

T.C.
MARMARA ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ
İSTATİSTİK ANABİLİM DALI

İSTATİSTİKSEL MAKİNE ÖĞRENME YÖNTEMLERİYLE
BANKALARIN DERECELENDİRİLMESİ:
TÜRKİYE BANKACILIK SİSTEMİ UYGULAMASI

Yüksek Lisans Tezi

AMURI BINAMAN

DANIŞMAN:
Prof. Dr. Birsen EYGİ ERDOĞAN

İSTANBUL, 2022

ÖNSÖZ

“İstatistiksel Öğrenme Yöntemleriyle Bankaların Derecelendirilmesi: Türkiye Bankacılık Sistemi Uygulaması” isimli bu tez Marmara Üniversitesi Fen Bilimleri Enstitüsü İstatistik Anabilim Dalı, Yüksek Lisans Programı’nda hazırlanmıştır.

Çalışma konusunun belirlenmesinde ve çalışmanın hazırlanma sürecinin her aşamasında bilgilerini, tecrübelerini ve değerli zamanlarını esirgemeyerek bana her fırsatta yardımcı olan değerli hocam Sayın Prof. Dr. Birsen EYGİ ERDOĞAN’a teşekkürü borç bilirim.

Teşekkürlerin az kalacağı diğer bölüm hocalarım Prof. Dr. Müjgan TEZ ve Prof. Dr. Deniz İNAN’a Marmara Üniversitesi’nde yüksek lisans hayatım boyunca kazandırdıkları her şey için ve beni gelecekte söz sahibi yapacak bilgilerle donattıkları için hepsine teker teker teşekkürlerimi sunuyorum.

Son olarak, tüm hayatım boyunca maddi ve manevi her zaman beni destekleyen, her adımda arkamda duran aileme sonsuz teşekkürlerimi sunarım.

Bu tezin, bundan sonraki çalışmalara katkı sağlamasını temenni ederim.

İÇİNDEKİLER

ÖNSÖZ.....	i
İÇİNDEKİLER.....	ii
ÖZET	iv
ABSTRACT.....	v
SEMBOLLER	vi
KISALTMALAR.....	vii
ŞEKİL.....	viii
TABLO LİSTESİ.....	ix
1. BÖLÜM: GİRİŞ.....	1
1.1. LİTERATÜR TARAMASI.....	2
1.2. BANKACILIKTA BAŞARISIZLIK KAVRAMI	5
1.3. AMAÇ VE MOTİVASYONUMUZ	6
2. BÖLÜM: FİNANSAL ORANLAR	8
3. BÖLÜM: VERİ	11
3.1.1. TÜRK BANKALARI İÇİN BAŞARISIZLIK TANIMI	11
3.2. VERİ TEMİZLEME	12
3.2.1. EKSİK VERİLERİN YÖNETİMİ	13
3.2.2. AYKIRI DEĞERLER YÖNETİMİ	15
3.2.3. VERİ STANDARTLAŞTIRMASI.....	17
3.3. HEDEF SONUÇ	19
4. BÖLÜM: YÖNTEM VE ALGORİTMALAR.....	20
4.1. İSTATİSTİKSEL ÖĞRENME YÖNTEMLERİ.....	20
4.1.1. DENETİMLİ MAKİNE ÖĞRENMESİ.....	20
4.1.2. DENETİMSİZ MAKİNE ÖĞRENMESİ.....	21
4.2. MODEL DEĞERLENDİRME METRİKLERİ.....	21
4.3. EN ÇOK BİLİNER BAZI MAKİNE ÖĞRENME ALGORİTMALARI	23
4.3.1. PERCEPTRON	23
4.3.2. KARAR AĞACI.....	26
4.4. LOJİSTİK REGRESYON.....	28
4.5. DOĞRUSAL DİSKRİMİNANT ANALİZİ	29

4.6.	SUPPORT VECTOR MACHINES (SVM).....	30
4.7.	TEMEL BİLEŞENLER ANALİZİ.....	32
5.	BÖLÜM: ALGORİTMAMIZ	37
5.1.	BANKALARI BAŞARILI-BAŞARISIZ ŞEKLİNDE AYIRACAK ALGORİTMANIN SEÇİMİ	37
5.2.	OPTİMAL DEĞİŞKEN SAYISI BELİRLEME	37
5.3.	GENETİK ALGORİTMASI	39
5.4.	GENETİK ALGORİTMANIN UYGULAMASI	40
5.5.	UYGULAMA	41
5.6.	MODELLERİN KARŞILAŞTIRMASI.....	45
5.7.	DERECELENDİRME	46
6.	BÖLÜM: SONUÇ.....	50
	KAYNAKLAR.....	51

ÖZET

İSTATİSTİKSEL ÖĞRENME YÖNTEMLERİYLE BANKALARIN DERECELENDİRİLMESİ: TÜRKİYE BANKACILIK SİSTEMİ UYGULAMASI

Derecelendirme kuruluşları tarafından yapılan finansal derecelendirmeler, özel ya da kamu şirketlerine veya devlete verilen kredilerin geri ödenmemesi riskini değerlendirmek için kullanılan bir araçtır. İyi bir derecelendirmeye sahip olmak, söz konusu şirketin finansal sağlığını ve güvenilirliğini ortaya koymak açısından önemlidir. Bu tezde, Türkiye Cumhuriyeti'ndeki bankaların finansal oranları kullanılarak derecelendirilmesi amaçlanmaktadır. Bu amaçla istatistiksel öğrenme yöntemlerinden yararlandık.

Türkiye bankacılık sisteminde 2001 yılında varlık ve bilanço denetleme sisteminde yapılan değişikliklere kadar pek çok finansal kriz yaşanmıştır. Bu krizler neticesinde banka iflasları gözlemlenmiştir. Banka iflasları genel olarak ekonomik çöküntüye yol açabileceğinden bankaların düzenli olarak izlenmesi çok önemlidir. Genelde bu izleme bankaların yıllık bilançolarının takibi üzerinden yapılmaktadır. Yıllık bilançosuna bakarak, bir bankanın iflas edip etmediğini açıkça söyleyebiliriz ancak bankalar arasında hangi bankanın diğerinden daha fazla çökmeye yakın olduğunu söylemek zor olabilir. Bu çalışmanın amacı bankaların başarı düzeyleri açısından sıralanmalarını sağlayan bir derecelendirme sistemi önermektir.

Bu çalışmada kullanılan veriler Türkiye Bankalar Birliği'nin (TBB) web sitesinden elde edilmiştir. TBB'den gelen veriler iki türe ayrıldığı için (iki bin bir yılından önceki bilançolardan ve iki bin bir yılından sonra bilançolardan oluşan veriler) her şeyden önce bu verilerin uyumunu kontrol edip gereken düzeltmeleri yaptık. Sonra uygun veri düzenleme ve temizleme yöntemleri uygulayarak temel istatistiksel hesaplama işlemleri yaptık. Ardından, literatürde daha önce yapılmış bilimsel çalışmaları inceleyip alternatif bir derecelendirme modeli oluşturmak için istatistiksel öğrenme yöntemlerinden yararlandık.

ABSTRACT

BANKS RATING USING STATISTICAL LEARNING METHODS: AN IMPLEMENTATION WITH TURKEY BANKING SYSTEM

Financial ratings made by rating agencies are a tool used to assess the risk of non-repayment of loans given to private or public companies or to governments. Having a good rating is important to demonstrate the financial health and reliability of the company in question. In this thesis, it is aimed to rate banks of the Republic of Turkey using financial ratios. For this purpose, statistical learning methods were used.

Up to the balance sheet assets and audit changes made to the Turkey Banking System in 2001, there has been too much financial crisis. As a result of these crises, bankruptcies were observed. Regular monitoring of banks is very important, as bankruptcy can lead to economic downturn in general. In general, this monitoring is carried out by following the annual balance sheets of the banks. We can clearly say whether one bank went bankrupt or not by looking at their annual balance sheet, but it can be difficult to tell which bank is close to collapse among several banks. The aim of this study is to propose a rating system that allows to rank banks in terms of their success levels.

The data used in this study were obtained from the Banks Association of Turkey's (TBB) web site. Since the data from the TBB is divided into two types (balance sheets before the year two thousand one and balance sheets after the year two thousand one), first of all, the compliance of these data was checked and necessary corrections were made. Then, basic statistical calculations were performed by applying classical data editing and cleaning methods. Afterwards, statistical learning methods was used in order to form an alternative rating model by examining the previous scientific studies in the documentatons.

SEMBOLLER

- x : Bağımsız değişken ya da bağımsız değişkenlerin matrisi
- Y : Bağımlı değişken
- $\hat{y}$: Y değişkenin tahmini
- e : e sayısı veya Euler sayısı ($e = 2.718281\dots$)
- Σ : Toplam



KISALTMALAR

TBB	: Türkiye Bankalar Birliđi
TMSF	: Tasarruf Mevduatı Sigorta Fonu
CAMELS	: Sermaye Yeterliliđi, Aktif Kalitesi, Yönetim, Karlılık, Likidite, Duyarlılık
YP	: Yabancı Para
TP	: Türk Parası
SY	: Sermaye Yeterliliđi
BY	: Bilanço Yapısı
AK	: Aktif Kalitesi
L	: Likidite
K	: Karlılık
GG	: Gelir-Gider Yapısı
SP	: Sektör Payları
GP	: Grup Payları
SR	: Şube Rasyoları
FR	: Faaliyet Rasyoları
TG	: Target (hedef değışkeni, sınıflama problemindeki bağımlı değışken)
MCAR	: Missing Complete At Random (Tamamen Rastgele Eksik)
MAR	: Missing at Random (Rastgele Eksik)
NMAR	: Not Missing at Random (Rastgele Olmayan Eksik)
LDA	: Lineer Diskriminant Analizi
SVM	: Support Vector Machine
PC	: Principal Component (Temel Bileşen)

ŞEKİL

Şekil 1: Isolation Forest Algoritması.....	16
Şekil 2: Üç girişli Perceptron	24
Şekil 3: Perceptron'un Sınıflandırma Sonuçları.....	26
Şekil 4: 4 Değişkenli Karar Ağacının Çıktısı.....	27
Şekil 5: Karar Ağacının Sınıflandırma Sonuçları.....	28
Şekil 6: Lojistik Regresyonun Sınıflandırma Sonuçları.....	29
Şekil 7: LDA'nın Sınıflandırma Sonuçları	30
Şekil 8: SVM'in Sınıflandırma Sonuçları.....	31
Şekil 9: Yüzde 90 Varyansı Açıklayan Bileşenler	33
Şekil 10: PCA Dayalı Perceptron.....	34
Şekil 11: PCA Dayalı Karar Ağacı.....	34
Şekil 12: PCA Dayalı Lojistik Regresyon.....	35
Şekil 13: PCA Dayalı LDA	36
Şekil 14: PCA Dayalı SVM.....	36
Şekil 15: Genetik Algoritmasının Akış Şeması.....	40
Şekil 16: Perceptron'un Optimal Sonucu	42
Şekil 17: Lojistik Regresyonun Optimal Sonucu	43
Şekil 18: LDA'nın Optimal Sonucu	44
Şekil 19: SVM'in Optimal Sonucu	45
Şekil 20: Train Set için LDA'nın Gözlemleri Yansıttığı Eksen	47
Şekil 21: Train Seti İçin Derece Kümeleri	48

TABLO LİSTESİ

Tablo 1: Finansal Rasyo Değişkenlerinin kodlanması	9
Tablo 2: MCAR Rastgelelik için Test sonucu.....	14
Tablo 3: Aykırı Gözlemler	16
Tablo 4: Normal Dağılan Değişkenler	18
Tablo 5: Karşılaşma Tablosu.....	46
Tablo 6: Train Seti için LDA'nın Performansı	47
Tablo 7: 2021 Yılı için Bankaların Başarı Dereceleri	49

1. BÖLÜM: GİRİŞ

Bankalar ekonomimizde birkaç temel işlevi yerine getirir. Her şeyden önce, ödemeler ve finansal işlemler sisteminin merkezinde yer alırlar. Merkezi olmayan alışverişi kolaylaştırdığından, bu sistem piyasa ekonomisinin işleyişi için esastır (Hanh, Hang, Huy, & Nuong, 2021). Ayrıca bankalar, varlık ve yükümlülüklerin vadelerini dönüştürmede önemli bir işlev üstlenmektedir. Normalde mevduat şeklinde kısa vadeli borçları kabul ederler ve bunları ipotek veya şirket kredileri gibi uzun vadeli varlıklara dönüştürürler. Bunun dışında bankalar, aynı kısa vadeli varlıklara hızlı erişim sağlayarak müşterilerine likidite sağlar. Aslında bankalar çok çeşitli vadelerde operasyonlar gerçekleştirerek finansal piyasalarda daha fazla etkinliği teşvik eden arbitraja izin verir. Bankalar ayrıca, fonları tasarruf sahiplerinden yatırımcılara kanallara ederek kredi aracılığının temel rolünü üstlenirler. Tasarruf sahiplerinin risklerini çeşitlendirmesine ve hepimizin zaman içinde tüketimimizi yumuşatmasına izin verirler (Allen & Carletti, 2008). Böylece genç aileler bir ev satın almak için borç alabilir, üniversite öğrencileri öğrenim ücretlerini ödeyebilir ve işletmeler sermayelerini ve yatırımlarını finanse edebilir.

Bankaların modern ekonomimizde oynadığı önemli rolü göz önünde bulundurarak, herhangi bir iflas riski karşısında bankaların sağlamlığını mümkün olduğunca bilmemiz gerekir. Aksi halde bütün ekonominin zarara düşmesine neden olur (Wilson, 2006). Örneğin, 15 Eylül 2008'de New York yatırım bankası Lehman Brothers iflas etti. Lehman Brothers'ın iflası, etkileri bugün hala hissedilen 2008 küresel mali krizinin başlangıcı oldu (De Haas & Van Horen, 2012). Dünyanın dört bir yanındaki hükümetler, bankaları ve genel olarak finansal sistemi yeniden ayağa kaldırmak için çok fazla kamu parası kullanmak zorunda kaldı, bu da ulusal borçların patlamasına yol açtı.

Lehman Brothers'ın iflası öngörülebilmiş olsaydı gerek ulusal gerek global anlamda finansal sisteme bu derece kuvvetli negatif etkileri olmayabilirdi. Dolayısıyla bir bankanın mali sağlığını ve iflas riski karşısındaki sağlamlığını tahmin etmek önemlidir ve bunu tahmin etmemizi sağlayan araçlardan biri finansal derecelendirmedir.

Bu çalışmanın amacı Türkiye bankalarının 2018, 2019 ve 2020 yıllarına ait finansal oranlarından oluşan veri setini kullanarak 2021 yılı için bankaların başarısız olup olmadığını

tahmin etmek ve belli bir algoritmaya uygun olarak bir derecelendirme sistemi önermektir. Bu amaçla 2018, 2019 ve 2020 yılları baz alınarak elde edilen kümeleme analizi ve diskriminant analizinin sonuçları 2021 yılı bankalarının derecelendirmesi için kullanılmıştır.

1.1. LİTERARTÜR TARAMASI

Finansal derecelendirme elde etmemizi sağlayacak bir model oluşturmak için, daha önce bu konuda yapılmış çalışmaların literatürüne göz atmamız zorunludur. Gerekli incelemelerden sonra literatüre geçen çalışmaların genel olarak üç kategoriye ayrıldığını söyleyebiliriz:

İlk derecelendirme kategorisi, neredeyse sadece büyük derecelendirme kuruluşları tarafından yapılmakta olan hükümetlerin, bankaların ve diğer ticari işletmelerin finansal derecelendirmesidir. Bu derecelendirmeler yatırımcılar, ihracatçılar ve otoriteler için karar destek aracı olarak kullanılır (Tuzci, 2020) ve (De Haan & Amtenbrink, 2011).

Bir derecelendirme kuruluşu tarafından verilen finansal derecelendirme, bir şirketin kredi geri ödeme kapasitesine odaklanır. Bunun için ajanslar, derecelendirmeleri gerçekleştirmek için dikkate alınan verilere göre iki tür derecelendirme kullanırlar. İlk derecelendirme türü "stand-alone" olarak adlandırılır ve ikincisi "hepsi bir arada" derecelendirme olarak bilinir (King, Ongena, & Tarashev, 2016). 2000'li yılların sonundaki küresel mali krizden sonra, büyük zorluk çeken bankaların bir kısmı dış mali destek alıp iyileştiriler. Bu tür destekler hem bu bankaların ana kurumlarından hem de ait oldukları ülkelerin kamu yetkililerinden gelebilir.

Bir derecelendirme kuruluşu, şirkete hem iç öz kaynaklarını hem de potansiyel harici yardımı dikkate alarak bir derece verirse buna "hepsi bir arada" derecelendirme denir. Ve derecelendirme, sadece derecelendirilmiş şirketlerin dâhili varlıkları temelinde yapıldığında "bağımsız" bir derecelendirmeden bahsedilir.

Bu tür derecelendirme piyasada bulunan başlıca ajanslar Moody's, Standard & Poor's ve Fitch Ratings' gibi kurumlar tarafından yapılmaktadır. Bu üç ajans tek başına pazarın yaklaşık %90'ini elinde tutmaktadır (European Securities and Markets Authority (ESMA), 2021). Bu derecelendirme kuruluşları hakkında birçok eleştiri mevcut ancak en çok dile getirilen iki tanesi şu şekilde ifade edilmektedir:

İlk eleştiri, ajansların, notlarının kendi açılarından herhangi bir garanti veya taahhüt değil, bir görüş temsil ettiği konusunda çok ısrar etmeleridir (Frost, 2007), bu da birçok ülkenin kanunlarına göre ajansları yatırımcıların açabileceği davalarından korur.

İkinci eleştiri, derecelendirme kuruluşlarının 2007-2009 mali krizinin patlak vermesine iki nedenden dolayı katılmakla suçlanmalarıdır. Her şeyden önce, dayandıkları karmaşık finansal düzenlemelere rağmen sonradan “değersiz” olduğu ortaya çıkan senetleri (ipotek konulan konutlar vb.) derecelendirmeye geldiğinde subjektif davranma veya çıkar çatışması ile suçlandılar (White, 2009). Amerika Birleşik Devletleri'nde emlak piyasası koşulları bozulmaya başladığında, birçok finansal varlığın derecesini aniden düşürerek (“downgrading”) tepki gösterdiler ve piyasanın içine çekildiği aşağı yönlü spirale katkıda bulundular.

İkinci derecelendirme kategorisi ticari banka müşterileri için kredi derecelendirmedir. Bu kategori de bir kredi puanı ile ilgilidir ve bu kredi puanı, kredi geçmişinin bir değerlendirmesi ve bir kişinin kendisine verilen krediyi yönetme yeteneğidir. Risk seviyesini ölçmek için kullanılan matematiksel bir formülün sonucudur. Bu puan, kredi geçmişinin (kredilerin türleri ve süreleri, ödeme vadelerine uyulması vb.) ve kişisel bilgilerinin bulunduğu kredi raporunun ayrıntılı bir analizinden sonra elde edilir. Buradaki kredi notu genellikle bir kredi derecelendirme kuruluşu tarafından verilir. Bu şirketler, bir bireyin finansal taahhütlerini yerine getirebilme olasılığını değerlendirmek için veri toplar.

Bu amaçla son zamanlarda yapılan bir çalışmada (Ceylan M. , 2022), bankaların müşterileri için bir kredi derecelendirme veya kredi riski değerlendirme modeli önerilmiştir. Önerilen modelde eğitim verileri üzerindeki kredi riski değerlendirme ile ilgili doğruluk performansı %52, test verileri üzerindeki performans ise %41 çıkmıştır.

Üçüncü kategori ise genelde akademik çalışmalarda karşılaştığımız finansal başarısızlık öngörü modelleridir. Bu çalışmaların bir kısmı özellikle bankacılık faaliyetlerini konu almakta ve bankaların başarı ya da başarısızlığını belirlemek için yılsonunda toplanan yıllık finansal bilanço verilerini kullanmaktadır. Bilanço verilerinden başarısızlık tahmini yapmak isteyen araştırmacılar genellikle istatistiksel öğrenme yöntemlerini uygulamaktadır.

Bir bankanın borçlarını ödemek için yeterli fonu olmadığında bankanın başarısızlığı meydana

gelir. Kişiler hesaplarına para yatırırken genellikle para yatırımlarında ve farklı amaçlarla kullanılır. Ancak onu elinde bulunduran banka genellikle gerektiğinde iade edebileceğini garanti eder. Bir bankanın mali durumu, fon taleplerini işleme koyamayacak kadar kötüleşirse, banka kapanmak zorunda kalır. Daha sonra iflas ettiği söylenir. Bankaların ve diğer kurumların mali başarısızlığı belirlemek veya önceden öngörmek için birkaç çalışmalar yapılmıştır. Erdoğan (Erdoğan, 2008) Türkiye Cumhuriyeti'nde faaliyet gösteren 42 ticari bankayı 20 finansal oran ve lojistik regresyonu kullanarak başarılı ve başarısız olmak üzere iki sınıfa ayırıp finansal başarısızlığı iki sene önceden tahmin etmeye çalışmıştır. Uyguladığı model başarılı ve başarısız bankaları yüzde doksan beş (%95) doğrulukla tahmin edebilmiştir.

Altunöz (Altunoz, 2013) Yapay sinir ağlarının algoritması sayesinde bankaların 36 tane finansal rasyo kullanarak bankaların finansal başarısızlığını 1 ve 2 yıl önceden tahmin etmeye çalıştı. Yaptığı çalışmada 1 yıl önceden bankanın başarısızlığını öngörme başarısı yüzde seksen sekiz (%88) olup 2 yıl önceden ise öngörme başarısı yüzde yetmiş yedi (%77) çıkmıştır.

Aynı fikirle ve seneler önce Altman (Altman, 1968), mali başarısızlığı birkaç yıl önceden tahmin edecek bir model önermiştir. Modelinde 5 oran kullanmıştır: *İşletme sermayesi/Toplam varlıklar*, *Geçmiş Yıllar Karları/Toplam varlıklar*, *Faiz ve vergi öncesi kazanç/Toplam varlıklar*, *Piyasa değeri özkaynak/Toplam borcun defter değeri* ve *Satışlar/Toplam varlıklar*. Onun çıkardığı modelin 1 ve 2 sene önce öngörme doğruluğu %95 ve %72 iken modelin 3 yıl önce öngörme doğruluğu %50'nin altına düşmüştür. Seneler artıkça da doğruluk yüzdesi de azalarak gitmiştir.

Karşılaştırmak amacıyla Wu, Gaunt, & Gray bu konuyla ilgili çeşitli çalışmaları ve tahmin modellerini derlediler. Buldukları sonuca göre, örneğin, ana bilanço bilgilerinin karlılık, likidite ve kaldıraçla ilgili olduğu sonucuna vardılar. Spesifik olarak, şirketlerin toplam aktiflerinden nispeten daha düşük faiz ve vergi öncesi karlılık varsa, şirketlerin başarısız olma olasılığı daha yüksektir (Wu, Gaunt, & Gray, 2010).

Odom ve Sharda, Altman'ın (Altman, 1968) kullandığı değişkenleri, "A neural network model for Bankruptcy prediction" başlıklı (Odom & Sharda, 1990) ve yapay sinir ağlarına dayalı çalışmalarının modeline girdi olarak kullanmışlardır. Sonuç olarak, yapay sinir ağları kullanan modellerin basit Diskriminant modellerden daha iyi tahminler sunduğunu ifade etmişlerdir.

Son yıllarda makine öğreniminde, "Ensembles" denilen yöntemlerden ve "Hybrid" olarak adlandırılan yöntemlerden bahsedilmektedir. Ensembles yöntemleri, makine öğrenme modellerinin birleşik gücünün sınıflandırma problemi veya regresyon problemi gibi bir öğrenme probleminde kullanıldığı bir makine öğrenme kavramıdır. Bu yaklaşımda, birkaç homojen makine öğrenme modeli “weak learners” (zayıf öğrenenler) olarak alınır ve birlikte gruplandırılır. Probleme uygulandığında, zayıf öğrenenlerin her biri ya tüm eğitim setinde ya da tüm eğitim setinin bir kısmında kendi sonucunu gösterir. Son olarak, nihai sonucu elde etmek için her zayıf öğrenenin sonuçları birleştirilir. Ensembles yöntemlerinin iki popüler kategorisi vardır, biri Bagging, diğeri ise Boosting olarak adlandırılır. Rastgele orman, Bagging kategorisine giren popüler Ensemble öğrenme modelidir. AdaBoost, Boosting kategorisine giren bir başka popüler Ensemble öğrenme modelidir. Bagging modelleri, tüm veri kümesinin bir kısmı üzerinde çalışırken, Boosting modelleri tüm veri kümesi üzerinde çalışır. Hybrid yöntemler ise daha yüksek performans ve optimal sonuçlar için iki veya daha fazla makine öğrenme yöntemini birleştirir. Aslında Hybrid yöntemler, iki veya daha fazla yöntemin avantajından yararlanarak daha iyi performansa ulaşır. Aykut Ekinci ve Halil İbrahim Erdal (Ekinci & Erdal, 2017), bankaların başarısızlığını tahmin etmek için farklı temel, Ensembles ve Hybrid yöntemlerinin performanslarını karşılaştırmışlardır. Buldukları sonuçlara göre Hybrid öğrenme modellerinin banka başarısızlıkları için güvenilir bir tahmin modeli olarak kullanılabileceği gösterilmiştir.

Bu konuda son yıllarda yapılan çalışmalar hakkında daha fazla bilgi edinmek için Alberto Citterio'nun 2000 yılından sonra yapılan birçok çalışmayı derlediği yayınına başvurulabilir. (Citterio, 2020). Citterio, çalışmasında farklı başarısızlık tanımları, tahmin modellerinin seçimi ve değişken seçiminin farklı yöntemlerinden bahsetmiştir.

Yukarıda belirtilen tüm çalışmaların ortak noktası başarılı bankaları veya şirketleri başarısız bankalardan veya şirketlerden ayırma amacını taşımalarıdır. Bu tezde de önerilen yaklaşım ile bankaların sınıflandırılması ve derecelendirilmesi konusunda daha önce yapılan çalışmalara katkı sunmak amaçlanmaktadır.

1.2.BANKACILIKTA BAŞARISIZLIK KAVRAMI

Derecelendirme kuruluşları tarafından yapılan finansal derecelendirmenin yanı sıra,

çalışmalarımızı ilgilendiren özel araştırma alanı, banka mali başarısızlığı ile ilgili çalışmalardır. Finansal başarısızlık kavramı, kişinin içinde bulunduğu faaliyet sektörüne ya da konuştuğu alandaki uzmanlara göre farklı tanımlara tabidir. Bu nedenle bazıları, bir bankanın iflas ettiğinde veya ardışık üç yıl boyunca zarar etmesi ya da başka bir banka tarafından satın alınması durumunda da finansal başarısızlığa uğradığını söyleyebilir (Ceylan & Korkmaz, 2001). Bankalar için en yaygın finansal başarısızlık kriterleri aşağıdaki gibidir:

- Ödenecek borçların ödenememe durumuna uğramak
- İflas etmek
- Konkordato ilan etmek
- Merkez Bankası tarafından faaliyetlerin durdurulması
- Bankaya kayyum atanması
- Hisse senedi temettüsünü ödememek
- Bankanın Faaliyetlerine son verilmesi ya da durdurulması
- Türkiye’de TMSF kurumuna devir edilmek

Yukarıdaki kriterlerin birçoğu, bilanço verileri ve istatistiksel öğrenme teknikleri kullanılarak incelenebilir. Literatürde en çok kullanılan yöntemler arasında Yapay Sinir Ağları, Lojistik Regresyon ve Diskriminant Analizi’ni gösterebiliriz (Mansouri, Nazari, & Ramazani, 2016).

Hedefimiz sadece bir derecelendirme modeli oluşturmak olduğunda bankaların finansal başarısızlığı ile ilgili çalışmalara neden ilgi duyduğumuz sorusuna cevabımız, bu derecelendirme modelinin kurulmasının amacı başarısızlığı önlemek için bankaların iyi finansal sağlığa sahip olduklarına emin olmaktır. Geliştirmek istediğimiz derecelendirme sistemi finansal sağlığı ölçmek için bir ölçek olarak görülebilir. Dolayısıyla, bir bankanın derecesi ne kadar kötü olursa finansal başarısızlığa düşmeye o kadar yakın olur ve benzer şekilde bir bankanın derecesi ne kadar iyi olursa finansal başarısızlıkla karşılaşma riski o kadar az olur.

1.3.AMAÇ VE MOTİVASYONUMUZ

Finansal derecelendirme kuruluşları farklı ve çeşitlidir ancak hepsinin ortak bir noktası vardır: Derecelendirilen şirketin kurulduğu ülkeye bakılmaksızın, tüm şirketler için aynı derecelendirme yöntemlerini uygularlar. Bankacılık sektöründe bu yaklaşım, gelişmiş bir

ülkeden bir şirkete ve geliřmekte olan bir ülkeden bir şirkete aynı řekilde uygulandıęında aynı başarıyı göstermeyebilir. Bunun nedeni sermaye ve banka bilançolarına ilişkin yasaların gelişmemiş ülkelerde yeterince titizlikle ele alınmamasıdır. Bazen bu ülkelerde bilanço düzenlenmesi yasaları neredeyse hiç yoktur. Ayrıca bu ülkeler bankacılık hisse senedi piyasalarının eksiklięini göstermektedir. Bu çalışmanın amacı, bir ülkenin durumuna uyarlanmış bir model önermektir.

Biraz araştırma yaptıktan sonra Türkiye'deki ve dięer birçok ülkedeki bankaların derecelendirmesi yalnızca yabancı kuruluşlar tarafından yapılmaktadır. Finansal derecelendirme yapan yerel ajanslar bile sadece bankalar dışındaki şirketleri derecelendirmektedir. Bu ajanslar oldukça iyi bir iş çıkarsalar bile, yabancı kuruluşlar tarafından yapılan derecelendirmelerin objektiflięi zaman zaman sorgulanmaktadır.

Finansal derecelendirme kuruluşu Fitch'in yanı sıra Moody's veya Standard ve Poors gibi dięer ünlü kuruluşlar, sanayileşmiş ülkelerde oluşturulan büyük çok uluslu şirketlerdir. Bu ajansların dięer bir ortak özellięi, farklı ülkelerde var olabilecek yasa ve denetleme sistemlerindeki farklılıkları dikkate almamalarıdır. Bu nedenle Fitch gibi bir finansal derecelendirme kuruluşu, derecelendirdięi tüm şirketler ve bankalar için aynı metodolojiyi kullanır. Çoęu zaman büyük derecelendirme kuruluşları tarafından uygulanan metodolojiler, sanayileşmiş ülkelerin ekonomileri üzerine kurulan ekonomik göstergelere dayanmaktadır ve bu göstergeler, geliřmekte olan ülkelerin şirketlerine uygulanması söz konusu olduęunda zayıf performanslar vermektedir. Latin Amerika ve Doęu Asya'daki bazı ülkelerde bu durum gözlenmektedir (Rojas-Suarez, 2002).

Bu tezde, bir ülkenin bankaları için, finansal verilere dayanan bir derecelendirme metodu oluşturmak istedik, uygulama örneęi olarak Türkiye'de faaliyet gösteren bankaları ele aldık.

2. BÖLÜM: FİNANSAL ORANLAR

Bankalar faaliyet alanı özel olan işletmelerdir. Kredi verir, müşterilerinin varlıklarını yönetir, finansal piyasalara katılır, kısaca ekonominin işleyişinde önemli bir bağlantı oluştururlar. Bu nedenle faaliyetleri denetlenir ve belirli muhasebe standartlarının kullanılması gerekir.

Kasım 1979'da Amerika Birleşik Devletleri Federal Finansal Kurumlar İnceleme Konseyi (FFIEC: The Federal Financial Institutions Examination Council), yaygın olarak CAMEL derecelendirme sistemi olarak bilinen finansal kurumlar için tek tip derecelendirme sistemini kabul etmiştir (Christopoulos, Mylonakis, & Diktapanidis, 2011). Sistem, finansal durumlarına göre yaklaşık 8.500 bankayı sıralamak için kullanılıyor. Daha sonra sistem, başta Amerikan Federal Rezervi (FED) olmak üzere diğer birçok kuruluş tarafından benimsendi (Dang, 2011).

CAMELS kısaltmasıyla temsil edilen altı faktöre dayalı olarak bankacılık denetçilerinin finansal kurumları değerlendirmek için kullandıkları tanınmış bir Amerikan derecelendirme sistemidir (GÖKMEN, 2007). CAMELS kısaltmasının anlamı:

C: Capital adequacy (Sermaye Yeterliliği)

A: Assets (Aktif Kalitesi)

M: Management Capability (Yönetim)

E: Earnings (Karlılık)

L: Liquidity (Likidite)

S: Sensitivity (Duyarlılık)

Bankalar faaliyetlerini ölçmek ve bilançolarını değerlendirmek için senelik finansal göstergeleri kullanırlar. Bu göstergeleri ve değerlerini içeren tablolara finansal rasyolar tabloları denir. Bu çalışma için TBB'nin web sitesinden elde edilmiş veriler CAMELS sistemine uygun olarak rasyolar tablosu şeklinde olup başka rasyolar da dâhil edilmiştir. Bu çalışmada kullandığımız rasyolar ve değişken olarak kodlanmaları aşağıdaki tabloda sunulmaktadır.

Tablo 1: Finansal Rasyo Değişkenlerinin kodlanması

Grup	Kod	Rasyo	Grup	Kod	Rasyo
Sermaye Yeterliliği	SY1	Sermaye Yeterliliği Oranı	Grup Payları	GP1	Toplam Aktifler
	SY2	Özkaynaklar / Toplam Aktifler		GP2	Toplam Krediler
	SY3	(Özkaynaklar-Duran Aktifler) / Toplam Aktifler		GP3	Toplam Mevduat
	SY4	Özkaynaklar / (Mevduat + Mevduat Dışı Kaynaklar)		GG1	Özel Karşılıklar Sonrası Net Faiz Geliri / Toplam Varlıklar
	SY5	Bilanço içi Döviz Pozisyonu / Özkaynaklar	Gelir-Gider Yapısı	GG2	Özel Karşılıklar Sonrası Net Faiz Geliri / Faaliyet Brüt Karı
	SY6	Net Bilanço Pozisyonu / Özkaynaklar		GG3	Faiz Dışı Gelirler (Net) / Toplam Varlıklar
	SY7	(Net Bilanço Pozisyonu + Net Nazım Hesap Pozisyonu) / Özkaynaklar		GG4	Faiz Dışı Gelirler (Net) / Diğer Faaliyet Giderleri
Bilanço Yapısı	BY1	TP Varlıklar / Toplam Varlıklar		GG5	Diğer Faaliyet Giderleri / Faaliyet Brüt Karı
	BY2	YP Varlıklar / Toplam Varlıklar		GG6	Kredi Karşılıkları / Toplam Varlıklar
	BY3	TP Yükümlülükler / Toplam Yükümlülükler		GG7	Faiz Gelirleri / Faiz Giderleri
	BY4	YP Yükümlülükler / Toplam Yükümlülükler		GG8	Toplam Gelirler / Toplam Giderler
	BY5	YP Varlıklar / YP Yükümlülükler		GG9	Faiz Gelirleri / Toplam Varlıklar
	BY6	TP Mevduat / Toplam Mevduat		GG10	Faiz Giderleri / Toplam Varlıklar
	BY7	TP Krediler / Toplam Krediler		GG11	Faiz Gelirleri / Toplam Gelirler
	BY8	Toplam Mevduat / Toplam Varlıklar		GG12	Faiz Giderleri / Toplam Giderler
	BY9	Alınan Krediler / Toplam Varlıklar	Şube Rasyoları	SR1	Şube Başına Toplam Aktif
Aktif Kalitesi	AK1	Finansal Varlıklar (Net) / Toplam Varlıklar		SR2	Şube Başına Toplam Mevduat
	AK2	Toplam Krediler / Toplam Varlıklar		SR3	Şube Başına TL Mevduat
	AK3	Toplam Krediler / Toplam		SR4	Şube Başına YP Mevduat

Sektör Payları	Likidite	Mevduat			
		AK4	Donuk Alacaklar / Toplam Krediler	SR5	Şube Başına Krediler ve Alacaklar
		AK5	Duran Varlıklar / Toplam Varlıklar	SR6	Şube Başına Personel (kişi)
		AK6	Tüketici Kredileri / Toplam Krediler	SR7	Şube Başına Net Kar
		L1	Likit Aktifler / Toplam Aktifler	FR1	(Personel Gideri + Kıdem Tazminatı) / Toplam Varlıklar
		L2	Likit Aktifler / Kısa Vadeli Yükümlülükler	FR2	(Personel Gideri + Kıdem Tazminatı) / Personel Sayısı (Bin TL)
	Karlılık	L3	TP Likit Aktifler / Toplam Aktifler	FR3	Kıdem Tazminatı / Personel Sayısı (Bin TL)
		L4	Likit Aktifler / (Mevduat + Mevduat Dışı Kaynaklar)	FR4	Personel Gideri / Diğer Faaliyet Giderleri
		L5	YP Likit Aktifler / YP Pasifler	FR5	Diğer Faaliyet Giderleri / Toplam Varlıklar
		K1	Ortalama Aktif Karlılığı	FR6	Faaliyet Brüt Karı / Toplam Varlıklar
	Sektör Payları	K2	Ortalama Özkaynak Karlılığı	FR7	Net Faaliyet Karı(Zararı) / Toplam Varlıklar
		K3	Vergi Öncesi Kar / Toplam Aktifler		
		K4	Net Dönem Karı (Zararı) / Ödenmiş Sermaye		
		SP1	Toplam Aktifler		
		SP2	Toplam Krediler ve Alacaklar		
		SP3	Toplam Mevduat		

Bu çalışmada bağımsız değişken olarak toplam 63 tane finansal rasyo kullanılmıştır.

3. BÖLÜM: VERİ

Verinin istatistiksel analizine girmeden önce dikkatle incelenmesi, eğer varsa, eksik ya da aykırı gözlemlerinin tespit edilmesi uygun olacaktır.

3.1.1. TÜRK BANKALARI İÇİN BAŞARISIZLIK TANIMI

2001 yılının Türkiye Cumhuriyeti bankacılık tarihinde bir dönüm noktası olarak görebiliyoruz. Nitekim Ekim 2000 ve Şubat 2001'de Türk bankacılık sektörü tarihinin en büyük krizlerinden birini yaşamıştır (Samırkaş, Evcı, & Ergün, 2014). Bu büyük tarihi krizden önce birçok banka iflas nedeniyle Tasarruf Mevduatı Sigorta Fonu'na devredilmiştir.

- Türk Ticaret Bankası 6 Kasım 1997,
- Bank Ekspres 12 Aralık 1998,
- Interbank 7 Ocak 1999,
- Egebank, Yurtbank, Sümerbank, Eskişehir Bankası, Yaşarbank 21 Aralık 1999,
- Etibank, Bank Kapital 27 Ekim 2000,
- Demirbank 6 Aralık 2000,
- Ulusal Bank 28 Şubat 2001,
- İktisat Bankası 15 Mart 2001,
- Sitebank, Milli Aydın Bankası, Bayındırbank, Kentbank, EGS Bank 9 Temmuz 2001,
- Toprakbank 30 Kasım 2001

Tarihinde TMSF'ye devredilmiştir.

Şimdiye kadar literatürde tespit ettiğimiz Türk banka başarısızlıklarının tahmini ile ilgili tüm çalışmalarda yukarıda listelediğimiz bankaların senelik bilanço verileri kullanılmıştır. Başka bir deyişle, bankaların başarısı veya başarısızlığını tahmin etmeye yönelik çalışma, bu verileri modellerinde kullanılacak başarısız bankaların örneği olarak kullanmıştır. Ancak bankacılık sisteminde yapılan reformlar ve 2001 yılındaki büyük kriz sonrasında oluşturulan kontrol politikası sayesinde artık kriz öncesinde olduğu gibi iflas eden banka vakalarını gözlenmemiştir. Mutlak anlamda bankaların iflas ettiğini görmemek ekonomi için çok iyi bir şey fakat bizimki gibi bir bankanın finansal sağlığını güncel verilerle belirlemeyi amaçlayan bilimsel çalışmalar yapabilmek için başarısız banka örneklerine ihtiyaç vardır. Bariz iflas

vakalarının olmaması nedeniyle başarısız bankaların örneklerini oluşturacak veri setini belirlemek için tekrar literatüre bakmamıza fayda vardır.

2018 yılında Yakıcı Ayan ve Değirmenci'nin "Firma finansal başarısızlık öngörüsü için bir lojistik regresyon modeli" başlıklı (Yakıcı Ayan & Değirmenci, 2018) çalışmasında ve Yusuf Gör' ün "Kurumsal yönetimin finansal başarısızlığı önlemedeki yeri üzerine bir araştırma" yazısında (Gör, 2018), bir bankanın 3 yıl üst üste zarar beyan etmesi halinde başarısız sayılabileceğini ifade ettiler. "Yapay Zekâ Yöntemleri ile İşletmelerin Finansal Başarısızlığının Tahmin Edilmesi" (Yürük & Ekşi, 2019) ve "Teknoloji yoğunluğuna göre finansal başarısızlık tahmin modelleri değişir mi? imalat sanayi sektörü üzerine bir uygulama" (Giray Yakut & Bacaksız, 2020) çalışmalarında, Yürük & Ekşi ve Giray Yakut & Bacaksız sırasıyla 2 ve 3 yıl arka arkaya bir şirketin zarar beyan etmesi durumunda başarısız olarak kabul edileceğini de belirtmişlerdir.

Bu bilgiler sayesinde, modelimiz için başarısız bankalar kategorisini oluşturacak tüm bankaları belirledik. Bankaların 2021 yılındaki durumunu tahmin edebilmek için 2018, 2019 ve 2020 yıllarına ait finansal verileri kullanmak temel amacımız olan bu çalışmada, başarısız kategorisinde yer alan bankalar 2018, 2019 ve 2020 yıllarında art arda zarar eden (aktif karlılık oranı üç yıl art arda ortalamanın altında kalanlar) bankalar olacaktır.

3.2.VERİ TEMİZLEME

Veri Temizleme, veri kalitesini iyileştirmeyi amaçlayan çeşitli süreçleri kapsar. Bir veri setindeki sorunları ortadan kaldırmak için birçok araç ve uygulama mevcuttur. Bu işlemler, bir veritabanı veya veri setindeki hatalı kayıtları düzeltmek veya kaldırmak için kullanılır. Genel olarak konuşursak, veri temizleme işlemleri, eksik, yanlış, bozuk veya alakasız veri veya kayıtların belirlenmesini ve değiştirilmesini içerir. Veri Temizleme uygun bir şekilde gerçekleştirildikten sonra, tüm veri setleri tutarlı ve hatasız olmalıdır. Bu işlem, bu verilerin kullanımı ve işlenmesi için gereklidir.

Temizleme yapılmadan analiz sonuçları bozulabilir (CHU, ILYAS, KRISHNAN, & WANG, 2016). Benzer şekilde, kötü veriler üzerinde eğitilmiş bir makine öğrenme veya yapay zekâ modeli önyargılı olabilir veya düşük performans sağlayabilir.

Veri Temizleme birçok avantaj sunar. Başlıca faydalarından biri, daha iyi veriye dayalı karar vermeyi mümkün kılmaktır. Veriler daha yüksek kaliteli olunca, veriye dayanan tüm faaliyetleri olumlu yönde etkiler

Veri temizleme, yeni verilere yer açmak için satırları yeniden düzenlemek veya bilgileri silmek kadar basit bir işlem değildir. Bizim çalışmamızda veri temizleme süreci aşağıdaki işlemleri kapsamaktadır:

- Yazım ve sentaks hatalarını düzeltme
- Veri kümelerini standartlaştırma
- Boş alanlar gibi hataları düzeltme
- Kayıtlardaki yinelenenleri belirleme, vs.

3.2.1. EKSİK VERİLERİN YÖNETİMİ

Belirli bir birey için belirli bir değişken için gözlem olmadığında eksik veriden bahsediyoruz. Veri yönetimi sorunu geniş bir konudur. İstatistiksel analiz sırasında eksik veriler göz ardı edilemez. Ancak oranlarına ve türlerine göre farklı çözümler seçilebilir. Değişkenleri veya verileri eksik olan bireyleri kaldırabilir veya eksik verilere değer atanabilir, hatta eksik verilerin varlığında analizlerin yapılmasını mümkün kılan yöntemler (veya algoritmalar) geliştirilebilir.

Eksik veriler, sadece veri biliminde değil, istatistiksel modellemelerde de çok yaygın olan bir problemdir. Eksik veri problemi ile uğraşma aşamasında, temel veri setini önemli ölçüde değiştirmeden, bu eksik olan bilgileri doldurmak ya da kaldırmak seçeneklerinden doğru olana karar vermek gerekmektedir..

Eksik verilerle uğraşmanın zorluğu, eksik verilerin kalıpları hakkında yaptığımız varsayımlarda yatmaktadır. Genel olarak üç tip eksik veri tanımlanır (Alpar, 2017) ve (Kwak & Kim, 2017) .

Bir X açıklayıcı değişkenin değeri eksik olduğunda, bu eksikliğin MCAR (Missing Complete At Random) olduğunu söyleyebilmemiz için ancak ve ancak değer eksik olma olasılığının diğer açıklayıcı değişkenler tarafından alınan değerlerden bağımsız olmalıdır.

Örneğin, İki değişkene sahip bir veri setimiz olsun: değişkenlerin biri kişinin adı ve diğeri kişinin cinsiyeti. Diyelim ki bir gözlem için adı Kübra olsun ama cinsiyet değeri belirtilmemiş. “Kübra” adının değerinden onun bir kadın olduğu sonucunu çıkarabiliriz. Bu durumda, açıklayıcı değişken Cinsiyet MCAR değildir.

Bir X değişkeninin eksik verilerinin Rastgele Eksik (MAR) olduğunu söyleyebilmemiz için ancak ve ancak X değişkenin değerinin eksik olma olasılığının başka değerleri bildirilen Y açıklayıcı değişkenlerin varlığına bağlı olup fakat X değişkenine bağlı olmamalıdır.

Örneğin bir gözlemin “Cinsiyet” değişkeninin değeri “Kadın” olup “Yaş” değişkeninin Eksikliğine MAR (Missing At Random) olduğunu söyleyebiliriz çünkü kadınların yaşlarını gizleme olasılığı daha yüksektir. Yani Yaş değerinin eksik olma olasılığı “Cinsiyet” değişkenine bağlı olup “Yaş” değişkenine bağlı değildir.

Rastgele Eksik Olmayan (NMAR, Not Missing At Random), Açıklayıcı bir değişkenin bir değerinin eksik olma veya eksik olmama olasılığı, yalnızca kendisine bağlıdır ve diğer açıklayıcı değişkenlerin değerlerinden herhangi biri ile ilişkili değildir. Başka bir deyişle, bir değerin yokluğunu motive eden bir şablon vardır.

Reha Alpar’a göre (Alpar, 2017) Eksik veriler MCAR (Missing Completely At Random) olunca uygunluk durumuna göre bilinen bir atama(Imputation) yöntemi kullanarak eksik değerleri doldurulabilir.

Aşağıdaki SPSS yazılımından alınmış test sonucuna göre, çalışmamızdaki eksik verilerin MCAR (H_0 : Eksik veriler rastgele dağılmaktadır, $P \text{ val.} > 0.05$, H_0 reddedilemez.) olduğu söylenebilir. Bu sonuç göz önüne alınarak veri setimizdeki eksik değerler için Median Imputation yöntemi ile atama yapılmıştır.

Tablo 2: MCAR Rastgelelik için Test sonucu

a. Little's MCAR test: Chi-Square = 12,061, DF = 271, Sig. = 1,000
--

3.2.2. AYKIRI DEĞERLER YÖNETİMİ

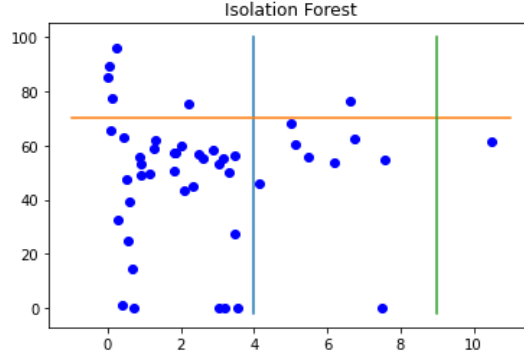
Aykırı değerler, veri setimizdeki olağandışı değerlerdir. Diğer veri noktalarından önemli ölçüde farklıdırlar ve analizimizi bozabilir, varsayımları ihlal edebilirler (Osborne & Overbay, 2004). Bunları analizden çıkarmak öznel bir uygulamadır ve neyi analiz etmeye çalıştığımıza bağlıdır. Genel olarak konuşursak, istenmeyen aykırı değerlerin kaldırılması, birlikte çalıştığımız verilerin performansını artırmaya yardımcı olabilir.

3.2.2.1.ISOLATION FOREST (YALITIM ORMANI) ALGORİTMASI

En yaygın olarak kullanılan aykırılık tespit yöntemleri, önce “normal” olarak adlandırılan verilerin bir profilini oluşturarak ve ardından bu normallikten sapan tüm gözlemleri aykırı değerler olarak sınıflandırarak çalışır. Isolation Forest algoritması farklı şekilde çalışır. Bu algoritma, büyük miktarda veri ile çalışma söz konusu olduğunda ve ayrıca değişken sayısı çok fazla olduğunda diğer yöntemlerden daha iyi performans sunar (Liu, Ting, & Zhou, 2008).

Bir anormallik tespit algoritmasının temel amacı, diğer verilere uymayan atipik verileri tanımlamaktır. Sorun basit değil çünkü bir aykırılığı neyin karakterize ettiğini önceden bilmemiz zordur. Daha sonra, bunları veriler içinde tespit etmek için uygun bir metriği öğrenmek algoritmaya bağlıdır.

Isolation Forest algoritması, Tony Liu, Kai Ming Ting ve Zhi-houa Zhou tarafından geliştirilmiştir (Liu, Ting, & Zhou, 2008). Veri kümesindeki her gözlem için bir anormallik puanı hesaplar. Bu puan, veri setine dayalı olarak her gözlemin normalliğinin bir ölçüsünü verir. Algoritma, bu puanı hesaplamak için söz konusu verileri özyinelemeli bir şekilde yalıtır: rastgele bir değişken seçer ve rastgele bir eşik belirler, ardından bunun belirli bir gözlemi izole etmeyi mümkün kılıp kılmadığını değerlendirir.



Şekil 1: Isolation Forest Algoritması

Hata! Başvuru kaynağı bulunamadı.'deki grafik aykırı değerlerini bulamak istediğimiz veriler olsun. Önce birinci değişkende rastgele bir dikey doğru çizilip (şekildeki mavi doğru) izole noktalar ortaya çıkıp çıkmadığına bakılır (Bizim durumumuzda izole nokta ortaya çıkmadı). Ardından rastgele ikinci bir doğru, bu sefer yatay bir doğru(sarı), çizilip tekrar izole noktaları kontrol edilir. Bu sefer birinci aykırı değerimiz ortaya çıktı (Grafikte, orta üst köşede). Aynı şekilde 3. Bir dikey doğru çizilip izole noktalara bakılır, bu sefer de ikinci izole bir nokta ortaya çıkmıştır(en sağda). Kesmeler ne kadar az ise o kadar da bulunmuş izole noktaların aykırı olma ihtimali yüksek olur.

Rastgele bölme işlemiyle, gözlemlerden birinin veri yığınınından ("normal" veri) yanlış bir şekilde izole edilmesinin mümkün olduğuna dikkat edilmelidir. Bu çok düşük bir ihtimal ama olabilir. Normal bir gözlem yanlış bir şekilde izole etme riskinin üstesinden gelmek için çözüm, tahmin ediciler olarak birkaç karar ağacı oluşturmaktan ibarettir. Bu ağaçların her biri bir dizi rastgele bölmeler gerçekleştirecektir. Ardından, tüm sonuçların ortalaması dikkate alınır. Çalışmamızdaki 2018-2020 yılı verilerine Isolation Forest algoritması uygulandığında çıkan aykırı bankalar Tablo 3' de verilmektedir.

Tablo 3: Aykırı Gözlemler

Adabank A.Ş.
Bank of America Yatırım Bank A.Ş.
Bank of China Turkey A.Ş.
JPMorgan Chase Bank N.A.
Standard Chartered Yatırım Bankası Türk A.Ş.
Türk Eximbank

Bir analizde veriler yanlış girilmişse veya ortak bir olgunun izole aykırı değerleriyse, aykırı değerler verilerden çıkarılabilir (Dhaka, 2017). Çalışmamızdaki 2018-2020 senelerin verilerine Isolation Forest algoritması uygulandığında çıkan aykırı bankalar bu kategoriye düşmediği için analizden çıkarılmamıştır. Aykırı olarak gözlenmiş bankaların 5'i büyük sermayeli yabancı bankalar (Adabank dâhil çünkü Kuveytli şirket The International Investor Company'na aittir) ve kalan Türk Eximbank sadece ihracat işleriyle ilgilendiği için aykırılık göstermesi normaldir. Diğer yıllar için de inceleme yapıldığında aynı bankaların aykırı değer olarak tespit edildiğini gözlemledik. Aykırı bankaların dâhil edilmediği bir senaryo ile çalışma tekrarlanabilir.

3.2.3. VERİ STANDARTLAŞTIRMASI

Standartlaştırma, bir değişkenin verilerini -1,0 ila 1,0 veya 0,0 ila 1,0 gibi daha küçük bir aralığa düşecek şekilde ölçeklendirmek için kullanılır. Genellikle sınıflandırma algoritmaları için kullanışlıdır.

Standartlaştırma, genellikle farklı bir ölçekteki değişkenlerle uğraşırken gereklidir, aksi takdirde, daha büyük ölçek değerlerine sahip olan diğer değişkenler nedeniyle eşit derecede önemli bir değişkenin (daha düşük bir ölçekte) etkinliğinin azalmasına yol açabilir.

Basitçe söylemek gerekirse, birden fazla öznelik mevcut olduğunda ancak öznelikler farklı ölçeklerde değerlere sahip olduğunda, veri madenciliği veya makine öğrenme işlemleri yapılırken zayıf veri modellerine yol açabilir. Bu nedenle, tüm nitelikleri aynı ölçeğe getirmek için standartlaştırılırlar.

Çeşitli standartlaştırma yöntemler vardır. Bunların arasında en çok bilinen z-standartlaştırma yöntemidir.

3.2.3.1. Z-standartlaştırması

Aşağıdaki gibi n gözlemlili sayısal bir değişken düşünelim:

$$(x_1 x_2 \dots x_n), x_i \in R$$

Sınırlı sayıda reel değere sahip olduğumuz için, çeşitli istatistiksel bilgiler elde edebiliriz:

min, maks, ortalama ve standart sapma.

Z-standartlaştırması, değerlerin ortalamasını 0'a ve standart sapmasını 1'e getirmeyi amaçlayan bir dönüşümdür. Bir x değerinin z-standart skoru aşağıdaki formül ile hesaplanır.

$$x_{std} = \frac{x - \mu}{\sigma}$$

x_{std} : x'in standartlaştırılmış değeri

μ : Dağılımın ortalama değeri

σ : Dağılımın standart sapması

3.2.3.2. Min-Max standartlaştırması

Bu standartlaştırma yöntemindeki işlem yalnızca minimum ve maksimum değerlere ihtiyaç duyar. Bu yaklaşım sonucunda, değerler arasındaki mesafe korunarak, değişkenin tüm değerleri 0 ile 1 arasına getirilir. Bunu yapmak için aşağıdaki formül kullanılır:

$$x_{std} = \frac{x - x_{min}}{x_{max} - x_{min}} \in [0, 1]$$

Veriler çok değişkenli normal dağılıma sahipse değişkenler de tek tek normal dağılıma sahiptir ancak tersi her zaman doğru değildir (Alpar, 2017). Çalışmamızda Min-Max standartlaştırma yöntemini kullandık çünkü uyguladığımız Shapiro-Wilk normallik test sonucuna göre verilerimizin değişkenlerinin çoğu tek tek normal dağılmadığı gözlenmiştir. 63 değişken üzerinde sadece 4 tanesi normal dağılıma sahiptir.

Tablo 4: Normal Dağılan Değişkenler

BY3: stat=0.962, p=0.119
BY4: stat=0.962, p=0.119
BY7: stat=0.962, p=0.116
GG10: stat=0.957, p=0.080

3.3. HEDEF SONUÇ

Bizimki gibi bir problemi çözmek için var olan algoritmalarından bahsetmeden önce, bu çalışmada izlediğimiz amacı ve özellikle sonunda nasıl bir sonuç elde etmek istediğimizi hatırlamak istiyoruz. Elimizde bulunan verilerden yola çıkarak bir bankanın bir sonraki yıl başarısız olup olmayacağını tahmin edebilecek bir algoritma geliştirmek istiyoruz. Ek olarak, algoritma, listelenen tüm bankalar için bir sonraki yıl için performans sırasına göre sıralama yeteneğine sahip olmalıdır.

Öncelikle başarısız bankaları diğer bankalardan net bir şekilde ayıracak bir yöntem bulmaya çalışarak başlayacağız. Söz konusu yöntem, mümkün olduğunca, her başarılı bankanın tüm başarısız bankalarla karşılaştırılması arasındaki uzaklığı net bir şekilde görebilmemiz için bize bir ayrım sağlamalıdır. Bu uzaklık, yıl içinde hangi bankaların diğerlerine göre en iyi performansı gösterdiğini belirlememizi sağlayacaktır.

4. BÖLÜM: YÖNTEM VE ALGORİTMALAR

4.1. İSTATİSTİKSEL ÖĞRENME YÖNTEMLERİ

İstatistiksel öğrenme veya Makine öğrenmesi, insanların öğrenme şeklini taklit etmek için öncelikle veri ve algoritmaları kullanan ve doğruluğunu kademeli olarak artıran bir yapay zekâ (AI) ve bilgisayar bilimi dalıdır.

Makine öğrenme, büyüyen veri bilimi alanının önemli bir bileşenidir. İstatistiksel yöntemlerin kullanılması yoluyla, algoritmalar, veri madenciliği projeleri için kritik bilgileri ortaya çıkaran sınıflandırmalar veya tahminler gerçekleştirmek üzere eğitilir. Bu bilgi daha sonra uygulamalar ve işletmeler arasında kararları yönlendirir ve ideal olarak temel büyüme ölçümelerini etkiler.

Makine öğrenme algoritmalarını 2 kategoriye ayırabiliriz (Bi, Goodman, Kaminsky, & Lessler, 2019):

- Denetimli öğrenme (supervised learning),
- Denetimsiz öğrenme (unsupervised learning)

Ancak Yarı denetimli ve Pekiştirmeli öğrenme (Reinforcement learning) gibi yukarıda bahsettiğimiz iki çeşitten türetilmiş başka yöntemler de mevcuttur. (İnceişçi, 2021)

4.1.1. DENETİMLİ MAKİNE ÖĞRENMESİ

Denetimli öğrenmede, operatörler bilgisayara girdi ve çıktı örnekleri sunar, daha sonra bu çıktıları girdilere göre elde etmek için aktif olarak çözümler arar. Amaç, makinenin kural eşleme girdilerini ve çıktılarını öğrenmesidir. Daha somut olarak, denetimli makine öğrenmesinde, makineye sunulan veriler, araması gereken kalıpları belirtmek için çoğunlukla etiketlenir. Söz konusu verilerin, eğitimden sonra sistemin yapması gereken şekilde zaten sınıflandırılması nadir değildir.

Denetimli öğrenme, daha az eğitim verisi gerektirir. Hâlihazırda etiketlenmiş verilerle karşılaştırma nedeniyle süreç de kolaylaştırılmıştır. Özellikle, gelecekteki veriler üzerinde

tahminler yapmak için denetimli makine öğrenmesi kullanılır. Bu durumda, girdi değişkenlerinden gelen algoritma, kesin olarak tahmin edebilen bir fonksiyon geliştirir.

Öğrenme algoritmaları arasında başlıca iki tür algoritma vardır: sınıflandırma ve regresyon algoritmaları. Sınıflandırma, bir kategori biçimindeki çıktı değişkenleri için belirli bir örneğin sonucunu tahmin etmeye yardımcı olur. Öte yandan regresyon, çıktı değişkeni belirli bir gerçek değer olduğunda bir örneğin sonucunu tahmin etmek için kullanılır.

4.1.2. DENETİMSİZ MAKİNE ÖĞRENMESİ

Denetimsiz öğrenme veya özellik öğrenme durumunda, giriş verilerinin yapısını kendi başına belirlemekten algoritma sorumludur. Bu yaklaşım, insanların dikkatinden kolayca kaçabilecek ilişkileri belirlemeye yardımcı olur.

Denetimsiz öğrenmenin farklı türleri de vardır: kümeleme, ilişkilendirme, boyut indirgeme. Kümeleme, örneklerin homojen olarak farklı gruplar veya kümeler halinde gruplandırılmasını sağlar. Her birinin içerdiği nesneler, diğer grupların nesnelerinden daha çok birbirine benzer. Veri biliminde ilişkilendirme, büyük veri gruplarını tanımlayan kuralları belirlemek için kullanılır. Boyut azaltma, adından da anlaşılacağı gibi, önemli bilgilerin iletilmesini sağlarken bir veri kümesindeki değişkenlerin sayısını azaltmak için kullanılır.

4.2. MODEL DEĞERLENDİRME METRİKLERİ

Bu çalışmada kullandığımız algoritmalar, bir dizi açıklayıcı değişkeniniz olduğunda ve bir yanıt değişkenini iki veya daha fazla sınıfa sınıflandırmak istediğinizde kullanabileceğiniz bir yöntemdir. İkili sınıflandırma yapıldıktan sonra bazı gözlemler gerçekte başarısız olup model tahmininde başarılı çıkabilir ve bazıları gerçekte başarılı olup başarısız olarak kestirilebilir. Karşımıza çıkabilen durumları dikkate alarak modelimizin performansını ölçmek istediğimizde bilmemiz gereken değerleri listeleyelim:

tp: True positive (gerçekte pozitif olup model tarafından pozitif kestirilen gözlem sayısı)

tn: True negative (gerçekte negatif olup model tarafından negatif kestirilen gözlem sayısı)

fn: False negative (gerçekte pozitif olup model tarafından negatif kestirilen gözlem sayısı)

fp: False positive (gerçekte negatif olup model tarafından pozitif kestirilen gözlem sayısı)

Listenlenmiş bu 4 değerlerini kullanarak performans metrikleri hesaplanabilir.

4.2.1. Doğruluk (Accuracy)

Doğruluğun anlaşılması oldukça basit bir anlamı vardır ve basitçe modelimizin doğru olarak sınıflandırdığı gözlemlerin oranıdır.

$$\text{Doğruluk} = \frac{tp + tn}{fp + tp + fn + tn}$$

4.2.2. Kesinlik (Precision)

Kesinlik, aslında pozitif olan tüm tahmin edilen pozitif değerlerimizin oranını temsil eder. Farklı sınıflandırma modellerini karşılaştırırken, yanlış pozitiflerden kaçınmak istediğimizde kesinlik iyi bir önlemdir.

$$\text{Kesinlik} = \frac{tp}{tp + fp}$$

Örneğin, spam e-postaları tespit ederken, yüksek Kesinliğe sahip bir model büyük olasılıkla tercih edilir. İstenmeyen e-posta durumunda, e-posta kullanıcısı, spam olmayan önemli bir e-postayı kaçırmaktan (yanlış pozitif) ziyade ara sıra istenmeyen e-posta almayı (yanlış negatif) tercih eder.

4.2.3. Recall (Hatırlama)

Recall, bir modelin pozitif gözlemleri gerçekten pozitif olarak ne kadar iyi sınıflandırdığını temsil eder. Farklı sınıflandırma modellerini karşılaştırırken, yanlış negatiflerden kaçınmak istediğimizde recall iyi bir önlemdir.

$$\text{Recall} = \frac{tp}{tp + fn}$$

4.2.4. F1-skoru

Son olarak F1-skoru inceleyelim. Bu F1 skoru, kesinlik ve hatırlamanın harmonik ortalamasıdır ve şu şekilde tanımlanır:

$$F1 - skoru = 2 \left(\frac{Kesinlik * Recall}{Kesinlik + Recall} \right)$$

Modelleri karşılaştırırken, Kesinlik ve Recall arasında bir denge kurmak istediğimizde F1-skoru kullanışlıdır. Yani, hem yanlış pozitiflerden (spam e-posta sınıflandırmasında olduğu gibi) hem de yanlış negatiflerden (anormallik tespitinde olduğu gibi) kaçınmak istediğimizde f1-skoru metriği tercih edilmelidir. Dolayısıyla bu çalışmada performans değerlendirme ölçüsü olarak f1-skor metriğini tercih ettik.

4.3. EN ÇOK BİLİLEN BAZI MAKİNE ÖĞRENME ALGORİTMALARI

Makine öğrenmesinde çok çeşitli ihtiyaçları karşılamak için farklı algoritmalar tasarlanmıştır. Hepsine hâkim olmak veya hepsini anlamak kolay değilse de, klasik bir veri bilimi eğitiminde öğretilen ve en yaygın uygulanan ana konuları bilmek önemlidir. Hepsi yukarıda belirtilen algoritma kategorilerinden birine aittir.

Bu çalışmada, elimizdeki tüm değişkenleri kullanarak bazı makine öğrenme algoritmalarını uygulamakla başlayacağız. Bir sonraki bölümde, daha iyi bir sınıflandırma sonucu elde etmek için mevcut makine öğrenme algoritmalarını bir optimizasyon yöntemiyle birleştirmekten oluşacak algoritmamız uygulanıp karşılaştırmalar yapılacaktır.

Bu çalışmanın amacı 2018, 2019 ve 2020 yıllarına ait veriler kullanılarak 2021 yılı için bankaların başarısız olup olmadığını tahmin etmek olup, bu çalışmada gerçekleştirilen tüm algoritma uygulamaları bu amaç doğrultusunda yapılmıştır. Yani her bir algoritma için sunulan sonuçlar, 2018, 2019 ve 2020 yıllarına ait verileri kullanarak 2021 yılı için tahmini sınıflandırma sonuçlarıdır.

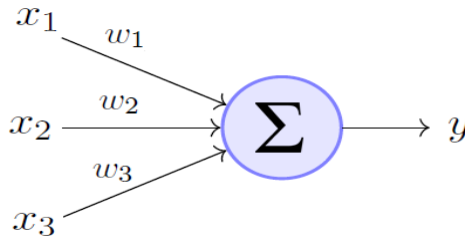
4.3.1. PERCEPTRON

Perceptron (Algılayıcı), denetimli Makine Öğrenme algoritmaları sınıfına aittir. Perceptron

algoritması 1950'lerin sonlarında Frank Rosenblatt tarafından (Rosenblatt, 1958) icat edildi ve doğrusal sınıflandırıcı (istatistiksel sınıflama için bir algoritma) olarak adlandırılır. Perceptron en eski makine öğrenme algoritmalarından biridir. Bir Perceptron, verileri yalnızca iki farklı sınıfa ayırmayı mümkün kılar. Koşul, bu verilerin doğrusal olarak ayrılabilir olması gerektiğidir (sorunu görüntüleyebilmek için bu, onları düz bir çizgi, bir düzlem vb. gibi doğrusal bir nesneyle ayırabilmemiz gerektiği anlamına gelir).

Perceptron, biyolojik nöronun bir taklidi olarak tasarlanmıştır. İnsan beyni bilgisinin günümüzdeki kadar kapsamlı olmadığı bir zamanda geliştirilmiş olup, oldukça basit ilkeler üzerine inşa edilmiştir.

Veri biliminde ve Makine Öğrenme projelerinde Perceptron önemli bir rol oynar. Normal ağırlık katsayılarını otomatik olarak öğrenmesini sağlayan Perceptron Öğrenme Kuralına göre çalışır (Noriega, 2005). Birkaç giriş sinyali alır ve bunların toplamı belirli bir eşiği aşıyorsa bir sonuç gönderilir veya aşmıyorsa başka bir sonuç gönderilir. Perceptron'nun diyagramı



Şekil 2: Üç girişli Perceptron

x_1, x_2, x_3 : Giriş sinyalleri

w_1, w_2, w_3 : Girişlerin ağırlıkları

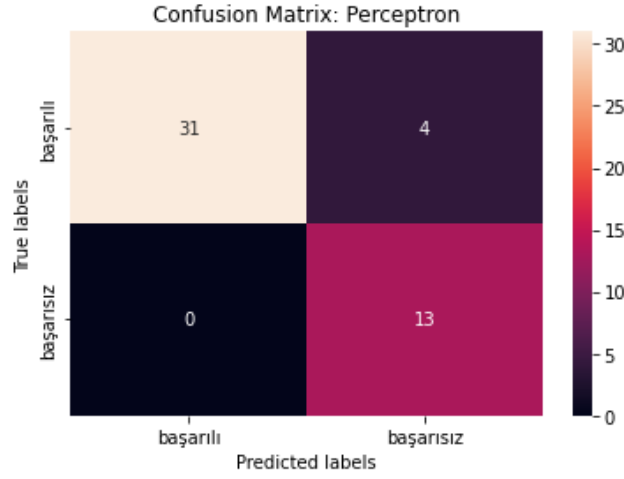
y : Çıkış değeri

X'ler girdi olarak, ardından bir çıktı değeri üretmekten sorumlu bir etkinleştirme fonksiyonu buluruz. Girişler ve fonksiyon arasındaki bağlantılar yukarıdaki grafikte w 'ler ile gösterilen ağırlıklardır. Bu ağırlıklar aslında sözde sinaptik bağlantıların bir analogisidir. Aktivasyon fonksiyonu, geçici bir girdi değeri hesaplamak için bu çarpılan ağırlıkları karşılık gelen girdilere ekler. Bu değer, bu geçici toplamın değerine bağlı olarak sinyali "1" veya "0"a çevirecek bir aktivasyon fonksiyonundan geçer. Hesaplanan çıktı ile beklenen çıktı arasında

bir fark varsa, bu hata payını kullanarak ağırlıkların değerini güncellenir, ardından tatmin edici bir sonuç elde edilene kadar çalışma yeniden başlar. Algılayıcıda neler olduğunu anlamak için, aşağıda verilen eğitim prosedürünün detayı incelenebilir:

- I. Ağırlıklar rastgele değerlerle (genellikle 0 ile 1 arasında) oluşturulur. Girdi başına değişken sayısı kadar ağırlık vardır. Örneğin, her biri iki değişkenli üç girdimiz varsa, iki ağırlık olacaktır.
- II. Eğitim için kullanılacak girdilerin çıktı değerlerini kayıt altına alıyoruz. İlk geçişte, ağırlıkların çarpımının sonuçlarını girdilerin değerleriyle ekliyoruz.
- III. Bu hesaplamanın sonucunu 1 veya 0 değerine dönüştürmek için kullanılacak olan bir aktivasyon fonksiyonuna geçer.
- IV. Perceptron tarafından hesaplanan bu çıktı değerinin, söz konusu girdinin gerçek çıktı değerine karşılık geldiğini doğrularız. Eğer öyleyse, ağırlık değerleri korunur. Değilse, hata payının bir kısmını onlardan kaldırarak ağırlıkların değerini güncelleriz.
- V. Tüm girdiler için tatmin edici bir sonuç elde edilene kadar süreç tekrarlanır.

2018, 2019 ve 2020 yıllarına ait verilerimize Perceptron algoritmasını uygulayarak 2021 yılı için her bankanın hangi sınıfa ait olacağını tahmin etmeye çalıştığımızda, elde edilen sonuç aşağıdaki gibidir.



	0	1	accuracy
precision	1	0,76	0,92
recall	0,89	1	0,92
f1-score	0,94	0,87	0,92
support	35	13	0,92

Şekil 3: Perceptron'un Sınıflandırma Sonuçları

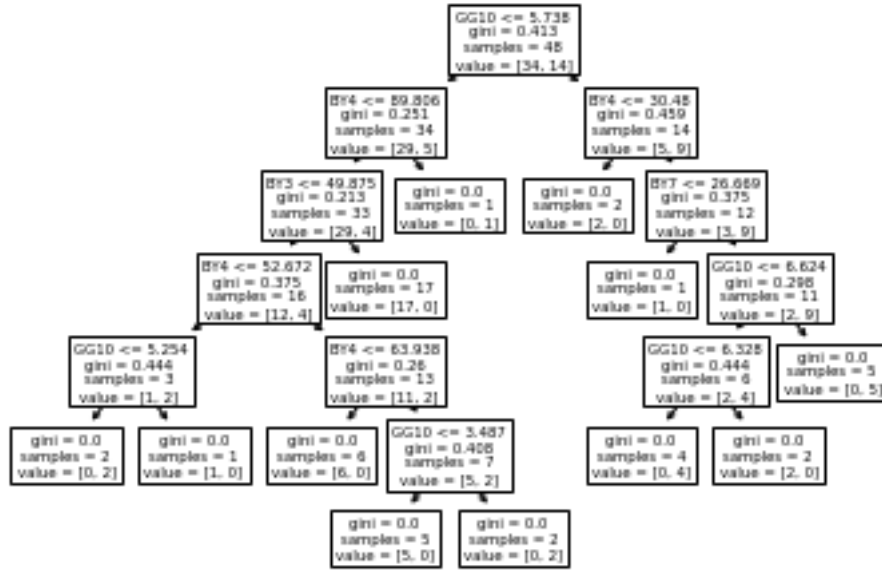
Şekil 3'teki çıktılara göre, algoritma, 13 başarısız gözlemde hepsini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 4 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.87** verdi. Modelin genel sınıflandırma doğruluğu 0.92'dir.

4.3.2. KARAR AĞACI

Karar ağacı algoritması, bir grafik modeline dayanan çok popüler bir algoritmadır. Ağaç, tüm gözlemleri sunan bir kökle başlar. Daha sonra kesişimlerine düğüm adı verilen bir dizi dalı ayırt ederiz. Bunlar, tahmin edilecek farklı sınıflara karşılık gelen yapraklarla sona erer. Her ağaç, bir yaprağa ulaşmadan önce karşılaşılan maksimum düğüm sayısının derinliği ile karakterize edilir. Bir karar ağacında, veriler iki veya daha fazla homojen kümeye bölünür. Benzer şekilde, tüm düğümlerin bir koşulu vardır. Bu şekilde, ne kadar aşağı inerseniz, koşullar o kadar çok olur. Problemimiz için değişken sayısının nispeten yüksek olması, algoritmanın zaman ve kaynak açısından çok maliyetli olması nedeniyle bu algoritmanın kullanılmasına büyük bir engel teşkil etmektedir.

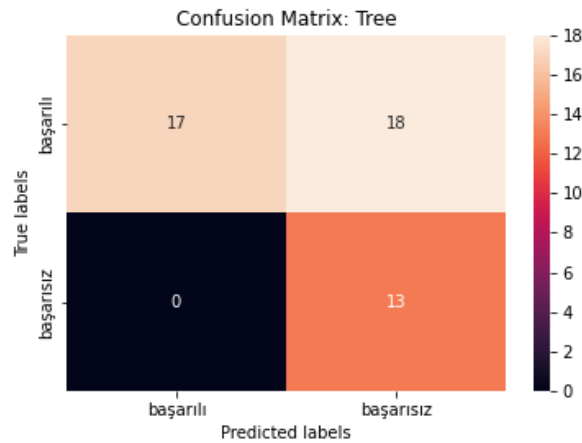
Aşağıdaki Şekil 4'e bakalım, bu algoritmanın sadece 4 değişken (BY3, BY4, BY7, GG10)

üzerinde bir uygulamasıdır, ancak modelin karmaşıklığı zaten kafa karıştırıcıdır. Bu modeli 63 değişkenin tamamıyla uygulamak, bundan daha da okunamaz bir sonuç üreteceği aşikârdır.



Şekil 4: 4 Değişkenli Karar Ağacının Çıktısı

Yine de bütün 63 değişkene göre gözlemleri sınıflandırmaya çalıştık. Söz konusu sınıflandırmanın sonucu Şekil 5'te verilmektedir.



	0	1	accuracy
precision	1	0,42	0,63
recall	0,49	1	0,63
f1-score	0,65	0,59	0,63
support	35	13	0,63

Şekil 5: Karar Ağacının Sınıflandırma Sonuçları

Şekil 3'teki çıktılara göre, algoritma, 13 başarısız gözlemde hepsini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 18 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.59** verdi. Modelin genel sınıflandırma doğruluğu 0.63'dir.

4.4.LOJİSTİK REGRESYON

Lojistik regresyon algoritmaları, ikili sınıflandırma yapmak için çok kullanışlıdır. Girdi olarak nitel ve/veya ordinal tahmin değişkenlerini alır ve ardından sigmoid işlevini kullanarak çıktı değerinin olasılığını ölçerler. Daha somut olarak, ana değişken ile açıklayıcı değişkenler arasındaki ilişkiyi incelerler. Lojistik regresyon çok sınıflı bir sınıflandırma için kullanılabilir. Örneğin giysileri gömlek, pantolon, şort gibi üç olanakta sınıflandırmak istediğimizde de kullanabiliriz. Bir gözlemin hangi sınıfa(y'nin değeri) koyulacağını belirlemek için aşağıdaki lojistik regresyon denkleminde yararlanılır.

$$y = \frac{e^{b_0+b_1x}}{1 - e^{b_0+b_1x}}$$

x : Bağımsız değişken

b_1, b_{20} : Parametre veya ağırlıklar

y : Çıkış değeri

Tüm değişkenlerle bir Lojistik regresyon modeli kurduğumuzda modelin sınıflandırma sonuçları Şekil 6'da gösterilmiştir.



	0	1	accuracy
precision	0,83	0,75	0,81
recall	0,94	0,46	0,81
f1-score	0,88	0,57	0,81
support	35	13	0,81

Şekil 6: Lojistik Regresyonun Sınıflandırma Sonuçları

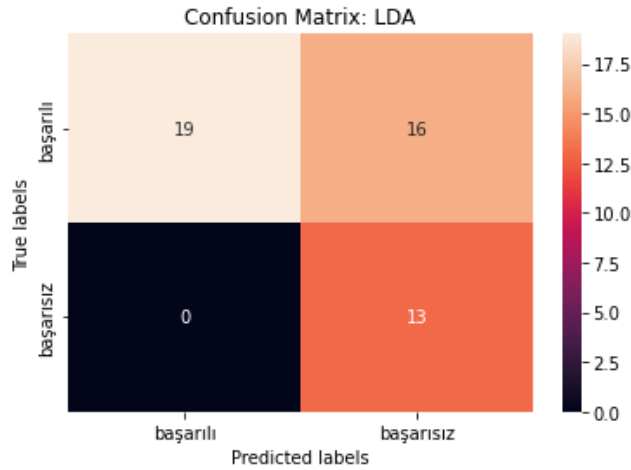
Şekil 3'teki çıktılarına göre, algoritma, 13 başarısız gözlemde 6 tanesini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 2 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.57** verdi. Modelin genel sınıflandırma doğruluğu 0.81'dir.

4.5.DOĞRUSAL DİSKRİMİNANT ANALİZİ

Doğrusal veya Lineer diskriminant analizi (LDA) veya normal diskriminant analizi veya diskriminant fonksiyon analizi, denetimli sınıflandırma problemleri için yaygın olarak kullanılan bir boyut indirgeme tekniğidir. Gruplardaki farklılıkları modellemek, yani iki veya daha fazla sınıfı ayırmak için kullanılır. Daha yüksek boyutlu uzaydaki özellikleri daha düşük boyutlu uzaya yansıtmak için kullanılır (Tang, Peng, Bi, Shan, & Hu, 2014).

LDA, tahmin edicilerin doğrusal kombinasyonlarını kullanarak verilen gözlemlerin sınıfını tahmin etmeye çalışır. Doğrusal diskriminant analizinin rolü, optimal olanı bulmak için birkaç olası ayırma eksenini kombinasyonunu test etmektir. 63 Değişkenlerimizle bir LDA modeli

kurduğumuzda elde edilen sonuç aşağıdaki gibidir.



	0	1	accuracy
precision	1	0,45	0,67
recall	0,54	1	0,67
f1-score	0,70	0,62	0,67
support	35	13	0,67

Şekil 7: LDA'nın Sınıflandırma Sonuçları

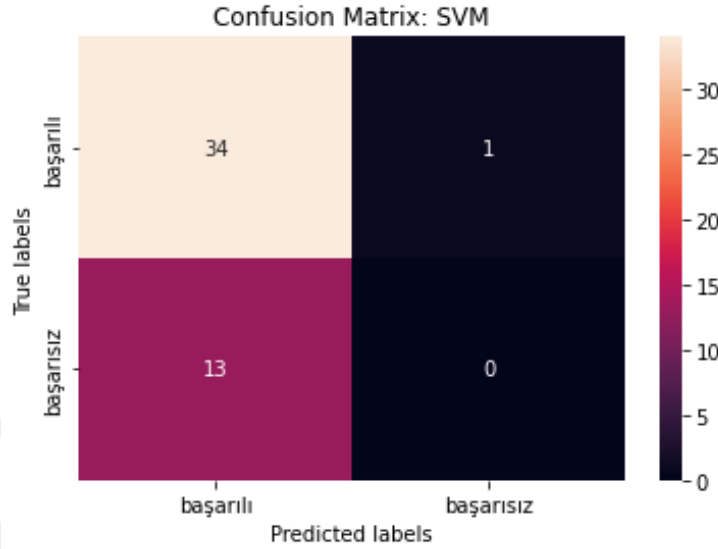
Şekil 3'teki çıktılarına göre, algoritma, 13 başarısız gözlemden hepsini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemden, algoritma 16 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.62** verdi. Modelin genel sınıflandırma doğruluğu 0.67'dir.

4.6.SUPPORT VECTOR MACHINES (SVM)

SVM, lojistik regresyon gibi bir ikili sınıflandırma algoritmasıdır. Bu iki algoritma, bir dizi öğeyi iki sınıfa ayırmayı mümkün kılarsa, destek vektör makinesi mümkün olan en net ayrımı seçer. Bunun için "maksimum marj" sayesinde veriler birkaç sınıfa ayrılır. Bu nedenle Large Margin Classifier (Büyük Marjlar sınıflandırıcısı) olarak adlandırılır (Xue, Chen, & Yang, 2011).

Bir kişinin etnik kökenini sınıflandırmak için iki kriterin vücuttaki melanin seviyesi ve burnunun uzunluğu olduğunu varsayarsak, bir grup insanı iki boyutta temsil edebilir ve onları

iki ana kategoriye ayırabiliriz: Afrikalılar ve Kafkasyalılar. SVM, böyle bir ayrımın elde edilmesini mümkün kılabilir. Benzer şekilde 63 değişkenli SVM modelin sonuçları aşağıda verilmiştir.



	0	1	accuracy
precision	0,72	0	0,71
recall	0,97	0	0,71
f1-score	0,83	0	0,71
support	35	13	0,71

Şekil 8: SVM'in Sınıflandırma Sonuçları

Şekil 3'teki çıktılarına göre, algoritma, 13 başarısız gözlemde hiçbirini doğru bir şekilde sınıflandırmadı. Öte yandan, 35 başarılı gözlemde, algoritma 1 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0** verdi. Modelin genel sınıflandırma doğruluğu 0.71'dir.

Bunlar en popüler makine öğrenme algoritmalarının bazı örnekleriydi ancak liste tam değil çünkü daha birçokları mevcuttur.

Bu çalışmadaki amaçlarımızdan biri, bir algoritmanın uygulamasında yer alan değişken sayısını azaltarak en iyi gözlem sınıflama tahmini bulmaktır. Boyut indirgeme yöntemleri arasında en çok bilinen Temel Bileşen Analizidir (PCA, Principal Components Analysis). Bu nedenle PCA algoritması ile gerçekleştirilen boyut indirgemesinden sonra elde edilen

bileşenleri kullanarak modeller oluşturduk.

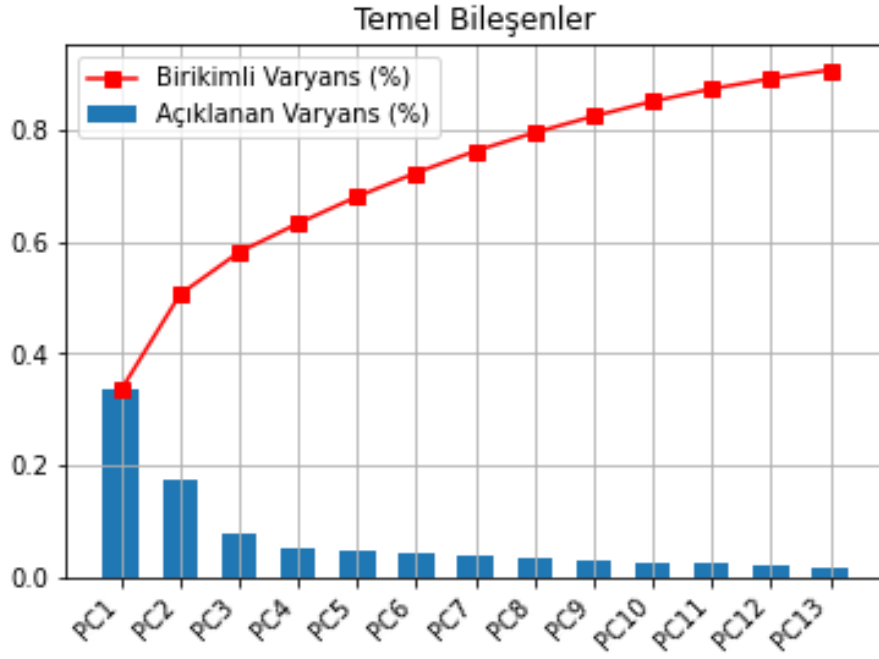
4.7.TEMEL BİLEŞENLER ANALİZİ

Temel bileşen analizi veya Principal Components Analysis (PCA), çok sayıda değişken tarafından tanımlanan gözlemleri içeren bir veri setini analiz etmeyi ve görselleştirmeyi mümkün kılar. Veri kümesinde 3'ten fazla değişkeniniz varsa, verileri çok boyutlu bir "hiper uzayda" görselleştirmek çok zor olabilir. Her değişken farklı bir boyut olarak kabul edilebilir. Bu durumda bir boyut indirgeme yaklaşımı olarak da kullanılan temel bileşenler analizinden yararlanılabilir.

Temel bileşen analizi, çok değişkenli bir veri tablosunda yer alan önemli bilgileri çıkarmak ve görselleştirmek için kullanılır. PCA, bu bilgiyi temel bileşenler adı verilen birkaç yeni değişkende sentezler. Bu yeni değişkenler, orijinal değişkenlerin doğrusal bir kombinasyonuna karşılık gelir. Temel bileşenlerin sayısı, orijinal değişkenlerin sayısından az veya ona eşittir.

Bir veri kümesinde yer alan bilgiler, içerdiği toplam varyansa karşılık gelir. PCA'nın amacı, veri varyasyonunun maksimum olduğu yönleri (yani ana eksenler veya ana bileşenler) belirlemektir. Başka bir deyişle, PCA, çok değişkenli bir verinin boyutlarını mümkün olduğunca az bilgi kaybederek grafiksel olarak görselleştirilebilen iki veya üç ana bileşene indirger.

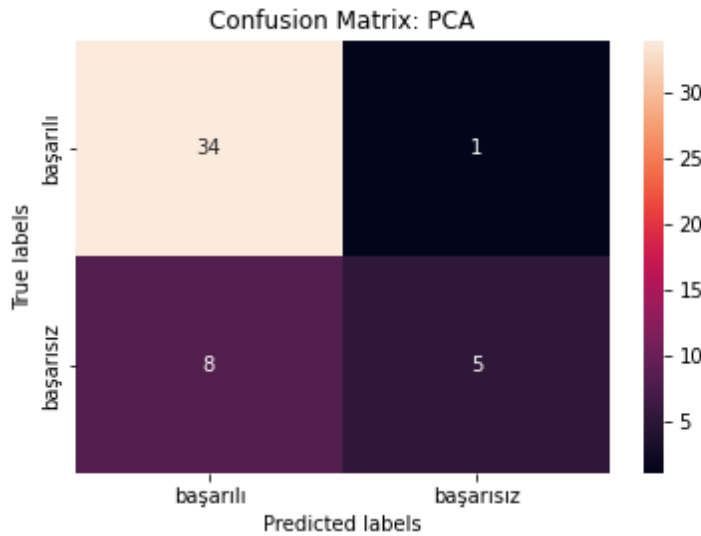
Problemimiz için temel bileşen analiz algoritmasını tek başına kullanmamız mümkün değildir. Temel bileşenlerin yanı sıra orijinal verilerimizin dönüştürülmesinden elde edilen skorları da bulduktan sonra, biraz önce belirtilen makine öğrenme yöntemlerinden bir yöntemi yardımıyla ve elde edilen temel bileşen skorları kullanarak bir model kurulabilir.



Şekil 9: Yüzde 90 Varyansı Açıklayan Bileşenler

Şekil 9'deki grafik, verilerimizdeki toplam varyansın %90'ını oluşturan temel bileşenleri göstermektedir. Bu varyans yüzdesini açıklayan toplam 13 temel bileşen vardır. Orijinal verilerimizden giderek temel bileşen skorları hesaplamak için bu bileşenleri kullanırız. Skorlar elde edildikten sonra çeşitli makine öğrenme algoritmaları kullanarak modelleri kurup karşılaştırma yapılacaktır.

4.7.1. TEMEL BİLEŞENLER ANALİZİNE DAYALI PERCEPTRON

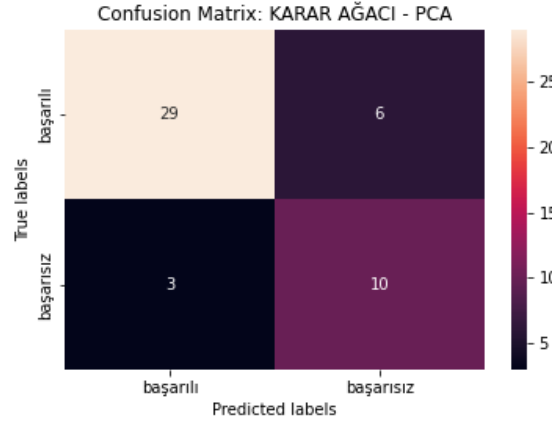


	0	1	accuracy
precision	0,81	0,83	0,81
recall	0,97	0,38	0,81
f1-score	0,88	0,53	0,81
support	35	13	0,81

Şekil 10: PCA Dayalı Perceptron

Şekil 10'teki çıktılara göre, algoritma, 13 başarısız gözlemde 5 tanesini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 1 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru 0.53 verdi. Modelin genel sınıflandırma doğruluğu 0.81'dir.

4.7.2. TEMEL BİLEŞENLER ANALİZİNE DAYALI KARAR AĞACI

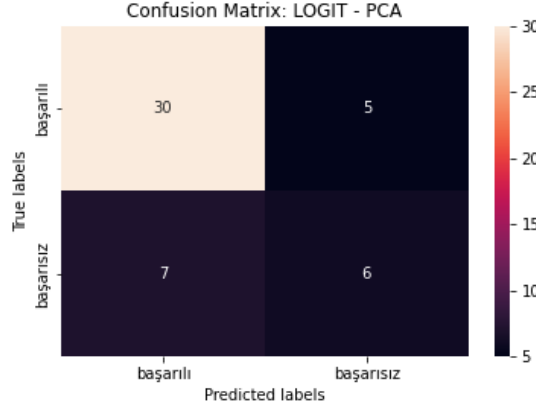


	0	1	accuracy
precision	0,91	0,63	0,81
recall	0,83	0,77	0,81
f1-score	0,87	0,69	0,81
support	35	13	0,81

Şekil 11: PCA Dayalı Karar Ağacı

Şekil 3'teki çıktılara göre, algoritma, 13 başarısız gözlemde 10 tanesini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 6 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.69** verdi. Modelin genel sınıflandırma doğruluğu 0.81'dir.

4.7.3. TEMEL BİLEŞENLER ANALİZİNE DAYALI LOJİSTİK REGRESYON

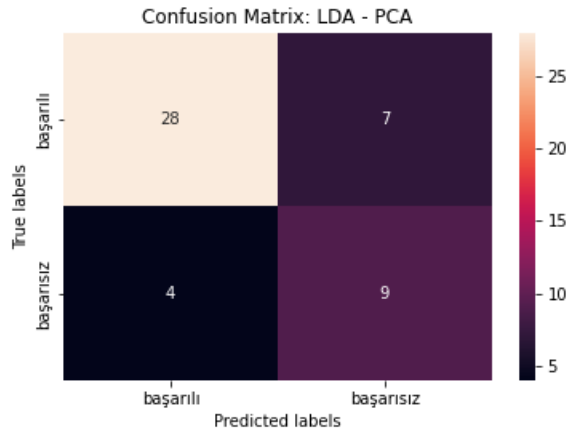


	0	1	accuracy
precision	0,81	0,55	0,75
recall	0,86	0,46	0,75
f1-score	0,83	0,50	0,75
support	35	13	0,75

Şekil 12: PCA Dayalı Lojistik Regresyon

Şekil 12’teki çıktılarına göre, algoritma, 13 başarısız gözlemde 6 tanesini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 5 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru 0.50 verdi. Modelin genel sınıflandırma doğruluğu 0.75’dir.

4.7.4. EMEL BİLEŞENLER ANALİZİNE DAYALI DİSKRİMİNANT ANALİZİ

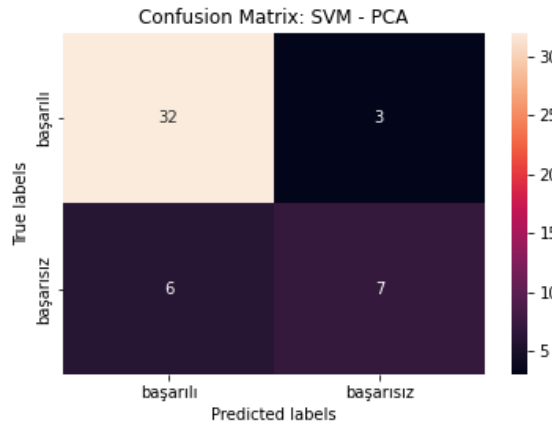


	0	1	accuracy
precision	0,88	0,56	0,77
recall	0,80	0,69	0,77
f1-score	0,84	0,62	0,77
support	35	13	0,77

Şekil 13: PCA Dayalı LDA

Şekil 3'teki çıktılarına göre, algoritma, 13 başarısız gözlemde 9 tanesini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 7 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.62** verdi. Modelin genel sınıflandırma doğruluğu 0.77'dir.

4.7.5. EMEL BİLEŞENLER ANALİZİNE DAYALI SVM



	0	1	accuracy
precision	0,84	0,7	0,81
recall	0,91	0,54	0,81
f1-score	0,88	0,61	0,81
support	35	13	0,81

Şekil 14: PCA Dayalı SVM

Şekil 3'teki çıktılarına göre, algoritma, 13 başarısız gözlemde 7 tanesini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 3 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.61** verdi. Modelin genel sınıflandırma doğruluğu 0.81'dir.

5. BÖLÜM: ALGORİTMAMIZ

Algoritma, belirli bir sorunu çözmek için izlenecek bir dizi adımdır. Algoritmamızdan bahsetmeden önce, takip ettiğimiz amacı ve sahip olduğumuz araçları hatırlayarak başlayalım.

Verileri temizledikten ve kullanılacak değişkenleri seçtikten sonra, bu seviyede elimizde Türkiye'de faaliyet gösteren 48 bankanın yıllık finansal tablo verileri bulunmaktadır. Bu bankaların 13 tanesi 2020 yılına kadar son üç sene üst üste zarar etmiş başarısız sayılmış bankalardır ve kalan 35 banka başarılı sayılmış bankalardır. Test seti olarak kullandığımız 2021 senesi veri setinde is 14 tane başarısız ve 34 tane başarılı banka bulunmaktadır. Bu veri kullanılarak bir yılda finansal performanslarına göre bankaları en iyiden en kötüye olacak şekilde sıralamak (Derecelendirmek) istiyoruz.

5.1. BANKALARI BAŞARILI-BAŞARISIZ ŞEKLİNDE AYIRACAK ALGORİTMANIN SEÇİMİ

Bu tezde bankaların başarılı-başarısız biçiminde sınıflandırılması amacıyla çeşitli istatistiksel ve makine öğrenme yöntemlerinden yararlanılmaktadır. göz önünde bulundurduğumuz yöntemler literatürde sıklıkla kullanılan sınıflandırma yöntemlerinden perceptron, karar ağacı, lojistik regresyon, lineer diskriminant analizi, destek vektör makineleridir. Bunun yanında, değişken seçimine alternatif olarak, boyut indirgeme yöntemi olarak Temel Bileşenleri analizinden yararlanılmaktadır.

Makine öğrenmesindeki sınıflandırma algoritmaları arasında bu iki banka kategorisini en iyi ayıracak bir algoritma bulma sorunudur. Seçebileceğimiz algoritmalar arasında Perceptron, Karar Ağacı, lojistik regresyon, lineer diskriminant analizi, destek vektör makineleri, K-means veya Temel Bileşenleri bulunmaktadır.

5.2.OPTİMAL DEĞİŞKEN SAYISI BELİRLEME

Daha önce yapılmış çalışmalardaki modellerde kullanılmış değişkenlerin seçimi ile ilgili kurallar pek yapılmamıştır. Ekici ve Erdal'ın (Ekinci & Erdal, 2017) çalışmasında Bankaların başarısızlığını tahmin etmek için değişken olarak 35 tane finansal oran kullanılmıştır.

Türkcan, Bozcuk ve Kemal Türkcan (Türkcan, Bozcuk, & Türkcan, 2018) bankaların mali başarısızlığını tahmin etmek için hem finansal hem de finansal olmayan bir sürü değişkeni modeline katmışlardır.

Bu çalışmada, birkaç sınıflandırma algoritması ile modeller oluşturduk, ancak her seferinde elimizdeki tüm değişkenleri modele dâhil ettik. Elde edilen sonuçlar doğruluk açısından başarılı olmadığından dolayı en iyi sonuçları elde edebilmek için sadece faydalı değişkenleri tutmak istedik.

Modeller için gerçekten faydalı olan değişkenleri belirlemek için çeşitli yöntemler geliştirilmiştir. Kumar ve Minz (Kumar & Minz, 2014) değişkenlerin bir modelde nasıl seçildiğine dair kapsamlı bir açıklama vermek için literatürü taramışlardır. James, Witten, Hastie ve Tibshirani (James, Witten, Hastie, & Tibshirani, 2013) tarafından ve literatürdeki diğer istatistiksel kaynaklarda sıklıkla önerilen Forward Stepwise (Adım Adım İleriye) ve Backward Stepwise (Adım Adım Geriye) gibi değişkenleri seçme yöntemlerinden yararlanabilir.

Geldiğimiz noktada elimizde 63 açıklayıcı değişken var. Amacımız, modelin başarısız bankalar sınıfını başarılı bankalardan optimal bir şekilde ayırmasını sağlayacak değişkenleri içeren bir model oluşturmaktır. Bunun için değişkenlerin alt kümeleri oluşturulabilir, ardından her alt küme için bir model oluşturularak ve sonra yalnızca iyi performans gösterenleri tutmak için modelleri karşılaştırabiliriz. Mesela rastgele sekizer değişkenli alt kümeleri oluşturalım:

SY1	SY2	SY3	SY4	SY5	SY6	SY7	BY1
BY2	BY3	BY4	BY5	BY6	BY7	BY8	BY9
AK1	AK2	AK3	AK4	AK5	AK6	L1	L2
L3	L4	L5	K1	K2	K3	K4	GG1
GG2	GG3	GG4	GG5	GG6	GG7	GG8	GG9
GG10	GG11	GG12	SP1	SP2	SP3	GP1	GP2
GP3	SR1	SR2	SR3	SR4	SR5	SR6	SR7
FR1	FR2	FR3	FR4	FR5	FR6	FR7	

Bu yarattığımız alt kümelerle ilgili endişe, bu tür alt kümeleri oluşturmak için çok sayıda olası kombinasyonun olmasıdır. 63 değişkenden 8'şer değişkenli bütün alt kümelerin sayısı aşağıdaki formülle hesaplanabilir.

$$\binom{63}{8} = \frac{63!}{(63-8)!8!} = 3.872.894.697$$

Görüldüğü gibi, tüm bu kombinasyonları denemek zaman ve kaynak açısından çok maliyetli olacaktır. Dolayısıyla yaklaşık bir çözüme ulaşmak amacıyla bir optimizasyon algoritmasına ihtiyacımız vardır. Optimizasyon algoritmalar arasında problemimize en uygun Genetik Algoritmasının olduğunu düşündük.

5.3.GENETİK ALGORİTMASI

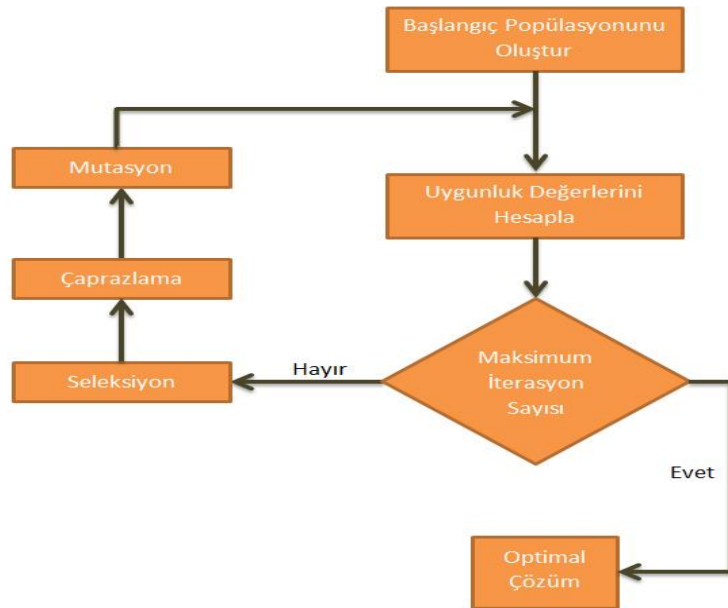
1859'da doğa bilimci Charles Darwin, evrim olgusunu açıklamaya yönelik bir teori sunan ünlü kitabı Türlerin Kökeni'ni yayınladı (Darwin, 1909). Bu devrimci teori, o zamanlar birkaç tartışmaya neden oldu, ancak şimdi yaygın olarak kabul ediliyor. Yaklaşık 100 yıl sonra, doğal seçim sürecine dayalı evrimsel sistemleri yapay olarak uygulamaya çalışan Profesör John Holland'a ilham verdi (Holland, 1975). Bu çalışma daha sonra bazı zor optimizasyon problemleri için yaklaşık çözümler elde etmek için kullanılan genetik algoritmalara yol açtı.

Doğal seleksiyon süreci ile genetik algoritma arasındaki analogiyi anlamak için önce evrimin temelindeki en önemli mekanizmaları hatırlayalım.

- 1) Evrim, bir popülasyondaki bireylerin her birini temsil eden kromozomlar üzerinde gerçekleşir.
- 2) Doğal seçim süreci, en uygun kromozomların daha sık üremesine ve gelecekteki popülasyonlara daha fazla katkıda bulunmasına neden olur.
- 3) Üreme sırasında, ebeveynlerin kromozomlarında bulunan bilgiler birleştirilir ve çocukların kromozomlarını üretmek için karıştırılır (“crossover”).
- 4) Geçişin sonucu, rastgele bozulmalar (mutasyonlar) ile değiştirilebilir.

Bu mekanizmalar, aşağıdaki adımları izleyerek yapılır.

- I. N çözümlerinden oluşan bir başlangıç popülasyonu oluşturulur.
- II. Popülasyondaki her bir bireyi değerlendirilir.
- III. Ebeveynleri puanlarına göre seçerek yeni çözümler üretilir. Üreme sırasında çaprazlama ve mutasyon operatörleri uygulanır.
- IV. N yeni birey üretildiğinde, eski popülasyonun yerini alırlar. Yeni popülasyonun bireyleri sırayla değerlendirilir.
- V. Tahsis edilen süre aşılmadıysa (veya maksimum nesil sayısına ulaşılmadıysa), 3. adıma dönülür.

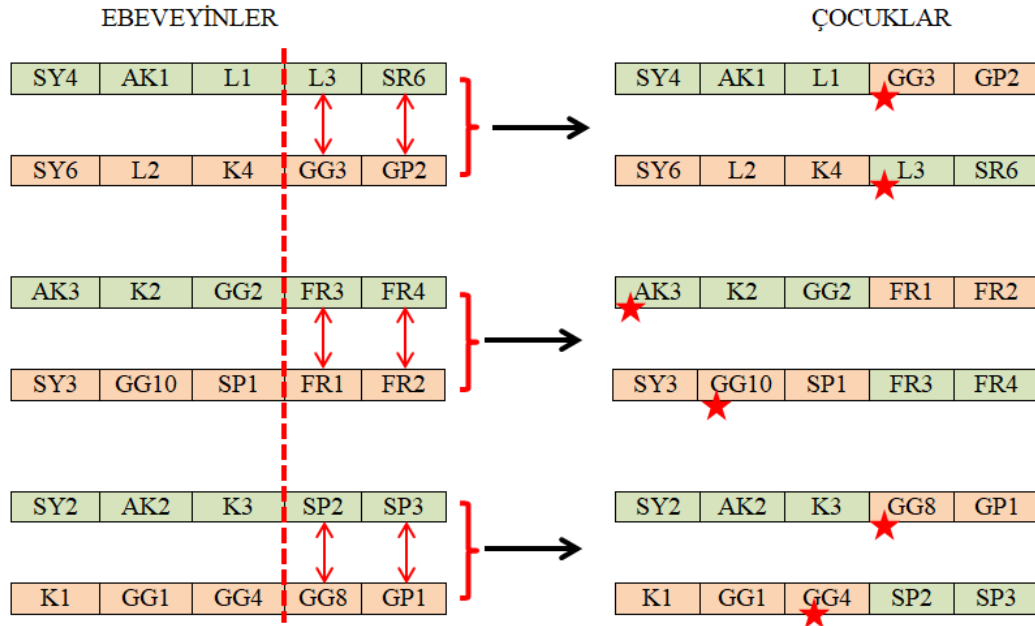


Şekil 15: Genetik Algoritmasının Akış Şeması

5.4.GENETİK ALGORİTMANIN UYGULAMASI

Örneğin 30 değişkenimiz var ve $n = 6$ (kromozom) çözümlerinden oluşan bir başlangıç

popülasyonu oluşturup ikişer kromozom oluşan çiftleri çıkaralım. Bu çiftler arasında çaprazlama yaparak yeni nesil kromozomlar üretilir.



Yeni üretilmiş bireylerde rastgele bir gende(kırmızı yıldızın işaretlediği) mutasyon yapılır yeni oluşan popülasyonumuzun her bir bireyi için ayrı bir model kurulursa modellerin f1 skorlarını elde edip değerlendirebiliriz.

Genetik Algoritma sürecindeki başlangıç kuşaktan f1-skoruna göre en iyi performans göstermiş popülasyonun yarısı sonraki aşamaya dahil edip diğer yarısı dışlıyoruz.

Benzer şekilde ikişer kromozomla oluşan çiftleri yeniden oluşturup yeni neslin üremesini sağlanır. Yeni popülasyon oluştuktan sonra değerlendirmeler yapılır en iyi bireyler kalır ve düşük performanslı bireyler süreçten atılır. Bu işlemleri 100. Kuşağa kadar devam ettirdikten sonra Algoritmamız durur.

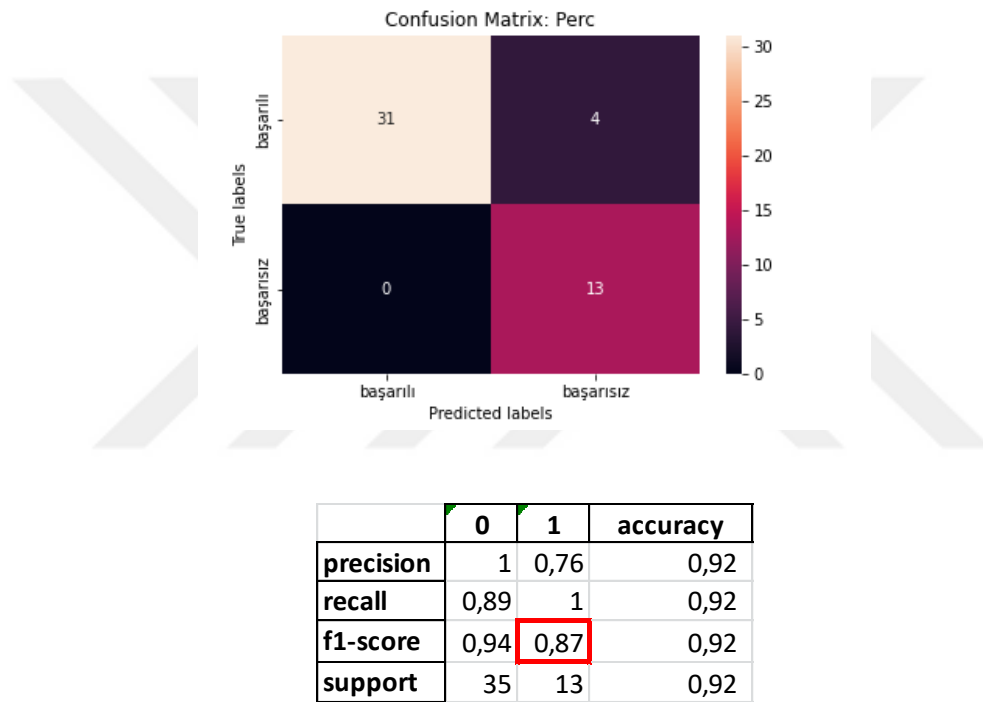
5.5. UYGULAMA

Bu çalışmada ele aldığımız algoritmalar arasında en iyi performansı üretecek sınıflandırma algoritmasını bulmak için her bir algoritmayı genetik algoritma ile birlikte çalıştırıp model için en kullanışlı değişkenleri bulmuştuk. Her bir algoritma için elde edilen sonuçlar aşağıdaki

gibidir:

Bu çalışmanın amacı 2018, 2019 ve 2020 yıllarına ait veriler kullanılarak 2021 yılı için bankaların başarısız olup olmadığını tahmin etmek olup, bu çalışmada gerçekleştirilen tüm algoritma uygulamaları bu amaç doğrultusunda yapılmıştır. Yani her bir algoritma için sunulan sonuçlar, 2018, 2019 ve 2020 yıllarına ait verileri kullanarak 2021 yılı için tahmini sınıflandırma sonuçlarıdır.

5.5.1. Perceptron



Şekil 16: Perceptron'un Optimal Sonucu

Bu sonuca varmak için Genetik algoritmanın seçtiği değişkenler:

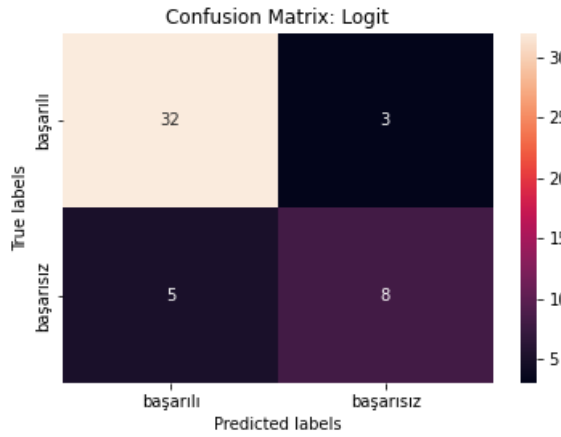
['BY2', 'AK2', 'AK5', 'L1', 'K1', 'K2', 'GG3', 'GG9', 'GG10', 'SR6', 'SR7', 'FR6', 'FR7']

Şekil 3'teki çıktılara göre, algoritma, 13 başarısız gözlemde hepsini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 4 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.87** verdi. Modelin genel sınıflandırma doğruluğu 0.92'dir.

5.5.2. Karar Ağacı

Karar ağacı için değişkenleri seçmeye çalışırken bir tutarsızlık sorunuyla karşılaştık. Normalde, genetik algoritma gibi bir optimizasyon algoritması, rastgele bir çözümden başlayıp, onu en uygun çözüme ulaşmak için geliştirerek ilerlerdir. Ancak bizim durumumuzda değişken seçim süreci, neredeyse her zaman %0 ile %15 arasında doğruluk oranıyla sonuçlar verdi. Son karşılaştırma için sadece tüm değişkenlerle oluşturulmuş karar ağacı modeli dikkate alacağız.

5.5.3. Lojistik Regresyon



	0	1	accuracy
precision	0,86	0,73	0,83
recall	0,91	0,62	0,83
f1-score	0,89	0,67	0,83
support	35	13	0,83

Şekil 17: Lojistik Regresyonun Optimal Sonucu

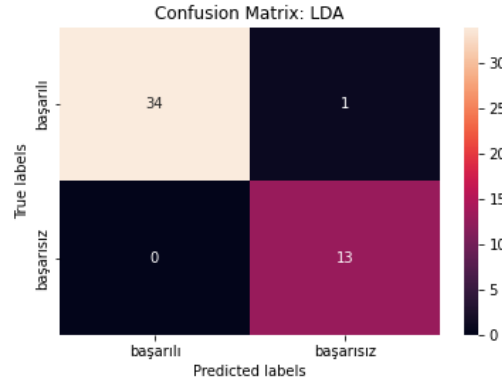
Şekil 3'teki çıktılarına göre, algoritma, 13 başarısız gözlemde 8 tanesini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 3 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.67** verdi. Modelin genel sınıflandırma doğruluğu 0.83'dir.

Seçilmiş değişkenler:

['SY2', 'SY4', 'BY2', 'BY5', 'BY9', 'AK1', 'AK3', 'AK6', 'L1', 'L3', 'K1', 'K2', 'K3', 'K4',

'GG1', 'GG2', 'GG3', 'GG5', 'GG6', 'GG8', 'GG9', 'GG11', 'SP1', 'SP2', 'SR2', 'SR3', 'SR5', 'SR7', 'FR1', 'FR2', 'FR5', 'FR6']

5.5.4. Lineer Diskriminant



	0	1	accuracy
precision	1	0,93	0,98
recall	0,97	1	0,98
f1-score	0,99	0,96	0,98
support	35	13	0,98

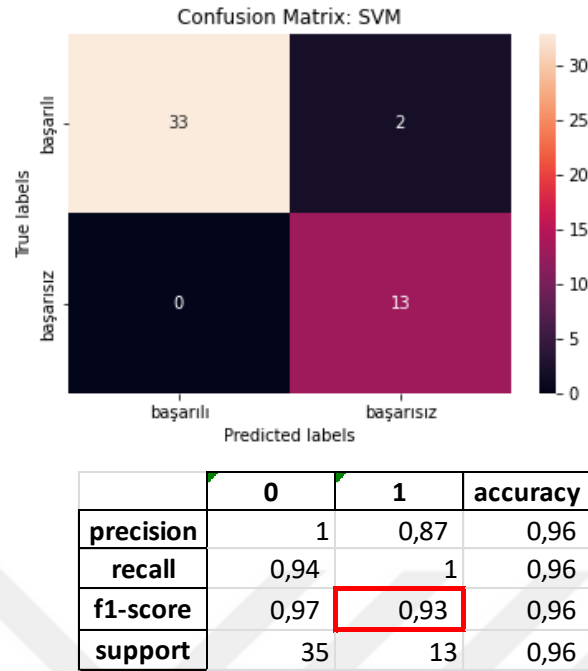
Şekil 18: LDA'nın Optimal Sonucu

Seçilmiş değişkenler:

['SY4', 'BY4', 'BY5', 'AK1', 'AK5', 'K2', 'GG3', 'GG4', 'GG11', 'SR1', 'SR3', 'SR6', 'FR4', 'FR6', 'FR7']

Şekil 3'teki çıktılarına göre, algoritma, 13 başarısız gözlemde hepsini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 1 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.96** verdi. Modelin genel sınıflandırma doğruluğu 0.98'dir.

5.5.5. SVM



Şekil 19: SVM'ın Optimal Sonucu

SVM'de kullanılmış değişkenler:

['SY4', 'SY7', 'BY6', 'AK3', 'AK5', 'L1', 'L3', 'L4', 'K1', 'K3', 'K4', 'GG8', 'SR4', 'SR7', 'FR7']

Şekil 3'teki çıktılarına göre, algoritma, 13 başarısız gözlemde hepsini doğru bir şekilde sınıflandırdı. Öte yandan, 35 başarılı gözlemde, algoritma 2 tanesini başarısız olarak sınıflandırdı ve bu da bize başarısız gözlemlerin sınıflandırılmasında F1-skoru **0.93** verdi. Modelin genel sınıflandırma doğruluğu 0.96'dir.

5.6. MODELLERİN KARŞILAŞTIRMASI

Aynı değişken seçim yöntemini Perceptron, Karar Ağacı, Lojistik Regresyon, Lineer Diskriminant, Destek Vektör Makinesi, K-ortalamlar ve Temel Bileşenler Analizine dayalı algoritmalarına uyguladık. Elde ettiğimiz eğitim (train) ve test setine ait sonuçlar aşağıdaki tabloda özetlenmiştir:

Tablo 5: Karşılaştırma Tablosu

	Train Set		Test Set	
	Accuracy (Doğruluk)	f1-score	Accuracy (Doğruluk)	f1-score
Perceptron	0,92	0,88	0,92	0,87
Karar Ağacı	0,69	0,65	0,94	0,89
Lojistik Regresyon	0,79	0,44	0,81	0,57
LDA	0,98	0,96	0,98	0,96
SVM	0,75	0,25	0,96	0,93
PCA-PERC	0,77	0,35	0,81	0,53
PCA-TREE	1	1	0,77	0,56
PCA-LOGIT	0,79	0,55	0,75	0,50
PCA-LDA	0,88	0,77	0,77	0,62
PCA-SVM	0,85	0,70	0,81	0,61

Bu seçenekler arasında favorimiz Lineer Diskriminant sınıflandırmasıdır çünkü Trevor Hastie'ye göre (James, Witten, Hastie, & Tibshirani, 2013):

1. Elimizde çok sayıda değişken içeren verilerimiz var ve tam olarak Lineer Diskriminant Analizi yöntemi boyut azaltılmasına dayanıyor.
2. Farklı sınıflar iyi ayrıldığında, lojistik regresyon modelinin parametre tahminleri şaşırtıcı bir şekilde dengesiz olabiliyor. Doğrusal Diskriminant analizinde bu sorun yaşanmamaktadır.
3. Gözlem sayısı n küçükse ve X bağımsız değişken dağılımı sınıfların her birinde yaklaşık olarak normal ise, gene Diskriminant Analizi modeli, lojistik regresyon modelinden daha istikrarlı sonuçlar vermektedir.

Ayrıca uygulamalarımın sonuçlarına göre LDA modelinin performansının daha iyi olduğu görülmüştür.

5.7.DERECELENDİRME

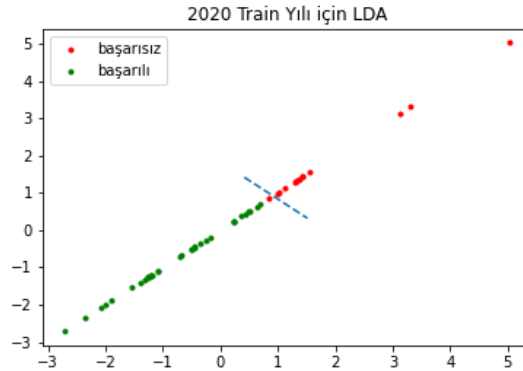
Tablo 5'e göre en iyi performansı sunan model LDA modelidir. Bu nedenle bankaların derecelendirmesini yapmak için LDA modelini seçmiştik. LDA modelinin train seti

üzerindeki performansı **Tablo 6**'da özetlenmiştir.

Tablo 6: Train Seti için LDA'nın Performansı

	0	1	accuracy
precision	0,97	1	0,98
recall	1	0,93	0,98
f1-score	0,99	0,96	0,98
support	34	14	0,98

Daha önce de belirttiğimiz gibi, Lineer Diskriminant sınıflandırması sadece gözlemleri sınıflandırmaya değil, aynı zamanda başlangıç verilerinden daha düşük boyutlardaki yüzeylere yansıtmaya da izin verir. Problemimiz ikili sınıflandırma bir problemi olduğu için LDA algoritması verilerimizi n-1 boyutlu (yani 1 boyutlu) bir uzaya yansıtmaya çalışır. Bir boyutlu uzay aslında bir eksene denk gelir ve bir eksen üzerinde gösterilmiş noktalar sıralanabilir. Gösteriş için **Şekil 20**'daki gibi iki boyutta çizmiştik. Kırmızı noktalar başarısız bankalar olup yeşil noktalar başarılı bankalardır.

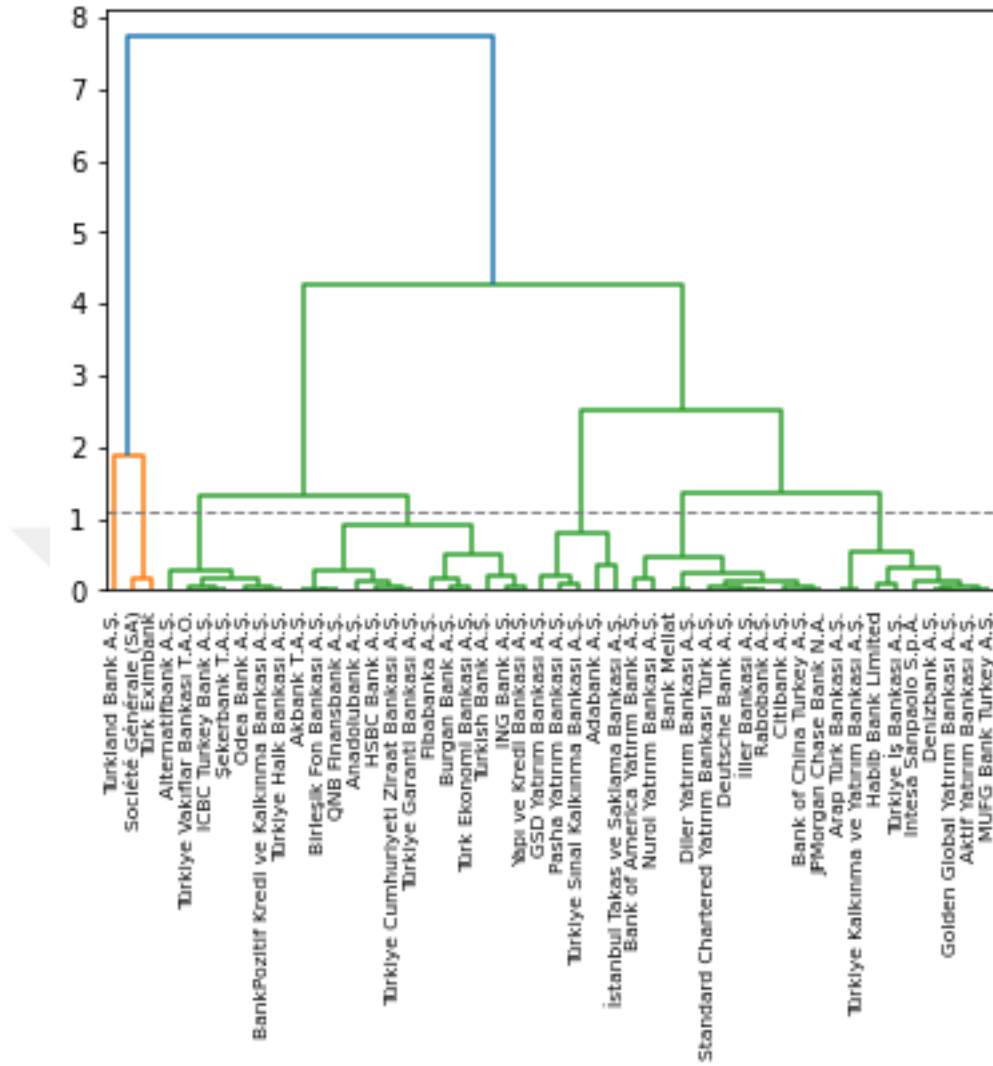


Şekil 20: Train Set için LDA'nın Gözlemleri Yansıttığı Eksen

Bankalara skor vermek için öncelikle bankanın verilerini modelin bulduğu eksene yansıtıyoruz. Bu eksende en az performans gösteren bankalar en yüksek değerlere sahip bankalar olurken, en iyi performansa sahip bankalar bu eksende en düşük değerlere sahiptir.

Train seti için LDA sayesinde elde edilen değerlere hiyerarşik kümeleme algoritması uygulayıp elimizde olan derecelendirme kuruluşu Fitch'in Türk bankalarına verdiği derece sınıfları gibi, bankaların performanslarının 7 sınıfa kümelenebileceğini çıkarmıştık. Hiyerarşik kümeleme sonucu **Şekil 21**'de görülmektedir. Nitekim elimizdeki Fitch'in 2019 sonunda Türk bankalarına verdiği derecelerin sınıfları şu şekildedir: BBB-, BB+, BB, BB-,

B+, B ve B-.



Şekil 21: Train Seti İçin Derece Kümeleri

Bu kümeleme algoritmasını 2021 test verilerine uygulayarak aşağıdaki Tablo 7'de gösterildiği gibi bankaları derecelendirebiliriz. Skor sütunu LDA analizi sonucunda elde edilen ayırma skorlarıdır ve bankalar bu ayırma skorlarına göre sıraya dizilmiştir. Başarısız bankalar 1 başarılı bankalar 0 olarak kodlandığı için en büyük negatif skor en başarılı bankayı en büyük pozitif skor ise en başarısız bankayı göstermektedir.

Tablo 7: 2021 Yılı için Bankaların Başarı Dereceleri

	Skor	Rate
Adabank A.Ş.	-7,65	BBB-
Bank of America Yatırım Bank A.Ş.	-2,45	BB+
İller Bankası A.Ş.	-2,28	BB+
Pasha Yatırım Bankası A.Ş.	-1,87	BB+
Rabobank A.Ş.	-1,77	BB+
Nurol Yatırım Bankası A.Ş.	-1,68	BB+
İstanbul Takas ve Saklama Bankası A.Ş.	-1,65	BB+
Citibank A.Ş.	-1,51	BB+
Diler Yatırım Bankası A.Ş.	-1,39	BB+
Golden Global Yatırım Bankası A.Ş.	-1,38	BB+
Standard Chartered Yatırım Bankası Türk A.Ş.	-1,36	BB+
Denizbank A.Ş.	-1,10	BB
GSD Yatırım Bankası A.Ş.	-0,89	BB
Deutsche Bank A.Ş.	-0,81	BB
Arap Türk Bankası A.Ş.	-0,70	BB
Türkiye Sınai Kalkınma Bankası A.Ş.	-0,61	BB
Aktif Yatırım Bankası A.Ş.	-0,44	BB
Türkiye İş Bankası A.Ş.	-0,26	BB
Bank of China Turkey A.Ş.	-0,17	BB
Anadolubank A.Ş.	-0,02	BB
Intesa Sanpaolo S.p.A.	0,17	BB
Türkiye Kalkınma ve Yatırım Bankası A.Ş.	0,21	BB
Akbank T.A.Ş.	0,34	BB
BankPozitif Kredi ve Kalkınma Bankası A.Ş.	0,48	BB-
Bank Mellat	0,54	BB-
QNB Finansbank A.Ş.	0,55	BB-
HSBC Bank A.Ş.	0,55	BB-
MUFG Bank Turkey A.Ş.	0,62	BB-
Türkiye Garanti Bankası A.Ş.	0,67	BB-
JPMorgan Chase Bank N.A.	0,78	BB-
Yapı ve Kredi Bankası A.Ş.	0,83	BB-
Birleşik Fon Bankası A.Ş.	0,84	BB-
Habib Bank Limited	0,86	BB-
ING Bank A.Ş.	0,88	BB-
Türkiye Cumhuriyeti Ziraat Bankası A.Ş.	0,90	BB-
Burgan Bank A.Ş.	0,94	BB-
Turkish Bank A.Ş.	1,11	BB-
Türk Ekonomi Bankası A.Ş.	1,19	BB-
Fibabanka A.Ş.	1,42	B+
Şekerbank T.A.Ş.	1,48	B+
Odea Bank A.Ş.	1,78	B+
Alternatifbank A.Ş.	1,79	B+
Türkiye Halk Bankası A.Ş.	1,81	B+
Türkiye Vakıflar Bankası T.A.O.	1,90	B+
ICBC Turkey Bank A.Ş.	1,93	B+
Turkland Bank A.Ş.	4,05	B
Türk Eximbank	4,12	B
Société Générale (SA)	5,77	B-

6. BÖLÜM: SONUÇ

Literatür taraması sonucunda bir bankanın başarısızlığı ile ilgili çeşitli tanımlamalar incelendikten sonra, bu tezde, bankanın arka arkaya 3 yıl zarar etmesi (aktif karlılık oranının arka arkaya üç yıl ortalamasının altında kalması) halinde bankanın başarısız sayılacağına karar verilmiştir. Bu kabulden yola çıkarak, bankaların tahmin edilecek yıldan (2021) önceki son 3 yıla ait (2018, 2019, 2020) finansal verileri ile bankanın başarısız olup olmadığını gösteren etiketlerin listelendiği bir veri seti oluşturduk. Bu verileri kullanarak başarılı ve başarısız bankaları ayırmak için Perceptron, Karar Ağacı, Lojistik Regresyon, Lineer Diskriminant Analizi ve SVM (Destek Vektör Makinaları) sınıflandırma yöntemlerini uyguladık. Ayrıca aynı yöntemleri temel bileşenler analizine dayalı olarak uygulayarak sınıflandırma modelleri oluşturduk ve birbiriyle karşılaştırdık. Karşılaştırma ölçütü olarak kullandığımız f1 skoru değerlendirildiğinde en iyi performansı verdiği için çalışmamızın sınıflandırma modeli olarak Lineer Diskriminant Analizi yöntemini seçtik. Son aşamada, bankaların 2021 yılı için derecelendirilmesi amacıyla, tahmin etmek istediğimiz yıldan (2021) bir önceki yıl olan 2020 yılının LDA ile elde edilen skorlarına hiyerarşik kümeleme analizi uygulanmıştır. Bankaların bu eğitim setinden (2020) elde edilen küme sayısı baz alınarak 2021 yılına ait LDA skorları için kümeler ve karşılık gelen dereceleri belirlenmiştir.

KAYNAKLAR

- Allen, F., & Carletti, E. (2008). The Roles of Banks in Financial Systems. *Oxford handbook of banking*, 32-57.
- Alpar, R. (2017). *Çok Değişkenli İstatistiksel Yöntemler*. Ankara: Detay Yayıncılık.
- Altman, E. I. (1968). Financial Ratios, Discriminant Analysis and the Prediction of Corporate Bankruptcy. *The Journal of Finance*, 23(4), 589-609.
- Altunoz, U. (2013). Bankaların Finansal Başarısızlıklarının Yapay Sinir Ağları Modeli Çerçevesinde Tahmin Edilebilirliği. *Dokuz Eylül Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi*, 28(2), 189-217.
- Bi, Q., Goodman, K. E., Kaminsky, J., & Lessler, J. (2019). What is machine learning? A primer for the epidemiologist. *American journal of epidemiology*, 188(12), 2222-2239.
- Ceylan, A., & Korkmaz, T. (2001). *İşletmelerde finansal Yönetim*. Bursa: Ekin Kitabevi yayınları.
- Ceylan, M. (2022). Bankacılıkta Kredi Derecelendirme modeli. *Marmara Üniversitesi Sosyal Bilimler Enstitüsü Yüksek Lisans Tezi*.
- Christopoulos, A. G., Mylonakis, J., & Diktapanidis, P. (2011). Could Lehman Brothers' collapse be anticipated? An examination using CAMELS rating system. *International Business Research*, 4(2), 11.
- CHU, X., ILYAS, I. F., KRISHNAN, S., & WANG, J. (2016). Data cleaning: Overview and emerging challenges. *Proceedings of the 2016 international conference on management of data.*, (s. 2201-2206).
- Citterio, A. (2020). Bank failures: review and comparison of prediction models. *Available at SSRN 3719997*.
- Dang, U. (2011). rating system in banking supervision. A case study.
- Darwin, C. (1909). *The origin of species*. New York: PF Collier & son.
- De Haan, J., & Amtenbrink, F. (2011). *Credit rating agencies*. In *Handbook of Central Banking, Financial Regulation and Supervision*. . Edward Elgar Publishing.
- De Haas, R., & Van Horen, N. (2012). International shock transmission after the Lehman Brothers collapse: Evidence from syndicated lending. *American Economic Review*, 102(3), 231-237.
- Dhokal, C. P. (2017). Dealing with outliers and influential points while fitting regression. *Journal of Institute of Science and Technology*, 22(1), 61-65.
- Ekinci, A., & Erdal, H. İ. (2017). Forecasting Bank Failure: Base Learners, Ensembles and

- Hybrid Ensembles. *Comput Econ*, 49, 677-686.
- Erdoğan, B. E. (2008). Bankruptcy Prediction of Turkish Commercial Banks Using Financial Ratios. *Applied Mathematical Sciences*, 2(60), 2973-2982.
- European Securities and Markets Authority (ESMA). (2021). *Report on CRA Market Share*.
- Frost, A. C. (2007). Credit Rating Agencies in Capital Markets: A Review of Research Evidence on Selected Criticisms of the Agencies. *Journal of Accounting, Auditing & Finance*, 22(3), 469-492.
- Giray Yakut, S., & Bacaksız, N. E. (2020). Teknoloji yoğunluğuna göre finansal başarısızlık tahmin modelleri değişir mi? imalat sanayi sektörü üzerine bir uygulama. *Marmara Üniversitesi İktisadi ve İdari Bilimler Dergisi*, 41(2), 536-555.
- GÖKMEN, B. (2007). Bankalarda Finanssal toblolar analizi. 20-23. İstanbul, Türkiye: İ.Ü. Sosyal bilimler enstitüsü, Yüksek lisans tezi.
- Gör, Y. (2018). Kurumsal yönetimin finansal başarısızlığı önlemedeki yeri üzerine bir araştırma. *Avrasya Sosyal ve Ekonomi Araştırmaları Dergisi*, 5(12), 689-697.
- Hanh, P. T., Hang, N. T., Huy, D. T., & Nuong, L. N. (2021). Enhancing Roles of Banks and the Comparison of Market Risk and Risk Policy Implications in Group of Listed Vietnam Banks During 2 Stages:. (11), 1723-1735. *Revista geintec-gestao Inovacao e Tecnologias*.
- Holland, J. (1975). Adaptation in natural and artificial systems. *Ann Arbor*.
- İnceişçi, F. K. (2021). Makine Öğrenme ile gemi makinelerinin hata analizi. *Marmara Üniversitesi, Bilgisayar Mühendisliği Programı Yüksek Lisans Tezi*.
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2013). *An introduction to statistical learning (Vol. 112, p. 18)*. New York: Springer.
- King, M. R., Ongena, S., & Tarashev, N. A. (2016). Bank Standalone Credit Ratings.
- Kumar, V., & Minz, S. (2014). Feature selection: a literature review. *Smart Computing Review*, 4(3), 211-229.
- Kwak, S. K., & Kim, J. H. (2017). Statistical data preparation: management of missing values and outliers. *Korean journal of anesthesiology*, 70(4), 407-411.
- Liu, F. T., Ting, K. M., & Zhou, Z.-H. (2008). Isolation forest. *2008 eighth ieee international conference on data mining.*, 413-422.
- Mansouri, A., Nazari, A., & Ramazani, M. (2016). A comparison of artificial neural network model and logistics regression in prediction of companies' bankruptcy (A case study of Tehran stock exchange). *International Journal of Advanced Computer Research*, 6(24).

- Noriega, L. (2005). Multilayer perceptron tutorial.
- Odom, M. D., & Sharda, R. (1990). A neural network model for bankruptcy prediction. 1990 *IJCNN International Joint Conference on neural networks*. (s. 163-168). IEEE.
- Osborne, J. W., & Overbay, A. (2004). The power of outliers (and why researchers should always check for them). *Practical Assessment, Research, and Evaluation*, 9(1).
- Rojas-Suarez, L. (2002). Rating banks in emerging markets: What credit rating agencies should learn from financial indicators. *Ratings, rating agencies and the global financial system*, 177-201.
- Rosenblatt, F. (1958). The perceptron: a probabilistic model for information storage and organization in the brain. *Psychological review*, 65(6), 386.
- Samırkaş, M. C., Evcı, S., & Ergün, B. (2014). Türk Bankacılık Sektöründe Karlılığın belirleyicileri. *Kafkas Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi*, 5(8).
- Tang, L., Peng, S., Bi, Y., Shan, P., & Hu, X. (2014). A New Method Combining LDA and PLS for Dimension Reduction. *PLoS ONE*, 9(5).
- Tuzci, D. Ş. (2020). *Bankacılıkta Kredi Derecelendirme sistemleri*. İstanbul: Bahçeşehir Üniversitesi, Sosyal Bilimler Enstitüsü, İşletme Yüksek Lisans Tezi.
- Türkcan, Z., Bozcuk, A. E., & Türkcan, K. (2018, Ekim). Türk Bankalarında Mali Başarısızlığın Tahmin Edilmesine Yönelik Ampirik Bir Çalışma. *Muhasebe ve Finansman Dergisi*(80), 253-274.
- White, L. J. (2009). The credit-rating agencies and the subprime debacle. *Critical Review*, 21(2-3), 389-399.
- Wilson, I. (2006). Regulatory and institutional challenges of corporate governance in post banking consolidation Nigeria. *Economic and Policy Review* 12, 2.
- Wu, Y., Gaunt, C., & Gray, S. (2010). A comparison of alternative bankruptcy prediction models. *Journal of Contemporary Accounting & Economics*, 6(1), 34-45.
- Xue, H., Chen, S., & Yang, Q. (2011). Structural regularized support vector machine: a framework for structural large margin classifier. *IEEE Transactions on Neural Networks*, 22(4), 573-587.
- Yakıcı Ayan, T., & Değirmenci, N. (2018). Firma finansal başarısızlık öngörüsü için bir Lojistik Regresyon modeli. *Uluslararası İktisadi ve İdari İncelemeler Dergisi*, 18, 77-88.
- Yürük, M. F., & Ekşi, İ. H. (2019). Yapay Zeka Yöntemleri ile İşletmelerin Finansal Başarısızlığının Tahmin Edilmesi. *Bist İmalat Sektörü Uygulaması*, 10(1), 393-422.