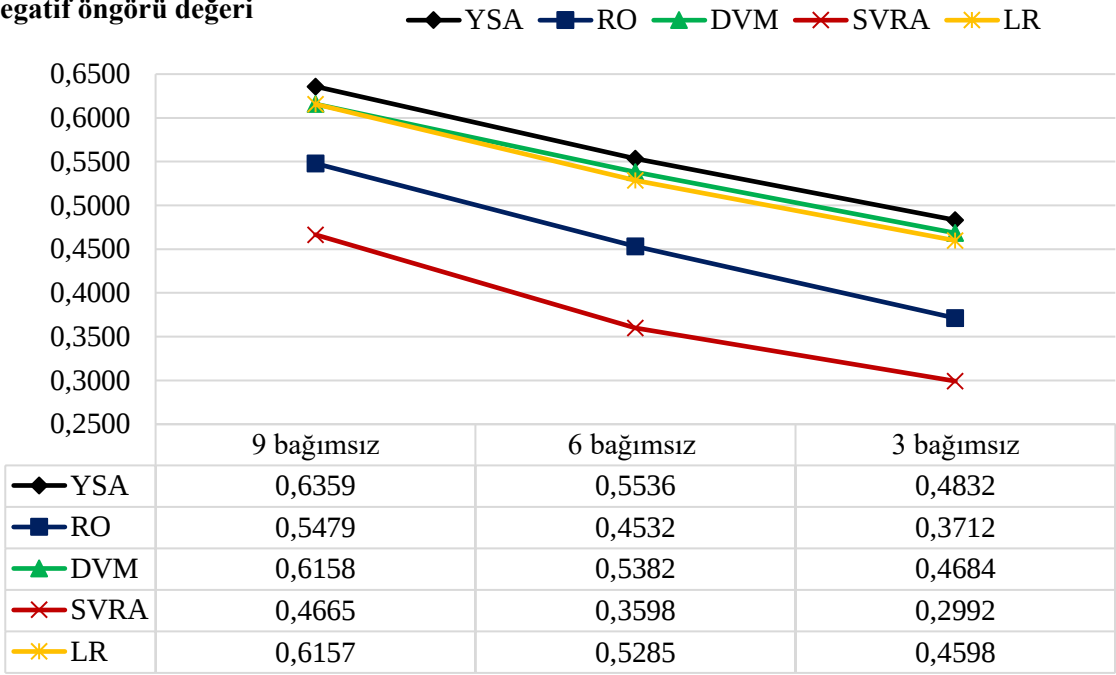
	10 bağımsız	7 bağımsız	4 bağımsız
YSA	0,8670	0,8564	0,8547
RO	0,8054	0,7989	0,7960
DVM	0,8821	0,8750	0,8652
SVRA	0,6308	0,6266	0,6199
LR	0,8676	0,8667	0,8664

EK 9. Bağımsız değişken sayısının 9, 6 ve 3 olması durumunda yöntemlerin değerlendirme kriterlerinin değerlerine ait grafikler.

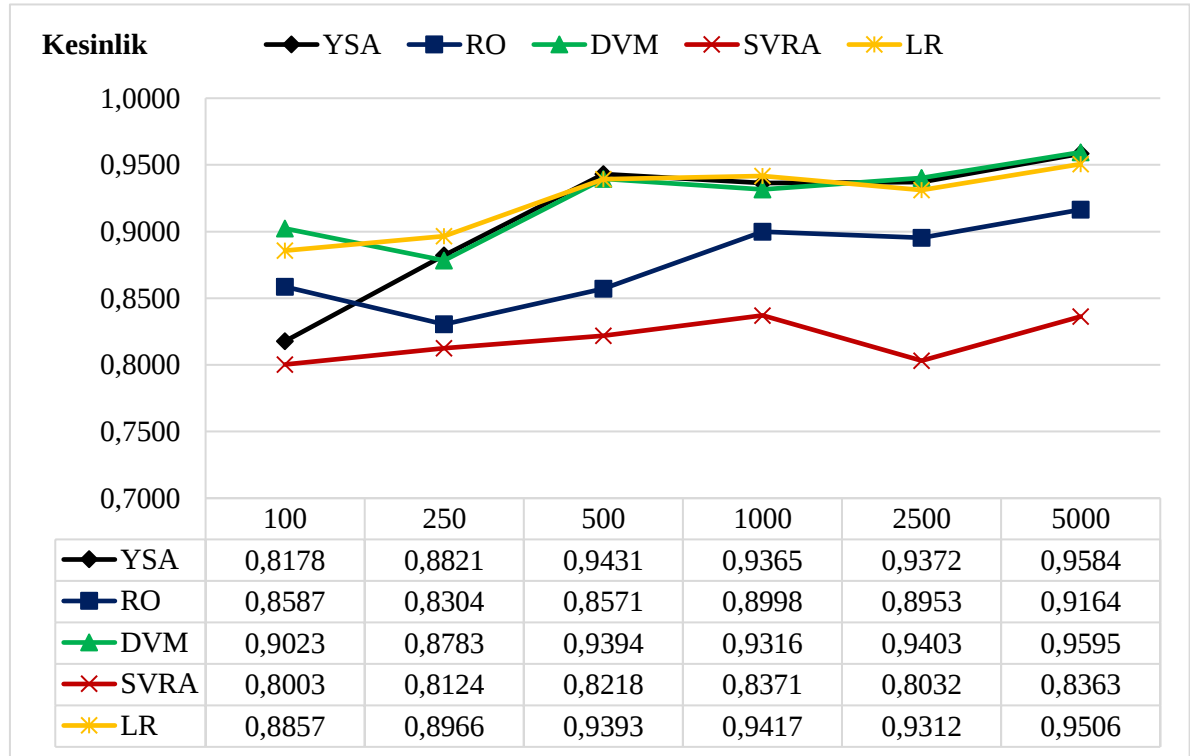
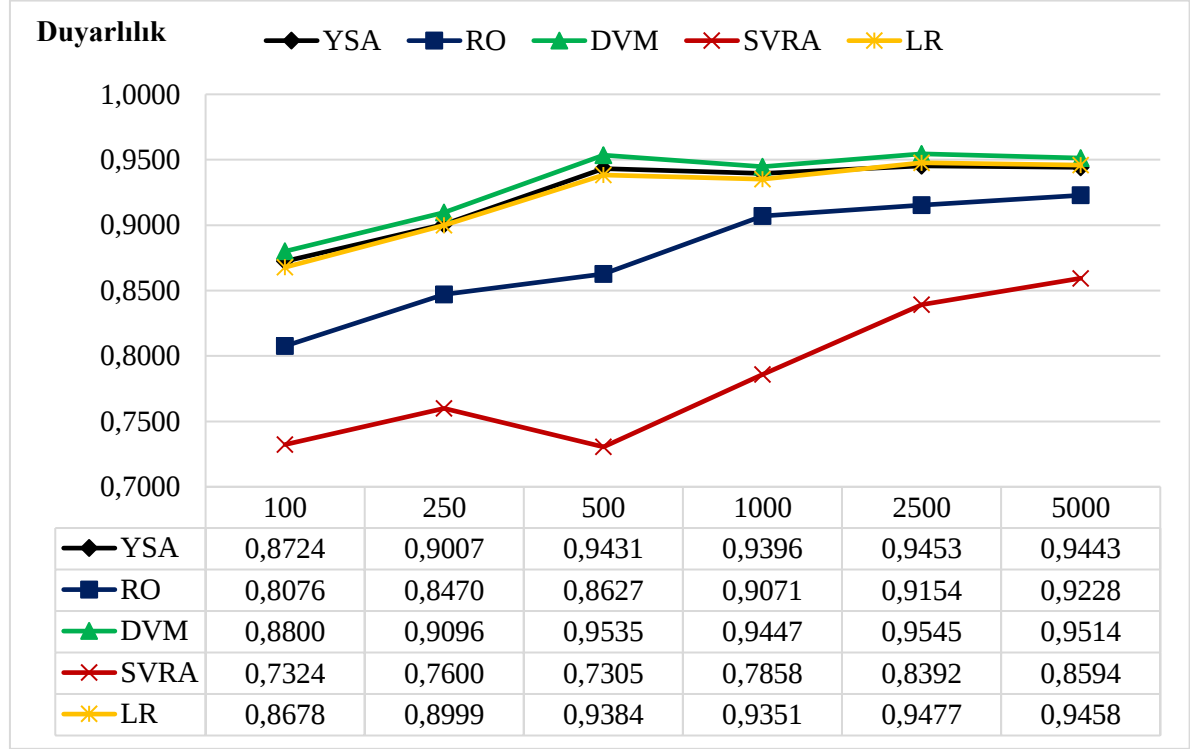


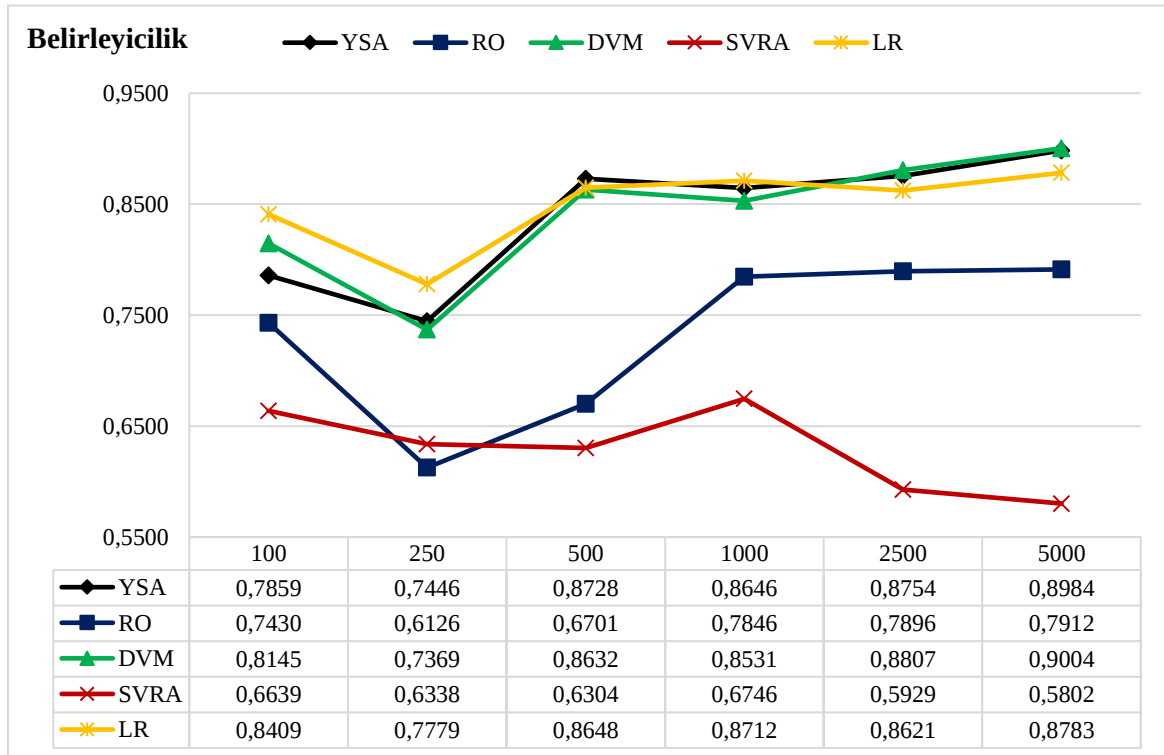
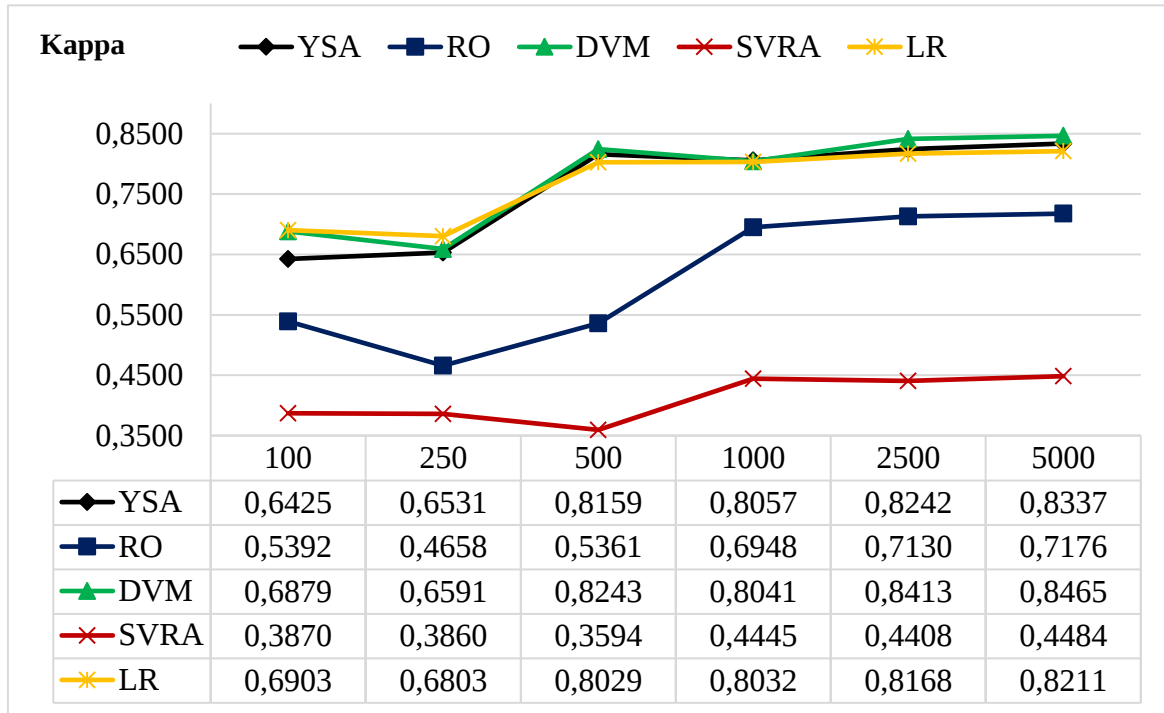


Negatif öngörü değeri



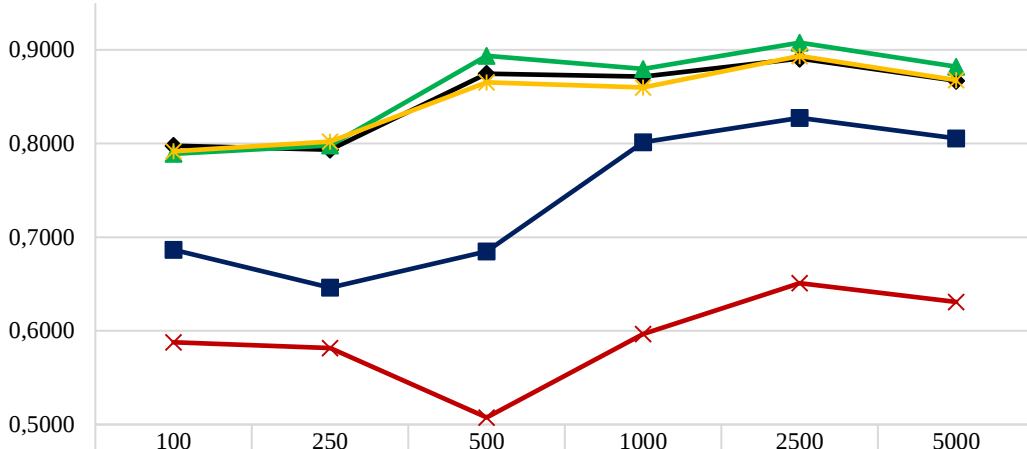
EK 10. Örneklem büyüklüğüne göre yöntemlerin değerlendirme kriterlerinin değerlerine ait grafikler.





Negatif öngörü değeri

YSA RO DVM SVRA LR



	100	250	500	1000	2500	5000
YSA	0,7977	0,7934	0,8745	0,8716	0,8907	0,8670
RO	0,6867	0,6461	0,6850	0,8013	0,8274	0,8054
DVM	0,7890	0,7981	0,8938	0,8796	0,9077	0,8821
SVRA	0,5878	0,5817	0,5075	0,5968	0,6509	0,6308
LR	0,7917	0,8022	0,8653	0,8599	0,8935	0,8676

EK 11. Etik Kurul Onayı

Evrak Tarih ve Sayısı: 27.04.2020-E.51153



T.C.
GAZİ ÜNİVERSİTESİ
Ölçme Değerlendirme Etik Alt Çalışma Grubu



Sayı : 91610558-302.08.01-
Konu : Bilimsel ve Eğitim Amaçlı

EĞİTİM BİLİMLERİ ENSTİTÜSÜ MÜDÜRLÜĞÜNE

İlgi : 10.03.2020 tarihli ve E.37115 sayılı yazı.

İlgi yazınız ile göndermiş olduğunuz, Eğitim Bilimleri Ana Bilim Dalı **Doktora öğrencisi Görkem CEYHAN'ın, Prof. Dr. İsmail KARAKAYA'nın** danışmanlığında yürüttüğü **"Eğitim Bilimleri Alanına Ait Veriler Üzerinde Veri Madenciliğinde Kullanılan Yöntemlerin Sınıflandırma Performanslarının Değerlendirilmesi"** adlı tez çalışması ile ilgili konu Kurulumuzun **07.04.2020** tarih ve **04** sayılı toplantısında görüşülmüş olup,

İlgilinin çalışmasının, yapılması planlanan yerlerden izin alınması koşuluyla yapılmasında etik açıdan bir sakınca bulunmadığına oybirliği ile karar verilmiş ve karara ilişkin imza listesi ekte gönderilmiştir.

Bilgilerinizi ve gereğini rica ederim.

e-imzalıdır
Prof. Dr. İsmail KARAKAYA
Kurul Başkanı

Araştırma Kod No: 2020-160

Ek: 1 Liste



Emniyet Mahallesi Bandırma Caddesi No :6/1 06560 Yenimahalle/ANKARA
Tel:0 (312) 202 20 57 - 0 (312) 2... Faks:0 (312) 202 38 76
İnternet Adresi :http://etikkomisyon.gazi.edu.tr/

Bilgi için :Esengül BOŞNAK
Birim Evrak Sorumlusu
Telefon No:03122022666



GAZİLİ OLMAK AYRICALIKTIR...