



T.C.
NECMETTİN ERBAKAN
ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ



KALP YETMEZLİĞİNDE ÖLÜM RİSKİNİN
TAHMİN EDİLMESİNDE EN ETKİLİ
ÖZNİTELİKLERİN BULUNMASI

Buse Nur PAYDAŞ

YÜKSEK LİSANS TEZİ

Elektrik Elektronik Mühendisliği Anabilim Dalı

Haziran-2026
KONYA
Her Hakkı Saklıdır

TEZ KABUL VE ONAYI

Buse Nur PAYDAŞ tarafından hazırlanan “*Kalp Yetmezliğinde Ölüm Riskinin Tahmin Edilmesinde En Etkili Özniteliklerin Bulunması*” adlı tez çalışması 11/06/2026 tarihinde aşağıdaki jüri tarafından oy birliği / oy çokluğu ile Necmettin Erbakan Üniversitesi Fen Bilimleri Enstitüsü Elektrik – Elektronik Mühendisliği Anabilim Dalı’nda YÜKSEK LİSANS TEZİ olarak kabul edilmiştir.

Jüri Üyeleri

İmza

Başkan

Doç. Dr. Güzin ÖZMEN

.....

Danışman

Doç. Dr. Sabri ALTUNKAYA

.....

Üye

Dr. Öğretim Üyesi Ali Osman ÖZKAN

.....

Fen Bilimleri Enstitüsü Yönetim Kurulu’nun/.../20.. gün ve sayılı kararıyla onaylanmıştır.

Prof. Dr. Havvanur UÇBEYİAY
FBE Müdürü

TEZ BİLDİRİMİ

Bu tezdeki bütün bilgilerin etik davranış ve akademik kurallar çerçevesinde elde edildiğini ve tez yazım kurallarına uygun olarak hazırlanan bu çalışmada bana ait olmayan her türlü ifade ve bilginin kaynağına eksiksiz atıf yapıldığını bildiririm.

DECLARATION PAGE

I hereby declare that all information in this document has been obtained and presented in accordance with academic rules and ethical conduct. I also declare that, as required by these rules and conduct, I have fully cited and referenced all material and results that are not original to this work.

Buse Nur PAYDAŞ

Tarih:11/06/2026

ÖZET

YÜKSEK LİSANS TEZİ

KALP YETMEZLİĞİNDE ÖLÜM RİSKİNİN TAHMİN EDİLMESİNDE EN ETKİLİ ÖZNİTELİKLERİN BULUNMASI

Buse Nur PAYDAŞ

NECMETTİN ERBAKAN ÜNİVERSİTESİ FEN BİLİMLERİ ENSTİTÜSÜ
ELEKTRİK ELEKTRONİK MÜHENDİSLİĞİ ANABİLİM DALI

Danışman: Doç. Dr. Sabri ALTUNKAYA

2026, 53 Sayfa

Jüri

Doç. Dr. Sabri ALTUNKAYA

Doç. Dr. Güzin ÖZMEN

Dr. Öğretim Üyesi Ali Osman ÖZKAN

Bu çalışmada, kalp yetmezliği hastalarına ait klinik veriler kullanılarak hasta sağkalımının tahmin edilmesinde en uygun parametrelerin belirlenmesine yönelik makine öğrenmesi tabanlı öznitelik seçimi yaklaşımları kapsamlı bir şekilde incelenmiştir. Farklı filtre, sarmal ve gömülü öznitelik seçimi yöntemleri çeşitli makine öğrenmesi sınıflandırıcılarıyla entegre edilerek birlikte değerlendirilmiştir.

Deneyssel analizler; takip süresi değişkeninin modele dâhil edildiği ve hariç tutulduğu iki farklı senaryo altında, modellerin genellenebilirliğini artırmak amacıyla Monte Carlo çapraz doğrulama yöntemiyle gerçekleştirilmiştir. Bu kapsamda klinik veri seti, her bir iterasyonda rastgele %70 eğitim ve %30 test olarak parçalanmış; bu işlem 100 kez tekrarlanarak (%70-30, x100) elde edilen performans sonuçlarının ortalaması alınmıştır. Uygulanan bu strateji, model başarısının veri bölünmesine olan duyarlılığını azaltmış, aşırı uyum riskini önlemiş ve sonuçların istatistiksel olarak daha kararlı ve güvenilir olmasını sağlamıştır. Model başarımı; doğruluk, F1-skoru, duyarlılık, özgüllük ve dengesiz veri setlerinde kararlı bir performans göstergesi olan Matthews korelasyon katsayısı (MCC) metrikleri ile ölçülmüştür.

Elde edilen sonuçlar, takip süresi değişkeninin modele dâhil edilmesinin F1-skoru ve MCC değerlerini belirgin şekilde artırdığını ve sınıf dengesizliği üzerinde olumlu bir etki yarattığını göstermektedir. Özellikle ReliefF tabanlı öznitelik seçimi yöntemi ile k-NN sınıflandırıcısının kombinasyonundan oluşan model, %85,56 doğruluk, %74,56 F1-skoru ve %65,94 MCC değeri ile en yüksek ve en dengeli performansı sunmuştur. Takip süresi değişkeni hariç tutulduğunda modellerin genel performansında düşüş gözlenmiş, ancak TreeBagger tabanlı öznitelik seçimi yöntemiyle eğitilen modellerin hâlâ tatmin edici ve literatürün üzerinde sonuçlar sağladığı saptanmıştır. Bu çalışma, makine öğrenmesi tabanlı öznitelik seçimi ve sınıflandırma yöntemlerinin kalp yetmezliği hasta sağkalımını yüksek güvenilirlikle tahmin edebileceğini ortaya koymakta; klinisyenler için tıbbi pratik süreçlerinde uygulanabilir ve kararlı bir klinik karar destek sistemi çerçevesi sunmaktadır.

Anahtar Kelimeler: Kalp yetmezliği, makine öğrenmesi, öznitelik seçimi, sağkalım analizi.

ABSTRACT

MS THESIS

FINDING THE MOST EFFECTIVE FEATURES FOR PREDICTING THE RISK OF DEATH IN HEART FAILURE

Buse Nur PAYDAŞ

**THE GRADUATE SCHOOL OF NATURAL AND APPLIED SCIENCES OF
NECMETTİN ERBAKAN UNIVERSITY THE DEGREE OF MASTER OF
SCIENCE IN ELECTRICAL ELECTRONICS ENGINEERING**

Advisor: Doç. Dr. Sabri ALTUNKAYA

2026, 53 Pages

Jury

Doç. Dr. Sabri ALTUNKAYA

Doç. Dr. Güzin ÖZMEN

Dr. Öğretim Üyesi Ali Osman ÖZKAN

In this study, machine learning-based feature selection approaches aimed at determining the most appropriate parameters for predicting patient survival using clinical data of heart failure patients were comprehensively investigated. Different filter, wrapper, and embedded feature selection methods were integrated and jointly evaluated with various machine learning classifiers.

Experimental analyses were carried out using the Monte Carlo cross-validation method to improve the generalizability of the models under two different scenarios, where the follow-up time variable was included in and excluded from the model. In this context, the clinical dataset was randomly partitioned into 70% training and 30% testing in each iteration; this process was repeated 100 times (70-30, x100), and the average of obtained performance results was taken. This applied strategy reduced the sensitivity of the model success to data splitting, prevented the risk of overfitting, and ensured that the results were statistically more stable and reliable. Model performance was evaluated using accuracy, F1-score, sensitivity, specificity, and the Matthews correlation coefficient (MCC) metrics, the latter being a robust indicator of performance in imbalanced datasets.

The obtained results indicate that the inclusion of the time variable in the model significantly increases the F1-score and MCC values and exerts a positive effect on class imbalance. In particular, the model consisting of the combination of the ReliefF-based feature selection method and the K-NN classifier provided the highest and most balanced performance with 85,56% accuracy, 74,56% F1-score, and 65,94% MCC value. A decrease in the overall performance of the models was observed when the follow-up time variable was excluded; however, it was determined that the models trained with the Tree Bagger-based feature selection method still provided satisfactory results that are above literature. This study demonstrates that machine learning-based feature selection and classification methods can predict heart failure patient survival with high reliability, offering a practicable and stable clinical decision support system framework for clinicians in medical practice processes.

Keywords: Feature selection, heart failure, machine learning, survival analysis.

ÖNSÖZ

Bu tez çalışmasının her aşamasında bilgi, destek ve sabrıyla bana her zaman yol gösteren saygıdeğer danışman hocam Doç. Dr. Sabri ALTUNKAYA'ya içtenlikle teşekkür ederim.

Hayatımın her aşamasında olduğu gibi, bu eğitim sürecinde de maddi ve manevi destekleriyle daima yanımda olan sevgili aileme teşekkürlerimi borç bilirim.

Son olarak; yüksek lisans eğitimim boyunca desteğini her zaman hissettiğim ve bu süreçte vefat eden sevgili abim Serkan'ı rahmet ve minnetle yâd ediyorum.

Buse Nur PAYDAŞ
KONYA-2026



İÇİNDEKİLER

ÖZET	iv
ABSTRACT.....	v
ÖNSÖZ	vi
ŞEKİLLER LİSTESİ	ix
TABLolar LİSTESİ	x
SİMGELER VE KISALTMALAR	xi
1. GİRİŞ	1
2. KAYNAK ARAŞTIRMASI	4
3. MATERYAL VE YÖNTEM.....	9
3.1. Veri Seti	9
3.2. Deneysel Tasarım	11
3.3. Veri Ön İşleme.....	11
3.4. Öznitelik Seçimi Yöntemleri	11
3.4.1. Komşuluk bileşeni analizi (Neighborhood component analysis, NCA).....	13
3.4.2. ReliefF	14
3.4.3. Minimum tekrar - maksimum ilişki (Minimum redundancy-maximum relevance, mRMR).....	14
3.4.4. Ki-kare (Chi-Square)	15
3.4.5. Ardışık öznitelik seçimi (Sequential feature selection, SFS)	16
3.4.6. LASSO (Least absolute shrinkage and selection operator)	16
3.4.7. Tree tabanlı öznitelik seçimi.....	17
3.4.8. Topluluk tabanlı ağaç yöntemi (TreeBagger).....	18
3.5. Kullanılan Makine Öğrenmesi Algoritmaları	18
3.5.1. k-En yakın komşular (k-Nearest neighbors, k-NN).....	18
3.5.2. Destek vektör makineleri (Support vector machine, SVM)	19
3.5.3. Naive bayes.....	20
3.5.4. Karar ağacı (Decision tree).....	21
3.5.5. Ayrımcı analiz (Discriminant analysis)	21
3.5.6. Lojistik regresyon (Logistic regression)	22
3.5.7. Çekirdek yaklaştırma (Kernel approximation)	23
3.5.8. Topluluk yöntemleri (Ensemble)	24
3.5.9. Yapay sinir ağları (Artificial neural networks, YSA).....	25
3.5.10. Doğrusal sınıflandırıcı (Linear classifier).....	26
3.6. Performans Değerlendirme Metrikleri.....	26
3.6.1. Doğruluk (Accuracy)	27
3.6.2. Duyarlılık (Sensitivity)	27
3.6.3. Özgüllük (Specificity).....	27
3.6.4. F1-skoru	27
3.6.5. Matthews korelasyon katsayısı (Matthews Correlation Coefficient, MCC).28	
4. ARAŞTIRMA SONUÇLARI VE TARTIŞMA.....	29

4.1. Araştırma Sonuçları	31
4.1.1. Takip süresi özniteliği dahil edilmeden elde edilen sonuçlar	31
4.1.2. Takip süresi özniteliği dahil edildiğinde elde edilen sonuçlar.....	32
4.2. Tartışma	33
4.2.1. Sonuçların Analizi	33
4.2.2. Literatürle Karşılaştırma	34
4.2.3. Çalışmanın Sınırlılıkları.....	42
5. SONUÇLAR VE ÖNERİLER	43
5.1. Sonuçlar	43
5.2. Öneriler	44
KAYNAKLAR	45
EKLER	50



ŞEKİLLER LİSTESİ

Şekil 3.1. Veri tipine dayalı öznitelik seçimi ve model doğrulama sürecinin genel akışı 9



TABLÖLAR LİSTESİ

Tablo 3.1. Veri setindeki özniteliklerin anlamları, ölçüm birimleri ve aralıkları	10
Tablo 3.2. Öznitelik seçimi yöntemlerinin kullanılan veri türlerine göre dağılımı	12
Tablo 4.1. Takip süresi özniteliği hariç tutulduğunda elde edilen sonuçlar.....	31
Tablo 4.2. Takip süresi özniteliği dahil edildiğinde elde edilen sonuçlar.....	32
Tablo 4.3. Literatürde kalp yetmezliği sağkalım tahmini üzerine takip süresi dahil edilerek yapılan çalışmaların karşılaştırılması.....	38
Tablo 4.4. Literatürde kalp yetmezliği sağkalım tahmini üzerine takip süresi dahil edilmeden yapılan çalışmaların karşılaştırılması	41



SİMGELER VE KISALTMALAR

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

SİMGELER

A	Projeksiyon (Dönüşüm) matrisi
b	Yanlılık (bias) terimi
B	Toplam karar ağacı sayısı
C	Sınıf etiketi veya kovaryans matrisi
d	Öklidyen mesafenin karesi
E	Beklenen değer (Expected value)
f(x)	Karar fonksiyonu veya aktivasyon fonksiyonu
H(S)	Entropi (Belirsizlik ölçüsü)
I(X; C)	Karşılıklı bilgi (Mutual information)
K	En yakın komşu sayısı
m	İterasyon sayısı veya örneklem büyüklüğü
mcg/L	Mikrogram/Litre (Kreatin kinaz vb. konsantrasyonu)
mL	Mililitre (Hacim ölçüsü)
mEq/L	Miliekivalan/Litre (Serum sodyum vb. elektrolit ölçüsü)
n	Toplam veri örneği sayısı
O	Gözlemlenen değer (Observed value)
p	Olasılık değeri veya öznitelik sayısı
P(C)	Öncel olasılık (Prior probability)
S	Seçilen öznitelikler kümesi
w/W	Ağırlık vektörü veya matrisi
x	Öznitelik vektörü (Giriş verisi)
y	Hedef değişken (Sınıf etiketi)
α (Alpha)	Lagrange çarpanı
β (Beta)	Model katsayıları
λ (Lambda)	Düzenleme (regülerizasyon) parametresi
μ (Mu)	Ortalama vektörü
ξ (Xi)	Gevşeme değişkeni (Slack variable)
σ (Sigma)	Aktivasyon fonksiyonu (Sigmoid vb.)
ϕ (Phi)	Çekirdek (Kernel) dönüşüm fonksiyonu
χ^2	Ki-kare (Chi-square) istatistiği
Δ (Delta)	Safılık azalması (Gain)
Σ (Sigma)	Toplam sembolü veya Kovaryans matrisi

KISALTMALAR

ANOVA	Varyans Analizi (Analysis of Variance)
BRF	Dengelenmiş Rastgele Orman (Balanced Random Forest)
CNN	Evrişimsel Sinir Ağı (Convolutional Neural Network)
CPK	Kandaki kreatinin fosfokinaz enzimi düzeyi (Creatinine phosphokinase)
DA	Ayrımcı Analiz (Discriminant Analysis)
DT	Karar Ağaçları (Decision Tree)
EF	Ejeksiyon fraksiyonu (Ejection fraction)
ETC	Ekstra Ağaçlar Sınıflandırıcısı (Extra Trees Classifier)

F1	F1-Skoru (Hassasiyet ve Duyarlılığın Harmonik Ortalaması)
GBM	Gradyan Artırma Makineleri (Gradient Boosting Machines)
GLM	Genelleştirilmiş Doğrusal Modeller (Generalized Linear Models)
GLMNET	Düzenleştirilmiş Genelleştirilmiş Doğrusal Modeller
HBP	Hipertansiyon (High blood pressure)
k-NN	k-En Yakın Komşu (k-Nearest Neighbors)
LASSO	En Az Mutlak Seçme ve Daralma Operatörü
LDA	Doğrusal Diskriminant Analizi (Linear Discriminant Analysis)
LR	Lojistik Regresyon (Logistic Regression)
MATLAB	Matris Laboratuvarı (Matrix Laboratory)
MCC	Matthews Korelasyon Katsayısı (Matthews Correlation Coefficient)
MÖ	Makine Öğrenmesi (Machine Learning)
MLP	Çok Katmanlı Algılayıcı (Multi-Layer Perceptron)
mRMR	Minimum Artıklık Maksimum İlginlik
NB	Naive Bayes
NCA	Komşu Bileşen Analizi (Neighborhood Component Analysis)
NHi	En Yakın Hit (Nearest-Hit)
NMi	En Yakın Miss (Nearest-Miss)
PCA	Temel Bileşenler Analizi (Principal Component Analysis)
RF	Rastgele Orman (Random Forest)
RFE	Özyinelemeli Öznitelik Eleme (Recursive Feature Elimination)
SC	Serum kreatin (Serum creatinine)
SFS	Ardışık Öznitelik Seçimi (Sequential Feature Selection)
SMOTE	Sentetik Azınlık Aşırı Örnekleme Tekniği
SS	Serum sodyumu (Serum sodium)
SVM	Destek Vektör Makineleri (Support Vector Machines)
UCI	University of California, Irvine - Machine Learning Repository
YSA	Yapay Sinir Ağları (Artificial Neural Network)

1. GİRİŞ

Kalp yetmezliği, kalbin vücudun metabolik ihtiyaçlarını karşılayacak düzeyde kan pompalayamaması sonucu ortaya çıkan klinik bir durumdur. Bu durum, kalbin yeterli miktarda kanla dolamaması veya mevcut kanı etkili bir şekilde dolaşıma gönderememesi nedeniyle gelişebilmektedir (National Heart, Lung, and Blood Institute, 2022).

Kalp hastalıkları ve inme için başlıca davranışsal risk faktörleri arasında sağlıklı beslenme, fiziksel hareketsizlik, tütün kullanımı ve alkol tüketimi yer almaktadır. Bunun yanı sıra hava kirliliği gibi çevresel faktörler de kardiyovasküler hastalık riskini artırmaktadır. Bu risk faktörleri bireylerde yüksek tansiyon, hiperglisemi, dislipidemi, aşırı kilo ve obezite gibi klinik bulguların ortaya çıkmasına yol açabilmektedir (World Health Organization, 2025).

İngiliz Kalp Vakfı (British Heart Foundation) tarafından yayımlanan veriler, Birleşik Krallık'ta 75 yaşından önce kalp ve dolaşım hastalıkları nedeniyle gerçekleşen ölümlerin son yıllarda artış gösterdiğini ortaya koymaktadır. 2014 yılında 41.042 olan bu sayı, 2017 yılında 42.384'e yükselmiştir. Ayrıca 2017'ye kadar olan beş yıllık dönemde, 65 yaşından önce kalp ve dolaşım hastalıklarına bağlı ölümlerde %4 oranında bir artış gözlenmiştir; bu durum, önceki beş yılda kaydedilen %19'luk düşüşle karşılaştırıldığında dikkat çekici bir değişime işaret etmektedir (Siddique, 2019). Kalp yetmezliği tedavisi, yüksek gelirli ülkelerde önemli düzeyde finansal yükler oluşturmaktadır (Cook ve Collins, 2015). Literatürde, geleneksel tanısal yöntemlerin kalp yetmezliğinin patofizyolojisini anlamada değerli katkılar sunduğu; ancak hastalığın çok boyutlu yapısını tam olarak yansıtmada konusunda bazı sınırlılıklar taşıdığı belirtilmektedir. Bu bağlamda, H2FPEF ve HFA-PEFF gibi tanısal skor sistemlerindeki son gelişmelerin, tanısal belirsizliklerin azaltılması ve hasta bakımının optimize edilmesi açısından umut verici olanaklar sunmaktadır. Bu bağlamda makine öğrenmesi (MÖ), kalp yetmezliğinde daha doğru risk tahmini, prognoz öngörüsü ve tedavi optimizasyonunu mümkün kılmak üzere zengin klinik verileri kullanabilen dönüştürücü bir araç olarak öne çıkmaktadır. Karmaşık algoritmalar ve hesaplamalı tekniklerden yararlanan MÖ yöntemleri, büyük veri kümeleri içindeki karmaşık örüntüleri ortaya çıkarma potansiyeline sahiptir ve bu sayede klinisyenleri kişiselleştirilmiş hasta bakımı için uygulanabilir ve yol gösterici bilgilerle destekleyebilir (Averbuch vd., 2022). Kalp yetmezliği verilerinin toplanması ve analizinde; hasta sağkalımını tahmin etmeye yönelik MÖ sınıflandırıcıları, denetimli

derin öğrenme uygulamaları ve çeşitli MÖ algoritmalarının yaygın biçimde kullanıldığı bildirilmektedir (J. Wang, 2021).

Araştırmanın amacı

Literatürde kalp yetmezliği hastalarına ait klinik veriler kullanılarak farklı yapay zekâ ve makine öğrenmesi tabanlı sınıflandırma algoritmalarının performanslarının karşılaştırıldığı çok sayıda çalışma bulunmaktadır. Ancak, bu çalışmalarda öznelite seçme yöntemlerinin sistematik olarak ele alındığı çalışmaların sınırlı olduğu görülmektedir. Bu doğrultuda, bu çalışmanın amacı, kalp yetmezliği hastalarına ait klinik veriler üzerinde farklı öznelite seçme yöntemleri ile makine öğrenmesi tabanlı sınıflandırma algoritmalarının birlikte değerlendirilerek performanslarının karşılaştırılması ve hasta sağkalımının tahminine yönelik daha doğru ve etkili bir öngörü modelinin tasarlanmasıdır.

UCI Kalp Yetmezliği veri seti üzerine literatürde çok sayıda çalışma bulunmasına rağmen, bu çalışmaların büyük bir kısmı veri setindeki takip süresi öznelitesinin sızıntı (data leakage) riskini veya model başarısındaki dominant etkisini göz ardı etmiştir. Bu tez çalışması, söz konusu veri setini sadece bir sınıflandırma problemi olarak ele almanın ötesine geçerek; farklı öznelite seçimi yöntemlerinin takip süresi faktörü ile nasıl etkileşime girdiğini ve bu durumun klinik öngörü değerini nasıl değiştirdiğini irdelemektedir. Ayrıca, sınırlı örneklem sayısına sahip bu veri setinde aşırı uyum (overfitting) riskini bertaraf etmek adına uygulanan 100 tekrarlı Monte Carlo simülasyonu, çalışmamızı metodolojik olarak daha sağlam bir zemine oturtmaktadır.

Araştırmanın önemi

Bu çalışmanın önemi, kalp yetmezliği gibi yüksek mortaliteye sahip bir hastalıkta erken risk sınıflandırması ve sağkalım tahmini yapılmasına katkı sağlayabilecek makine öğrenmesi tabanlı bir yaklaşım sunmasından kaynaklanmaktadır. Elde edilecek bulguların, klinik karar destek sistemlerinin oluşturulmasına ve hasta yönetiminin iyileştirilmesine yönelik bilimsel bir zemin oluşturması beklenmektedir. Ayrıca çalışma, sağkalım tahmini üzerinde en yüksek belirleyiciliğe sahip olan klinik parametrelerin öznelite seçimi (feature selection) yöntemleriyle analiz edilmesini amaçlamaktadır. Bu sayede, hangi değişkenlerin modelin öngörü başarısını daha fazla etkilediği ortaya konularak, klinik pratikte öncelik verilmesi gereken risk faktörlerine dair bir perspektif sunulması hedeflenmektedir. Son olarak çalışma, mevcut literatürde yer alan yöntemlerin

karşılaştırmalı analizi yoluyla hangi algoritmaların daha yüksek öngörü başarısı sağladığını ortaya koymayı amaçlamaktadır.



2. KAYNAK ARAŞTIRMASI

Kalp yetmezliği hastalarının sağkalım tahmini üzerine literatürde çok sayıda çalışma bulunmaktadır. Bu çalışmalar, makine öğrenmesi yöntemlerinin klinik veriler üzerinde etkin bir biçimde uygulanabildiğini göstermektedir. Bu çalışmada aynı veri seti üzerinden gerçekleştirilen çalışmalara yer verilmiştir (Ahmad vd., 2017).

Bu çalışmalar arasında yer alan ve kalp yetmezliği hastalarının sağkalım analizi ve risk sınıflandırması üzerine kapsamlı bir inceleme sunan bir çalışmada, makine öğrenmesi algoritmalarının bu alandaki etkinliği ve pratik klinik uygulama potansiyeli değerlendirilmiştir (Çakır, 2025). Araştırma kapsamında, Lojistik Regresyon (Logistic Regression, LR), Doğrusal Ayrımcı Analizi (Linear Discriminant Analysis, LDA), Gradyan Artırma Makineleri (Gradient Boosting Machines, GBM), Destek Vektör Makineleri (Support Vector Machines, SVM) ve Yapay Sinir Ağları (YSA) gibi geniş bir yelpazede yer alan modeller karşılaştırmalı olarak analiz edilmiştir. Elde edilen bulgular; GLM, LDA, GBM, Doğrusal SVM ve GLMNET algoritmalarının %96,6 doğruluk ve %97,6 F1-skoru değerlerine ulaşarak en yüksek başarıyı sergilediğini ortaya koymuştur.

Çalışmanın bir diğer önemli çıktısı ise değişkenlerin önem düzeylerine ilişkindir; Rastgele Orman (Random Forest) algoritmasıyla yürütülen öznitelik analizinde, mortalite tahmininde en kritik belirleyicinin takip süresi (time) olduğu, bu parametreyi ise serum kreatinin ve ejeksiyon fraksiyonu değerlerinin izlediği saptanmıştır. Teorik modellerin ötesine geçen araştırmanın özgün yanı, tasarlanan bu yüksek performanslı algoritmaların R-Shiny platformu üzerinden interaktif bir Klinik Karar Destek Sistemi (KKDS) arayüzüne dönüştürülmüş olmasıdır. Bu sistem, sağlık profesyonellerine klinik veriler üzerinden gerçek zamanlı risk tahmini yapma imkânı sunarak, yapay zekâ temelli çözümlerin tıbbi uygulama süreçlerine entegrasyonu konusunda somut bir çerçeve çizmektedir.

Souza ve Lima (2024) çalışmasında ise kalp yetmezliği veri seti üzerinde k-En Yakın Komşu (k-Nearest Neighbors, k-NN) algoritması ve korelasyon tabanlı öznitelik seçimi yöntemleri kullanılmıştır. Çalışmada farklı k değerleri ve korelasyon filtreleri test edilmiş ve en iyi sonuç, takip süresi özniteliği dahil edilerek $k=21$ ve $CF=0,1$ parametreleri ile %83,50 doğruluk değeri olarak elde edilmiştir. Bu en iyi performans elde edilirken kullanılan öznitelik kümesi; kreatin fosfokinaz, serum sodyumu, cinsiyet ve takip süresi değişkenlerinden oluşmaktadır. Bununla birlikte, çalışmada klinik açıdan en anlamlı özniteliklerin anemi, yüksek tansiyon, serum kreatinin ve cinsiyet olduğu

belirtilmiştir. Ayrıca, takip süresi özniteliğinin çıkarılması durumunda model performansında belirgin bir düşüş gözlemlenmiş olup, bu durum takip süresi bilgisinin kalp yetmezliği tahmininde önemli bir rol oynadığını göstermektedir. Bu bulgular, takip süresi özniteliğinin modele dahil edilmesinin performans üzerinde kritik bir etkisi olduğunu ortaya koymaktadır (Souza ve Lima, 2024).

Qadeer vd. (2025) tarafından yapılan çalışmada, kalp yetmezliği hastalarının sağkalım durumunun tahmini için makine öğrenmesi ve derin öğrenme yöntemleri karşılaştırılmıştır. Çalışmada UCI Kalp Yetmezliği Klinik Kayıtları (Heart Failure Clinical Records) veri seti kullanılmış olup, veri setinde yer alan 299 hastaya ait 12 klinik öznitelik modele dahil edilmiştir. Veri ön işleme aşamasında StandardScaler ile ölçekleme yapılmış ve sınıf dengesizliği problemini gidermek amacıyla Sentetik Azınlık Aşırı Örnekleme Tekniği (Synthetic Minority Over-sampling Technique, SMOTE) uygulanmıştır. Modelleme sürecinde k-En Yakın Komşular (k-NN), Naive Bayes (NB), Rastgele Orman, SVM ve Evrişimli Sinir Ağları (Convolutional Neural Network, CNN) gibi farklı algoritmalar kullanılmıştır. Model performansı %70 eğitim ve %30 test verisi olacak şekilde tek veri bölme (hold-out) yöntemi ile değerlendirilmiştir. Elde edilen sonuçlara göre en yüksek performans CNN modeli ile sağlanmış olup doğruluk değeri %99,75, F1 skoru %99,78 ve duyarlılık değeri %99,76 olarak raporlanmıştır. Makine öğrenmesi modelleri arasında ise Rastgele Orman ve k-NN algoritmaları sırasıyla yaklaşık %96 ve %95 doğruluk değerleri ile öne çıkmıştır (Qadeer vd., 2025).

Moreno-Sánchez (2023) çalışmasında, kalp yetmezliği hastalarının sağkalım tahmini için açıklanabilir yapay zekâ (Explainable Artificial Intelligence, XAI) tabanlı bir yaklaşım önerilmiştir. Çalışmada UCI Kalp yetmezliği (Heart Failure) veri seti kullanılmış olup, 299 hastaya ait 12 klinik öznitelik değerlendirilmiştir. Model geliştirme sürecinde veri hazırlama, öznitelik seçimi ve modelleme adımlarını otomatik olarak gerçekleştiren bir veri işleme Scikit-learn tabanlı bir veri işleme pipeline'ı (Scikit-learn-based data processing pipeline, SCI-XAI) kullanılmıştır.

Çalışmada geleneksel modellerin yanı sıra Rastgele Orman, Ekstra Ağaçlar (Extra Trees), AdaBoost (Adaptive Boosting), Gradyan Artırma (Gradient Boosting) ve XGBoost (Extreme Gradient Boosting) gibi topluluk (ensemble) yöntemleri ile test edilmiş ve model performansı 5-katlı çapraz doğrulama (5-fold cross validation) ile değerlendirilmiştir. Ayrıca modelin yorumlanabilirliğini artırmak amacıyla SHAP (SHapley Additive exPlanations), Kısmi Bağımlılık Grafiği (Partial Dependence Plot - PDP) ve öznitelik önemi (feature importance) gibi açıklanabilirlik teknikleri

uygulanmıştır. Elde edilen sonuçlara göre, en dengeli model Ekstra Ağaçlar sınıflandırıcısı ile elde edilmiş olup yaklaşık %84,4 doğruluk ve %79,5 dengelenmiş doğruluk (balanced accuracy) değeri elde edilmiştir. Bu model, 12 öz nitelik arasından seçilen 5 öz nitelik ile oluşturulmuştur. Çalışmada takip süresi öz niteliğinin en önemli değişken olduğu ve model performansı üzerinde belirleyici etkisi olduğu vurgulanmıştır (Moreno-Sánchez, 2023).

Zhang (2025) çalışmasında, kalp yetmezliği hastalarında mortalite tahmini amacıyla makine öğrenmesi yöntemlerinin performansı karşılaştırılmıştır. Çalışmada Kaggle üzerinden elde edilen Kalp Yetmezliği Klinik Kayıtları veri seti kullanılmış olup, 299 hastaya ait 12 klinik öz nitelik değerlendirilmiştir. Veri ön işleme aşamasında aykırı değerlerin temizlenmesi ve veri ölçekleme işlemleri gerçekleştirilmiştir.

Öz nitelik seçimi amacıyla Temel Bileşenler Analizi (Principal Component Analysis, PCA) yöntemi kullanılarak veri boyutu azaltılmış ve modelin karmaşıklığı düşürülmüştür. Modelleme aşamasında Lojistik Regresyon, Rastgele Orman ve k-En Yakın Komşular algoritmaları uygulanmıştır. Model performansı doğruluk, kesinlik (precision), duyarlılık ve F1-skoru metrikleri ile değerlendirilmiş, ayrıca çapraz doğrulama kullanılarak model genellenebilirliği incelenmiştir. Elde edilen sonuçlara göre en yüksek performans Rastgele Orman modeli ile elde edilmiş olup %90,7 doğruluk, %80 kesinlik, %57,1 duyarlılık ve %66,7 F1-skoru değerleri raporlanmıştır (Zhang, 2025).

Bir diğer çalışma ise, kalp yetmezliği hastalarının sağ kalım öngörüsünde sınıf dengesizliği probleminin çözümüne odaklanır (Ishaq vd., 2021). Yazarlar, mevcut veri setindeki ölümlü vakaların azınlıkta olmasının model performansını olumsuz etkilediğini saptayarak, bu sorunu aşmak için SMOTE kullanmışlardır. Çalışmada, Rastgele Orman, Karar Ağaçları (Decision Tree, DT), Lojistik Regresyon ve Gradyan Artırma gibi dokuz farklı makine öğrenmesi algoritması test edilmiştir.

Araştırmanın sonuçları, veri seti SMOTE tekniği ile dengelendiğinde modellerin tahmin gücünün belirgin ölçüde arttığını ortaya koymuştur. Özellikle Ekstra Ağaçlar Sınıflandırıcısı (Extra Trees Classifier, ETC) algoritmasının, öz nitelik seçimi yöntemleriyle birleştirildiğinde %92,62 doğruluk oranına ulaştığı rapor edilmiştir. Çalışmada ayrıca öz niteliklerin önem düzeyleri analiz edilmiş; yaş, ejeksiyon fraksiyonu, serum kreatinin ve takip süresi değişkenlerinin sağ kalım tahmini için en kritik parametreler olduğu doğrulanmıştır. Buna karşın; anemi, diyabet, cinsiyet ve sigara kullanımı gibi düşük öneme sahip 4 öz nitelik modelleme sürecinden elenmiştir. Düşük öneme sahip değişkenlerin elenmesiyle geriye kalan 9 öz nitelik üzerinden yapılan

testlerde, ETC algoritmasının %92,62 doğruluk oranına ulaştığı ve çalışmadaki en başarılı model olduğu rapor edilmiştir. Öznitelik eleme sürecinin, bilgi yoğunluğunu azaltarak modellerin en yüksek performansını korumasını sağladığı, hatta Karar Ağaçları ve Gradyan Artırma gibi belirli algoritmalarda doğrudan başarı artışına yol açtığı saptanmıştır. Ishaq vd. (2021), elde ettikleri sonuçların veri madenciliği teknikleri ve dengeleme yöntemlerinin entegrasyonu sayesinde klinik karar verme süreçlerinde yüksek güvenilirlikle kullanılabileceğini savunmaktadırlar.

Bu çalışma ise, kalp yetmezliği sağkalım tahmini üzerine literatürdeki en temel çalışmalardan biridir (Chicco ve Jurman, 2020a). Yazarlar, 299 hastadan oluşan veri setini kullanarak, makine öğrenmesi modellerinin karmaşık tıbbi kayıtlar içerisinde en belirleyici risk faktörlerini ayırt etme yeteneğini incelemişlerdir. Yapılan öznitelik sıralama (feature ranking) analizleri sonucunda, sadece serum kreatinin ve ejeksiyon fraksiyonu değişkenlerinin birlikte kullanılmasının, hastaların sağkalım durumunu tahmin etmek için yeterli ve güçlü birer gösterge olduğu saptanmıştır.

Çalışmada, Rastgele Orman ve Destek Vektör Makineleri (SVM) gibi algoritmalar kullanılarak yapılan deneylerde, takip süresi değişkeni dışarıda bırakıldığında bile bu iki parametrenin yüksek öngörü kapasitesine sahip olduğu vurgulanmıştır. Chicco ve Jurman (2020a), biyokimyasal bir ölçüm olan serum kreatinin ile kalp fonksiyonlarını temsil eden ejeksiyon fraksiyonunun, diğer 11 klinik değişkene kıyasla mortalite riskini çok daha yüksek bir doğrulukla yansıttığını savunmaktadır. Bu bulgular, klinik pratikte daha az değişkenle daha hızlı ve etkili karar destek sistemlerinin oluşturulabileceğini kanıtlaması bakımından büyük önem taşımaktadır.

Newaz vd. (2021)'nin çalışmasında ise, kalp yetmezliği hastalarının klinik kayıtlarını kullanarak karar destek sistemleri geliştirmeyi amaçlayan bir çalışmadır (Newaz vd., 2021). Yazarlar, Pakistan'daki Faisalabad Kardiyoloji Enstitüsü'nden toplanan 299 kişilik veri seti üzerinde, tahmin doğruluğunu artırmak amacıyla özyinelemeli öznitelik eleme (Recursive Feature Elimination, RFE) tekniğini uygulamışlardır. Çalışmada, öznitelik seçimi ile model karmaşıklığının azaltılması ve en anlamlı prognostik değişkenlerin belirlenmesi hedeflenmiştir.

Araştırma sonucunda, Rastgele Orman algoritmasının diğer sınıflandırıcılara kıyasla üstün bir performans sergilediği saptanmıştır. RFE yöntemiyle yapılan analizler, ejeksiyon fraksiyonu, serum kreatinin, yaş ve takip süresi (time) değişkenlerinin sağkalım tahmininde en etkili "ana sürücüler" (key drivers) olduğunu ortaya koymuştur. Newaz vd. (2021), modellerinin dengesiz veri kümesi üzerinde yüksek bir hassasiyetle çalıştığını

belirterek, elde edilen sonuçların klinisyenlerin tedavi yoğunluğunu belirleme süreçlerinde güvenilir bir rehber olabileceğini savunmuşlardır.

Qadri vd. (2024) çalışmasında, kalp yetmezliği sağkalım tahmininde yenilikçi yaklaşımları inceleyen güncel çalışmalardan biridir (Qadri vd., 2024). Çalışmada, geleneksel makine öğrenmesi yöntemlerinin ötesine geçerek, Transfer Öğrenme (Transfer Learning) tabanlı olasılıksal özniteliklerin kullanıldığı hibrit bir modelleme çerçevesi önermişlerdir. Çalışmada, önceden eğitilmiş modellerden elde edilen bilgilerin klinik veri setine aktarılmasıyla, modelin genelleştirme yeteneğinin ve tahmin doğruluğunun artırılması hedeflenmiştir.

Araştırma kapsamında, çeşitli denetimli öğrenme algoritmaları ve topluluk (ensemble) yöntemleri karşılaştırmalı olarak test edilmiştir. Elde edilen bulgular, önerilen transfer öğrenme destekli modelin %97,5 doğruluk ve %98 F1-skoru gibi oldukça yüksek performans değerlerine ulaştığını göstermiştir. Yazarlar, bu yaklaşımın özellikle küçük ölçekli klinik veri setlerinde karşılaşılan sınırlılıkları aşmada etkili olduğunu ve karmaşık hastalık paternlerini anlamada derin öğrenme tabanlı öznitelik çıkarımının klasik yöntemlere göre üstünlük sağladığını vurgulamışlardır. Bu çalışma, yapay zekadaki modern mimarilerin kalp yetmezliği gibi kritik tıbbi durumlarda son derece hassas karar destek mekanizmaları oluşturabileceğini kanıtlamaktadır.

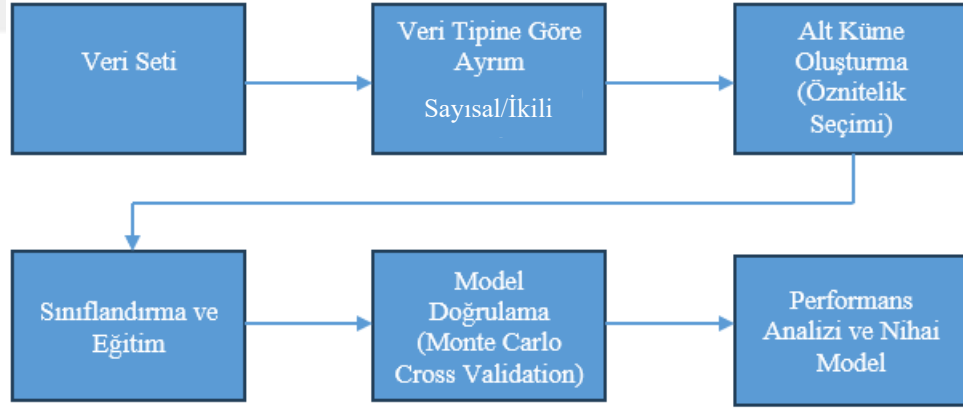
Literatürde yer alan bu çalışmaların bulguları, araştırma sonuçları ve tartışma bölümünde Tablo 4.3 ve Tablo 4.4'de çalışmamızdan elde edilen sonuçlarla birlikte karşılaştırmalı olarak sunulmuştur.

3. MATERYAL VE YÖNTEM

Bu bölümde, çalışmada kullanılan veri seti ve uygulanan makine öğrenmesi yöntemleri ayrıntılı olarak açıklanmaktadır. Öncelikle veri setinin genel özellikleri tanıtılmakta, ardından kullanılan sınıflandırma algoritmaları ve deneysel düzen sunulmaktadır.

Bu çalışmada izlenen metodolojik akış, kalp yetmezliği veri setindeki nümerik ve kategorik değişkenlerin bir arada bulunduğu heterojen yapıyı yönetmek ve sağkalım analizindeki klinik gereksinimleri karşılamak üzere kurgulanmıştır. Öznitelik seçimi sürecinde, nümerik ve kategorik alt uzaylardan en yüksek skorlu öznitelikler eş zamanlı olarak seçilerek 2, 4, 6 gibi simetrik artış gösteren boyutlu alt kümeler oluşturulmuştur. Modellerin genelleme yeteneğini ölçmek amacıyla, %70 eğitim ve %30 test ayrımı üzerinden 100 bağımsız iterasyonlu Monte Carlo doğrulama stratejisi uygulanmıştır. MATLAB ortamında eğitilen sınıflandırma algoritmalarından elde edilen ortalama performans metrikleri analiz edilerek, klinik kullanım potansiyeli en yüksek olan nihai model önerisi belirlenmiştir.

Şekil 3.1’de yapılan çalışmanın şekil işleyişi verilmiştir.



Şekil 3.1. Veri tipine dayalı öznitelik seçimi ve model doğrulama sürecinin genel akışı

3.1. Veri Seti

Bu çalışmada kullanılan veri seti, Ahmad ve çalışma arkadaşları tarafından toplanan ve kalp yetmezliği hastalarına ait klinik bilgileri içeren açık erişimli bir veri setidir. Veri seti, 2015 yılı içerisinde Pakistan’ın Faisalabad kentinde Faisalabad Kardiyoloji Enstitüsü (Faisalabad Institute of Cardiology) ve Allied Hastanesi’nde

(Allied Hospital) takip edilen toplam 299 kalp yetmezliği hastasına ait kayıtları kapsamaktadır (Ahmad vd., 2017).

Bu çalışmada kullanılan veri seti, kalp yetmezliği tanısı konulmuş 299 hastaya ait toplam 13 klinik ve demografik özniteligi içermektedir. Toplam 299 kaydın 194'ü erkek, 105'i kadındır. Tüm hastaların yaşları 40 yaş üzerindedir.

Tablo 3.1'de veri setinde yer alan tüm öznitelikler, değişken türleri ve kısa açıklamaları sunulmaktadır.

Tablo 3.1. Veri setindeki özniteliklerin anlamları, ölçüm birimleri ve aralıkları

	Öznitelik	Açıklama	Ölçüm Birimi	Aralık
1	Yaş (Age)	Hastanın yaşı	Yıl	[40 - 95]
2	Anemi (Anaemia)	Kırmızı kan hücreleri veya hemoglobindeki azalma	İkili (Binary)	0 - 1
3	Yüksek tansiyon (High blood pressure, HBP)	Hipertansiyon varlığı	İkili	0 - 1
4	Kreatin fosfokinaz (Creatinine phosphokinase, CPK)	Kandaki CPK enzimi düzeyi	mcg/L	[23 - 7861]
5	Diyabet (Diabetes)	Diyabet hastalığının varlığı	İkili	0 - 1
6	Ejektasyon fraksiyonu (Ejection fraction, EF)	Kalbin her kasılmada pompaladığı kan yüzdesi	Yüzde (%)	[14 - 80]
7	Cinsiyet (Sex)	Hastanın cinsiyeti (kadın/erkek)	İkili	0 - 1
8	Trombositler (Platelets)	Kandaki trombosit sayısı	kilotrombosit/mL	[25,01 - 850,00]
9	Serum kreatinin (Serum creatinine, SC)	Kandaki kreatinin düzeyi	mg/dL	[0,50 - 9,40]
10	Serum sodyumu (Serum sodium, SS)	Kandaki sodyum düzeyi	mEq/L	[114 - 148]
11	Sigara kullanımı (Smoking)	Sigara kullanımı durumu	İkili	0 - 1
12	Takip süresi (Time)	Takip süresi	Gün	[4 - 285]
13	Ölüm olayı (Death event/ hedef değişken)	Takip süresince ölüm durumu	İkili	0 - 1

mcg/L: mikrogram/litre; mL: mililitre; mEq/L: miliekivalan/litre.

Geçmişte sol ventrikül sistolik disfonksiyonu ve kalp yetmezliği öyküsü bulunan tüm 299 hasta, New York Kalp Cemiyeti (New York Heart Association) sınıflandırmasına göre sınıf III veya IV grubunda yer almaktadır. Öznitelikler, hastaların demografik bilgileri, eşlik eden klinik durumları ve laboratuvar ölçümlerini kapsayacak şekilde yapılandırılmıştır.

Bu öznitelikler, kalp yetmezliğinin klinik seyri ve hasta sağkalımı üzerinde etkili olduğu bilinen faktörleri kapsamaktadır. Özniteliklerin bu şekilde gruplandırılması, makine öğrenmesi modellerinin hem eğitilmesi hem de sonuçların klinik olarak yorumlanabilirliği açısından önem taşımaktadır.

3.2. Deneyisel Tasarım

Bu çalışmada, farklı öznitelik seçimi ve sınıflandırıcı kombinasyonlarının performanslarını değerlendirmek amacıyla kapsamlı bir deneyisel tasarım oluşturulmuştur. Modellerin genellenebilirliğini değerlendirmek için Monte Carlo doğrulama stratejisi kullanılmıştır. Veri seti her iterasyonda stratifiye rastgele örnekleme ile %70 eğitim ve %30 test olacak şekilde ayrılmıştır. Bu işlem 100 bağımsız iterasyon halinde tekrarlanmıştır. Nihai performans değerleri bu iterasyonlardan elde edilen aritmetik ortalama sonuçlar üzerinden raporlanmıştır.

Bu çalışmada sunulan tüm makine öğrenmesi deneyleri, MATLAB (MathWorks, Natick, MA, ABD) yazılım ortamı kullanılarak gerçekleştirilmiştir. Öznitelik seçimi yöntemleri ve sınıflandırma algoritmaları, MATLAB bünyesinde yer alan Statistics and Machine Learning Toolbox fonksiyonları aracılığıyla uygulanmıştır. Model eğitim ve test süreçlerinde hem tüm öznitelikler kullanılarak hem de time özniteliği çıkarılarak ayrı ayrı deneyler yürütülmüştür. Deneyler MATLAB ortamında tekrarlı olarak gerçekleştirilmiş ve performans ölçütleri otomatik olarak hesaplanmıştır. Bu yaklaşım, deneylerin tutarlı ve tekrarlanabilir biçimde yürütülmesini sağlamıştır (MathWorks, 2024).

3.3. Veri Ön İşleme

Bu çalışmada kullanılan veri seti, makine öğrenmesi modellerine doğrudan uygulanabilecek şekilde yapılandırılmıştır. Yapılan incelemeler sonucunda veri setinde eksik gözlem bulunmadığı belirlenmiş ve bu nedenle herhangi bir eksik veri tamamlama işlemi uygulanmamıştır. İkili değişkenler özgün hâllerıyla modele dâhil edilirken, sayısal değişkenler için herhangi bir ek dönüştürme veya normalizasyon işlemi gerçekleştirilmemiştir.

3.4. Öznitelik Seçimi Yöntemleri

Öznitelik seçimi, uygulayıcılar tarafından boyut indirgeme için yaygın olarak kullanılan bir tekniktir. Bu yöntem, belirli bir ilgililik değerlendirme ölçütüne göre özgün öznitelikler arasından ilgili olan küçük bir öznitelik alt kümesinin seçilmesini amaçlar. Bu yaklaşım genellikle daha iyi öğrenme performansı (örneğin, sınıflandırma için daha yüksek doğruluk), daha düşük hesaplama maliyeti ve daha iyi model yorumlanabilirliği sağlar (Tang vd., 2014).

Bu çalışmada, kalp yetmezliği veri seti üzerinde farklı öznitelik seçimi stratejilerinin sınıflandırma performansı üzerindeki etkilerini incelemek amacıyla filtre, sarmalayıcı ve gömülü yaklaşımları temsil eden çeşitli öznitelik seçimi yöntemleri uygulanmıştır. Elde edilen öznitelik alt kümeleri, farklı makine öğrenmesi tabanlı sınıflandırıcılar kullanılarak değerlendirilmiştir.

Ayrıca, öznitelik seçimi yöntemlerinin farklı öznitelik türleri üzerindeki davranışlarını daha sağlıklı değerlendirebilmek amacıyla, sayısal (numerical) ve ikili (binary) öznitelikler ayrı ayrı ele alınmıştır. Ön deneysel analizler sonucunda, bazı öznitelik seçimi yöntemlerinin sayısal öznitelikler üzerinde, bazılarının ise ikili öznitelikler üzerinde daha başarılı sonuçlar ürettiği gözlemlenmiştir. Bu doğrultuda, NCA gibi yöntemlerin sayısal öznitelikler üzerinde daha iyi performans sergilediği durumlarda yalnızca sayısal öznitelikler kullanılarak modeller eğitilmiş; Ki-kare (Chi-Square) gibi ikili öznitelikler üzerinde daha etkili olduğu bilinen yöntemler için ise yalnızca ikili öznitelikler dikkate alınmıştır. Hem sayısal hem de ikili öznitelikler üzerinde etkili performans gösterebilen yöntemler ise tüm öznitelikler birlikte kullanılarak değerlendirilmiştir. Bu yaklaşım, öznitelik seçimi ve sınıflandırma süreçlerinin yöntemlerin güçlü olduğu veri türleri üzerinde ele alınmasını sağlayarak, deneysel sonuçların daha güvenilir ve tutarlı biçimde analiz edilmesine olanak tanımıştır. Bu bağlamda, çalışmada yer verilen öznitelik seçimi yöntemlerinin kullanılan veri türlerine göre dağılımı Tablo 3.2'de özetlenmiştir.

Tablo 3.2. Öznitelik seçimi yöntemlerinin kullanılan veri türlerine göre dağılımı

Öznitelik Seçimi Yöntemi	Kullanılan Veri Türü
Ki-kare	İkili
ReliefF	Sayısal
NCA	Sayısal
LASSO	Sayısal
SFS	Sayısal
TreeBagger	Sayısal + İkili
Mrmr	Sayısal + İkili
Tree	Sayısal + İkili

Bu çalışmada öznitelik seçimi yöntemleri, çalışma prensipleri ve teorik varsayımları dikkate alınarak farklı veri türleri üzerinde uygulanmıştır. Ki-kare tabanlı öznitelik seçimi yöntemi, ikili değişkenler arasındaki istatistiksel bağımlılığı ölçmeye dayalı bir yaklaşım olduğundan yalnızca ikili öznitelikler kullanılarak değerlendirilmiştir (Liu ve Setiono, 1997).

Buna karşılık, ReliefF, Komşuluk bileşeni analizi (Neighborhood Component Analysis, NCA), LASSO ve Ardışık öznitelik seçimi (Sequential Feature Selection, SFS) yöntemleri; mesafe hesaplama, ağırlıklandırma ve optimizasyon temelli yapıları nedeniyle yalnızca sayısal öznitelikler üzerinde uygulanmıştır (Goldberger vd., 2004; Kononenko, 1994; Tibshirani, 1996).

Ağaç tabanlı öznitelik seçimi yöntemleri ile Minimum Tekrar - Maksimum İlişki (Minimum redundancy-maximum relevance, mRMR) yöntemi ise hem sayısal hem de ikili öznitelikleri birlikte değerlendirebilen esnek yapıları sayesinde tüm öznitelikler kullanılarak uygulanmıştır (Breiman, 2001; Peng vd., 2005).

Aşağıda çalışmada kullanılan öznitelik seçimi yöntemleri ayrıntılı olarak açıklanmaktadır.

3.4.1. Komşuluk bileşeni analizi (Neighborhood component analysis, NCA)

NCA, k-NN algoritmasından türetilmiş, parametrik olmayan (yani belirli bir parametre veya dağılıma bağlı olmayan) bir öznitelik seçimi yöntemidir. En iyi öznitelik alt kümelerini seçmek için bir öznitelik ağırlıklandırma şeması kullanır (Goldberger vd., 2004).

NCA algoritması, veriyi düşük boyutlu bir alt uzaya projekte eden bir matris A öğrenerek en yakın komşu sınıflandırıcı performansını artırmayı amaçlar. Bu matrisi kullanarak, yüksek boyutlu veriler daha az boyutlu bir uzaya taşınır ve sınıflandırma daha etkili hale gelir. Veri noktaları, birbirlerine olan mesafeleri ile ters orantılı bir olasılıkla komşu olarak atanır. Bu olasılık Denklem 3.1’de gösterilmiştir.

$$P_{ij} = \frac{\exp(-||Ax_i - Ax_j||^2)}{\sum_{k \neq i} \exp(-||Ax_i - Ax_k||^2)} \quad (3.1)$$

Bu eşitlikte x_i ve x_j , sırasıyla i’inci ve j’inci veri örneklerini (öznitelik vektörlerini) temsil etmektedir. A, NCA algoritması tarafından öğrenilen dönüşüm (projeksiyon) matrisidir ve yüksek boyutlu veriyi daha düşük boyutlu bir alt uzaya taşımaktadır. Ax_i ve Ax_j , veri örneklerinin dönüştürülmüş uzaydaki karşılıklarını göstermektedir. $||Ax_i - Ax_j||^2$, dönüştürülmüş uzaydaki iki örnek arasındaki Öklidyen mesafenin karesini ifade etmektedir. $\exp(.)$, üstel (exponential) fonksiyondur ve birbirine yakın örnekler daha yüksek olasılık atanmasını sağlamaktadır. P_{ij} , i’inci örneğin j’inci örneği komşu olarak seçme olasılığını ifade etmektedir. Paydadaki $\sum_{k \neq i}$ terimi, i’inci örnek dışındaki tüm örnekler üzerinden hesaplanan normalizasyon terimidir ve olasılıkların toplamının bire eşit olmasını sağlamaktadır. Projeksiyon matrisi

A, doğru sınıflandırma olasılığını maksimize etmek amacıyla gradyan tabanlı yöntemlerle optimize edilir. Ancak fonksiyonun konveks olmaması nedeniyle başlangıç değerleri ve optimizasyon teknikleri dikkatle seçilmelidir (Siddique, 2019).

3.4.2. ReliefF

ReliefF, Relief algoritmasının güncellenmiş bir versiyonudur ve gürültülü (yani hatalar veya parazit içeren) ve eksik verilerle başa çıkabilir. Algoritmanın temel fikri, özniteliklere, bu özniteliklerin sınıflandırma için ne kadar önemli olduğuna göre ağırlık atamaktır.

ReliefF algoritması ilk olarak, referans bir örnekten yola çıkarak, aynı sınıfa ait en yakın örnekleri (NHi) ve farklı sınıfa ait en yakın örnekleri (NMi) bulur. Bu işlem sırasında, örnekler arasındaki mesafeyi ölçmek için Manhattan uzaklığı kullanılır. Eğer referans örneğinden (Xi) NMi'ye olan mesafe, NHi'ye olan mesafeden daha küçükse, bu durumda ilgili özniteliğin önemi artar (Z. Wang vd., 2016).

ReliefF, çok sınıflı sınıflandırma problemleri için geliştirilmiş olup, örneklerin en yakın komşularıyla olan mesafelerini dikkate alan matematiksel bir formül aracılığıyla öznitelik ağırlıklarını hesaplayan bir algoritmadır.

ReliefF'in matematiksel ifadesi Denklem 3.2'de verilmiştir (Dou vd., 2022).

$$W_F(R_i) = W_F(R_i) + \sum_{j=1}^k \frac{\text{diff}(F, R_i, H_j)}{m \cdot k} + \sum_{C \neq \text{class}(R_i)} P(C) \left(1 - P(\text{class}(R_i))\right) \sum_{j=1}^k \frac{\text{diff}(F, R_i, M_j^C)}{m \cdot k} \quad (3.2)$$

Burada $W_F(R_i)$ öznitelik F için R_i örneğinin ağırlığı, H_j aynı sınıftan k en yakın komşu (hit), M_j^C farklı sınıflardan k en yakın komşu (miss) ve C sınıfı $\text{diff}(F, R_i, H_j)$ öznitelik F ile R_i örneği ve H_j arasındaki mesafe, $\text{diff}(F, R_i, M_j^C)$ öznitelik F ile R_i örneği ve M_j^C arasındaki mesafe, m kullanıcı tarafından belirlenen iterasyon sayısı k, aranan en yakın komşu sayısı P(C), C sınıfının prior (öncelikli) olasılığı $P(\text{class}(R_i))$, R_i örneğinin ait olduğu sınıfın olasılığı göstermektedir.

3.4.3. Minimum tekrar - maksimum ilişki (Minimum redundancy-maximum relevance, mRMR)

mRMR algoritması, seçilen özniteliklerin sınıfa olan ilişkisinin yüksek olmasını sağlarken, aynı zamanda bu öznitelikler arasındaki benzerlikleri (tekrarı) en aza

indirmeye çalışır (Peng vd., 2005). mRMR yöntemi, öz nitelikleri tek tek inceleyerek, her bir öz niteliğin diğerleriyle olan ilişkisini belirlemek için karşılıklı bilgiyi kullanır. Bu sayede, hangi öz niteliklerin sınıf etiketine (hedef değişken) daha fazla katkı sağladığını belirlemeye çalışır (Anuragi vd., 2022).

mRMR algoritması, önemli öz nitelikleri seçmek için iki kriteri birleştirir:

Maksimum İlgililik (Max-Relevance): Seçilen öz niteliklerin hedef sınıf C ile olan karşılıklı bilgisinin maksimize edilmesi hedeflenmektedir. Bu kriter, seçilen öz niteliklerin sınıf etiketini ayırt etme gücünü artırmayı amaçlamaktadır. Denklem 3.3'de tanımlanmaktadır.

$$\frac{1}{|S|} \sum_{f_i \in S} I(f_i; C) \quad (3.3)$$

Bu eşitlikte S seçilen öz nitelikler kümesini, f_i öz nitelik kümesindeki i'inci öz niteliği, C hedef sınıf değişkenini, $I(f_i; C)$ f_i öz niteliği ile hedef sınıf C arasındaki karşılıklı bilgiyi (mutual information) ifade etmektedir.

Minimum Tekrar (Min-Redundancy): Seçilen öz nitelikler arasındaki bağımlılığın ve fazlalığın en aza indirilmesi hedeflenmektedir. Bu kriter Denklem 3.4'de verilmiştir.

$$\frac{1}{|S|^2} \sum_{f_i, f_j \in S} I(f_i; f_j) \quad (3.4)$$

Bu eşitlikte $I(f_i; f_j)$, f_i ve f_j öz nitelikleri arasındaki karşılıklı bilgiyi ifade etmektedir. Bu ölçüt, birbirine benzer bilgi taşıyan öz niteliklerin birlikte seçilmesini engelleyerek daha ayırt edici bir öz nitelik kümesi oluşturulmasını sağlar.

mRMR Kriteri: Maksimum ilgililik ve minimum tekrar kriterlerinin birlikte kullanılmasıyla mRMR ölçütü Denklem 3.5'de tanımlanmaktadır.

$$\text{mRMR-MID: } \max_S \left(\frac{1}{|S|} \sum_{f_i \in S} I(f_i; C) - \frac{1}{|S|^2} \sum_{f_i, f_j \in S} I(f_i; f_j) \right) \quad (3.5)$$

Bu eşitlikte birinci terim, seçilen öz niteliklerin hedef sınıfla olan ilişkisini maksimize etmeyi, ikinci terim ise öz nitelikler arasındaki tekrar bilgisini minimize etmeyi amaçlamaktadır. Bu sayede mRMR algoritması, hedef değişkenle yüksek derecede ilişkili ve birbirleriyle düşük derecede bağımlı olan öz nitelikleri seçerek sınıflandırma performansını artırmayı hedeflemektedir (Peng vd., 2005).

3.4.4. Ki-kare (Chi-Square)

Ki-kare yaklaşımı, öne çıkan bir öz nitelik seçimi yöntemidir. Bu, gözlemlenen değerlerin beklenen sonuçlardan ne kadar önemli ölçüde farklılaştığını değerlendirmek için kullanılan bir istatistiksel testtir ve tahmin edici değişkeni belirlemek için kullanılır

(Shanmugarajeshwari ve Lawrance, 2016). Hesaplama sırasında, ki-kare testi yalnızca örnekler arasındaki gözlemlenen farklılıklar hakkında bilgi sağlamakla kalmaz, aynı zamanda belirli kategoriler arasındaki anlamlı farklılıklar hakkında da detaylı bilgi sunar. Özellikle, χ^2 tüm örnekler için hesaplanır ve bu özniteliklerin hedef üzerindeki etkisini keşfetmek için kullanılır. Bu etki, daha sonra özniteliklerin karar üzerindeki önemine, yani etkisinin büyüklüğüne göre haritalanır (Bryant ve Satorra, 2012; McHugh, 2013).

Ki-kare testi, gözlemlenen ve beklenen frekanslar arasındaki farkı ölçer. Bu denklemler Denklem 3.6 ve Denklem 3.7’de verilmiştir.

Formül:

$$\chi^2 = \sum \frac{(O-E)^2}{E} \quad (3.6)$$

O: Gözlemlenen değer, E: Beklenen değeri temsil eder.

Beklenen Değer:

$$E = \frac{MR \cdot MC}{n} \quad (3.7)$$

MR satır toplamı, MC sütun toplamı ve n veri kümesindeki toplam örnek sayısıdır.

Genel χ^2 : Tüm kategoriler için hesaplanan χ^2 değerlerinin toplamıdır ve hedef değişkenle olan ilişkinin gücünü gösterir (Hou vd., 2024).

3.4.5. Ardışık öznitelik seçimi (Sequential feature selection, SFS)

SFS, yinelemeli bir öznitelik seçimi yöntemidir ve boş bir küme ile başlayarak model için en iyi puanı veren bir özniteliği ekler. Ardından, başka bir öznitelik test edilip önceki alt kümeye eklenir. Sonrasında, yeni alt kümenin analizi yapılır. Bu öznitelik, sınıflandırma doğruluğunu maksimize ediyorsa eklenir. Bu adım, istenilen öznitelik kümesi elde edilene kadar ardışık olarak devam eder (Chandrashekar ve Sahin, 2014).

Bu işlem matematiksel olarak Denklem 3.8’de ifade edilmiştir.

$$S_{k-1} \cup \arg \max_{x \in X \setminus S_{k-1}} J(S_{k-1} \cup \{x\}) \quad (3.8)$$

S_k , k. adımda seçilen öznitelik kümesini, X, tüm öznitelikler kümesini, $J(\cdot)$, sınıflandırma performansını ölçen amaç (değerlendirme) fonksiyonunu göstermektedir (Pudil vd., 1994).

3.4.6. LASSO (Least absolute shrinkage and selection operator)

LASSO, gömülü (embedded) bir öznitelik seçimi yöntemi olup, model eğitimi sırasında öznitelik seçimini düzenleme (regularization) yoluyla gerçekleştirmektedir.

LASSO yöntemi, mutlak değer cezası (L1 normu) kullanarak bazı özniteliklerin katsayılarını sıfıra indirir ve böylece modele katkısı olmayan veya düşük olan özniteliklerin elenmesini sağlar. Bu yaklaşım, özellikle yüksek boyutlu veri setlerinde aşırı uyum (overfitting) problemini azaltmakta ve daha sade, yorumlanabilir modeller elde edilmesine katkı sağlamaktadır (Tibshirani, 1996).

LASSO optimizasyon problemi Denklem 3.9'da verilmiştir.

$$\min_{\beta} \left\{ \sum_{i=1}^n (y_i - X_i \beta)^2 + \lambda \sum_{j=1}^p |\beta_j| \right\} \quad (3.9)$$

Burada λ düzenleme parametresi olup, ceza teriminin etkisini kontrol eder. L_1 norm temelli bu ceza fonksiyonu, bazı katsayıların tam olarak sıfıra eşitlenmesine yol açarak otomatik değişken seçimi yapılmasını sağlar (Tibshirani, 1996).

3.4.7. Tree tabanlı öznitelik seçimi

Karar ağacı (Decision Tree) tabanlı yöntemler, sınıflandırma süreci sırasında öznitelik seçimini doğal olarak gerçekleştiren gömülü yaklaşımlar arasında yer almaktadır. Bu yöntemlerde, özniteliklerin önemi; düğüm bölünmelerinde kullanılan saflık ölçütleri (örneğin Gini indeksi veya entropi) üzerinden değerlendirilmektedir. Ağaç yapısında üst seviyelerde yer alan ve daha fazla bölünmeye katkı sağlayan öznitelikler, sınıflandırma sürecinde daha yüksek öneme sahip kabul edilmektedir. Bu sayede, model performansına en fazla katkı sağlayan özniteliklerin belirlenmesi mümkün olmaktadır (Salzberg, 1994).

Karar ağacı tabanlı yöntemlerde en yaygın kullanılan saflık ölçütleri Gini indeksi ve entropidir. Gini indeksi Denklem 3.10'da entropi ise Denklem 3.11'de verilmiştir.

$$Gini(t) = 1 - \sum_{k=1}^K p_k^2 \quad (3.10)$$

$$H(t) = - \sum_{k=1}^K p_k \log p_k \quad (3.11)$$

Burada p_k , t düğümündeki k . sınıfa ait örneklerin oranını ifade etmektedir. Bir öznitelik X_j ile yapılan bölünme sonucunda elde edilen saflık azalması denklem 3.12'de verilmiştir.

$$\Delta I(X_j) = I(t) - \sum_{m=1}^M \frac{N_m}{N_t} I(t_m) \quad (3.12)$$

Bu ifadede $I(t)$ ana düğümdeki saflık ölçütünü, t_m alt düğümleri, N_t ana düğümdeki örnek sayısını ve N_m alt düğümlerdeki örnek sayısını göstermektedir. En büyük saflık azalmasını sağlayan öznitelik bölünme için seçilmekte olup, bu sayede sınıflandırma performansına en fazla katkı sağlayan özniteliklerin belirlenmesi mümkün olmaktadır (Hastie vd., 2009).

3.4.8. Topluluk tabanlı ağaç yöntemi (TreeBagger)

TreeBagger, birden fazla karar ağacının rastgele örneklenmiş veri alt kümeleri üzerinde eğitilmesiyle oluşturulan topluluk (ensemble) tabanlı bir yaklaşımdır. Bu yöntemde, her bir karar ağacı farklı öznitelik alt kümeleri üzerinde eğitilmekte ve özniteliklerin önemi, tüm ağaçlar boyunca ortalama katkıları üzerinden hesaplanmaktadır. TreeBagger yaklaşımı, tekil karar ağaçlarına kıyasla daha kararlı ve genellenebilir öznitelik önem ölçümleri sunmakta, böylece daha güvenilir bir öznitelik seçimi sağlamaktadır (Breiman, 2001).

TreeBagger yöntemi, bootstrap örnekleme ile oluşturulan B adet karar ağacından oluşan bir topluluk modelidir. Her bir karar ağacı farklı eğitim örnekleri ve rastgele seçilmiş öznitelik alt kümeleri üzerinde eğitilir. Modelin genel çıktısı Denklem 3.13’de verilmiştir.

$$\hat{f}(x) = \frac{1}{B} \sum_{b=1}^B T_b(x) \quad (3.13)$$

Burada $T_b(x)$, b . karar ağacının tahminini ve B toplam ağaç sayısını göstermektedir.

Öznitelik önemi, tüm ağaçlarda ilgili özneliğin sağladığı ortalama saflık azalması üzerinden hesaplanır. Bu hesaplama Denklem 3.14’de verilmiştir.

$$VI(X_j) = \frac{1}{B} \sum_{b=1}^B \sum_{t \in T_b} \Delta I_t(X_j) \quad (3.14)$$

Burada $\Delta I_t(X_j)$, b . ağaçtaki t düğümünde X_j özneliği ile yapılan bölünme sonucu elde edilen saflık azalmasını ifade etmektedir. Bu yöntem, tekil karar ağaçlarına kıyasla daha kararlı ve genellenebilir öznitelik önem ölçümleri sunmaktadır (Breiman, 2001).

3.5. Kullanılan Makine Öğrenmesi Algoritmaları

Öznitelik seçimi aşamasının ardından elde edilen öznitelik alt kümeleri, hasta sağkalımını tahmin etmeye yönelik olarak çeşitli makine öğrenmesi tabanlı sınıflandırma algoritmaları kullanılarak değerlendirilmiştir. Bu çalışmada, farklı öğrenme prensiplerine dayanan ve literatürde yaygın olarak kullanılan sınıflandırma algoritmaları tercih edilmiştir.

3.5.1. k-En yakın komşular (k-Nearest neighbors, k-NN)

k-NN, sınıflandırma ve regresyon için kullanılan parametrik olmayan bir yöntemdir. Bu yöntemde girdi, her bir uzayda bulunan en yakın k örnekten oluşmaktadır.

Verilen N adet eğitim vektörü için, k -en yakın komşu (k -NN) algoritması en yakın k adet komşuyu belirler. Bu komşular arasındaki en yakın olanlar dikkate alınır ve buna bağlı olarak test örneği ilgili sınıfa atanır (Pandey ve Jain, 2017).

İki örnek arasındaki uzaklık genellikle Öklid mesafesi kullanılarak hesaplanır. Bu hesaplama Denklem 3.15’de ifade edilmiştir.

$$d(x_i, x_j) = \sqrt{\sum_{m=1}^p (x_{im} - x_{jm})^2} \quad (3.15)$$

Burada x_i ve x_j , p boyutlu öznitelik vektörlerini temsil etmektedir.

Bir test örneği x için en yakın k komşu kümesi $N_k(x)$ olarak tanımlanır. Sınıflandırma işlemi, bu komşuların sınıf etiketleri dikkate alınarak çoğunluk oylaması (majority voting) yöntemiyle gerçekleştirilir. Denklem 3.16’da bu hesaplama verilmiştir.

$$\hat{y}(x) = \arg \max_c \sum_{x_i \in N_k(x)} I(y_i = c) \quad (3.16)$$

Burada $I(.)$ gösterge fonksiyonunu, y_i eğitim örneklerinin sınıf etiketlerini ve c olası sınıf etiketlerini ifade etmektedir. En fazla sayıda komşuya sahip olan sınıf, test örneğinin tahmin edilen sınıfı olarak atanır (Cover ve Hart, 1967).

3.5.2. Destek vektör makineleri (Support vector machine, SVM)

Destek Vektör Makineleri, 1995 yılında keşfedilmiş popüler bir sınıflandırma tekniğidir. Bir sınıflandırma görevi genellikle veriyi eğitim ve test setlerine ayırmayı içerir. Eğitim setindeki her örnek bir “hedef değer” ve birkaç “öznitelik” içerir. SVM’nin amacı, yalnızca test verisinin öznitelikleri kullanılarak, test verisinin hedef değerlerini tahmin eden bir model üretmektir (Nugrahaeni ve Mutijarsa, 2016).

SVM, farklı sınıflara ait örnekleri en iyi şekilde ayıran ve marjini maksimize eden bir ayırıcı hiper düzlem elde etmektir. Lineer olarak ayrılabilir durumda SVM, Denklem 3.17’de ki optimizasyon problemini çözerek sınıflandırıcıyı belirler.

$$\min_{\omega, b, \xi} \frac{1}{2} \|\omega\|^2 + C \sum_{i=1}^N \xi_i \quad (3.17)$$

koşuluyla $y_i(\omega^T x_i + b) \geq 1 - \xi_i$, $\xi_i \geq 0$, $i = 1, 2, \dots, N$ burada ω ağırlık vektörüne, b bias terimini, x_i eğitim örneklerini ve $y_i \in \{-1, +1\}$ sınıf etiketlerini temsil etmektedir.

ξ_i gevşeme değişkenleri (slack variables) olup hatalı sınıflandırmaya izin verirken, C parametresi hata ile marjin genişliği arasındaki dengeyi kontrol etmektedir.

Doğrusal olarak ayrılamayan durumlarda ise SVM, veriyi daha yüksek boyutlu bir uzaya taşıyan çekirdek (kernel) fonksiyonları yardımıyla sınıflandırma yapmaktadır. Bu durumda karar fonksiyonu Denklem 3.18’de ifade edilmiştir.

$$f(x) = \sum_{i=1}^N \alpha_i y_i K(x_i, x) + b \quad (3.18)$$

Burada $K(x_i, x)$ kernel fonksiyonunu, α_i ise Lagrange çarpanlarını temsil etmektedir. Bu yaklaşım sayesinde SVM, doğrusal olmayan sınıflandırma problemlerinde de etkin sonuçlar üretebilmektedir (Hastie vd., 2009).

3.5.3. Naive bayes

Naive Bayes (NB), denetimli sınıflandırma problemlerinde yaygın olarak kullanılan, olasılıksal bir sınıflandırma yöntemidir. Yöntem, Bayes teoremine dayanmakta ve özniteliklerin sınıf koşullu olarak birbirinden bağımsız olduğu varsayımını kabul etmektedir. Bu basitleştirici varsayıma rağmen, NB algoritması birçok uygulamada beklenenden başarılı ve kararlı sonuçlar sunabilmektedir. NB’nin en önemli avantajları arasında kolay uygulanabilir olması, büyük veri kümelerinde düşük hesaplama maliyetiyle çalışabilmesi ve yüksek yorumlanabilirliğe sahip olması yer almaktadır. Bu özellikleri sayesinde, özellikle temel karşılaştırma modeli olarak ve klinik veri analizlerinde sıklıkla tercih edilmektedir (Domingos ve Pazzani, 1997; Hand ve Yu, 2001; Wu vd., 2008).

Buna göre, bir örneğin C_k sınıfına ait olma olasılığı Denklem 3.19’da ki Bayes teoremi ile hesaplanmaktadır.

$$P(C_k | x) = \frac{P(x | C_k) P(C_k)}{P(x)} \quad (3.19)$$

Burada $x = (x_1, x_2, \dots, x_n)$ öznitelik vektörünü, $P(C_k)$ sınıfın öncül olasılığını (prior), $P(x | C_k)$ sınıf koşullu olasılığı (likelihood) ve $P(x)$ normalizasyon terimini ifade etmektedir.

NB algoritmasında özniteliklerin koşullu olarak bağımsız olduğu varsayımı yapıldığından, ortak olasılık Denklem 3.20’de ki şekilde ayrıştırılmaktadır. Buna bağlı olarak sınıflandırma kuralı, maksimum ardıl olasılığa (Maximum A Posteriori – MAP) göre Denklem 3.21’de tanımlanmaktadır.

$$P(x|C_k) = \prod_{i=1}^n P(x_i|C_k) \quad (3.20)$$

$$\hat{C} = \arg \max_{C_k} P(C_k) \prod_{i=1}^n P(x_i | C_k) \quad (3.21)$$

Bu yaklaşım sayesinde NB, basitleştirici bağımsızlık varsayımına rağmen birçok pratik uygulamada başarılı ve kararlı sonuçlar sunabilmektedir. Düşük hesaplama

maliyeti, kolay uygulanabilirliği ve yüksek yorumlanabilirliği nedeniyle özellikle temel karşılaştırma modeli olarak ve klinik veri analizlerinde sıklıkla tercih edilmektedir (Bishop, 2006; Hastie vd., 2009; Murphy, 2012).

3.5.4. Karar ağacı (Decision tree)

Karar ağaçları, veriyi özniteliklere dayalı karar kurallarıyla hiyerarşik olarak bölen ve sınıflandırmayı yaprak düğümler üzerinden gerçekleştiren sezgisel sınıflandırma yöntemleridir. Doğrusal olmayan ilişkileri modelleyebilme yetenekleri ve yüksek yorumlanabilirlikleri nedeniyle klinik veri analizlerinde yaygın olarak kullanılmaktadır (Salzberg, 1994).

Bir karar ağacında en uygun bölme işlemi, genellikle bilgi kuramına dayalı ölçütler kullanılarak belirlenir. En yaygın kullanılan ölçüt entropi (Entropy) ve bilgi kazancı (Information Gain)'dır. Entropi Denklem 3.22'de hesaplanmaktadır.

$$H(S) = -\sum_{i=1}^c p_i \log_2(p_i) \quad (3.22)$$

Burada p_i , veri kümesi S içindeki i . Sınıfa ait örneklerin oranını, c ise sınıf sayısını ifade etmektedir.

Bir öznitelik A kullanılarak yapılan bölmenin sağladığı bilgi kazancı ise Denklem 3.23'de gösterilmektedir.

$$IG(S, A) = H(S) - \sum_{v \in \text{Values}(A)} \frac{|S_v|}{|S|} H(S_v) \quad (3.23)$$

Burada S_v , öznitelik A 'nın v değerine karşılık gelen alt veri kümesini temsil etmektedir. Karar ağacı algoritması, her düğümden bilgi kazancını maksimize eden özniteliği seçerek ağacı oluşturur.

Alternatif olarak CART algoritmasında bölme işlemi Gini indeksi kullanılarak yapılır. Denklem 3.24'de Gini indeksi belirtilmiştir. Amaç, her düğümden Gini değerini minimize eden bölmeyi seçmektir (Han vd., 2012).

$$Gini(S) = 1 - \sum_{i=1}^c p_i^2 \quad (3.24)$$

3.5.5. Ayrımcı analiz (Discriminant analysis)

Ayrımcı analiz, sınıflar arasındaki farkları maksimize eden doğrusal veya doğrusal olmayan karar sınırları oluşturarak sınıflandırma gerçekleştiren istatistiksel bir yöntemdir. Sınıfların dağılım özelliklerini dikkate alarak yeni gözlemlerin sınıf etiketlerini belirlemeyi amaçlamaktadır. En yaygın kullanılan türleri Doğrusal Ayrımcı

Analiz (Linear Discriminant Analysis, LDA) ve Kuadratik Ayrımcı Analiz (Quadratic Discriminant Analysis, QDA)'dır (Fisher, 1936; Hastie vd., 2009).

Ayrımcı analizde, her sınıfın çok değişkenli normal dağılıma sahip olduğu varsayılmaktadır. Bu durum Denklem 3.25'de gösterilmiştir.

$$p(x|C_k) = \frac{1}{(2\pi)^{d/2} |\Sigma_k|^{1/2}} \exp \left(-\frac{1}{2} (x - \mu_k)^T \Sigma_k^{-1} (x - \mu_k) \right) \quad (3.25)$$

Burada x gözlem vektörünü, μ_k k . sınıfın ortalama vektörünü ve Σ_k ise kovaryans matrisini temsil etmektedir.

Bayes karar kuralına göre, bir gözlem Denklem 3.26'da ki ayrım fonksiyonu maksimize edilerek sınıflandırılmaktadır.

$$\hat{C} = \arg \max_{C_k} \delta_k(x) \quad (3.26)$$

Doğrusal Ayrımcı Analiz (LDA) için, tüm sınıfların aynı kovaryans matrisine sahip olduğu $\Sigma_k = \Sigma$ varsayılır ve ayrım fonksiyonu Denklem 3.27'de tanımlanmıştır. Bu durumda sınıflar arasındaki karar sınırı doğrusal olmaktadır.

$$\delta_k(x) = x^T \Sigma^{-1} \mu_k - \frac{1}{2} \mu_k^T \Sigma^{-1} \mu_k + \ln P(C_k) \quad (3.27)$$

Kuadratik Ayrımcı Analiz yönteminde ise her sınıfın farklı bir kovaryans matrisine sahip olduğu kabul edilir ve ayrım fonksiyonu Denklem 3.28'de ki gibi ifade edilir.

$$\delta_k(x) = -\frac{1}{2} \ln |\Sigma_k| - \frac{1}{2} (x - \mu_k)^T \Sigma_k^{-1} (x - \mu_k) + \ln P(C_k) \quad (3.28)$$

Bu durumda elde edilen karar sınırı kuadratik (ikinci dereceden) bir yapı göstermektedir. Ayrımcı analiz yöntemleri, istatistiksel temellere dayanması ve yorumlanabilir olması nedeniyle özellikle tıbbi ve klinik veri analizlerinde sıklıkla tercih edilmektedir (Fisher, 1936; Hastie vd., 2009).

3.5.6. Lojistik regresyon (Logistic regression)

Lojistik regresyon, ikili sınıflandırma problemlerinde kullanılan olasılıksal bir modeldir. Bağımsız değişkenler ile hedef değişken arasındaki ilişkiyi lojistik fonksiyon aracılığıyla modelleyerek sınıf olasılıklarını tahmin eder. Basit yapısı ve yüksek yorumlanabilirliği nedeniyle temel karşılaştırma modeli olarak sıklıkla tercih edilmektedir. Lojistik regresyonun temel amacı, bir gözlemin belirli bir sınıfa ait olma olasılığını tahmin etmektir (Hosmer ve Lemeshow, 2000).

Bir gözlemin pozitif sınıfa ait olma olasılığı Denklem 3.29’da ki lojistik fonksiyon ile ifade edilmektedir. Burada x öznitelik vektörünü, β ağırlık vektörünü ve β_0 sabit terimi (bias) temsil etmektedir.

$$P(y = 1|x) = \frac{1}{1+e^{-(\beta_0+\beta^T x)}} \quad (3.29)$$

Lojistik regresyon modelinde log-olasılık oranı (logit fonksiyonu) doğrusal bir biçimde Denklem 3.30’da tanımlanmaktadır.

$$\log\left(\frac{P(y = 1|x)}{P(y = 0|x)}\right) = \beta_0 + \beta^T x \quad (3.30)$$

Model parametreleri genellikle maksimum olabilirlik (Maximum Likelihood Estimation, MLE) yöntemi ile tahmin edilmektedir. Bu amaçla negatif log-olabilirlik fonksiyonu minimize edilmektedir. Bu fonksiyon Denklem 3.31’de verilmiştir.

$$L(\beta) = -\sum_{i=1}^N [y_i \log(p_i) + (1 - y_i) \log(1 - p_i)] \quad (3.31)$$

Denklem 3.31’de $p_i = P(y_i = 1|x_i)$ olarak tanımlanmaktadır.

Basit yapısı, düşük hesaplama maliyeti ve yüksek yorumlanabilirliği nedeniyle lojistik regresyon, birçok çalışmada temel karşılaştırma modeli olarak sıklıkla tercih edilmektedir (Bishop, 2006; Hosmer ve Lemeshow, 2000).

3.5.7. Çekirdek yaklaşırma (Kernel approximation)

Çekirdek yaklaşırma yöntemleri, doğrusal olmayan veri yapılarını yaklaşık olarak daha yüksek boyutlu doğrusal uzaylara dönüştürmeyi amaçlamaktadır. Bu yaklaşım, çekirdek tabanlı yöntemlerin sağladığı esnekliği daha düşük hesaplama maliyetiyle elde etmeyi hedeflemektedir (Schölkopf ve Smola, 2002). Temel amaç, çekirdek fonksiyonunu açıkça hesaplamak yerine, veriyi düşük boyutlu bir öznitelik uzayına eşleyerek doğrusal yöntemlerle sınıflandırma veya regresyon gerçekleştirmektir (Hosmer ve Lemeshow, 2000).

Bir çekirdek fonksiyonu genel olarak Denklem 3.32’de tanımlanmaktadır.

$$K(x, x') = \langle \phi(x), \phi(x') \rangle \quad (3.32)$$

Burada $\phi(x)$, girdiyi yüksek boyutlu öznitelik uzayına eşleyen dönüşüm fonksiyonunu temsil etmektedir. Çekirdek yaklaşırma yöntemlerinde bu dönüşüm açık biçimde yaklaşık olarak tanımlanır ve yeni öznitelikler doğrudan hesaplanır.

En yaygın kullanılan yöntemlerden biri Rastgele Fourier Öznitelikleri (Random Fourier Features – RFF) yaklaşımıdır. Kaydırma değişmez (shift-invariant) çekirdekler için Gauss çekirdeği Denklem 3.33’de ifade edilmektedir.

$$K(x, x') = \exp\left(-\frac{\|x-x'\|^2}{2\sigma^2}\right) \quad (3.33)$$

Bu çekirdek fonksiyonu, Denklem 3.34'de yaklaşık olarak yazılabilmektedir.

$$K(x, x') \approx z(x)^T z(x') \quad (3.34)$$

Burada $z(x)$ Denklem 3.35'de tanımlanmaktadır. Burada ω_i vektörleri uygun bir olasılık dağılımından rastgele örneklenmekte, b_i ise $[0, 2\pi]$ aralığından seçilmektedir. Bu yaklaşım sayesinde, doğrusal olmayan çekirdek fonksiyonları doğrusal iç çarpım işlemleriyle yaklaşık olarak temsil edilebilmektedir.

$$z(x) = \sqrt{\frac{2}{D}} \begin{bmatrix} \cos(\omega_1^T x + b_1) \\ \cos(\omega_2^T x + b_2) \\ \vdots \\ \cos(\omega_D^T x + b_D) \end{bmatrix} \quad (3.35)$$

Alternatif bir yaklaşım olan Nyström yöntemi ise çekirdek matrisinin düşük-ranklı bir yaklaşımını kullanmaktadır. Bu yöntemde çekirdek matrisi yaklaşık olarak Denklem 3.36'da ifade edilmektedir.

$$K \approx CW^{-1}C^T \quad (3.36)$$

Burada C , çekirdek matrisinin seçilen örnekler üzerindeki sütunlarını W ise bu sütunların kesişiminden elde edilen alt matrisi temsil etmektedir. Bu yöntemle büyük veri kümeleri için hesaplama maliyeti önemli ölçüde azaltılmaktadır.

Çekirdek yaklaştırma yöntemleri, doğrusal olmayan problemlerde çekirdek tabanlı modellerin avantajlarını korurken, bellek ve zaman karmaşıklığını düşürmesi nedeniyle büyük ölçekli veri kümelerinde etkin bir çözüm sunmaktadır (Rahimi ve Recht, 2007; Williams ve Seeger, 2001).

3.5.8. Topluluk yöntemleri (Ensemble)

Topluluk yöntemleri, birden fazla temel sınıflandırıcının bir araya getirilmesiyle daha güçlü ve kararlı bir model oluşturmayı amaçlamaktadır. Bu yöntemler, aşırı uyum problemini azaltarak modelin genellenebilirliğini artırmaktadır (Dietterich, 2000).

Bir topluluk modelinde, M adet temel sınıflandırıcının çıktısı Denklem 3.37'de ki gibi birleştirilebilir.

$$\hat{y} = \arg \max_{c \in C} \sum_{m=1}^M \omega_m P_m(y = c|x) \quad (3.37)$$

Burada, C sınıf kümesini, M temel sınıflandırıcı sayısını, ω_m m . sınıflandırıcının ağırlığını, $P_m(y = c|x)$: m . sınıflandırıcının x örneği için c sınıfına ait olasılık tahminini ifade etmektedir. Eğer tüm sınıflandırıcılara eşit ağırlık verilirse ($\omega_m = 1$), bu ifade

çoğunluk oylaması (majority voting) yöntemine indirgenmektedir. Tahmin edilen sınıf Denklem 3.38’de ki şekilde hesaplanmaktadır.

$$\hat{y} = \arg \max_{c \in C} \sum_{m=1}^M \mathbb{I}(h_m(x) = c) \quad (3.38)$$

Burada $h_m(x)$, m . sınıflandırıcının tahmin ettiği sınıfı, $\mathbb{I}(\cdot)$ ise gösterge fonksiyonunu temsil etmektedir.

Literatürde yaygın olarak kullanılan topluluk yöntemleri arasında Torbalama (Bagging), Artırma (Boosting) ve Rastgele Orman (Random forest) yer almaktadır. Torbalama yöntemi, bootstrap örnekleme ile oluşturulan alt veri kümeleri üzerinde eğitilen modellerin çıktılarının birleştirilmesine dayanırken (Breiman, 1996). Artırma yöntemi, yanlış sınıflandırılan örneklerle daha fazla ağırlık vererek ardışık modellerin oluşturulmasını sağlamaktadır (Freund ve Schapire, 1997). Rastgele orman yöntemi ise çok sayıda karar ağacının rastgele alt öznelik kümeleri üzerinde eğitilmesiyle oluşturulan bir topluluk modelidir (Breiman, 2001).

Bu yöntemler, özellikle karmaşık ve gürültülü veri kümelerinde aşırı uyum problemini azaltarak modelin genellenebilirliğini artırmaktadır (Zhou, 2012).

3.5.9. Yapay sinir ağları (Artificial neural networks, YSA)

Yapay sinir ağları, biyolojik sinir sisteminden esinlenerek geliştirilmiş, katmanlı yapılar aracılığıyla karmaşık ve doğrusal olmayan ilişkileri öğrenebilen güçlü makine öğrenmesi modelleridir. Özellikle veri içerisindeki karmaşık örüntülerin keşfedilmesinde etkili olmaları nedeniyle klinik tahmin problemlerinde yaygın olarak kullanılmaktadır. Bir yapay sinir ağı, giriş katmanı, bir veya daha fazla gizli katman ve çıkış katmanından oluşmaktadır. Her bir katmandaki nöronlar, önceki katmandan gelen ağırlıklı toplamı doğrusal olmayan aktivasyon fonksiyonu yardımıyla dönüştürmektedir (Heaton, 2018).

Bir yapay sinir ağında l . katmandaki çıktı Denklem 3.39’da ki gibi ifade edilmektedir.

$$a^{(l)} = f(W^{(l)}a^{(l-1)} + b^{(l)}) \quad (3.39)$$

Burada $a^{(l)}$, l . katmanın çıktı vektörünü, $W^{(l)}$ ağırlık matrisini, $b^{(l)}$ bias vektörünü ve $f(\cdot)$ aktivasyon fonksiyonunu göstermektedir.

Model parametreleri, genellikle bir kayıp fonksiyonunun minimize edilmesi amacıyla geri yayılım (backpropagation) algoritması kullanılarak güncellenmektedir (Nielsen, 2016).

3.5.10. Doğrusal sınıflandırıcı (Linear classifier)

Doğrusal sınıflandırıcılar, sınıflar arasındaki ayrımı öznitelik uzayında doğrusal bir karar sınırı (hiper düzlem) oluşturarak gerçekleştiren temel makine öğrenmesi modelleridir. Bu yaklaşımlar, özellikle yüksek boyutlu veri kümelerinde düşük hesaplama maliyetleri ve kararlı performansları nedeniyle yaygın olarak tercih edilmektedir. Doğrusal modeller hem yorumlanabilir olmaları hem de analitik çözümler sunabilmeleri sayesinde birçok sınıflandırma probleminde temel karşılaştırma yöntemi olarak kullanılmaktadır (Bishop, 2006).

Bu modellerde karar fonksiyonu Denklem 3.40'da tanımlanmaktadır.

$$f(x) = w^T x + b \quad (3.40)$$

Burada x giriş öznitelik vektörünü, w ağırlık vektörünü ve b sabit terimini (bias) göstermektedir.

Sınıf etiketi, karar fonksiyonunun işaretine göre belirlenmektedir. Denklem 3.41'de matematiksel ifadesi gösterilmektedir.

$$\hat{y} = \text{sign}(f(x)) \quad (3.41)$$

Bu yapı, doğrusal ayrılabilir veri kümelerinde etkili sonuçlar üretmekte olup, birçok sınıflandırma algoritmasının temelini oluşturmaktadır (Bishop, 2006; Hastie vd., 2009).

3.6. Performans Değerlendirme Metrikleri

Bu çalışmada, makine öğrenmesi modellerinin hasta sağkalımını tahmin etme performansları; doğruluk (accuracy), duyarlılık (sensitivity), özgüllük (specificity), F1-skoru ve Matthews Korelasyon Katsayısı (Matthews Correlation Coefficient, MCC) kullanılarak değerlendirilmiştir. Farklı metriklerin birlikte kullanılması, modellerin sınıflandırma başarımının daha kapsamlı ve dengeli biçimde analiz edilmesini sağlamaktadır.

Duyarlılık ve özgüllük metrikleri, özellikle klinik uygulamalarda yanlış negatif ve yanlış pozitif sınıflandırmaların etkilerini değerlendirmek açısından önem taşımaktadır. Ancak bu çalışmada, sınıf dağılımı ve problem yapısı göz önünde bulundurularak, modellerin genel performanslarının karşılaştırılmasında F1-skoru ve MCC metriklerine öncelik verilmiştir.

Model performanslarının karşılaştırılmasında doğruluk, duyarlılık, özgüllük, F1-skoru ve Matthews Korelasyon Katsayısı (MCC) metrikleri kullanılmıştır. Bununla

birlikte, özellikle sınıf dengesizliği ve klinik karar verme bağlamı göz önünde bulundurularak, F1-skoru ve MCC metrikleri temel değerlendirme ölçütleri olarak ele alınmıştır.

3.6.1. Doğruluk (Accuracy)

Doğruluk, doğru sınıflandırılan örneklerin toplam örnek sayısına oranını ifade etmektedir. Bununla birlikte, sınıf dağılımının dengesiz olduğu veri setlerinde tek başına yeterli bir performans ölçütü sunmayabileceğinden, diğer metriklerle birlikte değerlendirilmiştir (Fawcett, 2006).

Matematiksel ifadesi Denklem 3.42’de verilmiştir.

$$\text{Doğruluk} = \frac{DP+DN}{DP+DN+YP+YN} \quad (3.42)$$

Doğru Pozitif (DP), Doğru Negatif (DN), Yanlış Pozitif (YP), Yanlış Negatif (YN) örnek sayılarını göstermektedir.

3.6.2. Duyarlılık (Sensitivity)

Duyarlılık, gerçek pozitif örneklerin ne kadarının doğru şekilde sınıflandırıldığını gösteren bir ölçüttür. Klinik bağlamda, özellikle riskli hastaların doğru şekilde tespit edilmesi açısından kritik öneme sahiptir (Fawcett, 2006; Sokolova ve Lapalme, 2009).

Matematiksel ifadesi Denklem 3.43’te verilmiştir.

$$\text{Duyarlılık} = \frac{DP}{DP+YN} \quad (3.43)$$

3.6.3. Özgüllük (Specificity)

Özgüllük, gerçek negatif örneklerin doğru sınıflandırılma oranını ifade etmektedir. Yanlış pozitif sınıflandırmaların azaltılması açısından önemli bir değerlendirme ölçütüdür ve duyarlılık metriğiyle birlikte yorumlanmaktadır (Fawcett, 2006; Sokolova ve Lapalme, 2009).

Matematiksel ifadesi Denklem 3.44’te verilmiştir.

$$\text{Özgüllük} = \frac{DN}{DN+YP} \quad (3.44)$$

3.6.4. F1-skoru

F1-skoru, kesinlik (precision) ve duyarlılığın harmonik ortalaması olarak tanımlanmakta olup, özellikle dengesiz veri setlerinde daha dengeli bir performans

değerlendirmesi sunmaktadır. Bu nedenle, bu çalışmada modellerin karşılaştırılmasında temel ölçütlerden biri olarak kullanılmıştır (Powers, 2007).

Matematiksel ifadesi Denklem 3.45 ve Denklem 3.46'da verilmiştir.

$$F1 - skoru = \frac{2 * Kesinlik * Duyarluluk}{Kesinlik + Duyarluluk} \quad (3.45)$$

veya eşdeğer biçimde,

$$F1 - skoru = \frac{2 * DP}{2 * DP + YP + YN} \quad (3.46)$$

3.6.5. Matthews korelasyon katsayısı (Matthews Correlation Coefficient, MCC)

Matthews Korelasyon Katsayısı (MCC), ikili sınıflandırma problemlerinde tüm karmaşıklık matrisi bileşenlerini birlikte dikkate alan korelasyon temelli bir performans ölçütüdür. Literatürde yapılan karşılaştırmalı çalışmalar, MCC'nin özellikle sınıf dengesizliğinin bulunduğu veri setlerinde doğruluk ve F1-skoruna kıyasla daha dengeli ve güvenilir bir değerlendirme sunduğunu göstermektedir. Bu nedenle MCC, klinik veri analizleri ve hasta sağkalımı tahmini gibi problemlerde model karşılaştırmaları için önerilen temel metriklerden biri olarak kabul edilmektedir (Chicco ve Jurman, 2020b, 2023; Jurman vd., 2012).

Matematiksel ifadesi Denklem 3.47'de verilmiştir.

$$MCC = \frac{DP * DN - YP * YN}{\sqrt{(DP + YP)(DP + YN)(DN + YP)(DN + YN)}} \quad (3.47)$$

4. ARAŞTIRMA SONUÇLARI VE TARTIŞMA

Bu çalışmada, güncel makine öğrenmesi modellerinin kalp yetmezliği hastalarının sağkalım verilerine uyarlanması ve tahmin performanslarının değerlendirilmesi amaçlanmıştır. Bu doğrultuda, yüksek ayırt edici güce sahip klinik özniteliklerin tespit edilmesi ve bu özniteliklerle en optimum şekilde çalışan sınıflandırma algoritmalarının belirlenmesi için kapsamlı bir analiz süreci yürütülmüştür. Sonuç olarak, klinik kullanım potansiyeli en yüksek olan ideal öznitelik alt kümesi ve sınıflandırma algoritması belirlenerek nihai bir model önerisinde bulunulmuştur.

Çalışmada kullanılan ham veri seti toplam 12 öznitelikten oluşmaktadır. Bu değişkenlerden biri olan “*takip süresi*” (time), hastanın hastaneye ilk başvuru anındaki klinik durumunu değil, sonradan elde edilen izlem sürecini yansıtmaktadır. Modelin yalnızca başvuru anındaki klinik bulgularla, henüz takip süresi bilinmeyen yeni hastalar için öngörü yapabilmesini sağlamak ve literatürde veri sızıntısı olarak bilinen, geleceğe dair bilgilerin modele sızması problemini engellemek amacıyla, temel modelleme “*takip süresi*” özniteliği hariç tutularak gerçekleştirilmiştir. Bununla birlikte, takip süresinin sağkalım üzerindeki doğrudan istatistiksel etkisini ve diğer klinik parametrelerle olan etkileşimini gözlemleyebilmek adına, “*takip süresi*” özniteliğinin dahil edildiği ikinci bir analiz daha kurgulanmıştır. Dolayısıyla, izleyen bölümlerde detaylandırılan tüm öznitelik seçimi, sınıflandırma ve model doğrulama prosedürleri; “*takip süresi*” özniteliğinin varlığı ve yokluğu senaryoları altında iki bağımsız kol halinde tekrarlanmıştır.

Öznitelik seçim algoritmalarının matematiksel temelleri, girdi verisinin sayısal (numerical) veya ikili (binary) olma durumuna göre yapısal farklılıklar göstermektedir. Bu çalışmada analiz edilen kalp yetmezliği kohortuna ait veri seti, 7'si nümerik ve 5'i kategorik olmak üzere toplam 12 bağımsız değişkenden oluşan heterojen bir yapıya sahiptir. Bölüm 3.4'te teorik altyapıları detaylandırılan 8 farklı öznitelik seçim algoritmasının; dördü spesifik olarak sayısal uzaylar, biri ikili uzaylar, kalan üçü ise her iki veri tipini (karma) işleyebilen hibrit yaklaşımlar için tasarlanmıştır. Veri setinin bu karmaşık yapısı göz önüne alınarak, ayırt edici özniteliklerin tespitinde iki farklı metodolojik strateji kurgulanmış ve karşılaştırmalı olarak analiz edilmiştir:

İlk stratejide, veri tipleri arasında herhangi bir ayrım gözetilmeksizin, tüm sayısal ve ikili öznitelikler tek bir girdi matrisi halinde birleştirilerek ilgili öznitelik seçim algoritmalarına topluca beslenmiştir. İkinci stratejide, veri seti doğasına uygun olarak sayısal ve ikili olmak üzere iki alt uzaya ayrıştırılmıştır. İkili değişkenler bağımsız olarak

ikili öznitelik algoritmalarıyla; sayısal değişkenler ise sayısal algoritmalarla değerlendirilerek, her iki grup için kendi içinde hiyerarşik önem sıralamaları elde edilmiştir. Modellerin sınıflandırma performansını test etmek amacıyla, bu iki listeden en yüksek önem skoruna sahip öznitelikler eş zamanlı olarak çekilmiştir. Her iterasyonda, her iki gruptan birer öznitelik kümeyle dahil edilerek simetrik bir artış (sırasıyla 2, 4, 6... boyutlu öznitelik vektörleri) sağlanmıştır. Oluşturulan her bir yeni alt küme için sınıflandırma algoritmaları tamamen sıfırdan ve bağımsız olarak yeniden eğitilmiş, elde edilen performans metrikleri kümülatif etkiyi gözlemlemek üzere kaydedilmiştir. Çalışmada kullanılan iki adet hibrit öznitelik seçim algoritması, her iki veri tipini de işleyebilme kapasiteleri göz önünde bulundurularak, metodolojik açıdan hem birer nümerik hem de birer kategorik yöntem olarak kabul edilmiştir. Bu varsayım doğrultusunda, ayrıştırılmış öznitelik seçimi yaklaşımında söz konusu hibrit yöntemler her iki algoritma havuzuna da dahil edilmiş; böylece hem nümerik hem de kategorik veri gruplarının kendi içlerindeki analizinde doğrudan kullanılmıştır.

Sınıflandırma modellerinin performans analizi ve optimal öznitelik alt kümelerinin değerlendirilmesi kapsamında, sağlam bir doğrulama stratejisi olan tekrarlı bekletme (repeated hold-out / Monte Carlo cross-validation) yaklaşımı benimsenmiştir. Bu metodolojik çerçevede, modellerin aşırı öğrenme riskini minimize etmek ve yeni verilere karşı genelleştirilebilirlik kapasitesini istatistiksel bir güvenilirlikle ölçmek üzere, hasta kohortu rastgele alt örnekleme yöntemiyle %70 eğitim ve %30 test alt kümelerine ayrıştırılmıştır.

Tekil bir veri bölütleme işleminin yaratabileceği rastgelelik sapmasını (selection bias) ve performans varyansını ortadan kaldırmak, ayrıca daha kararlı (stable) değerlendirme metrikleri elde etmek adına, bu ayırma ve test işlemi her bir sınıflandırma algoritması ve öznitelik kombinasyonu için 100 bağımsız iterasyon halinde tekrarlanmıştır. İlgili algoritmaların nihai eğitim ve test performansları, bu 100 bağımsız iterasyondan elde edilen sonuçların aritmetik ortalaması alınarak hesaplanmış ve karşılaştırmalı analizler bu ortalama değerler üzerinden yürütülmüştür.

Çalışma kapsamında kullanılan makine öğrenmesi modellerinin uygulama sürecinde, yöntemler arası karşılaştırılabilirliği korumak adına temel olarak MATLAB ortamının varsayılan parametre ayarları benimsenmiştir. Bununla birlikte, model mimarilerini tanımlayan bazı kritik hiper parametreler araştırmacı tarafından spesifik olarak belirlenmiştir. Bu doğrultuda; k-En Yakın Komşu (k-NN) algoritmasında komşu sayısı $k=5$ olarak set edilmiş, Destek Vektör Makineleri (SVM) modelinde ise doğrusal

olmayan ilişkileri yakalayabilmek amacıyla Radyal Tabanlı Fonksiyon (RBF) çekirdeği tercih edilmiştir. Topluluk öğrenme modellerinden TreeBagger yönteminde toplam ağaç sayısı 50 olarak belirlenirken, Lojistik Regresyon modeli lojistik öğrenici çerçevesinde yapılandırılmıştır. Naive Bayes, Karar Ağaçları ve Ayrımcı Analizi gibi diğer algoritmalar ise herhangi bir dış müdahale olmaksızın standart yazılım ayarlarıyla çalıştırılmıştır.

4.1. Araştırma Sonuçları

Bu bölümde deneylerden elde edilen sonuçlar sunulmaktadır. Sonuçların daha anlaşılır şekilde değerlendirilebilmesi amacıyla yalnızca performansı en yüksek ilk 10 model tablolar halinde verilmiştir. Her bir öznitelik seçimi yöntemi ve sınıflandırıcı kombinasyonu için elde edilen en iyi ilk beş performans sonucu, deney senaryolarına göre EK-1, EK-2, EK-3 ve EK-4'te sunulmuştur.

4.1.1. Takip süresi özniteliği dahil edilmeden elde edilen sonuçlar

Tablo 4.1'de, takip süresi (time) özniteliğinin veri setinden çıkarılması durumunda elde edilen sınıflandırma sonuçları sunulmaktadır. Elde edilen bulgular incelendiğinde, modellerin performans değerlerinin birbirine oldukça yakın olduğu ve genel olarak orta düzeyde bir başarı sergilediği görülmektedir.

Tablo 4.1. Takip süresi özniteliği hariç tutulduğunda elde edilen sonuçlar

Öznitelik Seçici	Sınıflandırıcı	ÖS	F1 (%) (ort±ss)	MCC (%) (ort±ss)	Doğruluk (%) (ort±ss)	Duyarlılık (%) (ort±ss)	Özgüllük (%) (ort±ss)
SFS + TreeBagger	Doğrusal	6 ^a	59±6	45,27±8	77,33±3	51,36±8	89,57±4
mRMR + Tree	Doğrusal	6 ^b	58,49±7	44,41±8	76,88±3	51,29±8	88,98±4
mRMR	Doğrusal	4 ^c	57,11±9	43,99±9	76,93±4	49,14±10	89,97±5
Tree + Tree	Doğrusal	6 ^d	58,34±7	43,95±9	76,69±3	51,34±8	88,67±4
TreeBagger+ TreeBagger	Doğrusal	6 ^e	57,54±6	43,60±7	76,75±3	50,08±8	89,26±3
ReliefF + Tree	Doğrusal	6 ^f	57,11±8	43,55±9	76,69±3	49,39±9	89,55±4
TreeBagger	Lojistik	3 ^g	57,06±7	43,30±8	76,60±3	49,17±9	89,51±4
Ki-kare	Doğrusal	3 ^h	56,56±8	43,22±9	76,58±3	48,36±10	89,93±4
Tree	Lojistik	4 ⁱ	57,11±7	43,15±9	76,53±3	49,38±8	89,28±5
TreeBagger	Lojistik	5 ⁱ	56,60±7	43,08±9	76,61±5	48,13±8	90,01±4

ÖS: Öznitelik Sayısı, ort: ortalama, ss: standart sapma ^a Ejeksiyon fraksiyonu, Yaş, Serum kreatinin, Sigara kullanımı, Cinsiyet, Yüksek tansiyon ^b Serum kreatinin, Ejeksiyon fraksiyonu, Yaş, Cinsiyet, Sigara kullanımı, Diyabet ^c Serum kreatinin, Ejeksiyon fraksiyonu, Yaş, Yüksek tansiyon ^d Serum kreatinin, Ejeksiyon fraksiyonu, Yaş, Cinsiyet, Sigara kullanımı, Diyabet ^e Ejeksiyon fraksiyonu, Serum kreatinin, Yaş, Sigara kullanımı, Cinsiyet, Yüksek tansiyon ^f Ejeksiyon fraksiyonu, Yaş, Serum kreatinin, Cinsiyet, Sigara kullanımı, Diyabet ^g Ejeksiyon fraksiyonu, Serum kreatinin, Yaş ^h Ejeksiyon fraksiyonu, Serum kreatinin, Yaş ⁱ Serum kreatinin, Ejeksiyon fraksiyonu, Yaş, Serum sodyumu ⁱ Ejeksiyon fraksiyonu, Serum kreatinin, Yaş, Serum sodyumu, Cinsiyet

En yüksek performans, SFS + TreeBagger tabanlı öznitelik seçimi ve Doğrusal sınıflandırıcının birlikte kullanıldığı modelde elde edilmiştir. Bu model, 6 öznitelik (ejeksiyon fraksiyonu, yaş, serum kreatinin, sigara kullanımı, cinsiyet, yüksek tansiyon) kullanarak %77,33 doğruluk, %59 F1-skoru ve %45,27 MCC değeri ile diğer yöntemlere kıyasla daha dengeli ve üstün bir performans göstermiştir. Buna en yakın sonuç mRMR + Tree öznitelik seçme ve yine Doğrusal sınıflandırıcı ile elde edilmiştir. Bu modelde ise 6 öznitelik (serum kreatinin, ejeksiyon fraksiyonu, yaş, cinsiyet, sigara kullanımı, diyabet) kullanarak %76,88 doğruluk, %58,49 F1-skoru ve %44,41 MCC değeri elde edilmiştir.

Genel olarak değerlendirildiğinde, takip süresi özneliği hariç tutulduğunda modeller arasında büyük performans farkları oluşmadığı görülmektedir. Bununla birlikte, ejeksiyon fraksiyonu, serum kreatinin ve yaş değişkenlerinin hemen hemen tüm başarılı modellerde ortak olarak yer aldığı dikkat çekmektedir. Bu durum, ilgili özneliklerin kalp yetmezliği hasta sağkalım tahmininde önemli belirteçler olduğunu göstermektedir.

4.1.2. Takip süresi özneliği dahil edildiğinde elde edilen sonuçlar

Tablo 4.2’de takip süresi bilgisi dahil edilerek elde edilen sınıflandırma performansları sunulmaktadır. Tablo 4.2 incelendiğinde, bütün modellerde en iyi sonucun k-NN sınıflandırıcı ile elde edildiği görülmektedir. En yüksek performans k-NN sınıflandırıcı ile birlikte ReliefF öznitelik seçme algoritmasının kullanıldığı modelde elde edilmiştir.

Tablo 4.2. Takip süresi özneliği dahil edildiğinde elde edilen sonuçlar

Öznitelik Seçici	Sınıflandırıcı	ÖS	F1 (%) (ort±ss)	MCC (%) (ort±ss)	Doğruluk (%) (ort±ss)	Duyarlılık (%) (ort±ss)	Özgüllük (%) (ort±ss)
ReliefF	k-NN	5 ^a	74,56±5	65,94±7	85,56±3	66,59±7	94,51±3
TreeBagger	k-NN	4 ^b	74,26±5	65,75±6	85,47±3	66,07±8	94,61±3
Ki-kare	k-NN	2 ^c	74,37±6	65,71±7	85,49±3	66,43±8	94,47±3
Tree	k-NN	3 ^d	74,24±6	65,59±7	85,45±3	66,16±8	94,53±3
Ki-kare	k-NN	3 ^e	73,84±6	65,28±7	85,34±3	65,70±8	94,53±3
TreeBagger + Tree	Lojistik	2 ^f	72,20±6	62,24±8	84,01±3	65,07±7	92,94±4
ReliefF + TreeBagger	k-NN	4 ^g	73,34±7	62,13±9	83,71±4	70,47±9	89,95±4
TreeBagger + Tree	k-NN	4 ^h	73,23±5	62,09±8	83,69±3	69,7±8	90,31±4
ReliefF + Tree	k-NN	4 ⁱ	73,01±6	62,08±8	83,75±3	68,84±8	90,8±4
SFS + Ki-kare	Lojistik	4 ⁱ	71,43±6	61,71±8	83,8±3	63,82±8	93,2±4

ÖS: Öznitelik Sayısı, ort: ortalama, ss: standart sapma ^a Takip süresi, Ejeksiyon fraksiyonu, Diyabet, Anemi, Cinsiyet ^b Takip süresi, Serum kreatinin, Ejeksiyon fraksiyonu, Serum sodyumu ^c Takip süresi, Ejeksiyon fraksiyonu, ^d Takip süresi, Serum kreatinin, Ejeksiyon fraksiyonu ^e Takip süresi, Ejeksiyon fraksiyonu, Serum kreatinin ^f Takip süresi, Cinsiyet ^g Takip süresi, Ejeksiyon fraksiyonu, Sigara kullanımı, Cinsiyet ^h Takip süresi, Ejeksiyon fraksiyonu, Cinsiyet, Sigara Kullanımı ⁱ Takip süresi, Ejeksiyon fraksiyonu, Sigara kullanımı, Cinsiyet ⁱ Takip süresi, Serum sodyumu, Yüksek Tansiyon, Anemi

Bu modelde 5 öznitelik (takip süresi, ejeksiyon fraksiyonu, diyabet, anemi, cinsiyet) kullanarak %85,56 doğruluk, %74,56 F1 skoru, %65,94 MCC değeri elde edilmiştir. Bu modele en yakın sonuç k-NN sınıflandırıcı ile birlikte TreeBagger öznitelik seçme algoritmasının kullanıldığı modelde elde edilmiştir. Bu modelde 4 öznitelik (takip süresi, serum kreatinin, ejeksiyon fraksiyonu, serum sodyumu) kullanarak %85,47 doğruluk, %74,26 F1 skoru ve %65,75 MCC değeri elde edilmiştir.

Tablo genel olarak değerlendirildiğinde, takip süresi özneliğinin modele dahil edilmesiyle birlikte sınıflandırma performanslarında belirgin bir iyileşme elde edildiği görülmektedir. Özellikle başarılı modellerde takip süresi, ejeksiyon fraksiyonu ve serum kreatinin değişkenlerinin sıklıkla birlikte yer alması, bu özniteliklerin hasta sağkalım tahmini açısından yüksek ayırt edici güce sahip olduğunu göstermektedir. Elde edilen bulgular, takip süresi bilgisinin kalp yetmezliği hastalarının sağkalım durumunun tahmin edilmesinde önemli bir belirleyici olduğunu ortaya koymaktadır.

4.2. Tartışma

4.2.1. Sonuçların Analizi

Bu çalışmada elde edilen deneysel bulgular; takip süresi özneliğinin modele etkisi, öznitelik seçimi stratejilerinin başarısı ve performans metrikleri arasındaki dağılım çerçevesinde bütüncül bir yaklaşımla değerlendirilmiştir. Analiz sonuçları incelendiğinde, takip süresi özneliğinin modele dahil edilmesinin performans metriklerinde dramatik bir artış sağladığı, doğruluk oranlarını %84–%86, F1 skorlarını ise %71–%75 aralığına yükselttiği görülmektedir. Bu senaryolarda k-NN algoritması, çoğu öznitelik seçimi yöntemleriyle en uyumlu ve en yüksek performansı sergileyen sınıflandırıcı olarak öne çıkmıştır. Takip süresi değişkeninin sağkalım tahminindeki bu baskın rolü, literatürde veri sızıntısı riskine işaret etse de modelin teknik kapasitesini ölçmek açısından kritiktir. Bununla birlikte, çalışmanın gerçek dünya klinik senaryolarına uygunluğu açısından, takip süresi değişkeni hariç tutulduğunda elde edilen %77,33'lük doğruluk başarısı, modelin hastaneye ilk başvuru anındaki öngörü kapasitesi bakımından daha gerçekçi bir bulgu olarak değerlendirilmektedir.

Takip süresi özniteliğinin sisteme dahil edilmesi, sadece genel doğruluğu artırmakla kalmamış, aynı zamanda modellerin pozitif sınıfı (ölüm gerçekleşen vakalar) yakalama yeteneğinde de belirgin bir iyileşme sağlamıştır. Tablo 4.1'de %48-%51 bandında kalan duyarlılık değerlerinin, takip süresi değişkeniyle birlikte %66-%70 seviyelerine yükselmesi, takip süresinin mortalite riskini ayırt etmedeki yüksek bilgi kazancını doğrulamaktadır. Buna karşın özgüllük değerlerinin %94 seviyelerine ulaşması, modelin hayatta kalan hastaları ayırt etmedeki kararlılığını koruduğunu göstermektedir. Bu performans artışı, sınıfsal dengesizliğin yarattığı yanlılığın, güçlü özniteliklerin modele dahil edilmesiyle belirli ölçüde dengelenebildiğini kanıtlamaktadır.

Modelin ayırt edici gücü öznitelik sayısı bağlamında ele alındığında, ReliefF yöntemiyle seçilen 5 öznitelikli (takip süresi, ejeksiyon fraksiyonu, diyabet, anemi, cinsiyet) kümenin en yüksek performansı vermesi, ancak Ki-kare yöntemiyle seçilen sadece 2 öznitelikli (takip süresi, ejeksiyon fraksiyonu) modelin bile %85,49 doğruluk gibi oldukça yakın bir başarı sergilemesi dikkat çekicidir. Bu bulgu, kalp yetmezliği sağkalım tahmininde doğru öznitelik kombinasyonunun seçilmesinin, modelin genellenebilirliği ve hesaplama verimliliği açısından çok sayıda değişken kullanmaktan daha avantajlı olduğunu göstermektedir.

Tüm bu değerlendirmeler ışığında, farklı öznitelik seçimi algoritmaları ve sınıflandırma senaryoları genelinde bir kararlılık analizi yapıldığında; takip süresi, serum kreatinin, ejeksiyon fraksiyonu ve yaş değişkenlerinin neredeyse tüm senaryolarda en belirleyici faktörler olarak öne çıktığı tespit edilmiştir. Bu değişkenlerin farklı algoritmalar tarafından sürekli olarak en yüksek önem skorlarıyla seçilmesi, kalp yetmezliğine bağlı mortalite riskinin tahmininde bu parametrelerin evrensel ve en kritik biyobelirteçler olduğu gerçeğini metodolojik bir titizlikle doğrulamaktadır.

4.2.2. Literatürle Karşılaştırma

Bu çalışmada elde edilen bulguların literatürde raporlanan sonuçlarla karşılaştırılabilmesi amacıyla, kalp yetmezliği hasta sağkalım tahminine yönelik takip süresi özniteliğini kullanan önceki çalışmalar Tablo 4.3'de özetlenmiştir. Tablo 4.3. incelendiğinde yüksek başarı elde etmiş çalışmalar görülmektedir. Bu çalışmaların yüksek performans değerlerine ulaşmasında kullanılan veri işleme, modelleme ve sınıflandırıcı doğrulama yaklaşımlarının belirleyici olduğu görülmektedir.

Tablo 4.3'te time özniteliğini dahil eden çalışmalar metodolojik olarak incelendiğinde, öznitelik yönetimi ve sınıflandırma stratejilerinde geniş bir çeşitliliğe gidildiği görülmektedir. Bu doğrultuda Chicco ve Jurman (2020a); Pearson Korelasyon Katsayısı (Pearson Correlation Coefficient, PCC), Ki-Kare Testi, Mann-Whitney U Testi, Shapiro-Wilk Testi, ReliefF, Rastgele Orman Gini Önemi (Random Forest Gini Importance), Rastgele Orman OOB Permütasyon Önemi (Random Forest Out-of-Bag Permutation Importance), Tek Kural (One Rule, OneR), Yinelemeli Bölümleme ve Regresyon Ağaçları (Recursive Partitioning and Regression Trees, RPART), Doğrusal SVM ANOVA (Linear Support Vector Machine ANOVA) ve XGBoost öznitelik Skorlaması olmak üzere 11 farklı filtre ve gömülü yöntemi sınamış, bu listeleri Borda Sayımı algoritmasıyla birleştirerek Lojistik Regresyon (LR) ve Rastgele Orman (RF) modellerini eğitmiştir. Moreno-Sánchez (2023); Özyinelemeli Öznitelik Eleme (RFE), Karşılıklı Bilgi, ANOVA F-Testi ve Ki-Kare algoritmalarını ayrı ayrı test ederek öznitelik seçimi yapmış ve toplu öğrenme odaklı RF, Gradyan Artırma, AdaBoost, Aşırı Gradyan Artırma (eXtreme Gradient Boosting, XGBoost) ile Ekstra Ağaçlar sınıflandırıcılarını değerlendirmiştir. Ishaq vd. (2021), gömülü bir yaklaşım olan Rastgele Orman Önemi (Random Forest Importance) yöntemini kullanarak öznitelikleri filtrelemiş ve LR, SVM, k-NN, Karar Ağacı (DT), RF, Gradyan Artırma, AdaBoost, Ekstra Ağaçlar Sınıflandırıcısı (ETC) ve Işık Gradyan Artırma Sınıflandırıcısı (LightGBM) algoritmalarını kullanmıştır. Öznitelik azaltmak yerine genişletmeyi hedefleyen Qadri vd. (2024), RF tabanlı Transfer Öğrenme ve Gradyan Artırma Makinesi (Gradient Boosting Machine, GBM) mimarileri üzerinden öznitelik mühendisliği ile yeni olasılıksal alt öznitelikler türeterek LR, SVM, k-NN, DT, NB ve RF modellerini eğitmiştir. Buna karşın, doğrusal ilişkilere odaklanan Souza ve Lima (2024) farklı varyasyonlardaki Korelasyon Filtrelerini (Correlation Filters, CF) $CF \geq 0,1$ eşik değeriyle test edip sadece k-NN algoritmasını konumlandırırken; (Çakır, 2025) herhangi bir öznitelik seçimi yapmadan 12 değişkenin tamamıyla LR, Doğrusal Ayrımcı Analizi (LDA), GBM, Düzenleştirilmiş Genelleştirilmiş Doğrusal Modeller (GLMNET) ve Doğrusal Çekirdekli Destek Vektör Makineleri (SVM Linear) algoritmalarını sınamıştır. Benzer şekilde (Qadeer vd., 2025)'de öznitelik elemesi yapmaksızın tam boyutlu veri üzerinde LR, SVM, k-NN, RF, DT, NB, AdaBoost, Çok Katmanlı Algılayıcı (MLP) ve derin öğrenme tabanlı CNN mimarilerini test etmiştir. Literatürdeki diğer çalışmalardan yapısal olarak ayrışan Zhang (2025) ise öznitelikler arası çoklu doğrusallığı denetlemek için Varyans Şişirme Faktörü (Variance Inflation Factor, VIF) analizini, geometrik boyut

indirgeme için de PCA algoritmasını koşturarak Lojistik Regresyon, Rastgele Orman ve k-NN modellerini değerlendirmiştir.

Ishaq vd. (2021), Qadri vd. (2024) ve Qadeer vd. (2025) çalışmalarında sınıf dengesizliği problemini gidermek amacıyla SMOTE TEKNİĞİ kullanılmış ve bu sayede model performansında belirgin artışlar elde edilmiştir. Ancak SMOTE ile üretilen sentetik örnekler, özellikle küçük veri setlerinde modelin bu verilere aşırı uyum sağlamasına ve genellenebilirliğin azalmasına neden olabilmektedir. Benzer şekilde, Qadri vd. (2024) çalışmasında kullanılan transfer öğrenme (transfer learning) yaklaşımı ve türetilmiş olasılıksal öznitelikler performansı artırmakla birlikte, modelin gerçek veri dağılımına dayalı öğrenme kapasitesini ve yorumlanabilirliğini sınırlayabilmektedir. Bu nedenle, bu tür yöntemlerle elde edilen yüksek performans değerlerinin dikkatli değerlendirilmesi gerekmektedir.

Öznitelik seçimi açısından bakıldığında; literatürdeki çalışmaların çoğunda 12 özneliğin tamamı kullanılırken, bu çalışmada ReliefF algoritması ile belirlenen sadece 5 kritik öznitelik (takip süresi, ejeksiyon fraksiyonu, diyabet, anemi, cinsiyet) ile sonuca gidilmiştir. Literatürde az sayıda öznitelik (3 ile 5 arası) kullanarak sağkalım tahmini yapan modeller incelendiğinde, bu çalışmaların genellikle laboratuvar ortamında ölçülmesi gereken parametreler gerektirdiği görülmektedir. Örneğin, Souza ve Lima (2024) tarafından 4 öznitelik ile kurulan model; Serum sodyumu ve Kreatin fosfokinaz (CPK) gibi kan tahlili ve spesifik laboratuvar analizi gerektiren verilere ihtiyaç duymaktadır. Benzer şekilde Chicco ve Jurman (2020a) tarafından önerilen 3 öznelikli modelde yer alan serum kreatinin verisi, hastanın böbrek fonksiyonlarının takibi için ek biyokimyasal testleri zorunlu kılmaktadır. Aynı şekilde Ishaq vd. (2021) tarafından kullanılan öznitelik seti de kreatin fosfokinaz ve trombositler gibi tamamen laboratuvar analizi ve kan tahlili gerektiren parametrelere bağımlıdır. Bu tür öznitelikler, model boyutunu küçültse de klinik veri toplama sürecini laboratuvar imkanlarına bağımlı ve zahmetli hale getirebilmektedir. Bu çalışmada ise ReliefF öznitelik seçimi algoritması sonucunda en yüksek başarıyı veren 5 özneliğin (takip süresi, ejeksiyon fraksiyonu, diyabet, anemi, cinsiyet); literatürdeki diğer modellerin aksine, klinik ortamda en hızlı ve en düşük maliyetle elde edilebilecek veriler olduğu görülmüştür. Modelin başarısını borçlu olduğu bu öznitelik seti; fiziksel muayene, hasta anamnezi veya rutin bir ekokardiyografi raporuyla anlık olarak temin edilebilmektedir. Algoritmanın laboratuvar bağımlılığı olan karmaşık parametreler yerine bu erişilebilir öznitelikleri en belirleyici olarak seçmesi, uygulanan modelin pratik değerini artırmaktadır. Minimalist ve

maliyetsiz bir öznitelik yapısına sahip modelimizle elde edilen %85,56 doğruluk ve sınıf dengesizliğine rağmen yakalanan %65,94 MCC değeri, modelin laboratuvar odaklı karmaşık parametrelere ihtiyaç duymadan da yüksek bir genelleme kabiliyetine sahip olduğunu kanıtlamaktadır. Özellikle özgüllük değerindeki belirgin başarı, düşük maliyetli bu modelin klinik karar destek süreçlerinde güvenilir bir filtre görevi görebileceğini göstermektedir

Literatürdeki birçok çalışmada model performansının değerlendirilmesinde kullanılan doğrulama stratejilerinin sonuçların varyansını ölçmede sınırlı kaldığı dikkat çekmektedir. Takip süresi özneliğinin kullanılarak sınıflama performansının değerlendirildiği çalışmaların verildiği Tablo 4.3'deki çalışmalar içerisinde sadece birinci satırda verilen Chicco ve Jurman (2020a)'ın ve bu çalışma dışında bütün çalışmalarda model performansı ya sabit bir veri bölmesi ya da tekrarlı olmayan doğrulama stratejileri üzerinden raporlanmıştır. Bu durum, diğer çalışmaların elde ettiği yüksek performansın verinin belirli bir kombinasyonuna dayalı tesadüfi bir başarı olmasına yol açabilmekte ve modelin farklı veri dağılımlarındaki genel performansını tam olarak yansıtmamaktadır. Buna karşılık, bu çalışmada benimsenen 100 tekrarlı Monte Carlo çapraz doğrulama yaklaşımı, modelin başarısını farklı rastlantısal alt kümeler üzerinde test ederek sonuçların istatistiksel güvenilirliğini artırmış ve veri bölünmesine dayalı olası yanlışlıkları minimize etmiştir. Takip süresi özneliğinin dahil edildiği ve 100 tekrarlı Monte Carlo çapraz doğrulama yöntemiyle test edilen modeller kıyaslandığında; Chicco ve Jurman (2020a)'ın üç temel öznitelik (ejeksiyon fraksiyonu, serum kreatinin, takip süresi) ve Lojistik Regresyon (LR) ile elde ettiği sonuçlara karşın, bu çalışmada beş öznitelik (takip süresi, ejeksiyon fraksiyonu, diyabet, anemi, cinsiyet) ve ReliefF tabanlı k-En Yakın Komşu (k-NN) algoritması kullanılmıştır. Elde edilen bulgular, önerilen modelin MCC (%65,94), F1-skoru (%74,56) ve doğruluk (%85,56) metriklerinde Chicco ve Jurman (2020a)'ın çalışmasını geride bıraktığını göstermektedir. Özellikle özgüllük değerindeki belirgin artış (%86'dan %94,51'e), modelin negatif vakaları ayırt etme konusundaki üstünlüğünü kanıtlamaktadır. Ancak duyarlılık değerinde gözlemlenen nispi düşüş (%78,5'ten %66,59'a), eklenen özniteliklerin ve k-NN algoritmasının modelin "yakalama" kapasitesinden ziyade doğru sınıflandırma hassasiyetine odaklandığını düşündürmektedir. Sonuç olarak, daha geniş bir öznitelik kümesiyle desteklenen ReliefF k-NN yaklaşımının, çapraz doğrulama altında daha stabil ve yüksek korelasyonlu (yüksek MCC) sonuçlar ürettiği değerlendirilmektedir.

Tablo 4.3. Literatürde kalp yetmezliği sağkalım tahmini üzerine takip süresi dahil edilerek yapılan çalışmaların karşılaştırılması

Çalışma	Öznitelik Sayısı	Öznitelik Seçici	Veri Dengeleme	Sınıflandırıcı	Veri Bölme	MCC, F1, Doğruluk (Acc), Duyarlılık (Sen), Özgüllük (Spe)
(Chicco ve Jurman, 2020a)	3 (EF + SC + Takip süresi)	RF, Ekstra Ağaçlar + Pearson korelasyonu, tek değişkenli lojistik regresyon, t-testi	-	LR	Monte Carlo CV (70/30×100)	MCC: %61,6, F1: %71,9 Acc: %83,8, Sen: %78,5, Spe: %86
(Çakır, 2025)	12	-	-	LR, LDA, GBM, GLMNET, Doğrusal SVM	Hold-out (90/10) + 10-fold CV	MCC: -, F1: %97,6 Acc: %96,6, Sen: 1, Spe: %88,9
(Ishaq vd., 2021)	9 (Yaş + CPK + Diyabet + EF + Cinsiyet + Trombositler + SC + SS + Takip süresi)	Rastgele Orman Önemi (RF Importance)	SMOTE	ETC	Hold-out (70/30)	MCC: -, F1: %93 Acc: %92, Sen: %93, Spe: -
(Qadri vd., 2024)	12	Feature Engineering (RF-based)	SMOTE	RF	Hold-out (80/20) + 10-fold CV	MCC: -, F1: %98 Acc: %97,5, Sen: %98, Spe: -
(Souza ve Lima, 2024)	4 (Takip süresi + CPK + SS + Cinsiyet)	Korelasyon Filtresi (CF)	-	k-NN	Hold-out (70/30) + CV (x100)	MCC: -, F1: %88,9 Acc: %83,5, Sen: %97, Spe: %54,7
(Qadeer vd., 2025)	12	-	SMOTE	CNN	Hold-out (70/30)	MCC: -, F1: %99,78 Acc: %99,75, Sen: %99,76, Spe: -
(Moreno-Sánchez, 2023)	5 (Takip süresi + SC + EF + Yaş + Diyabet)	RFE, Karşılıklı Bilgi, ANOVA, Ki-kare	-	Ekstra Ağaçlar	Hold-out (70/30) + 5-fold CV	MCC: -, F1: %79,9 Acc: %87,1, Sen: %79,3, Spe: %90,8
(Zhang, 2025)	12	PCA	-	Rastgele Orman	CV+Test	MCC: -, F1: %66,7 Acc: %90,7, Sen: %57,1, Spe: -
Bu Çalışma	5 (Takip süresi + EF + Diyabet + Anemi + Cinsiyet)	ReliefF	-	k-NN	Monte Carlo CV (70/30 x100)	MCC: %65,94, F1: %74,56, Acc: %85,56, Sen: %66,59, Spe: %94,51

Takip süresi özniteliğinin dahil edildiği modeller yüksek performans sergilese de klinik uygulama sırasında takip süresinin önceden bilinmemesi bu modellerin gerçek zamanlı kullanımını kısıtlayabilmektedir. Literatürde yer alan çalışmalar incelendiğinde, takip süresi özniteliğinin model performansı üzerindeki etkisinin çoğu çalışmada sistematik olarak ele alınmadığı görülmektedir. Bu kısıtı aşmak ve modelin hastaneye başvuru anındaki öngörü yeteneğini test etmek amacıyla çalışmamız takip süresi özniteliği dahil edilmeden tekrarlanmıştır. Takip süresi özniteliği dahil edilmeden elde edilen bulguların literatürle karşılaştırılabilmesi amacıyla, kalp yetmezliği hasta sağkalım tahminine yönelik takip süresi özniteliğini kullanmayan önceki çalışmalar Tablo 4.4’de özetlenmiştir. Tablo 4.4’te yer alan çalışmalar metodolojik olarak incelendiğinde, öznitelik seçimi ve sınıflandırma aşamalarında geniş bir algoritma çeşitliliğinin denendiği görülmektedir. Literatürdeki bu çeşitlilik doğrultusunda Chicco ve Jurman (2020a) ile Souza ve Lima (2024) tarafından yürütülen öznitelik seçimi ve filtreleme temelli yaklaşımların yanı sıra, hibrit ve ardışık bir yaklaşım benimseyen Newaz vd. (2021) de dikkat çeken bir diğer çalışmadır. Newaz vd. (2021), Ki-Kare testi ile Özyinelemeli Öznitelik Eleme (RFE) algoritmalarını birlikte kullanarak öznitelikleri filtrelemiş; sınıflandırma aşamasında ise Lojistik Regresyon (LR), k-En Yakın Komşu (k-NN), Destek Vektör Makineleri (SVM), Karar Ağacı (DT), Rastgele Orman (RF), Gradyan Artırma, AdaBoost, MLP ve Dengelenmiş Rastgele Orman (BRF) sınıflandırıcılarını eğitmiş ve test etmiştir.

Tablo 4.4 incelendiğinde bilimiz dahilinde takip süresi özniteliği dahil edilmeden yapılan çalışma sayısının oldukça sınırlı olduğu görülmektedir. Takip süresi özniteliği hariç tutulan bu çalışmalar karşılaştırıldığında, takip süresi bileşeninin modelden çıkarılmasının literatürdeki diğer çalışmalarda da ciddi performans kayıplarına yol açtığı ve doğruluk değerlerinin %50-%70 bandına gerilediği açıkça görülmektedir. Örneğin, Chicco ve Jurman (2020a) takip süresi özniteliğini çıkardığında iki öznitelik (EF + SC) ve Rastgele Orman sınıflandırıcısı ile %58,5 doğruluk ve %41,8 MCC değerine düşmüştür. Newaz vd. (2021) ise Ki-kare ve RFE gibi öznitelik seçme yöntemleriyle öznitelik sayısını 3’e (EF, SC, yaş) düşürmüş, BRF sınıflandırıcısı ve 5-katlı çapraz doğrulama ile %76,25 doğruluk ve %52,53 MCC elde etmiştir.

Bu çalışmada ise takip süresi özniteliği dahil edilmeden, SFS + TreeBagger algoritması ile belirlenen 6 öznitelik (EF, yaş, SC, sigara kullanımı, cinsiyet, yüksek tansiyon) ve Doğrusal sınıflandırıcı (Linear classifier) kullanılarak yürütülen analizlerde %77,33 doğruluk ve %45,27 MCC değerlerine ulaşılmıştır. Elde edilen doğruluk değeri,

literatürde takip süresi özniteliği kullanmayan diğer çalışmaların üzerindedir. Daha da önemlisi, bu çalışmada ulaşılan %89,57'lük özgüllük değeri, Newaz vd. (2021) çalışmasının (%74,45) belirgin şekilde üzerindedir. Bu yüksek özgüllük oranı, takip süresi kısıtı olmaksızın hastanın ilk başvuru anında bile modelin düşük riskli hastaları (yaşayacak vakaları) yüksek bir güvenilirlikle ayırt edebildiğini ve klinik kaynakların gereksiz kullanımını önleyebileceğini göstermektedir.

Yöntemsel açıdan bakıldığında; Newaz vd. (2021) çalışmasında performans değerlendirmesi varyansı tam olarak yansıtamayan standart bir 5-katlı çapraz doğrulama (5-fold CV) üzerinden yapılmıştır. Buna karşın bu çalışmada, takip süresi özniteliği barındıran modellerde olduğu gibi, yine 100 tekrarlı Monte Carlo çapraz doğrulama protokolü korunmuştur. Katı ve rastlantısal test kümelerine dayanan bu protokole rağmen elde edilen stabil başarı, modelin takip süresinden bağımsız genelleme kabiliyetinin tesadüfi olmadığını doğrulamaktadır.

Sonuç olarak, takip süresi özniteliğinin klinik takip süreçlerindeki belirsizliği göz önüne alındığında; çalışmamızda sunulan takip süresinden bağımsız modelin, literatürdeki benzerlerine oranla hem daha yüksek bir sınıflandırma doğruluğu sunduğu hem de en katı doğrulama süreçleri altında bile istatistiksel olarak kararlı sonuçlar ürettiği değerlendirilmektedir.

Tablo 4.4. Literatürde kalp yetmezliği sağkalım tahmini üzerine takip süresi dahil edilmeden yapılan çalışmaların karşılaştırılması

Çalışma	Öznitelik Sayısı	Öznitelik Seçici	Veri Dengeleme	Sınıflandırıcı	Veri Bölme	MCC F1 Doğruluk (Acc) Duyarlılık (Sen) Özgüllük (Spe)
(Chicco ve Jurman, 2020a)	2 (EF + SC)	RF, Ekstra Ağaçlar + Pearson korelasyonu, tek değişkenli lojistik regresyon, t-testi	-	RF	Monte Carlo CV (70/30 ×100)	MCC: %41,8 F1: %75,4 Acc: %58,5 Sen: %54,1 Spe: %85,5
(Newaz vd., 2021)	3 (EF + SC + Yaş)	Ki-kare + RFE	-	BRF	5-fold CV	MCC: %52,53 F1: - Acc: %76,25 Sen: %80,21 Spe: %74,45
(Souza ve Lima, 2024)	4 (Anemi + HBP + SC + Cinsiyet)	Korelasyon Filtresi (CF)	-	k-NN	Hold-out (70/30) + CV (x100)	MCC: - F1: %81,4 Acc: %73,06 Sen: %86,6 Spe: %44,2
Bu Çalışma	6 (EF + Yaş + SC + Sigara Kullanımı + Cinsiyet + Yüksek Tansiyon)	SFS + TreeBagger	-	Doğrusal	Monte Carlo CV (70/30 ×100)	MCC: %45,27 F1: %59 Acc: %77,33 Sen: %51,36 Spe: %89,57

4.2.3. Çalışmanın Sınırlılıkları

Bu çalışmada elde edilen bulgular bazı sınırlılıklar çerçevesinde değerlendirilmelidir. Araştırma yalnızca kullanılan veri setinde mevcut olan klinik değişkenler üzerinden yürütülmüştür. Verilerin daha önce toplanmış klinik kayıtlardan elde edilmesi ve belirli bir hasta grubunu temsil etmesi, sonuçların genellenebilirliğini sınırlayan bir faktördür. Bununla birlikte, çalışmada kullanılan sınıflandırma algoritmaları ve performans değerlendirme ölçütleri model başarımını doğrudan etkilemektedir. Farklı algoritmalar, öznelik seçme yöntemleri ve değerlendirme kriterleri kullanıldığında farklı sonuçlar elde edilmesi mümkündür.

Çalışmada kullanılan veri seti 299 hasta kaydından oluşmaktadır. Veri setinin görece küçük olması bir sınırlılık olarak değerlendirilebilir. Ancak bu veri seti kalp yetmezliği tahmini üzerine yapılan makine öğrenmesi çalışmalarında literatürde yaygın olarak kullanılan veri setlerinden biridir ve farklı yöntemlerin karşılaştırmalı olarak değerlendirilmesine olanak sağlamaktadır.

Ayrıca bu çalışmada veri dengeleme yöntemleri (SMOTE, ADASYN vb.) uygulanmamıştır. Bunun temel nedeni, veri setinin doğal dağılımının korunmak istenmesidir. Sentetik veri üretimi içeren yöntemler bazı durumlarda model performansını yapay olarak artırabilmekte ve elde edilen sonuçların gerçek veri üzerindeki performansını tam olarak yansıtmayabilmektedir. Bu nedenle çalışmada modeller veri setinin doğal yapısı üzerinde değerlendirilmiştir.

Bu çalışmada kullanılan veri seti, makine öğrenmesi modellerine doğrudan uygulanabilecek şekilde yapılandırılmıştır. Yapılan incelemeler sonucunda veri setinde eksik gözlem bulunmadığı belirlenmiş ve bu nedenle herhangi bir eksik veri tamamlama işlemi uygulanmamıştır. İkili değişkenler özgün hâllerıyla modele dâhil edilmiştir. Sayısal değişkenler için ise literatürde aynı veri seti üzerinde gerçekleştirilen referans çalışmalarla metodolojik tutarlılığı korumak amacıyla herhangi bir ölçekleme, standardizasyon veya normalizasyon işlemi uygulanmamıştır.

5. SONUÇLAR VE ÖNERİLER

5.1. Sonuçlar

Bu tez çalışmasında, klinik karar destek sistemlerinde kullanılmak üzere kalp yetmezliği hastalarının sağkalım durumunu tahmin eden makine öğrenmesi tabanlı, optimize edilmiş bir mimari önerilmiştir. 299 hastadan oluşan heterojen kohort üzerinde filtre, sarmal ve gömülü öznitelik seçimi algoritmaları ile farklı matematiksel tabanlı sınıflandırıcıların kombinasyonları sistematik olarak test edilmiştir. Deneysel analizler, hastanın hastanede izlendiği gün sayısını ifade eden takip süresi (time) değişkeninin model başarısı üzerinde doğrudan belirleyici olduğunu doğrulamıştır. Takip süresi bilgisinin sisteme dâhil edildiği senaryolarda; ReliefF öznitelik seçimi yöntemi ile k-En Yakın Komşu (k-NN) sınıflandırıcısının entegrasyonu en yüksek ve kararlı performansı sunmuştur. Bu model yalnızca beş öznitelik (takip süresi, ejeksiyon fraksiyonu, diyabet, anemi, cinsiyet) kullanarak %85,56 doğruluk, %74,56 F1-skoru ve %65,94 Matthews Korelasyon Katsayısı (MCC) değerlerine ulaşmıştır. Özellikle modelin %94,51 gibi son derece yüksek bir özgüllük oranına ulaşması, sistemin hayatta kalma olasılığı yüksek olan düşük riskli hastaları ayırt etmedeki hassasiyetini kanıtlamakta ve gereksiz tıbbi müdahale ile klinik kaynak israfını engelleme potansiyeli taşımaktadır. Öte yandan, hastanın takip süresinin henüz bilinmediği ilk muayene anını simüle eden senaryoda, SFS + TreeBagger öznitelik seçimi ile Doğrusal sınıflandırıcı kombinasyonu öne çıkmıştır. Bu model; ejeksiyon fraksiyonu, yaş, serum kreatinin, sigara kullanımı, cinsiyet ve yüksek tansiyondan oluşan altı öznitelikli alt küme ile %77,33 doğruluk ve %89,57 özgüllük üreterek, takip süresinden bağımsız koşulda literatürdeki benzer kısıtlı modellerin üzerine çıkmıştır. Literatürde genellikle tekil veri bölmelerine dayalı tesadüfi yüksek başarılar raporlanırken, bu çalışmada uygulanan 100 tekrarlı Monte Carlo çapraz doğrulama yöntemi, elde edilen başarı metriklerinin istatistiksel olarak kararlı olduğunu ve aşırı uyum riskinden arındığını kanıtlamıştır. Tüm analizler genelinde ejeksiyon fraksiyonu, serum kreatinin ve yaş değişkenlerinin sürekli olarak en yüksek önem skorlarını alarak seçilmesi, yapay zekâ modelinin mekanik seçimlerinin kardiyovasküler tıp literatüründeki klinik bulgularla tam bir uyum içinde olduğunu doğrulamaktadır. Sonuç olarak bu çalışma, veri setinin doğal yapısını bozacak sentetik veri üretme tekniklerine başvurmadan, veri doğasına uygun güçlü öznitelik seçimi ve kararlı

doğrulama protokolleriyle klinikte uygulanabilir, güvenilir ve düşük maliyetli bir karar destek mimarisi sunulabileceğini ortaya koymuştur.

5.2. Öneriler

Bu tez çalışmasında önerilen modellerin klinik başarımlarını artırmak ve gerçek dünya tıbbi uygulama süreçlerine entegrasyonunu sağlamak amacıyla gelecekte yapılacak araştırmalara yönelik çeşitli öneriler sunulmaktadır. Uygulanan model mimarisi retrospektif veriler üzerinde test edildiğinden, modelin gerçek dünya koşullarındaki genellenebilirliğini artırmak adına farklı coğrafi bölgelerdeki hastaneleri kapsayan ileriye dönük ve çok merkezli klinik çalışmalarda doğrulanması gerekmektedir. Takip süresi değişkeninin performansa olan belirgin etkisi göz önüne alındığında, gelecek çalışmalarda makine öğrenmesi algoritmalarının takip süresi kısıtlı verileri doğrudan işleyebilen Cox regresyonu, Kaplan-Meier analizi veya derin öğrenme tabanlı sağkalım ağırları ile hibrit hale getirilerek dinamik bir risk skorlaması yapılması önerilmektedir. Ayrıca, önerilen modellerin pratik hekimlik süreçlerinde yer bulabilmesi için hekimin muayene anında klinik parametreleri girerek anlık mortalite risk yüzdesi ve triyaj önerisi alabileceği; hekimin laboratuvar sonuçları ve klinik bulguları pratik bir şekilde girerek saniyeler içinde anlık risk skorlaması elde edebileceği kullanıcı dostu bir Klinik Karar Destek Sistemi (KKDS) yazılım arayüzü hazırlanmalıdır.

KAYNAKLAR

- Ahmad, T., Munir, A., Bhatti, S. H., Aftab, M., ve Raza, M. A. (2017). Survival analysis of heart failure patients: A case study. *PLOS ONE*, 12(7), e0181001. <https://doi.org/10.1371/journal.pone.0181001>
- Anuragi, A., Singh Sisodia, D., ve Pachori, R. B. (2022). EEG-based cross-subject emotion recognition using Fourier-Bessel series expansion based empirical wavelet transform and NCA feature selection method. *Information Sciences*, 610, 508-524. <https://doi.org/10.1016/j.ins.2022.07.121>
- Averbuch, T., Sullivan, K., Sauer, A., Mamas, M. A., Voors, A. A., Gale, C. P., Metra, M., Ravindra, N., ve Van Spall, H. G. C. (2022). Applications of artificial intelligence and machine learning in heart failure. *European Heart Journal - Digital Health*, 3(2), 311-322. <https://doi.org/10.1093/ehjdh/ztac025>
- Bishop, C. M. (2006). *Pattern recognition and machine learning*. Springer.
- Breiman, L. (1996). Bagging predictors. *Machine Learning*, 24(2), 123-140. <https://doi.org/10.1007/BF00058655>
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5-32. <https://doi.org/10.1023/A:1010933404324>
- Bryant, F. B., ve Satorra, A. (2012). Principles and practice of scaled difference chi-square testing. *Structural Equation Modeling: A Multidisciplinary Journal*, 19(3), 372-398. <https://doi.org/10.1080/10705511.2012.687671>
- Chandrashekar, G., ve Sahin, F. (2014). A survey on feature selection methods. *Computers & Electrical Engineering*, 40(1), 16-28. <https://doi.org/10.1016/j.compeleceng.2013.11.024>
- Chicco, D., ve Jurman, G. (2020a). Machine learning can predict survival of patients with heart failure from serum creatinine and ejection fraction alone. *BMC Medical Informatics and Decision Making*, 20(1), 16. <https://doi.org/10.1186/s12911-020-1023-5>
- Chicco, D., ve Jurman, G. (2020b). The advantages of the Matthews correlation coefficient (MCC) over F1 score and accuracy in binary classification evaluation. *BMC Genomics*, 21(1), 6. <https://doi.org/10.1186/s12864-019-6413-7>
- Chicco, D., ve Jurman, G. (2023). A statistical comparison between Matthews correlation coefficient (MCC), prevalence threshold, and Fowlkes–Mallows index. *Journal of Biomedical Informatics*, 144, 104426. <https://doi.org/10.1016/j.jbi.2023.104426>
- Cook, J. A., ve Collins, G. S. (2015). The rise of big clinical databases. *The British Journal of Surgery*, 102(2), e93-e101. <https://doi.org/10.1002/bjs.9723>
- Cover, T., ve Hart, P. (1967). Nearest neighbor pattern classification. *IEEE Transactions on Information Theory*, 13(1), 21-27. <https://doi.org/10.1109/TIT.1967.1053964>

- Çakır, M. (2025). Comparison of machine learning models in heart failure prediztion and their integration into clinical decision support systems. *International Journal of 3D Printing Technologies and Digital Industry*, 9(2), 272-282. <https://doi.org/10.46519/ij3dptdi.1724620>
- Dietterich, T. G. (2000). Ensemble methods in machine learning. In J. Kittler ve F. Roli (Ed.), *Multiple Classifier Systems* (C. 1857, s. 1-15). Springer Berlin Heidelberg. https://doi.org/10.1007/3-540-45014-9_1
- Domingos, P., ve Pazzani, M. (1997). On the Optimality of the Simple Bayesian Classifier under Zero-One Loss. *Machine Learning*, 29, 103-130.
- Dou, J., Song, Y., Wei, G., ve Zhang, Y. (2022). Fuzzy information decomposition incorporated and weighted Relief-F feature selection: When imbalanced data meet incomplection. *Information Sciences*, 584, 417-432. <https://doi.org/10.1016/j.ins.2021.10.057>
- Fawcett, T. (2006). An introduction to ROC analysis. *Pattern Recognition Letters*, 27(8), 861-874. <https://doi.org/10.1016/j.patrec.2005.10.010>
- Fisher, R. A. (1936). The use of multiple measurements in taxonomic problems. *Annals of Eugenics*, 7(2), 179-188. <https://doi.org/10.1111/j.1469-1809.1936.tb02137.x>
- Freund, Y., ve Schapire, R. E. (1997). A decision-theoretic generalization of on-line learning and an application to boosting. *Journal of Computer and System Sciences*, 55(1), 119-139. <https://doi.org/10.1006/jcss.1997.1504>
- Goldberger, J., Hinton, G. E., Roweis, S. T., ve Salakhutdinov, R. R. (2004). *Neighbourhood Components Analysis*. 513-520.
- Han, J., Kamber, M., ve Pei, J. (2012). *Data mining: Concepts and techniques* (3rd ed.). Morgan Kaufmann.
- Hand, D. J., ve Yu, K. (2001). Idiot's bayes—Not so stupid after all? *International Statistical Review*, 69(3), 385-398. <https://doi.org/10.1111/j.1751-5823.2001.tb00465.x>
- Hastie, T., Tibshirani, R., ve Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction* (2nd ed.). Springer.
- Heaton, J. (2018). Ian Goodfellow, Yoshua Bengio, and Aaron Courville: Deep learning: The MIT Press, 2016, 800 pp, ISBN: 0262035618. *Genetic Programming and Evolvable Machines*, 19(1-2), 305-307. <https://doi.org/10.1007/s10710-017-9314-z>
- Hosmer, D. W., ve Lemeshow, S. (2000). *Applied Logistic Regression* (2nd ed.). John Wiley ve Sons.
- Hou, X., Chen, Y., Wu, K., Zhou, Y., Lu, J., ve Weng, X. (2024). A fuzzy rough granular ensemble learning based on the feature selection with chi-square. *Journal of Intelligent ve Fuzzy Systems*, 46(6201-6217). <https://doi.org/10.3233/JIFS-234510>

- Ishaq, A., Sadiq, S., Umer, M., Ullah, S., Mirjalili, S., Rupapara, V., ve Nappi, M. (2021). Improving the prediction of heart failure patients' survival using SMOTE and effective data mining techniques. *IEEE Access*, 9, 39707-39716. <https://doi.org/10.1109/ACCESS.2021.3064084>
- Jurman, G., Riccadonna, S., ve Furlanello, C. (2012). A comparison of MCC and CEN error measures in multi-class prediction. *PLoS ONE*, 7(8), e41882. <https://doi.org/10.1371/journal.pone.0041882>
- Kononenko, I. (1994). Estimating attributes: Analysis and extensions of RELIEF. In F. Bergadano ve L. Raedt (Ed.), *Machine Learning: ECML-94* (C. 784, s. 171-182). Springer Berlin Heidelberg. https://doi.org/10.1007/3-540-57868-4_57
- Liu, H., ve Setiono, R. (1997). Feature selection via discretization. *IEEE Transactions on Knowledge and Data Engineering*, 9(3), 545-556. <https://doi.org/10.1109/69.617912>
- MathWorks. (2024). *MATLAB* (Versiyon R2024a) [Yazılım]. The MathWorks, Inc. <https://www.mathworks.com/>
- McHugh, M. L. (2013). The Chi-square test of independence. *Biochemia Medica*, 23(2), 143-149. <https://doi.org/10.11613/BM.2013.018>
- Moreno-Sánchez, P. A. (2023). Improvement of a prediction model for heart failure survival through explainable artificial intelligence. *Frontiers in Cardiovascular Medicine*, 10, 1219586. <https://doi.org/10.3389/fcvm.2023.1219586>
- Murphy, K. P. (2012). *Machine learning: A probabilistic perspective* (1st ed.). MIT Press.
- National Heart, Lung, and Blood Institute. (2022). *What is heart failure?* <https://www.nhlbi.nih.gov/health/heart-failure>
- Newaz, A., Ahmed, N., ve Haq, F. S. (2021). Survival prediction of heart failure patients using machine learning techniques. *Informatics in Medicine Unlocked*, 26, 100772. <https://doi.org/10.1016/j.imu.2021.100772>
- Nielsen, M. (2016). *Neural networks and deep learning* [Online book (web book)]. Deep Learning. <http://neuralnetworksanddeeplearning.com/>
- Nugrahaeni, R. A., ve Mutijarsa, K. (2016). Comparative analysis of machine learning KNN, SVM, and random forests algorithm for facial expression classification. *2016 International Seminar on Application for Technology of Information and Communication (ISEMantic)*, 163-168. <https://doi.org/10.1109/ISEMANTIC.2016.7873831>
- Pandey, A., ve Jain, A. (2017). Comparative analysis of KNN algorithm using various normalization techniques. *International Journal of Computer Network and Information Security*, 9(11), 36-42. <https://doi.org/10.5815/ijcnis.2017.11.04>
- Peng, H., Long, F., ve Ding, C. (2005). Feature selection based on mutual information criteria of max-dependency, max-relevance, and min-redundancy. *IEEE*

- Transactions on Pattern Analysis and Machine Intelligence*, 27(8), 1226-1238.
<https://doi.org/10.1109/TPAMI.2005.159>
- Powers, D. (2007). *Evaluation: From precision, recall and F-factor to ROC, informedness, markedness & correlation* (SIE-07-001). Flinders University of South Australia.
- Pudil, P., Novovičová, J., ve Kittler, J. (1994). Floating search methods in feature selection. *Pattern Recognition Letters*, 15(11), 1119-1125.
[https://doi.org/10.1016/0167-8655\(94\)90127-9](https://doi.org/10.1016/0167-8655(94)90127-9)
- Qadeer, M., Ayaz, R., ve Thohir, M. I. (2025). Heart Failure Prediction Through a Comparative Study of Machine Learning and Deep Learning Models. *The 7th International Global Conference Series on ICT Integration in Technical Education & Smart Society*, 61. <https://doi.org/10.3390/engproc2025107061>
- Qadri, A. M., Hashmi, M. S. A., Raza, A., Zaidi, S. A. J., ve Rehman, A. U. (2024). Heart failure survival prediction using novel transfer learning based probabilistic features. *PeerJ Computer Science*, 10, e1894. <https://doi.org/10.7717/peerj-cs.1894>
- Rahimi, A., ve Recht, B. (2007). *Random features for large-scale kernel machines*. 20, 1177-1184.
- Salzberg, S. L. (1994). C4.5: Programs for machine learning by J. Ross Quinlan. *Machine Learning*, 16(3), 235-240. <https://doi.org/10.1007/BF00993309>
- Schölkopf, B., ve Smola, A. J. (2002). *Learning with kernels: Support vector machines, regularization, optimization, and beyond*. MIT Press.
- Shanmugarajeshwari, V., ve Lawrance, R. (2016). Analysis of students' performance evaluation using classification techniques. *2016 International Conference on Computing Technologies and Intelligent Data Engineering (ICCTIDE'16)*, 1-7.
<https://doi.org/10.1109/ICCTIDE.2016.7725375>
- Siddique, H. (2019). *UK heart disease fatalities on the rise for first time in 50 years*. The Guardian. <https://www.theguardian.com/society/2019/may/13/heart-circulatory-disease-fatalities-on-rise-in-uk>
- Sokolova, M., ve Lapalme, G. (2009). A systematic analysis of performance measures for classification tasks. *Information Processing & Management*, 45(4), 427-437.
<https://doi.org/10.1016/j.ipm.2009.03.002>
- Souza, V. S., ve Lima, D. A. (2024). Cardiac disease diagnosis using k-nearest neighbor algorithm: A study on heart failure clinical records dataset. *Artificial Intelligence and Applications*, 3(1), 56-71. <https://doi.org/10.47852/bonviewAIA42022045>
- Tang, J., Alelyani, S., ve Liu, H. (2014). Feature selection for classification: A review. In *Data Classification: Algorithms and Applications* (s. 37-64). CRC Press.
- Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 58(1), 267-288.
<https://doi.org/10.1111/j.2517-6161.1996.tb02080.x>

- Wang, J. (2021). Heart failure prediction with machine learning: A comparative study. *Journal of Physics: Conference Series*, 2031(1), 012068. <https://doi.org/10.1088/1742-6596/2031/1/012068>
- Wang, Z., Zhang, Y., Chen, Z., Yang, H., Sun, Y., Kang, J., Yang, Y., ve Liang, X. (2016). Application of ReliefF algorithm to selecting feature sets for classification of high resolution remote sensing image. *2016 IEEE International Geoscience and Remote Sensing Symposium (IGARSS)*, 755-758. <https://doi.org/10.1109/IGARSS.2016.7729190>
- Williams, C. K. I., ve Seeger, M. (2001). *Using the Nyström method to speed up kernel machines*. 13, 682-688.
- World Health Organization. (2025). *Cardiovascular diseases (CVDs)*. World Health Organization. [https://www.who.int/news-room/fact-sheets/detail/cardiovascular-diseases-\(cvds\)](https://www.who.int/news-room/fact-sheets/detail/cardiovascular-diseases-(cvds))
- Wu, X., Kumar, V., Ross Quinlan, J., Ghosh, J., Yang, Q., Motoda, H., McLachlan, G. J., Ng, A., Liu, B., Yu, P. S., Zhou, Z.-H., Steinbach, M., Hand, D. J., ve Steinberg, D. (2008). Top 10 algorithms in data mining. *Knowledge and Information Systems*, 14(1), 1-37. <https://doi.org/10.1007/s10115-007-0114-2>
- Zhang, H. (2025). Comparison of prediction models for heart failure related data. *Proceedings of the 2nd International Conference on Innovations in Applied Mathematics, Physics, and Astronomy*, 324-330. <https://doi.org/10.5220/0013825200004708>
- Zhou, Z.-H. (2012). *Ensemble Methods: Foundations and Algorithms*. Chapman ve Hall/CRC.

EKLER

EK-1 Tüm Öznitelikleri İçeren Veri Kümesi İçin En İyi Model Başarım Sonuçları (Zaman Dahil)

Öznitelik Seçici	Sınıflandırıcı	ÖS	Doğruluk (%) ort±ss	Duyarlılık (%) ort±ss	Özgüllük (%) ort±ss	F1-skoru (%) ort±ss	MCC (%) ort±ss
mRMR	Topluluk	10	84,02±3	72,97±8	89,27±5	74,56±6	63,43±8
mRMR	Topluluk	11	83,93±2	70,95±7	90,06±3	73,80±4	62,72±6
mRMR	Topluluk	9	83,93±3	71,03±7	90,2±3	73,80±5	62,68±7
mRMR	Topluluk	6	83,74±3	73,42±8	88,58±4	74,14±6	62,65±8
mRMR	Lojistik	1	84,08±3	64,97±7	93,08±3	72,18±5	62,33±7
ReliefF	kNN	5	85,56±2	66,59±7	94,51±2	74,56±5	65,94±6
ReliefF	kNN	7	85,29±2	66,48±7	94,18±2	74,22±5	65,26±6
ReliefF	Topluluk	9	84,80±3	72,61±7	90,54±4	75,33±5	64,82±7
ReliefF	kNN	3	85,09±2	65,69±6	94,25±2	73,78±5	64,78±7
ReliefF	kNN	9	85,07±2	65,32±7	94,38±2	73,56±5	64,68±7
Ki-kare	kNN	2	85,49±2	66,43±7	94,47±2	74,37±5	65,71±7
Ki-kare	kNN	3	85,34±2	65,70±8	94,53±2	73,84±5	65,28±7
Ki-kare	Topluluk	9	84,15±3	71,69±7	90,00±4	74,26±5	63,36±7
Ki-kare	SVM	3	83,89±3	72,69±8	89,17±4	74,21±5	62,90±8
Ki-kare	Topluluk	3	83,92±3	72,21±8	89,43±4	74,06±5	62,85±7
NCA	kNN	4	85,28±2	65,62±7	94,53±2	73,85±5	65,22±7
NCA	Topluluk	9	85,00±3	71,90±8	91,16±3	75,23±5	65,00±7
NCA	kNN	5	85,10±2	65,69±7	94,26±2	73,69±5	64,78±7
NCA	Topluluk	6	84,11±3	71,28±8	90,16±4	74,11±5	63,23±8
NCA	Topluluk	5	83,82±3	72,62±8	89,09±4	74,09±5	62,81±7
SFS	Topluluk	12	84,10±3	71,90±7	89,86±3	74,25±5	63,16±7
SFS	DA	12	82,61±3	67,20±8	89,82±3	70,97±6	59,13±7
SFS	DT	12	78,18±4	64,65±10	84,54±5	65,22±7	49,92±9
SFS	YSA	12	77,88±4	56,31±13	88,01±5	61,12±11	47,30±12
SFS	NB	12	76,6±3	45,5±9	91,21±3	54,97±8	42,45±9
LASSO	DA	7	83,84±3	69,57±7	90,58±3	73,28±5	62,20±7
LASSO	Topluluk	7	83,47±3	69,99±8	89,83±3	72,91±6	61,54±7
LASSO	Doğrusal	7	82,73±3	65,28±7	90,97±4	70,70±6	59,42±8
LASSO	Lojistik	7	82,54±3	65,24±9	90,69±4	70,31±6	58,89±8
LASSO	YSA	7	81,34±3	63,89±11	89,58±5	68,30±8	56,35±9
Tree	kNN	3	85,45±2	66,16±7	94,53±2	74,24±5	65,59±7
Tree	SVM	3	84,13±3	73,45±7	89,16±3	74,67±5	63,38±8
Tree	Topluluk	6	84,15±3	71,08±7	90,32±4	74,12±5	63,28±7
Tree	Topluluk	5	83,85±3	72,23±9	89,34±4	73,93±6	62,77±8
Tree	Topluluk	11	83,94±3	70,63±7	90,23±3	73,72±5	62,63±8
TreeBagger	kNN	4	85,47±2	66,07±7	94,61±2	74,26±5	65,75±6
TreeBagger	SVM	3	84,47±2	74,34±7	89,25±3	75,26±4	64,32±6
TreeBagger	kNN	3	84,88±3	63,90±7	94,76±2	72,87±6	64,23±7
TreeBagger	Topluluk	11	84,30±3	71,77±7	90,19±3	74,43±5	63,56±7
TreeBagger	Topluluk	4	84,18±3	71,84±8	89,96±3	74,25±5	63,27±7

EK-2 Tüm Öznitelikleri İçeren Veri Kümesi İçin En İyi Model Başarım Sonuçları (Zaman Hariç)

Öznitelik Seçici	Sınıflandırıcı	ÖS	Doğruluk (%) ort±ss	Duyarlılık (%) ort±ss	Özgüllük (%) ort±ss	F1-skoru (%) ort±ss	MCC (%) ort±ss
mRMR	Doğrusal	4	76,93±3	49,14±10	89,97±4	57,11±8	43,99±9
mRMR	Doğrusal	3	76,34±3	50,52±8	88,49±4	57,46±6	42,92±8
mRMR	Topluluk	3	75,04±3	57,60±9	83,28±4	59,37±7	41,90±9
mRMR	Doğrusal	8	76,08±3	46,54±9	89,99±4	55,09±7	41,72±8
mRMR	DA	4	76,20±3	44,93±9	90,94±4	54,36±7	41,70±9
ReliefF	Doğrusal	9	75,69±3	46,19±9	89,56±3	54,44±7	40,59±7
ReliefF	Topluluk	10	73,80±3	50,11±9	84,95±4	54,73±7	37,48±8
ReliefF	DA	11	74,43±3	43,61±8	88,91±4	51,85±7	37,23±9
ReliefF	Lojistik	9	74,40±3	43,44±8	88,96±4	51,67±7	37,14±8
ReliefF	Topluluk	9	73,61±3	49,70±9	84,85±4	54,36±7	36,87±9
Ki-kare	Doğrusal	3	76,58±3	48,36±9	89,93±4	56,56±7	43,22±8
Ki-kare	Lojistik	3	76,60±3	48,42±8	89,83±3	56,67±7	43,02±9
Ki-kare	DA	3	76,12±3	44,19±9	91,17±3	53,82±8	41,31±9
Ki-kare	Lojistik	4	75,91±3	45,79±7	90,15±3	54,72±6	41,18±8
Ki-kare	Doğrusal	6	76,00±3	44,09±8	91,00±4	53,72±7	41,09±9
Ki-kare	Doğrusal	5	76,01±3	44,21±9	90,99±3	53,73±7	41,04±8
NCA	Doğrusal	4	75,83±2	45,30±7	90,26±4	54,29±5	41,04±6
NCA	Topluluk	2	74,33±4	57,60±8	82,17±5	58,79±6	40,51±9
NCA	DA	4	75,69±3	43,31±8	90,98±4	53,03±7	40,28±9
NCA	Lojistik	4	75,45±2	47,30±8	88,69±3	54,88±6	40,17±7
NCA	DA	7	75,04±3	42,65±8	90,30±4	51,89±8	38,50±9
SFS	Topluluk	10	74,16±4	51,29±9	84,95±4	55,74±7	38,48±10
SFS	Topluluk	11	73,96±3	48,61±9	85,89±4	54,09±7	37,30±8
SFS	Topluluk	9	73,90±3	48,77±9	85,75±4	54,19±7	37,19±9
SFS	DA	8	74,48±3	42,32±8	89,60±3	51,06±7	36,95±8
SFS	Doğrusal	6	74,57±2	39,91±13	90,96±4	48,57±14	36,33±11
LASSO	Topluluk	7	74,35±3	52,70±8	84,58±4	56,66±6	39,16±8
LASSO	Topluluk	6	74,24±3	52,26±9	84,57±4	56,18±6	38,78±8
LASSO	Topluluk	8	74,20±3	51,00±9	85,08±4	55,48±6	38,30±8
LASSO	Doğrusal	7	74,97±3	39,97±9	91,49±4	50,07±8	38,14±8
LASSO	Doğrusal	8	74,70±3	38,50±10	91,82±4	48,59±10	37,13±10
Tree	Lojistik	4	76,53±3	49,38±8	89,28±4	57,11±6	43,15±8
Tree	Lojistik	3	76,35±3	47,83±8	89,80±4	56,19±6	42,53±8
Tree	Doğrusal	3	76,15±3	49,21±9	88,92±4	56,66±7	42,31±9
Tree	Topluluk	3	74,70±3	56,96±8	83,10±5	58,92±6	41,18±8
Tree	Doğrusal	4	75,79±3	45,36±8	90,12±4	54,16±7	40,88±8
TreeBagger	DT	1	71,79±2	28,06±10	92,45±5	37,74±9	43,30±8
TreeBagger	DT	2	74,51±3	52,07±9	85,10±5	56,46±6	43,08±8
TreeBagger	DT	3	72,83±4	54,51±9	81,49±5	56,12±7	42,88±9
TreeBagger	DT	4	72,52±4	54,50±11	80,99±5	55,56±7	42,00±9
TreeBagger	DT	5	72,02±4	52,25±9	81,37±5	54,27±8	41,75±10

EK-3 Sayısal ve İkili Olarak Ayrıştırılmış Veri Kümesi İçin En İyi Model Başarım Sonuçları (Zaman Dahil)

Sayısal	İkili	Sınıflandırıcı	ÖS	Doğruluk (%) ort±ss	Duyarlılık (%) ort±ss	Özgüllük (%) ort±ss	F1-skoru (%) ort±ss	MCC (%) ort±ss
mRMR	mRMR	Lojistik	2	83,57±3	65,23±7	92,18±4	71,62±5	61,27±8
mRMR	Tree	Lojistik	2	83,54±3	64,83±7	92,35±4	71,59±6	61,19±8
mRMR	Ki-kare	Lojistik	2	83,07±3	64,71±7	91,71±4	70,93±6	60,05±8
mRMR	Tree	Topluluk	12	82,81±3	68,29±9	89,68±3	71,63±5	59,94±7
mRMR	mRMR	DA	2	82,55±3	67,33±7	89,74±5	71,22±6	59,41±8
ReliefF	mRMR	DT	4	80,96±3	67,78±8	87,18±5	69,44±6	62,13±8
ReliefF	mRMR	DT	6	80,12±4	65,96±10	86,78±6	67,79±6	62,08±8
ReliefF	mRMR	DT	8	79,09±3	65,77±9	85,35±5	66,63±6	61,09±9
ReliefF	mRMR	DT	10	79,49±3	65,31±10	86,17±4	66,74±7	60,99±9
ReliefF	mRMR	DT	12	78,56±3	65,22±10	84,84±5	65,81±6	60,78±8
NCA	mRMR	Topluluk	12	82,75±2	68,48±8	89,49±3	71,59±5	59,86±7
NCA	Tree	DA	12	82,53±3	67,19±7	89,73±4	71,01±5	59,20±7
NCA	Tree	Doğrusal	6	82,80±2	60,23±8	93,41±3	68,90±5	59,14±7
NCA	Tree	Topluluk	12	82,45±3	67,67±8	89,42±3	71,05±5	59,04±7
NCA	mRMR	Doğrusal	6	82,62±3	60,10±8	93,24±3	68,64±6	58,70±7
SFS	Ki-kare	Lojistik	4	83,80±3	63,82±8	93,20±3	71,43±5	61,71±7
SFS	Tree Bagger	Lojistik	2	83,67±3	65,79±7	92,11±4	72,06±6	61,68±8
SFS	Tree	Lojistik	4	83,69±3	63,61±8	93,13±3	71,21±6	61,41±8
SFS	Tree Bagger	Lojistik	4	83,40±3	63,69±7	92,68±3	70,96±5	60,69±7
SFS	Tree	Doğrusal	4	83,44±2	60,02±8	94,52±2	69,70±6	60,64±7
LASSO	mRMR	DA	12	82,91±3	68,07±8	89,90±3	71,72±6	60,05±8
LASSO	mRMR	Topluluk	12	82,75±3	68,43±8	89,50±4	71,66±5	59,84±7
LASSO	Ki-kare	Topluluk	12	82,80±3	67,61±8	89,92±4	71,37±6	59,75±8
LASSO	Tree Bagger	Topluluk	12	82,64±3	68,95±9	89,08±4	71,56±6	59,69±8
LASSO	Tree	DA	12	82,53±33	67,41±8	89,63±3	71,01±5	59,09±8
Tree	mRMR	Topluluk	10	83,44±3	69,24±8	90,16±3	72,68±6	61,37±8
Tree	Ki-kare	kNN	6	83,47±3	67,38±9	91,07±3	72,12±6	61,18±8
Tree	Tree	kNN	6	83,49±2	66,80±8	91,36±3	71,99±	61,17±6
Tree	Ki-kare	kNN	8	83,44±2	67,72±7	90,82±3	72,16±5	61,04±7
Tree	Tree Bagger	kNN	6	83,28±3	67,83±8	90,52±4	71,97±6	60,84±7
TreeBagger	Tree	Lojistik	2	84,01±3	65,07±7	92,94±3	72,20±5	62,24±8
TreeBagger	Tree	kNN	4	83,69±3	69,70±7	90,31±4	73,23±5	62,09±7
TreeBagger	Tree Bagger	kNN	4	83,45±3	69,52±8	90,00±4	72,77±6	61,51±8
TreeBagger	Ki-kare	Lojistik	2	83,56±3	66,75±7	91,47±4	72,15±5	61,39±8
TreeBagger	Tree	Doğrusal	10	83,36±2	69,52±8	89,88±3	72,67±5	61,26±7

EK-4 Sayısal ve İkili Olarak Ayrıştırılmış Veri Kümesi İçin En İyi Model Başarım Sonuçları (Zaman Hariç)

Sayısal	İkili	Sınıflan dırıcı	ÖS	Doğruluk (%) ort±ss	Duyarlılık (%) ort±ss	Özgüllük (%) ort±ss	F1-skoru (%) ort±ss	MCC (%) ort±ss
mRMR	Tree	Doğrusal	6	76,88±3	51,29±8	88,98±4	58,49±7	44,41±8
mRMR	Tree Bagger	Doğrusal	6	76,26±3	49,14±8	89,06±3	56,73±6	42,43±8
mRMR	mRMR	Doğrusal	8	76,25±2	46,51±9	90,28±3	55,23±7	42,04±7
mRMR	Tree Bagger	Doğrusal	8	76,16±2	46,68±8	90,12±3	55,37±7	41,89±7
mRMR	Ki-kare	Doğrusal	8	76,22±3	45,24±9	90,83±4	54,54±7	41,77±9
ReliefF	Tree	Doğrusal	6	76,69±3	49,39±9	89,55±4	57,11±8	43,55±9
ReliefF	Tree Bagger	Doğrusal	6	76,39±3	49,95±7	88,86±4	57,45±6	43,04±8
ReliefF	mRMR	Doğrusal	6	76,12±3	49,17±8	88,84±4	56,69±6	42,20±8
ReliefF	Tree	Lojistik	6	76,10±3	48,33±9	89,21±4	56,07±7	42,06±9
ReliefF	Ki-kare	Doğrusal	6	76,01±3	49,29±8	88,62±4	56,59±6	41,98±8
NCA	mRMR	DA	11	73,87±3	43,03±8	88,37±4	51,00±7	35,78±9
NCA	Tree	DA	11	73,82±3	42,70±9	88,47±4	50,65±7	35,54±9
NCA	Ki-kare	DA	11	73,84±3	40,99±8	89,37±4	49,79±7	35,48±9
NCA	Tree Bagger	DA	11	73,35±3	41,66±8	88,24±4	49,67±6	34,28±8
NCA	Ki-kare	Doğrusal	8	73,73±3	33,42±10	92,73±3	44,02±311	33,38±10
SFS	Tree Bagger	Doğrusal	6	77,33±3	51,36±8	89,57±4	59,00±6	45,27±8
SFS	Tree	Doğrusal	6	76,54±2	49,46±9	89,26±4	56,95±7	43,17±7
SFS	Tree	Lojistik	6	76,36±3	48,06±8	89,69±4	56,29±6	42,54±8
SFS	Ki-kare	Doğrusal	8	76,11±3	46,59±8	90,10±4	55,24±7	41,94±7
SFS	mRMR	Doğrusal	6	75,85±2	49,48±7	88,35±3	56,59±6	41,63±7
LASSO	Tree Bagger	DA	8	74,98±3	43,66±8	89,73±3	52,48±7	38,54±8
LASSO	mRMR	Doğrusal	8	75,00±3	43,58±11	89,83±4	51,78±11	38,35±11
LASSO	Tree	Doğrusal	6	74,73±3	43,67±9	89,37±4	52,01±8	38,21±8
LASSO	Tree Bagger	Doğrusal	8	74,84±3	43,51±10	89,59±4	51,77±10	38,04±9
LASSO	Ki-kare	Doğrusal	8	74,73±3	44,10±11	89,16±4	51,86±11	37,80±10
Tree	Tree	Doğrusal	6	76,69±3	51,34±8	88,67±4	58,34±7	43,95±9
Tree	Tree	Lojistik	6	76,66±3	48,13±8	90,11±3	56,61±6	43,21±8
Tree	mRMR	Doğrusal	6	76,27±2	48,81±7	89,18±3	56,55±6	42,36±7
Tree	Tree Bagger	Lojistik	6	76,15±3	48,23±7	89,32±4	56,24±36	42,19±8
Tree	Tree Bagger	Doğrusal	6	76,18±3	49,24±9	88,84±4	56,55±8	42,18±9
TreeBagger	Tree Bagger	Doğrusal	6	76,75±3	50,08±8	89,26±3	57,54±6	43,60±7
TreeBagger	Ki-kare	Doğrusal	6	76,08±3	48,44±8	89,12±4	56,20±37	41,96±8
TreeBagger	Tree Bagger	Lojistik	6	7603±3	47,94±7	89,34±3	56,08±6	41,78±8
TreeBagger	Tree	Doğrusal	6	75,67±3	48,87±9	88,39±4	56,03±7	41,22±9
TreeBagger	Tree	Doğrusal	8	75,81±3	46,74±9	89,56±4	54,81±8	41,11±8