



T.C.

BARTIN ÜNİVERSİTESİ

LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ

YÖNETİM BİLİŞİM SİSTEMLERİ ANABİLİM DALI

YÜKSEK LİSANS TEZİ

ÖĞRENCİNİN AKADEMİK PERFORMANSINI TAHMİN ETMEDE
AÇIKLANABİLİR YAPAY ZEKA VE MAKİNE ÖĞRENİMİ
KULLANIMI

MÜJDE DİLEK SOBUTAY

DANIŞMAN

PROF. DR. ÜMMÜHAN AVCI

BARTIN-2026



T.C.

BARTIN ÜNİVERSİTESİ

LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ

YÖNETİM BİLİŞİM SİSTEMLERİ ANABİLİM DALI

**ÖĞRENCİNİN AKADEMİK PERFORMANSINI TAHMİN ETMEDE
AÇIKLANABİLİR YAPAY ZEKA VE MAKİNE ÖĞRENİMİ KULLANIMI**

YÜKSEK LİSANS TEZİ

Müjde Dilek SOBUTAY

JÜRİ ÜYELERİ

Danışman : Prof. Dr. Ümmühan AVCI

Üye : Prof. Dr. Hatice YILDIZ DURAK

Üye : Dr. Öğr. Üyesi Gökhan KUTLUANA

BARTIN-2026

KABUL VE ONAY

Müjde Dilek SOBUTAY tarafından hazırlanan “ÖĞRENCİNİN AKADEMİK PERFORMANSINI TAHMİN ETMEDE AÇIKLANABİLİR YAPAY ZEKA VE MAKİNE ÖĞRENİMİ KULLANIMI” başlıklı bu çalışma, 27.03.2026 tarihinde yapılan savunma sınavı sonucunda oy birliği ile başarılı bulunarak jürimiz tarafından Yüksek Lisans Tezi olarak kabul edilmiştir.

Başkan : Prof. Dr. Hatice YILDIZ DURAK

Üye : Prof. Dr. Ümmühan AVCI

Üye : Dr. Öğr. Üyesi Gökhan KUTLUANA

Bu tezin kabulü Lisansüstü Eğitim Enstitüsü Yönetim Kurulu’nun/...../20... tarih ve 20...../.....-..... sayılı kararıyla onaylanmıştır.

Prof. Dr. Mustafa Sabri GÖK
Enstitü Müdürü

BEYANNAME

Bartın Üniversitesi Lisansüstü Eğitim Enstitüsü tez yazım kılavuzuna göre Prof. Dr. Ümmühan AVCI danışmanlığında hazırlamış olduğum “ÖĞRENCİNİN AKADEMİK PERFORMANSINI TAHMİN ETMEDE AÇIKLANABİLİR YAPAY ZEKA ve MAKİNE ÖĞRENİMİ KULLANIMI” başlıklı yüksek lisans tezimin bilimsel etik değerlere ve kurallara uygun, özgün bir çalışma olduğunu, aksinin tespit edilmesi halinde her türlü yasal yaptırımını kabul edeceğimi beyan ederim.

27.03.2026

Müjde Dilek SOBUTAY

YAPAY ZEKÂ KULLANIMINA İLİŞKİN BEYAN

☐ Tez yazım sürecinde yapay zekâ tabanlı araçlardan **faydalanmadım**.

☒ Tez yazım sürecinde yapay zekâ tabanlı araçlardan **faydalandım**. Bu kapsamda, yapay zekâ destekli araçların kullanımına ilişkin aşağıdaki hususları beyan ederim:

1. Bu çalışmada, yalnızca bilimsel üretim sürecini desteklemek amacıyla yapay zekâ tabanlı araçlardan yararlanılmış olup; kullanılan araçların işlevi, katkı düzeyi ve sınırları açık biçimde eklerde belirtilmiştir. Bu kapsamda kullanılan yapay zekânın kapsamı hipotez geliştirme, teorik tartışma ve yorumlama gibi üst düzey beceri, deneyim ve uzmanlık gerektiren aşamaları içermemektedir.
2. Yapay zekâ kullanımı, çalışmanın özgünlüğüne, bilimsel etik kurallarına ve ilgili mevzuata aykırı olmayacak biçimde gerçekleştirilmiştir.
3. Usulsüz yapay zekâ kullanımı (örneğin: kaynak göstermeksizin metin üretimi, otomatik içerik üretiminin özgün eser gibi sunulması, yanıltıcı kaynak gösterimi vb.) tarafımdan kesinlikle yapılmamıştır.
4. Çalışmanın hazırlanmasında kullanılan yapay zekâ araçları ve işlevleri ekte sunulan tabloda ayrıntılı olarak belirtilmiş olup bu kullanım danışman öğretim üyesinin bilgisi dahilinde gerçekleştirilmiştir.
5. Yukarıda yer alan beyanımın gerçeğe aykırı olduğunun tespiti hâlinde, yükseköğretim mevzuatı, ilgili etik kurallar ve Bartın Üniversitesi disiplin mevzuatı çerçevesinde hakkımda yaptırım uygulanabileceğini bildiğimi ve bu sonuçları kabul ettiğimi taahhüt ederim.

27.03.2026

Müjde Dilek SOBUTAY

ÖN SÖZ

Bu tezin hazırlanması sürecinde bilgi, birikim ve deneyimleriyle bana her zaman yol gösteren; değerli zamanını ayırarak çalışmamın her aşamasında sabırla destek olan danışman hocam Prof. Dr. Ümmühan AVCI'ya en içten teşekkürlerimi sunarım.

Zorlandığım her anda yanımda olarak bana güç veren sevgili eşim Utku SOBUTAY'a, bu yoğun süreçte varlığıyla bana huzur ve mutluluk veren minik oğlum Destan Alp'e, manevi desteğini her zaman hissettiren annem Gülay DANIŞMAN'a ve bana daima inanıp destek olan tüm aile üyelerime gönülden teşekkür ederim.

Müjde Dilek SOBUTAY

ÖZET

Yüksek Lisans Tezi

ÖĞRENCİNİN AKADEMİK PERFORMANSINI TAHMİN ETMEDE AÇIKLANABİLİR YAPAY ZEKA VE MAKİNE ÖĞRENİMİ KULLANIMI

Müjde Dilek SOBUTAY

Bartın Üniversitesi

Lisansüstü Eğitim Enstitüsü

Yönetim Bilişim Sistemleri Anabilim Dalı

Tez Danışmanı: Prof. Dr. Ümmühan AVCI

Bartın-2026, sayfa: 211

Bu çalışmanın amacı, öğrencilerin akademik performanslarının makine öğrenmesi yöntemleri kullanılarak tahmin edilmesi ve geliştirilen modellerin açıklanabilir yapay zekâ teknikleri ile yorumlanabilir hâle getirilmesidir. Günümüzde eğitim kurumlarında öğrenci bilgi sistemleri ve dijital öğrenme ortamları aracılığıyla büyük miktarda veri üretilmektedir. Bu verilerin makine öğrenmesi yöntemleriyle analiz edilmesi, öğrencilerin akademik performanslarının erken aşamada tahmin edilmesine ve risk altındaki öğrencilerin belirlenmesine olanak sağlamaktadır. Açıklanabilir yapay zekâ yaklaşımları ise bu tahminlerin hangi değişkenlere dayandığını ortaya koyarak model sonuçlarının daha şeffaf, yorumlanabilir ve güvenilir hâle gelmesini desteklemektedir.

Araştırma kapsamında öğrencilerin akademik performansını etkileyebileceği düşünülen çok boyutlu değişkenler kullanılarak bir veri seti oluşturulmuştur. Çalışmada kullanılan veriler öğrenci bilgi sistemi (UBYS) kayıtları ve öğrencilerden elde edilen anket verilerinden oluşmaktadır. Veri setinde öğrencilerin sınav notları, devam durumu, çevrimiçi öğrenme kaynaklarının kullanımı, sosyal medya kullanımı, teknoloji kullanımı ve sosyo-demografik özellikleri gibi çeşitli değişkenler yer almaktadır. Veri ön işleme sürecinde eksik verilerin kontrol edilmesi, kategorik değişkenlerin sayısallaştırılması ve veri setinin modelleme için

uygun hâle getirilmesi işlemleri gerçekleştirilmiştir.

Araştırmada öğrencilerin akademik performansını tahmin etmek amacıyla hem sınıflandırma hem de regresyon tabanlı makine öğrenmesi modelleri kullanılmıştır. Sınıflandırma analizlerinde Support Vector Machine (SVM) algoritması en yüksek performansı göstermiştir. Modelin test veri seti üzerinde elde ettiği doğruluk oranı 0.9071, F1 skoru 0.9129, dengeli doğruluk değeri 0.9062, ROC-AUC değeri 0.9389 ve PR-AUC değeri 0.9294 olarak hesaplanmıştır. Bu sonuçlar modelin öğrencilerin akademik performanslarını yüksek doğrulukla sınıflandırabildiğini ve sınıflar arasındaki ayrımı başarılı bir şekilde gerçekleştirebildiğini göstermektedir.

Regresyon analizlerinde ise öğrencilerin akademik performansları sürekli bir değişken olarak ele alınmış ve farklı regresyon modelleri karşılaştırılmıştır. Yapılan analizler sonucunda XGBoost regresyon modeli en başarılı performansı göstermiştir. Modelin test veri seti üzerinde elde ettiği R^2 değeri 0.6963, RMSE değeri 12.4614 ve MAE değeri 8.9904 olarak hesaplanmıştır. Bu sonuçlar modelin öğrencilerin akademik performansındaki varyansın önemli bir kısmını açıklayabildiğini ve gerçek değerlere yakın tahminler üretebildiğini göstermektedir.,

Çalışmanın önemli katkılarından biri geliştirilen makine öğrenmesi modellerinin açıklanabilir yapay zekâ yöntemleri kullanılarak yorumlanmasıdır. Bu amaçla SHAP ve LIME yöntemleri kullanılmıştır. SHAP analizi sonuçlarına göre öğrencilerin akademik performansını etkileyen en önemli değişkenlerin başında baba eğitim durumu, ders (çalışma) notları, çevrimiçi kurs kullanımı, sosyal medya kullanımının ders çalışmaya etkisi ve e-ders kaynaklarının kullanımı gelmektedir. LIME analizi sonuçları ise bireysel tahminler üzerinde derslere düzenli katılım, baba eğitim durumu, ders çalışma sırasında verilen mola sıklığı, çevrimiçi kurslara katılım ve öğrencinin yaşadığı yer gibi değişkenlerin önemli rol oynadığını göstermektedir.

Elde edilen bulgular öğrencilerin akademik performanslarının yalnızca akademik göstergelerle değil, aynı zamanda sosyo-demografik faktörler, öğrenme davranışları ve dijital öğrenme kaynaklarının kullanımı gibi çok boyutlu değişkenlerle ilişkili olduğunu ortaya koymaktadır. Ayrıca açıklanabilir yapay zekâ yöntemleri sayesinde makine

öğrenmesi modellerinin karar süreçleri daha anlaşılır hâle getirilmiş ve eğitim alanında veri temelli karar verme süreçlerinin desteklenebileceği gösterilmiştir.

Sonuç olarak bu çalışma, makine öğrenmesi ve açıklanabilir yapay zekâ yaklaşımlarının öğrencilerin akademik performansının tahmin edilmesinde etkili araçlar sunduğunu ve bu tür modellerin eğitim kurumlarında erken uyarı sistemlerinin geliştirilmesine katkı sağlayabileceğini ortaya koymaktadır.

Anahtar Kelimeler: Açıklanabilir Yapay Zekâ, Akademik Performans Tahmini, Eğitimde Yapay Zekâ, Makine Öğrenimi, Öğrenci Başarısı, Yapay Zekâ



ABSTRACT

M.Sc. Thesis

USE OF EXPLAINABLE ARTIFICIAL INTELLIGENCE AND MACHINE LEARNING IN PREDICTING STUDENT'S ACADEMIC PERFORMANCE

Müjde Dilek SOBUTAY

Bartın University

Graduate School

Department of Management Information Systems

Thesis Advisor: Prof. Dr. Ümmühan AVCI

Bartın- 2026, pp: 211

The aim of this study is to predict students' academic performance using machine learning methods and to make the developed models interpretable through explainable artificial intelligence (XAI) techniques. Today, large amounts of data are generated in educational institutions through student information systems and digital learning environments. The analysis of these data using machine learning methods enables early prediction of students' academic performance and the identification of at-risk students. Explainable artificial intelligence approaches support making model outcomes more transparent, interpretable, and reliable by revealing which variables these predictions are based on.

Within the scope of the study, a dataset was created using multidimensional variables that are considered to affect students' academic performance. The data used in the study consists of student information system (UBYS) records and survey data obtained from students. The dataset includes various variables such as students' exam scores, attendance, use of online learning resources, social media usage habits, technology use, and socio-demographic characteristics. During the data preprocessing phase, missing data were handled, categorical variables were transformed into numerical form, and the dataset was prepared for modeling.

In the study, both classification and regression-based machine learning models were used to predict students' academic performance. In classification analyses, the Support Vector Machine (SVM) algorithm demonstrated the highest performance. The model achieved an accuracy of 0.9071, an F1-score of 0.9129, a balanced accuracy of 0.9062, an ROC-AUC of 0.9389, and a PR-AUC of 0.9294 on the test dataset. These results indicate that the model can classify students' academic performance with high accuracy and effectively distinguish between classes.

In regression analyses, students' academic performance was treated as a continuous variable, and different regression models were compared. The results showed that the XGBoost regression model achieved the best performance. The model obtained an R^2 value of 0.6963, an RMSE of 12.4614, and an MAE of 8.9904 on the test dataset. These findings indicate that the model can explain a substantial portion of the variance in students' academic performance and produce predictions close to the actual values.

One of the key contributions of this study is the interpretation of the developed machine learning models using explainable artificial intelligence methods. For this purpose, SHAP and LIME techniques were employed. According to the SHAP analysis results, the most important variables affecting students' academic performance include father's education level, course (study) notes, use of online courses, the impact of social media usage on studying, and the use of e-learning resources. The results of the LIME analysis indicate that variables such as regular attendance, father's education level, frequency of breaks during studying, participation in online courses, and the place of residence play a significant role in individual predictions.

The findings reveal that students' academic performance is associated not only with academic indicators but also with multidimensional variables such as socio-demographic factors, learning behaviors, and the use of digital learning resources. Moreover, explainable artificial intelligence methods have made the decision-making processes of machine learning models more understandable and demonstrated their potential to support data-driven decision-making processes in education.

In conclusion, this study shows that machine learning and explainable artificial intelligence

approaches provide effective tools for predicting students' academic performance and can contribute to the development of early warning systems in educational institutions.

Keywords: Academic Performance Prediction, Artificial Intelligence, Artificial Intelligence in Education, Explainable Artificial Intelligence, Machine Learning, Student Achievement



İÇİNDEKİLER

KABUL VE ONAY.....	ii
BEYANNAME	iii
YAPAY ZEKÂ KULLANIMINA İLİŞKİN BEYAN.....	iv
ÖZET	vi
ABSTRACT	ix
İÇİNDEKİLER.....	xii
ŞEKİLLER DİZİNİ.....	xvi
TABLolar DİZİNİ.....	xvii
EKLER DİZİNİ	xxi
SİMGELER VE KISALTMALAR DİZİNİ.....	xxii
1. GİRİŞ.....	1
1.1 Problem Durumu	1
1.2 Araştırmanın Önemi.....	2
1.3 Araştırmanın Amacı ve Araştırma Problemleri	3
1.4 Varsayımlar	4
1.5 Araştırmanın Sınırlılıkları	4
1.6 Kuramsal Çerçeve.....	6
1.6.1 Yapay Zekâ, Makine Öğrenmesi ve Tahmin Yaklaşımları	6
1.6.1.1 Eğitim-Öğretim Süreçlerinde Öğrenci Akademik Performansında Makine Öğrenmesi Modelleri.....	9
1.6.1.2 Öğrenci Akademik Performansının Çok Boyutlu Yapısı	10
1.6.2 Açıklanabilir Yapay Zekâ ve Önemi.....	11
1.6.2.1 Açıklanabilir Yapay Zekâ Yöntemleri: LIME ve SHAP	12
1.6.2.2 Eğitimde Açıklanabilir Yapay Zekânın Kullanımı	13
2. ALANYAZIN	14
2.1 Uluslararası Alanyazında yapılan çalışmalar	14
2.1.1 Eğitimde Yapay Zekâ ve Performans Tahmini	14
2.1.2 Akademik Performansı Etkileyen Değişkenler ve Eğitim Verisinin Yapısı..	16
2.1.3 Açıklanabilir Yapay Zekâ (XAI) ve Eğitimde Şeffaflık	19
2.2 Ulusal Alanyazın	20

2.2.1 Ulusal Bağlamda Öğrenci Performansı Tahmini Çalışmaları	20
2.3 Alanyazının Sentezi ve Araştırma Gereksinimi.....	23
3. MATERYAL VE YÖNTEM.....	25
3.1 Materyal.....	25
3.1.1 Veri Seti	25
3.1.2 Veri Seti Özellikleri	26
3.1.3 Makine Öğrenmesi ve Açıklanabilir Yapay Zekâ Açısından Değişken Seçim Gerekçesi.....	30
3.2 Model Tasarımı.....	30
3.3 Veri İşleme Araçları	31
3.3.1 Python Programlama Dili	31
3.3.1.1 Veri Ön İşleme ve Düzenleme Araçları (Pandas ve NumPy)	32
3.3.1.2 Makine Öğrenmesi ve Modelleme Araçları (Scikit-learn).....	32
3.3.1.3 Veri Görselleştirme ve Analiz Ortamı	32
3.4 Yöntem.....	33
3.4.1 Karar Ağaçları (Decision Trees)	34
3.4.2 Destek Vektör Makineleri/ Destek Vektör Regresyonu (Support Vector Machines/ Support Vector Regression).....	35
3.4.3 Yapay Sinir Ağları (Artificial Neural Networks)	36
3.4.4 Rastgele Orman (Random Forest)	37
3.4.5 XGBoost (Extreme Gradient Boosting)	37
3.4.6 K-En Yakın Komşu (K-NN)	38
3.5 SHAP ve LIME	39
3.6 Veri Toplama Süreci.....	42
3.7 Veri Ön İşleme Süreci.....	43
3.7.1 Verilerin Okunması	43
3.7.2 Eksik Veri Analizi ve Düzenlenmesi	44
3.7.3 Sayısal Değişkenlerin Ölçeklendirilmesi.....	45
3.7.4 Verilerin Birleştirilmesi.....	46
3.7.5 Öznitelik Seçimi	46
3.7.6 Çoklu Öznitelik Ayırıştırması.....	48
3.7.7 Kategorik Verilerin Sayısallaştırılması	48
3.7.7.1 Ordinal Encoding Kullanılarak Kategorik Verilerin	

Sayısallaştırılması.....	49
3.7.7.2 One-Hot Encoding Kullanılarak Kategorik Verilerin	
Sayısallaştırılması.....	59
3.7.8 Yöntemsel Sorunlar: Dengesiz Veri, Metrikler ve Veri Sızıntısı	60
4. BULGULAR VE TARTIŞMA.....	62
4.1 Seçilen Öznitelikler ve Özellikleri	62
4.2 Öğrenci Performansının Sınıflandırılması	63
4.2.1 Karar Ağacı Modeli Parametreleri.....	66
4.2.2 Destek Vektör Makinesi Modeli Parametreleri	67
4.2.3 Rastgele Orman Modeli Parametreleri	67
4.2.4 K- En Yakın Komşu Modeli Parametreleri	68
4.2.5 Çok Katmanlı Perceptron (MLP) Modeli Parametreleri	69
4.2.6 XGBoost Modeli Parametreleri.....	69
4.3 Öğrenci Performansının Sınıflandırılma Sonuçları	70
4.3.1 Sınıflandırma Sonuçları (Deneme 1).....	71
4.3.1.1 XGBoost Sınıflandırma Modeli Sonuçları	78
4.3.2 Sınıflandırma Sonuçları (Deneme 2).....	82
4.3.2.1 Rastgele Orman Sınıflandırma Modeli Sonuçları	90
4.3.3 Sınıflandırma Sonuçları (Deneme 3).....	94
4.3.3.1 Destek Vektör Makinesi (SVM) Sınıflandırma Modeli Sonuçları	102
4.3.4 Sınıflandırma Sonuçları (Deneme 4).....	107
4.3.4.1 Destek Vektör Makinesi (SVM) Sınıflandırma Modeli Sonuçları	114
4.4 Öğrenci Performansının Tahmini	119
4.4.1 Regresyon Sonuçları (Deneme 1)	120
4.4.2 Regresyon Sonuçları (Deneme 2)	125
4.5 Açıklanabilir Yapay Zekâ.....	130
4.5.1 SHAP.....	130
4.5.2 LIME.....	137
5. SONUÇ, TARTIŞMA VE ÖNERİLER.....	144
5.1 Sınıflandırma Sonuçlarının Değerlendirilmesi	144
5.2 Regresyon Sonuçlarının Değerlendirilmesi	146
5.3 Model Karşılaştırması.....	147
5.4 Açıklanabilir Yapay Zekâ Bulgularının Değerlendirilmesi.....	148

5.4.1 SHAP Analizi Sonuçlarının Değerlendirilmesi	148
5.4.2 LIME Analizi Sonuçlarının Değerlendirilmesi	149
5.5 Eğitim Sistemine Yönelik Çıkarımlar	151
5.6 Öneriler	152
KAYNAKLAR	154
EKLER	163
ÖZGEÇMİŞ	212



ŞEKİLLER DİZİNİ

Şekil	Sayfa
No	No
3.1: Karar Ağaçları.....	35
3.2: Destek Vektör Makineleri	36
3.3: Yapay Sinir Ağları	36
3.4: Rastgele Orman.....	37
3.5: XGBoost.....	38
3.6: K-En Yakın Komşu.....	38
4.1: Seçilen öznitelikler ve özellikleri	62
4.2: Veriseti ve sınıf sayıları.....	63
4.3: SVM Confusion Matrix.....	116
4.4: SVM ROC Eğrisi	118
4.5: Orijinal değişken bazlı toplu SHAP ortalamaları grafiği.....	134
4.6: SHAP Arı Sürüsü (Beeswarm) grafiği	135
4.7: Modele en çok etki eden 20 özniteliğin ortalama LIME değerleri grafiği.....	137
4.8: Modele pozitif etki eden 20 özniteliğin ortalama LIME değerleri grafiği	141
4.9: Modele negatif etki eden 20 özniteliğin ortalama LIME değerleri grafiği	143

TABLÖLAR DİZİNİ

Tablo No	Sayfa No
3.1: Veri seti nitelikleri ve açıklamaları	28
3.2: Cinsiyet kategorik verisinin kodlanması	49
3.3: 'Aile gelir düzeyiniz' kategorik verisinin kodlanması	49
3.4: 'Anne/Baba eğitim durumunuz' kategorik verisinin kodlanması	50
3.5: 'Yaşadığınız yer' kategorik verisinin kodlanması	50
3.6: 'Liseden mezun olduğunuz okul türü' kategorik verisinin kodlanması	50
3.7: 'Öğrenim süreciniz boyunca barınma durumunuz nedir?' kategorik verisinin kodlanması	51
3.8: 'Ders çalışırken en çok hangi kaynakları kullanırsınız?' kategorik verisinin kodlanması	51
3.9: 'Haftalık ortalama ders çalışma süreniz kaç saattir?' kategorik verisinin kodlanması	52
3.10: 'Ders çalışmak için en çok kullandığınız ortam/mekân nedir?' kategorik verisinin kodlanması	52
3.11: 'Genel olarak çalışma ortamın seni çalışmaya ne kadar teşvik ediyor?' kategorik verisinin kodlanması	53
3.12: 'Ders çalışırken ne sıklıkla ara verirsiniz?' kategorik verisinin kodlanması	53
3.13: 'Sınavlara ne kadar süre önceden hazırlanmaya başlıyorsunuz?' kategorik verisinin kodlanması	54
3.14: 'Sınav kaygısı yaşıyor musunuz?' kategorik verisinin kodlanması	54
3.15: 'Derslere düzenli olarak katılıyor musunuz?' kategorik verisinin kodlanması	54
3.16: 'Öğrenme stilinizi nasıl tanımlarsınız?' kategorik verisinin kodlanması	55
3.17: 'Günlük ortalama internet kullanım süreniz kaç saattir?' kategorik verisinin kodlanması	56
3.18: 'Derslerinizle ilgili olarak interneti ne sıklıkla kullanırsınız?' kategorik verisinin kodlanması	56
3.19: 'Hangi çevrimiçi öğrenme araçlarını kullanıyorsunuz?' kategorik verisinin kodlanması	57
3.20: 'Çevrimiçi derslere katılımınız ne düzeyde?' kategorik verisinin kodlanması	57
3.21: 'Sosyal medya kullanımı ders çalışmanızı etkiliyor mu?' kategorik verisinin kodlanması	

kodlanması	58
3.22: 'Teknoloji kullanımının ders çalışmanıza yardımcı olduğunu düşünüyor musunuz?' kategorik verisinin kodlanması	58
3.23: 'Devam_x (Devam Zorunluluğu)' kategorik verisinin kodlanması	58
3.24: 'Başarı Durumu' kategorik verisinin kodlanması	59
4.1: SMOTE ve ADASYN karşılaştırması	64
4.2: %70 Eğitim %30 Test verileri ve sınıf sayıları (Deneme 1)	71
4.3: %60 Eğitim %40 Test verileri ve sınıf sayıları (Deneme 1)	72
4.4: %70 Eğitim %30 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 1)	72
4.5: %70 Eğitim %30 Test verileri ile tüm modellerin test performans sonuçları (Deneme 1)	74
4.6: %60 Eğitim %40 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 1)	75
4.7: %60 Eğitim %40 Test verileri ile tüm modellerin test performans sonuçları (Deneme 1)	76
4.8: XGBoost en iyi parametreler (Deneme 1)	78
4.9: XGBoost 10 Fold Cross Validation (CV) sonuçları (Deneme 1)	79
4.10: XGBoost Confusion Matrix (Deneme 1)	80
4.11: XGBoost % 60 Eğitim %40 Test verileri ile test sınıflandırma raporu (Deneme 1)	80
4.12: XGBoost genel test sınıflandırma sonuçları (Deneme 1)	82
4.13: Veri Arttırımı yapılmış, %70 Eğitim %30 Test verileri ve sınıf sayıları (Deneme 2)	83
4.14: Veri Arttırımı yapılmış, %60 Eğitim %40 Test verileri ve sınıf sayıları (Deneme 2)	84
4.15: %70 Eğitim %30 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 2)	85
4.16: %70 Eğitim %30 Test verileri ile tüm modellerin test performans sonuçları (Deneme 2)	86
4.17: %60 Eğitim %40 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 2)	87
4.18: %60 Eğitim %40 Test verileri ile tüm modellerin test performans sonuçları	

(Deneme 2).....	89
4.19: Rastgele Orman en iyi parametreler (Deneme 2).....	90
4.20: XGBoost 10 Fold Cross Validation (CV) sonuçları (Deneme 2)	91
4.21: Rastgele Orman Confusion Matrix (Deneme 2)	92
4.22: Rastgele Orman %70 Eğitim %30 Test verileri ile test sınıflandırma raporu (Deneme 2).....	93
4.23: Rastgele Orman genel test sınıflandırma sonuçları (Deneme 2).....	94
4.24: Veri Arttırımı yapılmış, %70 Eğitim %30 Test verileri ve sınıf sayıları (Deneme 3)	95
4.25: Veri Arttırımı yapılmış, %60 Eğitim %40 Test verileri ve sınıf sayıları (Deneme 3)	96
4.26: %70 Eğitim %30 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 3).....	97
4.27: %70 Eğitim %30 Test verileri ile tüm modellerin test performans sonuçları (Deneme 3).....	98
4.28: %60 Eğitim %40 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 3).....	100
4.29: %60 Eğitim %40 Test verileri ile tüm modellerin test performans sonuçları (Deneme 3).....	101
4.30: SVM en iyi parametreler (Deneme 3).....	102
4.31: SVM 10 Fold Cross Validation (CV) sonuçları (Deneme 3).....	103
4.32: SVM Confusion Matrix (Deneme 3).....	104
4.33: SVM %70 Eğitim %30 Test verileri ile test sınıflandırma raporu (Deneme 3).....	105
4.34: SVM genel test sınıflandırma sonuçları (Deneme 3).....	106
4.35: Veri Arttırımı yapılmış, %70 Eğitim %30 Test verileri ve sınıf sayıları (Deneme 4)	107
4.36: Veri Arttırımı yapılmış, %60 Eğitim %40 Test verileri ve sınıf sayıları (Deneme 4)	108
4.37: %70 Eğitim %30 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 4).....	109
4.38: %70 Eğitim %30 Test verileri ile tüm modellerin test performans sonuçları (Deneme 4).....	110
4.39: %60 Eğitim %40 Test verileri ile tüm modellerin Çapraz Doğrulama (CV)	

performans sonuçları (Deneme 4).....	112
4.40: %60 Eğitim %40 Test verileri ile tüm modellerin test performans sonuçları (Deneme 4).....	113
4.41: SVM en iyi parametreler (Deneme 4).....	114
4.42: SVM 10 Fold Cross Validation (CV) sonuçları (Deneme 4).....	115
4.43: SVM Confusion Matrix (Deneme 4).....	116
4.44: SVM %70 Eğitim %30 Test verileri ile test sınıflandırma raporu (Deneme 4).....	117
4.45: SVM genel test sınıflandırma sonuçları (Deneme 4).....	118
4.46: Çapraz Doğrulama (CV) performans sonuçları (Deneme 1)	120
4.47: Test performans sonuçları (Deneme 1)	121
4.48: XGBoost en iyi parametreler (Deneme 1)	123
4.49: XGBoost Çapraz Doğrulama performans sonuçları (Deneme 1)	124
4.50: XGBoost test performans sonuçları (Deneme 1)	124
4.51: Çapraz Doğrulama (CV) performans sonuçları (Deneme 2)	125
4.52: Test performans sonuçları (Deneme 2)	127
4.53: SVR en iyi parametreler (Deneme 2).....	128
4.54: SVR Çapraz Doğrulama performans sonuçları (Deneme 2).....	129
4.55: SVR test performans sonuçları (Deneme 2).....	129
4.56: SHAP ortalama.....	131
4.57: Modele pozitif etki eden 20 öz niteliğin ortalama LIME değerleri tablosu.....	140
4.58: Modele negatif etki eden 20 öz niteliğin ortalama LIME değerleri tablosu	142

EKLER DİZİNİ

Ek	Sayfa
No	No
Ek 1. Yapay zekâ kullanım tablosu.....	163
Ek 2. Sosyal ve Beşeri Bilimler Etik Kurulu Onay Belgesi.....	164
Ek 3. Sınıflandırma Sonuçları Tabloları	165
Ek 4. Regresyon Sonuçları Tabloları	206



SİMGELER VE KISALTMALAR DİZİNİ

R^2	: Determinasyon Katsayısı
Σ	: Toplam İşareti

KISALTMALAR

ADASYN	: Adaptive Synthetic Sampling
ANN	: Artificial Neural Network
CV	: Cross Validation
DT	: Decision Tree
KNN	: K- Nearest Neighbor
LIME	: Local Interpretable Model Agnostic Explanations
LMS	: Learning Management System
MAE	: Mean Absolute Error
Max	: Maximum
Min	: Minimum
ML	: Machine Learning
MLP	: Multi-Layer Perceptron
ÖYS	: Öğrenme Yönetim Sistemi
PR-AUC	: Precision-Recall Area Under The Curve
RBF	: Radial Basis Function
RF	: Random Forest
RMSE	: Root Mean Squared Error
ROC-AUC	: Receiver Operating Characteristic Area Under The Curve
SHAP	: Shapley Additive Explanations
SMOTE	: Synthetic Minority Over-Sampling Technique
SVM	: Support Vector Machine
SVR	: Support Vector Regression
UBYS	: Üniversite Bilgi Yönetim Sistemi
XAI	: Explainable Artificial Intelligence
XGBoost	: Extreme Gradient Boosting

1. GİRİŞ

Eğitim-öğretim süreçleri, öğrenenin bireysel özellikleri ile öğretim ortamının koşullarının etkileşimi içinde şekillenen, çok boyutlu ve dinamik bir yapıya sahiptir. Öğrencilerin akademik performansını belirleyen faktörler; bilişsel yeterlikler, sosyo-ekonomik koşullar, öğrenme stratejileri ve devam durumu gibi değişkenlerin yanında, çevrimiçi öğrenme ortamlarındaki etkileşim örüntülerini de kapsayacak biçimde giderek genişlemektedir (Sirin, 2005; Credé vd., 2010). Dijitalleşmeyle birlikte üniversitelerin Öğrenci Bilgi Yönetim Sistemleri (UBYS) ve Öğrenme Yönetim Sistemleri (LMS) üzerinden üretilen veriler artmış; bu veri yoğunluğu, öğrencinin öğrenme sürecini yalnızca dönem sonu çıktılarıyla değil süreç verileriyle de izlemeyi mümkün kılmıştır (Siemens ve Long, 2011).

Bu dönüşüm, yapay zekâ (YZ) ve özellikle makine öğrenmesi (ML) yaklaşımlarını eğitim araştırmalarında ve uygulamalarında merkezi bir konuma taşımıştır. Yapay zekâ, insan benzeri bilişsel süreçleri taklit edebilen sistemler geliştirmeyi amaçlayan disiplinlerarası bir alan olarak tanımlanırken; makine öğrenmesi, sistemlerin açıkça programlanmadan veriden örüntü öğrenerek tahmin üretmesini sağlayan yöntemler kümesini ifade etmektedir (Aggarwal, 2018; Goodfellow vd., 2016). Eğitim bağlamında bu yöntemler, öğrencilerin akademik risklerinin erken dönemde belirlenmesi, kişiselleştirilmiş destek mekanizmalarının tasarlanması ve kurumsal karar süreçlerinin veri temelli güçlendirilmesi açısından kritik fırsatlar sunmaktadır (Baker ve Siemens, 2014; Siemens ve Long, 2011).

Öte yandan, eğitim alanında tahmin modellerinin kullanımında yalnızca yüksek doğruluk değerleriyle yetinmek, kararların eğitim paydaşları açısından güvenilir ve eyleme dönük olmasını garanti etmemektedir. Özellikle karmaşık modellerin “kara kutu” niteliği, modelin hangi gerekçeyle belirli bir tahmini ürettiğinin anlaşılmasını zorlaştırmaktadır (Adadi ve Berrada, 2018). Bu durum, makine öğrenmesi modelleri nin karar süreçlerinin anlaşılabilir hale getirilmesi ve açıklanabilir yapay zekâ yaklaşımlarının kullanılmasının önemini ortaya koymaktadır.

1.1 Problem Durumu

Eğitim kurumlarında öğrenci akademik performansının düşüklüğü çoğu zaman dönem sonu

sonuçları ortaya çıktıktan sonra fark edilmekte; bu durum, zamanında destek sunmayı güçleştirmektedir. Öte yandan eğitim kurumları UBYS ve ÖYS üzerinden çok miktarda veri üretmesine rağmen bu veriler çoğu zaman bütünleşik biçimde analiz edilmemektedir; karar süreçleri ağırlıklı olarak deneyim ve sezgiye dayalı yürütülebilmektedir. Bu durum, erken uyarı ve hedefli müdahale mekanizmalarının sınırlı kalmasına yol açmaktadır (Siemens ve Long, 2011).

Alanyazın, makine öğrenmesi ile başarı tahmininin mümkün olduğunu göstermesine karşın; modellerin tek kaynaktan beslenmesi, bağlamsal farklılıklar nedeniyle genellenebilirlik sorunları, dengesiz veri ve yanlış veri göstergesi seçimi nedeniyle hatalı risk değerlendirmeleri, ayrıca veri sızıntısı gibi yöntemsel problemler nedeniyle sahada güvenilirliğin azalması gibi zorlukları da işaret etmektedir (He ve Garcia, 2009; Kaufman vd., 2012; Kapoor ve Narayanan, 2023; Rahman vd., 2025). Buna ek olarak, yüksek doğruluk üreten ancak karar gerekçesini açıklayamayan modeller, eğitim paydaşları açısından ‘neden’ sorusunu yanıtsız bırakmaktadır. Bu nedenle, açıklanabilir yapay zekâ teknikleriyle desteklenmeyen tahmin modellerinin kurum içinde benimsenmesi ve müdahale süreçlerine entegre edilmesi zorlaşabilmektedir (Adadi ve Berrada, 2018; Arrieta vd., 2020). Bu doğrultuda, makine öğrenmesi ile geliştirilen akademik performans tahmin modellerinin açıklanabilir yapay zekâ yaklaşımlarıyla desteklenmesi giderek daha fazla önem kazanmaktadır.

1.2 Araştırmanın Önemi

Öğrencilerin akademik performansını tahmin etmede makine öğrenmesi yöntemlerinin etkinliğinin incelenerek, risk altındaki öğrencilerin erken belirlenmesine yönelik karar destek mekanizmalarının geliştirilmesi önemlidir. Bu çalışmada öğrencilerin akademik performansını tahmin etmede makine öğrenmesi modeli geliştirilmesi ve açıklanabilir yapay zeka ile en iyi modelin özniteliklerinin belirlenmesi ve anlaşılıp yorumlanabilir hale getirilmesi hedeflenmektedir. Çalışmanın UBYS, ÖYS ve anket gibi çoklu veri kaynaklarını bir araya getirerek öğrencinin akademik sürecini daha bütüncül biçimde temsil etmeyi amaçlaması, alanyazındaki tek-kaynaklı yaklaşım sınırlılıklarına karşı önemli bir katkı sunmaktadır. Bununla birlikte, çalışmada kullanılan veri setinin hazır bir veri tabanından alınmamış olması, araştırmanın özgün yönlerinden birini oluşturmaktadır. Veri seti, Bartın

Üniversitesi öğrencilerine ait verilerin araştırma amacı doğrultusunda farklı kaynaklardan derlenmesi ve yapılandırılmasıyla oluşturulmuştur. Bu yönüyle çalışma, yalnızca model geliştirmeyi değil, aynı zamanda bağlama özgü ve özgün bir veri seti oluşturmayı da içermektedir.

Öğrenme sürecinin paydaşları olan öğretmenler ve yöneticiler, öğrencilerin akademik performansını değerlendirme süreçlerinde çok kaynaklı veriyi içeren yaklaşımları tercih etmektedirler (Dilbaz Sayın ve Arslan, 2017). Bu bulgu, eğitimde değerlendirme süreçlerinin çok boyutlu ve çok kaynaklı bir bakış açısıyla ele alınması gerektiğini göstermektedir. Dolayısıyla öğrenci performansının tahmin edilmesinde de birden fazla veri kaynağının kullanılması, performansı etkileyen farklı değişkenlerin birlikte ele alınmasını sağlayarak daha kapsamlı ve gerçekçi tahmin modellerinin geliştirilmesine katkı sunmaktadır.

1.3 Araştırmanın Amacı ve Araştırma Problemleri

Bu araştırmanın temel amacı, öğrencilerin akademik performansını tahmin etmede makine öğrenmesi temelli sınıflandırma ve regresyon modellerinin başarımını incelemek ve tahmin sonuçlarını SHAP ve/veya LIME gibi açıklanabilir yapay zekâ teknikleri ile yorumlanabilir hâle getirmektir.

Bu amaç doğrultusunda araştırmanın ana problemi aşağıdaki şekilde ifade edilebilir:

“UBYS, ÖYS ve anket verilerinden elde edilen çoklu değişkenler kullanılarak öğrencilerin akademik performansı ne ölçüde tahmin edilebilir ve model çıktıları açıklanabilir yapay zekâ yöntemleriyle nasıl yorumlanabilir?”

Ana problem çerçevesinde aşağıdaki alt problemlere yanıt aranmıştır:

1. Öğrencilerin akademik performans (başarılı/başarısız) durumu, farklı makine öğrenmesi sınıflandırma algoritmaları kullanılarak ne ölçüde tahmin edilebilmekte ve denenen modeller arasında en yüksek tahmin performansını hangi algoritma göstermektedir?

2. Çoklu veri kaynağından elde edilen değişkenler kullanılarak öğrencilerin akademik performansı regresyon modelleriyle ne ölçüde tahmin edilebilmektedir?
3. SHAP ve LIME açıklanabilir yapay zeka teknikleri kullanılarak model çıktılarında hangi değişkenler akademik performans tahmininde daha etkili görünmektedir ve bu etkiler eğitim-öğretim bağlamında nasıl yorumlanabilir?

1.4 Varsayımlar

Bu araştırma kapsamında;

- Araştırmada kullanılan UBYS, ÖYS ve anket verilerinin öğrencilerin akademik performansını etkileyen özellikleri doğru ve güvenilir biçimde yansıttığı,
- Anket verilerinin katılımcıların gerçek beyanlarını yansıttığı,
- Veri setinde yer alan değişkenlerin öğrencilerin gerçek öğrenme davranışlarını ve akademik durumlarını temsil ettiği,
- Araştırmada kullanılan makine öğrenmesi algoritmalarının veri setindeki örüntüleri ortaya çıkarabilecek ve akademik performans tahmini için uygun modeller üretebilecek nitelikte olduğu,
- Model değerlendirme sürecinde kullanılan performans ölçütlerinin (ör. doğruluk, F1, RMSE vb.) modellerin tahmin başarısını geçerli biçimde yansıttığı,
- Açıklanabilir yapay zekâ teknikleri olan SHAP ve/veya LIME yöntemlerinin model çıktılarındaki değişken etkilerini anlamlandırmada güvenilir açıklamalar sunduğu varsayılmaktadır.

1.5 Araştırmanın Sınırlılıkları

Bu araştırma tek bir kurumda yürütülmüştür; bu nedenle bulguların farklı kurum ve bölümlere genellenebilirliği sınırlı olabilir. Veri kaynakları UBYS, ÖYS ve anketle sınırlı tutulmuştur; ders dışı öğrenme etkinlikleri, bireysel psikolojik ölçekler veya daha ayrıntılı teknoloji kullanım verileri kapsama alınmamıştır. Bu çalışmada öznitelik seçimi, araştırmanın problem durumu, amacı ve eğitim bağlamındaki yorumlanabilirlik gereksinimi dikkate alınarak manuel biçimde gerçekleştirilmiştir. Herhangi bir otomatik öznitelik seçimi

yönteminin kullanılmamış olması çalışmanın sınırlılıklarından biri olarak değerlendirilebilir. Bununla birlikte, bu tercih rastlantısal değil, çalışmanın amacı doğrultusunda yapılmış bilinçli bir yöntemsel karardır. Nitekim öznitelik seçimi sürecinde kullanılabilecek tek bir “en iyi” yöntemin bulunmadığı, uygun yaklaşımın problemin yapısına, veri özelliklerine ve araştırma amacına göre değiştiği belirtilmektedir (Saeys vd., 2007). Ayrıca öznitelik seçiminin özgün anlamını koruması ve seçilen değişkenlerin alan bilgisi temelinde yorumlanabilmesi, yorumlanabilirlik açısından önemli görülmektedir (Suresh vd., 2022; Molnar, 2022). Bu doğrultuda, eğitim bağlamında değişkenlerin pedagojik olarak anlamlandırılabilmesi ve yorumlanabilirliğin korunması, otomatik öznitelik seçimine göre daha öncelikli kabul edilmiştir. Ancak burada belirtmelidir ki yorumlanabilirlik bağlama özgü bir kavramdır ve alan uzmanlığından bağımsız düşünülemez (Rudin, 2019). Ayrıca makine öğrenmesi modelleri, kullanılan değişkenlerin niteliği ve veri kalitesiyle sınırlıdır; özellikle dengesiz sınıf dağılımı ve veri sızıntısı riskleri titizlikle yönetilse de tamamen ortadan kaldırılması pratikte güç olabilir (Kaufman vd., 2012; Kapoor ve Narayanan, 2023). Son olarak, XAI yöntemleri (SHAP/LIME) model davranışını açıklamak için güçlü araçlar sunmakla birlikte, bu açıklamaların pedagojik yorumlanması uzman görüşü ve bağlamsal bilgi gerektirmektedir. Çalışmada bazı kategorik değişkenlerin sayısal kodlar ile temsil edilmesi, özellikle Support Vector Machine, K-Nearest Neighbors ve Neural Networks gibi mesafe temelli modellerde algoritmaların kategoriler arasında gerçek bir mesafe varmış gibi algılamasına yol açabilir. Çalışmada kategorik değişkenlerin sayısal olarak temsil edilmesine bağlı olarak ortaya çıkabilecek yapay mesafe algısı, one-hot encoding yöntemi kullanılarak büyük ölçüde azaltılmıştır. Bununla birlikte, bu yaklaşımın veri boyutunu artırması ve özellikle mesafe temelli modellerde seyreklik ve mesafe hesaplamalarına ilişkin sınırlılıklar oluşturabilmesi nedeniyle, bulguların yorumlanmasında bu durum dikkate alınması gereken bir sınırlılık olarak değerlendirilebilir. ADASYN veri artırma yönteminin verilerin eğitim/test olarak bölünmesinden önce yapılmış olması çalışmanın sınırlılığı olarak değerlendirilmektedir. Bu durum, test verisine ait bilgilerin dolaylı olarak veri sızıntısına neden olarak model performansının olduğundan daha yüksek tahmin edilmesine yol açmış olabilir. Bu nedenle elde edilen bulguların genellenebilirliği dikkatli bir şekilde yorumlanmalıdır.

1.6 Kuramsal Çerçeve

Bu bölümde yapay zeka, makine öğrenmesi ve tahmin yaklaşımları, eğitim-öğretim süreçlerinde öğrenci akademik performansında makine öğrenmesi modelleri, öğrenci akademik performansının çok boyutlu yapısı, açıklanabilir yapay zeka ve önemi, açıklanabilir yapay zekâ yöntemleri: LIME ve SHAP, eğitimde açıklanabilir yapay zekânın kullanımı konuları açıklanmıştır.

1.6.1 Yapay Zekâ, Makine Öğrenmesi ve Tahmin Yaklaşımları

Yapay zekâ, insan zekâsı gerektiren öğrenme, akıl yürütme, karar verme ve problem çözme gibi süreçlerin bilgisayar sistemleri aracılığıyla gerçekleştirilmesini amaçlayan geniş bir çalışma alanını ifade etmektedir. Makine öğrenmesi ise yapay zekânın bir alt alanı olarak, sistemlerin açık biçimde programlanmadan verilerden örüntü öğrenmesini ve bu örüntüler üzerinden tahmin ya da sınıflandırma yapmasını mümkün kılmaktadır (Alpaydin, 2020; Goodfellow vd., 2016).

Makine öğrenmesi temelli tahmin modelleri, eğitim alanında genellikle iki temel problem türü çerçevesinde ele alınmaktadır: sınıflandırma ve regresyon. Sınıflandırma yaklaşımları öğrencilerin “başarılı/başarısız” ya da “riskli/riskli değil” gibi kategorilere ayrılmasını hedeflerken; regresyon yaklaşımları not ortalaması veya başarı puanı gibi sürekli değişkenlerin tahmin edilmesine odaklanmaktadır (James vd., 2013). Eğitim verilerinde öğrencilerin öğrenme davranışları, katılım düzeyi ve çalışma alışkanlıkları gibi değişkenler arasında doğrusal olmayan ilişkiler sıklıkla görüldüğünden, ağaç tabanlı topluluk yöntemleri ve kernel tabanlı yaklaşımlar bu alanda yaygın biçimde kullanılmaktadır.

Karar ağaçları, sınıflandırma ve regresyon problemlerinde en temel ve en anlaşılır yöntemlerden biridir. Breiman vd. (1984), CART yaklaşımı ile karar ağaçlarının kuramsal temelini oluşturarak veri uzayının ardışık bölünmeler yoluyla nasıl anlamlı alt gruplara ayrılabilirliğini göstermiştir. Quinlan (1986) ise bilgi kazanımı gibi ölçütleri sistematikleştirerek karar ağacı üretimini daha açık bir yapıya kavuşturmuştur. Karar ağaçlarının en önemli avantajı, sonuçların yorumlanabilir olması ve değişkenler arasındaki doğrusal olmayan ilişkileri ortaya koyabilmesidir. Ancak tek bir karar ağacının veriye aşırı

uyum gösterebilmesi, bu yöntemin en önemli sınırlılıklarından biridir.

Bu sınırlılığı azaltmak amacıyla geliştirilen topluluk öğrenme yöntemleri, birden fazla modelin birlikte çalıştırılmasıyla daha güçlü tahminler üretmeyi hedeflemektedir. Breiman (2001), Random Forest yaklaşımıyla çok sayıda karar ağacını bir araya getirerek daha kararlı ve daha yüksek genelleme gücüne sahip modeller elde edilebileceğini göstermiştir. Liaw ve Wiener (2001) ise bu yaklaşımın uygulamalı kullanımını kolaylaştırmıştır. Benzer biçimde Friedman (2001), gradient boosting yöntemiyle zayıf öğrencilerin ardışık biçimde birleştirilmesi yoluyla güçlü tahminleyiciler kurulabileceğini ortaya koymuştur. Chen ve Guestrin (2016) tarafından geliştirilen XGBoost ise bu yaklaşımı ölçeklenebilirlik, hız ve düzenlileştirme gibi avantajlarla güçlendirmiştir. Bu nedenle Random Forest, Gradient Boosting ve XGBoost gibi yöntemler, özellikle çok değişkenli ve karmaşık veri yapılarında yüksek tahmin performansı sağlayan güçlü modeller arasında değerlendirilmektedir.

Destek Vektör Makineleri (SVM) de sınıflandırma ve regresyon problemlerinde önemli bir kuramsal temele sahiptir. Cortes ve Vapnik (1995), destek vektör makinelerinin temel mantığını ortaya koyarak gözlemleri en uygun biçimde ayıran karar sınırının oluşturulmasını esas alan bir yaklaşım geliştirmiştir. Vapnik (1995), bu yaklaşımı istatistiksel öğrenme teorisi içinde değerlendirerek modelin yalnızca mevcut veriye uyumunu değil, yeni veriler üzerindeki genelleme kapasitesini de ön plana çıkarmıştır. Schölkopf ve Smola (2002), çekirdek yöntemleri aracılığıyla doğrusal olmayan örüntülerin de etkili biçimde modellenilebileceğini göstermiştir. Smola ve Schölkopf (2004) ise Destek Vektör Regresyonu (SVR) ile bu yaklaşımın sürekli değişkenlerin tahmininde de kullanılabileceğini ortaya koymuştur. Bu özellikleri nedeniyle SVM ve SVR, özellikle yüksek boyutlu ve karmaşık veri yapılarında güçlü alternatif yöntemler olarak kabul edilmektedir.

Parametrik olmayan yaklaşımlar arasında yer alan k-en yakın komşu yöntemi, yeni bir gözlemi kendisine en çok benzeyen örnekler üzerinden değerlendirmektedir. Altman (1992), Cunningham ve Delany (2007), kNN yaklaşımının temel özelliğinin katı dağılım varsayımlarına ihtiyaç duymadan yerel örüntüler üzerinden sınıflandırma ve regresyon yapabilmesi olduğunu vurgulamıştır. Bu yönüyle yöntem basit, sezgisel ve uygulanabilir bir çerçeve sunmaktadır. Bununla birlikte kNN yöntemi, uzaklık ölçütüne, veri ölçeklendirmesine ve veri boyutuna duyarlıdır. Bentley (1975) tarafından geliştirilen KD-

tree yapıları ise komşuluk aramalarının hesaplama maliyetini azaltarak bu tür yöntemlerin daha verimli çalışmasına katkı sağlamıştır. Dolayısıyla kNN, özellikle karşılaştırma amacıyla yararlı bir temel model olarak değerlendirilebilir.

Yapay sinir ağları ve makine öğrenme yöntemleri ise çok katmanlı yapıları sayesinde karmaşık ve doğrusal olmayan ilişkileri modelleyebilme gücüne sahiptir. Aggarwal (2018), sinir ağlarının temel prensiplerini, optimizasyon mantığını ve düzenlileştirme tekniklerini kapsamlı biçimde ele almıştır. Goodfellow vd. (2016), derin öğrenmeyi temsil öğrenmesi bağlamında sistematikleştirerek, modellerin veriden giderek daha soyut özellikler çıkarabildiğini göstermiştir. Nair ve Hinton (2010), ReLU aktivasyon fonksiyonunun öğrenme sürecini iyileştirdiğini ortaya koyarken; Prechelt (2000), erken durdurmayı aşırı öğrenmeye karşı etkili bir strateji olarak tartışmıştır. Bergstra ve Bengio (2012) da hiperparametre ayarlamasında random search yaklaşımının birçok durumda verimli sonuçlar verdiğini göstermiştir. Bu çalışmalar birlikte değerlendirildiğinde, sinir ağlarının başarısının yalnızca model mimarisine değil, aynı zamanda eğitim süreci ve hiperparametre ayarına da bağlı olduğu anlaşılmaktadır.

Makine öğrenmesi yöntemlerinin uygulanmasında yalnızca algoritmalar değil, veri işleme süreci de belirleyici öneme sahiptir. Han, Kamber ve Pei (2012), veri madenciliği sürecini ön işleme, özellik çıkarımı, modelleme ve değerlendirme aşamalarıyla birlikte ele alarak veri kalitesinin ve iş akışının model performansı üzerindeki etkisini vurgulamıştır. James vd. (2013), Hastie vd. (2009) ise model karmaşıklığı, bias–variance dengesi ve model seçimi gibi temel ilkeleri açıklayarak uygun modelin yalnızca yüksek doğruluk sağlayan model olmadığını; aynı zamanda veri yapısına uygun, genellenebilir ve yorumlanabilir bir model olması gerektiğini ortaya koymuştur. Géron (2019) da bu kuramsal çerçeveyi uygulamalı makine öğrenmesi pratikleriyle birleştirerek model geliştirme sürecinin deneysel ve sistematik bir yapı gerektirdiğini göstermiştir.

Sonuç olarak sınıflandırma ve regresyon yaklaşımları, öğrenci akademik performansı gibi eğitim-öğretim problemlerinin modellenmesinde birbirini tamamlayan iki temel çerçeve sunmaktadır. Karar ağaçları ve topluluk yöntemleri doğrusal olmayan ilişkileri güçlü biçimde yakalayabilmekte; destek vektör makineleri yüksek boyutlu veri yapılarında etkili sonuçlar verebilmekte; kNN yerel benzerlikleri temel alan basit fakat işlevsel bir yaklaşım

sunmakta; yapay sinir ağıları ise daha karmaşık örüntülerin öğrenilmesine olanak tanımaktadır. Bu nedenle eğitim verilerinde model seçimi yapılırken, yalnızca performans ölçütleri değil; veri yapısı, problem türü, genelleme kapasitesi ve yorumlanabilirlik gereksinimi de birlikte dikkate alınmalıdır.

1.6.1.1 Eğitim-Öğretim Süreçlerinde Öğrenci Akademik Performansında Makine Öğrenmesi Modelleri

Makine öğrenmesi yöntemlerinin eğitim alanında kullanılmaya başlanması, öğrenci öğrenme süreçlerinin daha sistematik biçimde analiz edilmesine ve akademik performansın erken aşamalarda tahmin edilmesine olanak sağlamıştır. Makine öğrenmesi modelleri, öğrencilerin akademik başarı durumlarını tahmin ederek başarısızlık riski taşıyan öğrencilerin erken dönemde belirlenmesine katkı sağlayabilmektedir. Özellikle erken uyarı sistemleri, öğrencilerin akademik performanslarında ortaya çıkabilecek olumsuz eğilimlerin dönem sonu sonuçları ortaya çıkmadan önce tespit edilmesini mümkün kılmakta ve öğretim elemanlarına hedefli müdahale stratejileri geliştirme imkânı sunmaktadır.

Siemens ve Long (2011) makine öğrenmesinin, eğitim kurumlarının “erken uyarı ve müdahale” kapasitesini artıran bir karar destek yaklaşımı olarak ele almış; özellikle risk altındaki öğrencilerin dönem içinde belirlenmesine yönelik potansiyeline dikkat çekmiştir. Gašević vd. (2015), makine öğrenmesinin yalnızca ‘tahmin’ üretmekle sınırlı kalmaması; öğrenmeyi iyileştirmeye yönelik pedagojik yorum üretmesi gerektiğini savunmaktadır. Benzer şekilde Gašević vd. (2016), öğretim tasarımının hem demografik değişkenlerin hem de makine öğrenmesi metriklerinin öngörü gücünü etkileyebileceğini göstererek tek bir modelin farklı bağlamlarda aynı performansı göstermeyebileceğini tartışmıştır. Bu bulgular, kurum içi verilerle geliştirilen modellerin genellenebilirliği ve bağlamsal geçerliliğinin özellikle değerlendirilmesi gerektiğini ortaya koymaktadır. Romero ve Ventura (2010), eğitim veri madenciliğini öğrencilerin öğrenme süreçlerinden anlamlı örüntüler çıkarmayı amaçlayan bir alan olarak tanımlamakta ve bu alanın öğrenci performans tahmini, bırakma riski tespiti ve öğrenme davranışlarının analizi gibi konulara katkı sunduğunu vurgulamaktadır.

Bu çerçevede, UBYS ve ÖYS gibi sistemlerden gelen akademik ve davranışsal veriler,

anketler aracılığıyla elde edilen psiko-sosyal göstergelerle birlikte ele alındığında, akademik performansı daha bütüncül biçimde açıklama ve tahmin etme olanağı doğmaktadır. Örneğin performansı etkileyen değişkenlerin çok boyutlu olduğu; sosyo-ekonomik durumun akademik performansı ile anlamlı ilişkiler gösterdiği meta-analitik bulgularla desteklenmektedir (Kim vd., 2019). Devam durumunun başarıyla güçlü ilişkisi ise sınıf içi süreçlerin önemine işaret etmektedir (Credé vd., 2010). Pesovski vd. (2025) öğrencilerin akran etkileşimlerini ağ analiziyle modelleyerek başarı kümeleri üretmiş ve üretken yapay zekâ ile farklı kümelere yönelik kişiselleştirilmiş öğrenme içerikleri önermiştir. Bu yaklaşım, performans tahmininin yalnızca ‘tahmin’ olarak kalmayıp, tahmin çıktılarına pedagojik müdahaleye dönüştürme potansiyelini göstermesi açısından önemlidir.

Öğrenci akademik performansının tahmin edilmesi, eğitim kurumlarına yalnızca geriye dönük değerlendirme değil, ileriye dönük planlama yapma olanağı sunmaktadır. Erken uyarı sistemleri sayesinde risk altındaki öğrencilerin dönem içinde tespit edilmesi; rehberlik, akademik danışmanlık ve destek hizmetlerinin hedefli biçimde yürütülmesini sağlayabilir (Siemens ve Long, 2011). Bu durum, hem öğrenci başarısını artırmaya hem de kurumun kaynaklarını daha verimli kullanmasına katkıda bulunabilir. Rahman vd. (2025), öğrenci performansı tahminine yönelik makine öğrenmesi çalışmalarını sistematik biçimde incelemiş ve bu çalışmaların erken müdahale sistemlerinin tasarımında önemli bir rol oynadığını belirtmiştir. Bununla birlikte, alanyazında performans tahmini çalışmalarının bir kısmının tek veri kaynağına dayanması, çeşitli bağlamlarda genellenebilirlik sorunlarını gündeme getirmektedir. Bu nedenle çoklu veri kaynağı yaklaşımı, öğrencinin öğrenme sürecini yalnızca notlarla değil, süreç davranışları ve bağlamsal göstergelerle birlikte ele alması bakımından önemli bir avantaj sunmaktadır.

1.6.1.2 Öğrenci Akademik Performansının Çok Boyutlu Yapısı

Akademik performans, yalnızca bilişsel yeterliklerin veya dönem sonu notlarının bir yansıması değildir. Öğrenci başarısı; bireysel özellikler, sosyo-ekonomik koşullar, öğrenme davranışları, kurumsal katılım, devam durumu, öz-düzenleme becerileri ve öğretimsel bağlam gibi çok sayıda değişkenin etkileşimiyle şekillenmektedir. Bu nedenle akademik performansı açıklamaya ve tahmin etmeye yönelik yaklaşımların tek boyutlu değil, çok boyutlu bir çerçeveye dayanması gerekmektedir.

Alanyazında sosyo-ekonomik düzeyin öğrencinin akademik performansı ile anlamlı ilişkiler gösterdiği uzun süredir bilinmektedir. Gottfried (2010), sosyo-ekonomik durum ile akademik performans arasındaki ilişkinin güçlü ve sistematik olduğunu ortaya koymuştur. Benzer şekilde devam durumunun da öğrenci başarısının önemli belirleyicilerinden biri olduğu gösterilmiştir. Credé vd. (2010), ders devamı ile akademik performans arasında güçlü bir ilişki bulunduğunu vurgulayarak sınıf içi katılımın önemine dikkat çekmiştir. Bu bulgular, öğrenci başarısının yalnızca sınav performansı ile değil, öğrenme sürecine katılım ve bağlamsal koşullarla birlikte değerlendirilmesi gerektiğini ortaya koymaktadır.

Güncel derlemeler, öğrenci performansını tahmin etmeye yönelik modellerde demografik veriler, geçmiş akademik kayıtlar, Öğrenme yönetim sistemi (LMS) etkileşim verileri, katılım göstergeleri ve psikososyal ölçümlerin bir arada kullanılmasının daha kapsamlı bir açıklama zemini sunduğunu göstermektedir (Almalawi vd., 2024; Rodríguez-Ortiz vd., 2025; Kalita vd., 2025). Bu çerçevede UBYS, ÖYS ve anketlerden elde edilen verilerin birlikte kullanılması, öğrencinin yalnızca sonuç odaklı başarısını değil, bu başarıyı şekillendiren süreçleri de görünür kılmaktadır. Dolayısıyla akademik performansın tahmini, çoklu veri kaynaklarının bütünleştirildiği, öğrenciyi çok boyutlu biçimde ele alan bir yaklaşım gerektirmektedir.

1.6.2 Açıklanabilir Yapay Zekâ ve Önemi

Yapay zekâ ve makine öğrenmesi modellerinin eğitim alanında giderek daha fazla kullanılması, bu modellerin yalnızca yüksek tahmin başarımı göstermesini değil, aynı zamanda anlaşılabilir ve yorumlanabilir olmasını da gerekli kılmaktadır. Özellikle öğrenci başarısı, başarısızlık riski ve erken müdahale kararları gibi bireyin akademik sürecini etkileyen alanlarda, model çıktılarının nasıl üretildiğinin açıklanabilmesi önem taşımaktadır. Bu durum, açıklanabilir yapay zekâ kavramını eğitim araştırmalarında merkezi bir konuma taşımıştır.

Adadi ve Berrada (2018), açıklanabilir yapay zekâyı makine öğrenmesi modellerinin daha anlaşılır hâle getirilmesini amaçlayan yöntemler bütünü olarak ele almaktadır. Benzer biçimde Arrieta vd., (2020), açıklanabilirliğin yalnızca teknik bir özellik olmadığını; aynı zamanda sorumlu yapay zekâ, güven, şeffaflık ve etik karar verme ile doğrudan ilişkili

olduğunu vurgulamaktadır. Eğitim bağlamında bu durum daha da önem kazanmaktadır; çünkü model çıktıları öğretim elemanlarının, akademik danışmanların ve yöneticilerin öğrencilere yönelik müdahale kararlarını etkileyebilmektedir. Son dönem derlemelerde eğitim veri madenciliği çalışmalarında açıklanabilirliğin, güvenilir ve uygulanabilir yapay zekâ sistemleri geliştirmek açısından temel bir gereklilik hâline geldiğini göstermektedir (Gunasekara ve Saarela, 2024; Gunasekara ve Saarela, 2025; Türkmen, 2025).

1.6.2.1 Açıklanabilir Yapay Zekâ Yöntemleri: LIME ve SHAP

Açıklanabilir yapay zekâ alanında geliştirilen yöntemler, model kararlarının farklı düzeylerde yorumlanmasını mümkün kılmaktadır. Bu yöntemler genel olarak modelden bağımsız açıklama teknikleri ve belirli model türlerine özgü açıklama yaklaşımları biçiminde sınıflandırılabilir. Eğitim verilerinde yaygın olarak kullanılan yöntemlerden ikisi LIME ve SHAP'tır.

Ribeiro vd. (2016) tarafından geliştirilen LIME, herhangi bir sınıflandırıcının tekil tahminlerini yerel düzeyde yaklaşık doğrusal modeller aracılığıyla açıklamayı amaçlamaktadır. Bu yönüyle LIME, belirli bir öğrencinin neden “riskli” ya da “başarılı” olarak sınıflandırıldığını anlamaya yardımcı olmaktadır. Lundberg ve Lee (2017) tarafından önerilen SHAP ise her bir özelliğin model çıktısına katkısını Shapley değerleri aracılığıyla hesaplamakta ve hem yerel hem de küresel düzeyde açıklamalar sunmaktadır. Böylece modelin genel olarak hangi değişkenlere daha fazla ağırlık verdiği ve belirli bir tahminde hangi özelliklerin etkili olduğu birlikte değerlendirilebilmektedir.

Bu yöntemler, eğitim bağlamında devamsızlık, çevrim içi etkileşim, ders materyali kullanım sıklığı, çalışma süresi ya da önceki akademik performans gibi değişkenlerin öğrenci performansı üzerindeki etkisini görünür kılmaktadır. Özellikle SHAP ve LIME gibi tekniklerin öğretmenler ve karar vericiler için açıklamaları daha erişilebilir hâle getirdiği ve model sonuçlarının yorumlanmasını kolaylaştırdığı belirtilmektedir (Liu ve Khalil, 2024; Türkmen, 2025).

1.6.2.2 Eđitimde Açıklanabilir Yapay Zekânın Kullanımı

Öđrenci performansı tahmininin eđitimde etkili biçimde kullanılabilmesi için model açıklamalarının yalnızca teknik açıdan deđil, pedagojik açıdan da anlamlı olması gerekmektedir. Bir modelin hangi öđrencinin risk altında olduğunu göstermesi tek başına yeterli deđildir; aynı zamanda bu riskin hangi deđişkenlerle ilişkili olduğunu anlaşılması gerekir. Bu sayede öđretim elemanları ve akademik danışmanlar, tahmin sonuçlarını daha bilinçli biçimde yorumlayabilir ve daha hedefli müdahale stratejileri geliştirebilir.

Açıklanabilir yapay zekâ yaklaşımlarının eđitimdeki temel katkılarından biri, öđretmen güvenini ve sistem kabulünü artırmasıdır. Son dönem çalışmaları, açıklanabilir modellerin eđitimde makine öđrenmesi araçlarının daha şeffaf, güvenilir ve eyleme dönüştürülebilir olmasını desteklediğini göstermektedir (Türkmen, 2025; Alibayeva ve Allayarova, 2025). Ayrıca açıklanabilirlik, modelin yalnızca tahmin üretmesini deđil, aynı zamanda pedagojik kararları destekleyen bir araç hâline gelmesini sağlamaktadır. Bu bağlamda açıklanabilir yapay zekâ, öđrenci performansının tahmin edilmesinde teknik doğruluk ile eđitimsel yorumlanabilirlik arasında köprü kuran önemli bir yaklaşım olarak deđerlendirilebilir.

2. ALANYAZIN

Bu bölümde öğrencinin akademik performansını tahmin etmeye yönelik makine öğrenmesi yaklaşımları ile bu yaklaşımların açıklanabilir yapay zekâ (XAI) yöntemleriyle desteklenmesine ilişkin alanyazın, uluslararası ve ulusal çalışmalar olarak iki başlık altında sentezlenmiştir. Çalışmalar; (i) eğitimde yapay zekâ ve performans tahmini, (ii) akademik performansı etkileyen değişkenler ve eğitim verisinin yapısı, (iii) açıklanabilir yapay zekâ (xai) ve eğitimde şeffaflık, (iv) ulusal bağlamda öğrenci performansı tahmini çalışmaları ve (v) alanyazının sentezi ve araştırma gereksinimi temaları etrafında bütünleştirilmiştir.

2.1 Uluslararası Alanyazında yapılan çalışmalar

Bu bölümde; “Eğitimde Yapay Zekâ ve Performans Tahmini, Akademik Performansı Etkileyen Değişkenler ve Eğitim Verisinin Yapısı, Açıklanabilir Yapay Zekâ (XAI) ve Eğitimde Şeffaflık” başlıkları bulunmaktadır.

2.1.1 Eğitimde Yapay Zekâ ve Performans Tahmini

Uluslararası alanyazında yapay zekâ, eğitim bağlamında yalnızca otomasyon sağlayan bir teknoloji olarak değil, veriden örüntü çıkaran, tahmin üreten ve karar süreçlerini destekleyen bir yaklaşım olarak ele alınmaktadır. Morandín-Ahuerma (2022), yapay zekânın kavramsal çerçevesini tartışırken bu alanın temel gücünün büyük veri yapıları içerisinde anlamlı ilişkiler üretme kapasitesi olduğunu vurgulamaktadır. Bu değerlendirme, eğitim ortamlarında biriken çok boyutlu verilerin neden yapay zekâ temelli yöntemlerle analiz edilmeye çalışıldığını açıklamaktadır. Lexcellent (2019) ise yapay zekânın insan zekâsıyla karşılaştırmalı kullanımını tartışırken, özellikle insan-merkezli alanlarda teknik doğruluk kadar etik sorumluluğun da önemli olduğunu belirtmektedir. Bu vurgu, eğitimde yapay zekâ kullanımının yalnızca performans artışı üzerinden değil, kararların adilliği ve sorumluluğu üzerinden de değerlendirilmesi gerektiğini göstermektedir. Benzer biçimde Nagendraswamy ve Salis (2021), yapay zekânın eğitimde özellikle kişiselleştirme, öğrenci başarısını tahmin etme ve karar desteği sağlama alanlarında yoğunlaştığını belirterek, bu teknolojilerin öğrenme süreçlerinin izlenmesi ve yönlendirilmesinde giderek daha merkezi hâle geldiğini ortaya koymaktadır.

Gašević vd. (2016), dokuz lisans dersinde yürüttükleri çalışmada dersin öğretim tasarımı ve ÖYS kullanım biçiminin tahmin modellerinin performansını doğrudan etkilediğini göstermiştir. Araştırmanın önemli bulgusu, ders bazlı geliştirilen modeller ile genelleştirilmiş modeller arasında anlamlı farklar bulunması ve öğretim koşulları dikkate alınmadan geliştirilen genellemelerin öğrenci başarısını olduğundan fazla ya da eksik tahmin edebilmesidir. Bu sonuç, öğrenci performansı tahmininde “tek tip model” anlayışının sınırlı olduğunu ve bağlamın mutlaka hesaba katılması gerektiğini ortaya koymaktadır. Senin tezinde kurum içi verilerle model geliştirilmesi de bu nedenle önemlidir; çünkü bağlama duyarlı modeller, genel modellerden daha anlamlı sonuçlar verebilmektedir.

Almalawi vd. (2024) ise eğitim amaçlı tahmin modellerine ilişkin sistematik derlemelerinde, özellikle yükseköğretim ve çevrim içi öğrenme bağlamlarında topluluk öğrenmesi ve ÖYS temelli analitiklerin öne çıktığını, ancak modellerin veri kaynağına ve bağlama göre farklı performans gösterdiğini raporlamıştır. Bu sonuçlar, tahmin modellerinin yalnızca teknik üstünlük üzerinden değil, veri yapısı ve kullanım bağlamı içinde değerlendirilmesi gerektiğini desteklemektedir.

Rahman vd. (2024), öğrenci performansı tahmininde makine öğrenmesi kullanımını sistematik biçimde inceledikleri çalışmalarında, denetimli öğrenme yöntemlerinin alanda baskın olduğunu ve çalışmaların temel amaçlarının erken müdahale, karar verme süreçlerini iyileştirme ve tahmin doğruluğunu artırma olduğunu göstermiştir. Bununla birlikte derleme, aşırı öğrenme, ölçeklenebilirlik ve genellenebilirlik gibi teknik sorunların; adalet, hesap verebilirlik ve şeffaflık gibi etik sorunlarla birlikte düşünülmesi gerektiğini de vurgulamaktadır. Bu bulgu, eğitim alanında açıklanabilir yapay zekâ ve etik veri kullanımı tartışmalarının neden giderek daha merkezi hâle geldiğini açıklamaktadır.

Makine öğrenmesinde sınıflandırma ve regresyon, denetimli öğrenmenin iki temel yaklaşımını oluşturmaktadır. Sınıflandırma, gözlemlerin belirli kategorilere ayrılmasını amaçlarken; regresyon, sürekli bir değer tahmin edilmesine odaklanmaktadır. Eğitim verisi bağlamında sınıflandırma, öğrencilerin başarılı/başarısız ya da riskli/riskli değil biçiminde gruplandırılmasını sağlarken; regresyon, not ortalaması, sınav puanı veya başarı düzeyi gibi sürekli çıktıları tahmin etmek için kullanılmaktadır. Bu nedenle öğrenci performansının modellenmesinde hangi yaklaşımın tercih edileceği, araştırma probleminin

yapısına ve tahmin edilmek istenen çıktının türüne bağlı olarak değişmektedir (James vd., 2013; Hastie vd., 2009).

Sonuç olarak uluslararası alanyazın, eğitimde yapay zekâ ve performans tahmini probleminin çok katmanlı bir yapı sergilediğini göstermektedir. Bu alan yalnızca öğrencilerin başarılı ya da başarısız olacağını tahmin etmeyi amaçlamamakta; aynı zamanda bu tahminlerin pedagojik olarak anlamlandırılmasını, bağlama uygun biçimde yorumlanmasını ve öğrenci destek süreçlerine dönüştürülmesini de gerekli kılmaktadır. Dolayısıyla öğrenci performansı tahmini, teknik doğruluk kadar pedagojik açıklanabilirlik, bağlamsal geçerlilik ve etik sorumluluk gerektiren bütüncül bir araştırma alanı olarak değerlendirilmektedir.

2.1.2 Akademik Performansı Etkileyen Değişkenler ve Eğitim Verisinin Yapısı

Akademik performansı etkileyen değişkenler alanyazında tekil ve doğrusal bir yapı içinde değil; bireysel, bağlamsal, davranışsal ve ilişkisel boyutların etkileşimi çerçevesinde ele alınmaktadır. Bu nedenle öğrenci performansını açıklamaya ve tahmin etmeye yönelik modellerin yalnızca not, sınav puanı ya da demografik bilgilerle sınırlandırılması yeterli görülmemektedir. Alanyazındaki bulgular, akademik performansı etkileyen değişkenlerin öğrencinin içinde bulunduğu sosyo-ekonomik çevreden öğrenme sürecine katılımına, dijital sistemlerde bıraktığı izlerden akran etkileşim ağlarındaki konumuna kadar uzanan çok katmanlı bir yapı oluşturduğunu göstermektedir. Bu durum, eğitim verisinin de doğal olarak çok kaynaklı, çok boyutlu ve bağlama duyarlı bir yapı sergilediğine işaret etmektedir.

Sirin (2005), sosyo-ekonomik durum ile akademik başarı arasındaki ilişkiyi meta-analitik olarak incelemiş ve sosyo-ekonomik durumun farklı örneklerde ve farklı ölçüm biçimlerinde akademik performansın anlamlı bir yordayıcısı olduğunu ortaya koymuştur. Çalışmanın temel bulgusu, bu ilişkinin zayıf ve tesadüfi değil; orta ile güçlü düzey arasında değişen sistematik bir örüntü sergilemesidir. Bu bulgudan çıkarılabilecek temel sonuç, akademik başarının yalnızca bireysel bilişsel özelliklerle açıklanamayacağı; öğrencinin içinde bulunduğu sosyoekonomik bağlamın da başarı üzerinde belirleyici bir rol oynadığıdır. Dolayısıyla performans tahmini çalışmalarında sosyoekonomik göstergelerin modele dâhil edilmesi, yalnızca arka plan bilgisini genişletmek için değil, başarıyı yapısal boyutlarıyla

açıklayabilmek için de önem taşımaktadır.

Credé vd. (2010), derse katılım ile akademik performans arasındaki ilişkiyi meta-analitik düzeyde incelemiş ve devamın hem ders notları hem de genel not ortalaması ile güçlü ve tutarlı bir ilişki gösterdiğini ortaya koymuştur. Çalışmanın dikkat çekici sonucu, derse katılımın bazı durumlarda standart sınav puanları, çalışma alışkanlıkları ve çalışma becerileri gibi yaygın başarı göstergelerinden daha güçlü bir yordayıcı olabilesidir. Bu bulgu, öğrenme sürecine fiili katılımın başarı üzerindeki etkisinin göz ardı edilemeyecek düzeyde olduğunu göstermektedir. Buradan hareketle, devam verisinin yalnızca idari bir kayıt değil, öğrencinin öğrenme sürecine bağlılığını ve dersle kurduğu etkileşimi yansıtan güçlü bir davranışsal gösterge olarak değerlendirilmesi gerektiği söylenebilir.

OECD'nin (2019) PISA 2018 bulguları, başarı farklılıklarının yalnızca öğrenci düzeyindeki özelliklerle açıklanamayacağını; okul ve sistem düzeyindeki değişkenlerin de bu farklılıkların oluşumunda etkili olduğunu göstermektedir. PISA sonuçları, öğrencilerin performanslarının okul kaynakları, öğrenme ortamı, okul iklimi ve daha geniş eğitim sistemi içindeki eşitsizliklerle birlikte değerlendirilmesi gerektiğini ortaya koymaktadır. Bu bulgunun önemli çıkarımı, akademik başarının bireysel performansın ötesinde kurumsal ve yapısal koşulların da bir ürünü olduğudur. Dolayısıyla başarı tahmininde kullanılan değişkenlerin yorumlanması sürecinde, yalnızca mikro düzey bireysel verilerin değil, başarıyı şekillendiren makro düzey bağlamın da dikkate alınması gerekmektedir.

Güncel çalışmalar, akademik performansı etkileyen değişkenlerin yalnızca klasik göstergelerle sınırlı olmadığını; öğrencilerin dijital ortamlarda bıraktıkları izlerin de anlamlı bilgi taşıdığını göstermektedir. Canale vd. (2024), sınav performansını tahmin etmeye yönelik çalışmalarında sürüm kontrol sistemlerinden türetilen özellikleri kullanarak, dijital üretim süreçlerine dayalı verilerin de tahmin gücüne katkı sağlayabildiğini ortaya koymuştur. Çalışmanın temel bulgusu, öğrencinin yalnızca sonuçları değil, öğrenme ve üretim sürecindeki davranış örüntülerinin de akademik performansla ilişkili olduğudur. Bu bulgudan hareketle, eğitim verisinin yalnızca çıktı odaklı değil, süreç odaklı bir yapıya sahip olduğu sonucuna ulaşılmaktadır. Başka bir ifadeyle, öğrencinin ne kadar başarılı olduğu kadar, nasıl çalıştığına ilişkin dijital izler de başarı tahmininde anlamlı değişkenler üretmektedir.

Benzer biçimde Ilic vd. (2021), programlama öğretimi bağlamında yürüttükleri çalışmada öğretim sistemi içinde oluşan davranışsal verilerin öğrenci performansını tahmin etmede anlamlı katkılar sunduğunu göstermiştir. Etkileşim yoğunluğu, ilerleme düzeyi ve sistem içi kullanım örüntüleri gibi göstergelerin tahmin gücüne katkıda bulunduğu belirtilmektedir. Çalışmanın ortaya koyduğu temel sonuç, platform-içi davranışların öğrencinin öğrenme sürecine ilişkin doğrudan sinyaller taşıdığıdır. Bu bulgu, dijital öğrenme ortamlarından elde edilen verilerin yalnızca tamamlayıcı değil, öğrencinin öğrenme performansını açıklamada temel kaynaklardan biri olarak değerlendirilmesi gerektiğini göstermektedir. Dolayısıyla öğrenme yönetim sistemleri ve benzeri platformlardan elde edilen davranışsal veriler, performans tahmin modellerinde yüksek açıklayıcılık potansiyeline sahip değişkenler olarak ele alınmalıdır.

Pesovski vd. (2025) ise öğrenci başarısını yalnızca bireysel ya da davranışsal değişkenlerle değil, sosyal etkileşim ağları üzerinden de inceleyerek akran etkileşimlerinden türetilen merkezilik ölçütlerinin başarı tahmininde güçlü belirleyiciler olabileceğini göstermiştir. Bu çalışmanın temel bulgusu, öğrencinin öğrenme topluluğu içindeki konumunun ve akranlarıyla kurduğu ilişkilerin akademik performansla anlamlı biçimde ilişkili olduğudur. Bu bulgudan çıkarılabilecek önemli sonuç, başarının yalnızca bireyin kendi özelliklerinden türeyen bir çıktı olmadığı; aynı zamanda sosyal öğrenme dinamiklerinin de bu süreçte etkili olduğudur. Bu nedenle sosyal ağ yapılarından elde edilen ilişkisel değişkenlerin performans tahmini modellerine dâhil edilmesi, öğrencinin öğrenme deneyimini daha bütüncül biçimde temsil etme olanağı sunmaktadır.

Bu çalışmalar birlikte değerlendirildiğinde, akademik performansı etkileyen değişkenlerin en az üç temel düzeyde yapılandığı görülmektedir: bireysel ve bağlamsal değişkenler, öğrenme sürecine katılımı yansıtan davranışsal değişkenler ve akran etkileşimini yansıtan ilişkisel değişkenler. Bu durum, eğitim verisinin basit bir tablosal yapıdan ibaret olmadığını; zaman içinde üretilen, bağlama göre anlam kazanan ve farklı veri katmanlarını bir araya getiren çok boyutlu bir sistem olduğunu göstermektedir. Dolayısıyla öğrenci performansını tahmin etmeye yönelik modellerin gücü, yalnızca kullanılan algorithmadan değil, başarıyı çok boyutlu biçimde temsil eden uygun değişken setlerinin oluşturulmasından da doğrudan etkilenmektedir. Sonuç olarak, akademik performansı daha güçlü açıklayan ve daha güvenilir tahmin eden modeller geliştirebilmek için eğitim verisinin yapısal, davranışsal ve

ilişkisel boyutlarını birlikte dikkate alan bütüncül bir veri yaklaşımına ihtiyaç vardır.

2.1.3 Açıklanabilir Yapay Zekâ (XAI) ve Eğitimde Şeffaflık

Adadi ve Berrada (2018), yeni bir tahmin modeli geliştirmekten çok, açıklanabilir yapay zekâ alanındaki yöntemleri derleyen ve sistematik biçimde sınıflandıran bir anket çalışması sunmuştur. Çalışmanın teknik katkısı, açıklanabilirlik yaklaşımlarını model-içi ve model-dışı, yerel ve küresel, ayrıca yorumlama, görselleştirme ve örnek temelli açıklama gibi eksenlerde yapılandırarak alandaki yöntemlerin kavramsal haritasını ortaya koymasındır. Benzer şekilde Arrieta vd. (2020) de ağırlıklı olarak taksonomik ve kavramsal bir çerçeve geliştirmiş; açıklanabilir yapay zekâyı yalnızca açıklama üretme meselesi olarak değil, yorumlanabilirlik, şeffaflık, adillik, güvenilirlik ve sorumlu yapay zekâ boyutlarıyla birlikte ele almıştır. Doshi-Velez ve Kim (2017) ise deneysel bir model önermek yerine, yorumlanabilirliğin nasıl değerlendirileceğine ilişkin metodolojik bir zemin sunmuş; değerlendirme sürecini uygulama-temelli, insan-temelli ve işlev-temelli katmanlar üzerinden ele alarak “yorumlanabilirlik nasıl ölçülür?” sorusuna odaklanmıştır.

Yöntem öneren teknik çalışmalar arasında Ribeiro, Singh ve Guestrin’in (2016) geliştirdiği LIME yaklaşımı önemli bir yer tutmaktadır. Bu çalışmada, kara kutu modelin belirli bir gözlem etrafındaki davranışını açıklamak amacıyla ilgili gözlemin çevresinde pertürbe edilmiş örnekler üretilmiş, bu örnekler asıl modelden geçirilmiş ve yakın örneklerle daha yüksek ağırlık verilerek yerelde doğrusal ve yorumlanabilir bir vekil model kurulmuştur. Böylece LIME, herhangi bir sınıflandırıcının tekil kararını modelden bağımsız biçimde açıklayan yerel vekil modelleme yaklaşımı olarak öne çıkmıştır. Lundberg ve Lee (2017) tarafından geliştirilen SHAP yöntemi ise açıklamayı oyun teorisindeki Shapley değerleri üzerinden kurmakta ve her bir özneliliğin katkısını tüm olası öznelilik kombinasyonlarındaki marjinal katkıları dikkate alarak hesaplamaktadır. Bu sayede açıklamalar toplamsal, tutarlı ve karşılaştırılabilir hâle gelmekte; SHAP, LIME’daki gibi yaklaşık yerel doğrusal modelleme yerine daha güçlü kuramsal temele dayanan birleşik bir açıklama çerçevesi sunmaktadır.

Eğitim bağlamındaki çalışmalar ise daha çok açıklanabilirliğin neden gerekli olduğunu temellendirmektedir. Holmes vd. (2022) çalışmaları, doğrudan bir tahmin modeli

geliştirmekten çok, eğitimde kullanılan yapay zekâ sistemlerinin öğrenci hakları, şeffaflık, insan denetimi, hesap verebilirlik ve pedagojik uygunluk açısından nasıl değerlendirilmesi gerektiğini tartışmaktadır. Baker ve Hawn (2021) de benzer biçimde, eğitim teknolojilerinde algoritmik yanlılık, adillik ve veri temsiliyeti sorunlarını öne çıkarmakta; teknik olarak yeni bir XAI algoritması sunmaktan ziyade, eğitimde kullanılan algoritmaların hangi veri ve tasarım seçimleri nedeniyle eşitsizlik üretebildiğini ortaya koymaktadır. Buna karşılık Jang vd. (2022) gibi uygulamalı çalışmalar, öğrenci performansının erken tahmini için makine öğrenmesi modelleri kurmakta; veri ön işleme, sınıflandırma modeli oluşturma, performans ölçütlerine göre model seçimi ve ardından SHAP ya da LIME gibi yöntemlerle öznitelik katkılarının açıklanması adımlarını izlemektedir. Bu yönüyle söz konusu çalışmalar, yöntemsel olarak öğrenci performansı tahminine dayalı tez çalışmalarına daha yakındır. Basu ve Rao (2023) ise eğitim veri madenciliği bağlamında açıklanabilir yapay zekâyı daha çok uygulamaya dönük karar desteği aracı olarak ele almakta; açıklamaların öğretmen ve yöneticilerin müdahale kararlarına dönüştürülebilir olmasına vurgu yapmaktadır.

Sonuç olarak alanyazın, açıklanabilir yapay zekânın yalnızca teknik bir model açıklama yöntemi olmadığını; aynı zamanda güven, şeffaflık, adillik ve sorumlu yapay zekâ ilkeleriyle yakından ilişkili olduğunu göstermektedir. Bu nedenle SHAP ve LIME gibi açıklanabilirlik yöntemleri, öğrenci performansı tahmin modellerinin yorumlanması ve eğitimsel karar süreçlerinin desteklenmesi açısından önemli araçlar olarak değerlendirilmektedir.

2.2 Ulusal Alanyazın

Bu bölümde; “Ulusal Bağlamda Öğrenci Performansı Tahmini Çalışmaları, Alanyazının Sentezi ve Araştırma Gerekisini” başlıkları bulunmaktadır.

2.2.1 Ulusal Bağlamda Öğrenci Performansı Tahmini Çalışmaları

Ulusal bağlamındaki öğrenci performansı tahmini çalışmaları incelendiğinde, uluslararası alanyazına benzer biçimde makine öğrenmesi yöntemlerinin eğitim verileri üzerinde giderek daha yoğun biçimde uygulandığı görülmektedir. Bununla birlikte ulusal alanyazının öne çıkan yönü, yalnızca tahmin başarımına odaklanmakla kalmayıp, veri dengesizliği, özellik seçimi, model sadeleştirme ve pedagojik yorumlama gibi uygulamaya dönük sorunları da

tartışmaya açmasıdır. Bu durum, ulusal çalışmaların öğrenci performansı tahminini salt teknik bir problem olarak değil, aynı zamanda eğitimsel karar desteği ve yöntemsel sağlamlık çerçevesinde ele almaya başladığını göstermektedir.

Başer vd. (2020), ortaöğretim öğrencilerinin performansını tahmin etmek amacıyla okul raporları ve anket verilerini kullanarak Iterative Classifier, OneR, LogitBoost ve Yapay Sinir Ağları gibi farklı yöntemleri karşılaştırmıştır. Çalışmada OneR yönteminin yüksek doğruluk ve duyarlılık değerleriyle öne çıkması, daha basit ve yorumlanabilir modellerin bazı eğitim veri setlerinde karmaşık yöntemlerle rekabet edebildiğini göstermesi bakımından dikkat çekicidir. Bu bulgu, öğrenci performansı tahmininde model karmaşıklığının tek başına üstünlük sağlamadığını; veri yapısına uygun ve anlaşılabilir yöntemlerin de güçlü sonuçlar verebildiğini ortaya koymaktadır. Ulusal bağlamda bu tür sonuçlar, özellikle eğitim yöneticileri ve uygulayıcılar açısından yorumlanabilir modellerin neden önemli olduğunu göstermesi bakımından anlamlıdır.

Aghalarova ve Bozkurt Keser (2021), ortaokul öğrencilerinin Matematik ve Portekizce ders başarılarını tahmin etmeye yönelik çalışmalarında SMOTE tabanlı ön işleme, özellik seçimi ve normalizasyon süreçlerini hiperparametre optimizasyonu ile birlikte ele almıştır. Çalışmanın temel katkısı, dengesiz veri yapısının ve model ayar süreçlerinin tahmin başarımı üzerinde belirleyici rol oynadığını açık biçimde göstermesidir. Bu yönüyle araştırma, ulusal alanyazında veri ön işleme ile hiperparametre optimizasyonunu birlikte değerlendiren daha sistematik bir çizgiyi temsil etmektedir. Aynı zamanda, öğrenci başarı tahmininde yalnızca algoritma seçiminin değil, veri hazırlama ve ayarlama süreçlerinin de model performansını doğrudan etkilediğini ortaya koyması bakımından bu tez çalışmasının yöntemsel çerçevesiyle doğrudan ilişkilidir. Bu çalışma DergiPark üzerinde doğrulanabilmektedir.

Göreke (2024), öğrencilerin akademik başarısını etkileyen faktörleri analiz ederek Random Forest, Ridge ve Derin Sinir Ağları ile tahmin modelleri geliştirmiş; özellikle Random Forest üzerinden özellik önem düzeylerini hesaplayarak düşük önem düzeyli değişkenlerin elenmesinin performansa etkisini tartışmıştır. Çalışmanın önemli katkısı, yalnızca tahmin doğruluğuna odaklanmak yerine hangi değişkenlerin model için daha anlamlı olduğunu da incelemesidir. Bu yaklaşım, ulusal bağlamda özellik önem analizi ve model sadeleştirme konularının öne çıkmaya başladığını göstermektedir. Başka bir ifadeyle, ulusal alanyazında

artık yalnızca “hangi model daha iyi tahmin ediyor?” sorusu değil, “hangi değişkenler gerçekten belirleyici?” sorusu da daha görünür hâle gelmektedir. Bu çalışma da DergiPark üzerinde doğrulanabilmiştir.

Bakan ve Kanbay (2024), makine öğrenmesi yöntemleriyle eğitim başarısına etki eden faktörlerin modellenmesine odaklanarak ulusal alanyazında faktör analizi ile performans modellemesini buluşturan çalışmalardan biri olarak değerlendirilebilir. Bu tür çalışmaların temel önemi, öğrenci performansını yalnızca sonuç değişkeni olarak ele almamakta; bu sonuca etki eden belirleyicileri de görünür kılmaya çalışmaktadır. Böylece tahmin modelleri, yalnızca sınıflandırma ya da tahmin aracı olmaktan çıkıp, eğitim sürecinde hangi değişkenlerin daha yakından izlenmesi gerektiğine dair karar desteği de üretmektedir.

Çoban Kapuoğlu (2024) ise makine öğrenmesinin eğitimde kullanımını eğitim felsefesi perspektifinden tartışarak ulusal alanyazına farklı bir katkı sunmaktadır. Çalışmanın öne çıkardığı temel nokta, yapay zekâ uygulamalarının eğitimde yalnızca teknik verimlilik temelinde değil, pedagojik ve etik boyutlarıyla birlikte değerlendirilmesi gerektiğidir. Bu vurgu, Türkiye bağlamında gelişen öğrenci performansı tahmini çalışmalarının yalnızca yöntemsel doğrulukla sınırlı kalmaması; aynı zamanda eğitimsel anlam, adalet ve sorumluluk çerçevesinde de tartışılması gerektiğini göstermektedir. Bu yönüyle söz konusu çalışma, ulusal alanyazında teknik modelleme ile pedagojik değerlendirme arasında köprü kuran önemli bir yaklaşımı temsil etmektedir.

Bu çalışmalar birlikte değerlendirildiğinde, ulusal bağlamdaki öğrenci performansı tahmini araştırmalarının dört temel çizgide geliştiği söylenebilir. İlk olarak, farklı makine öğrenmesi algoritmalarının karşılaştırılması yoluyla hangi yöntemlerin eğitim verilerinde daha etkili olduğu araştırılmaktadır. İkinci olarak, veri dengesizliği, özellik seçimi ve hiperparametre optimizasyonu gibi yöntemsel sorunlar daha görünür hâle gelmektedir. Üçüncü olarak, özellik önem analizi ve model sadeleştirme gibi yaklaşımlarla tahmin modellerinin daha yorumlanabilir ve uygulanabilir hâle getirilmesi amaçlanmaktadır. Dördüncü olarak ise pedagojik ve etik boyutların tartışmaya açılması, ulusal alanyazının giderek daha bütüncül bir çerçeveye yöneldiğini göstermektedir. Dolayısıyla ulusal bağlamdaki çalışmalar, öğrenci performansı tahminini yalnızca teknik bir sınıflandırma problemi olarak değil, eğitimsel karar desteği üreten çok boyutlu bir araştırma alanı olarak konumlandırmaya başlamıştır.

2.3 Alanyazının Sentezi ve Araştırma Gereksinimi

Uluslararası alanyazın, öğrenci performansı tahmininde denetimli öğrenme yöntemlerinin yaygınlaştığını; davranışsal/dijital iz verilerinin (ÖYS, VCS, ağ verisi) giderek daha fazla kullanıldığını ve tahminin erken müdahale tasarımlarına bağlandığını göstermektedir. Buna karşın genellenebilirlik, dengesiz veri ve sızıntı gibi yöntemsel riskler, gerçek dünya uygulamalarında önemli sorunlar olarak kalmaktadır. Kaufman vd. (2012) ve Kapoor ve Narayanan (2023) çizgisinde tartışılan sızıntı problemleri; He ve Garcia (2009) ve Saito ve Rehmsmeier (2015) çizgisinde tartışılan dengesiz veri ve metrik seçimi sorunları, eğitim verilerinde ‘yüksek skor’ görünen modellerin sahada başarısız olabileceğini göstermektedir.

Ulusal alanyazında ise öğrenci başarı tahminine yönelik çalışmaların sayısı artmakla birlikte, çoğunlukla tek veri kaynağı ve sınırlı açıklanabilirlik düzeyinde kaldığı görülmektedir. SHAP ve LIME gibi modern XAI yöntemlerinin ulusal uygulamalarda daha sınırlı görünmesi, model çıktılarının paydaşlar tarafından yorumlanabilirliğini azaltabilmektedir. Bu noktada teziniz, UBYs + ÖYS + anket gibi çoklu veri kaynaklarını bütünleştirerek hem sınıflandırma hem regresyon modelleri kurmayı ve çıktıları XAI ile açıklayarak karar desteğe dönüştürmeyi hedeflemesi bakımından alanyazındaki önemli bir boşluğu doldurmaya adaydır.

Sonuç olarak alanyazındaki ihtiyaç, yalnızca yüksek doğruluklu tahmin modelleri değil; aynı zamanda

- Sızıntı kontrolü yapılmış,
- Dengesiz veri koşullarına duyarlı metriklerle değerlendirilmiş,
- Bağlama duyarlı ve genellenebilirliği tartışılmış,
- SHAP/LIME gibi yöntemlerle açıklanabilirliği sağlanmış modellerin geliştirilmesidir. Bu tez, söz konusu gereksinimleri tek bir araştırma tasarımında bütünleştirmeyi amaçlamaktadır.

Van Rossum ve Drake (2009), Python dilinin referans yapısını sunarak bilimsel hesaplama ekosisteminin temelini oluşturur. McKinney (2010), pandas veri yapılarıyla istatistiksel hesaplama süreçlerini kolaylaştırırken; Pedregosa vd. (2011) scikit-learn kütüphanesiyle

sınıflandırma, regresyon, model seçimi ve değerlendirme adımlarını standartlaştırmıştır. Bu yazılım ekosistemi, eğitim verilerinde tekrar üretilebilir (reproducible) analiz akışları kurmak için pratik bir altyapı sağlar ve sızıntı/ön işleme adımlarının kod düzeyinde izlenebilirliğini artırır.



3. MATERYAL VE YÖNTEM

Bu bölümde, çalışmada kullanılan veri kaynakları, araştırmanın yürütülmesinde izlenen yöntemsel yaklaşım ve uygulanan analiz süreçleri ayrıntılı olarak açıklanmaktadır. Araştırma kapsamında öğrencilerin akademik performanslarını etkileyen faktörlerin çok boyutlu biçimde incelenebilmesi amacıyla; anket uygulamaları, üniversite bilgi yönetim sistemi kayıtları ve e-öğrenme platformlarından elde edilen veriler birlikte kullanılmıştır. Elde edilen veriler, makine öğrenmesi ve açıklanabilir yapay zekâ yaklaşımlarına uygun şekilde işlenmiş; veri ön işleme, özellik seçimi, model eğitimi ve performans değerlendirme aşamalarından geçirilmiştir.

3.1 Materyal

Bu bölümde; “Veri Seti, Veri Seti Özellikleri, Makine Öğrenmesi ve Açıklanabilir Yapay Zeka Açısından Değişken Seçim Gerekçesi” başlıkları bulunmaktadır.

3.1.1 Veri Seti

Bu çalışmada yararlanılan veri seti, Bartın Üniversitesi’nde öğrenim gören lisans ve önlisans öğrencilerinden elde edilen çok kaynaklı verilerden oluşmaktadır. Araştırma kapsamında öğrencilerin akademik performanslarını etkileyen faktörlerin bütüncül biçimde analiz edilebilmesi amacıyla, anket uygulamaları, Üniversite Bilgi Yönetim Sistemi (UBYS) kayıtları ve e-öğrenme (e-ders) platformlarından sağlanan veriler birlikte kullanılmıştır. Böylece yalnızca akademik performans göstergelerine dayalı bir değerlendirme yerine, öğrencilerin demografik özellikleri, öğrenme alışkanlıkları ve dijital etkileşimlerini de kapsayan çok boyutlu bir veri yapısı oluşturulmuştur.

Anketler aracılığıyla öğrencilerin demografik bilgileri (yaş, cinsiyet, aile eğitim durumu vb.) ile ders çalışma alışkanlıkları ve öğrenme ortamlarına ilişkin veriler toplanmıştır. UBYS üzerinden elde edilen veriler, öğrencilerin vize-final notları ve genel not ortalamaları gibi akademik performans göstergelerini içermektedir. Buna ek olarak, e-ders platformlarından sağlanan veriler; ders materyallerine erişim süreleri, etkileşim sayıları ve çevrimiçi öğrenme etkinliklerine katılım düzeylerini yansıtmaktadır. Bu farklı veri kaynakları, her bir öğrenci

için tekil ve bütünleşik bir veri kümesi oluşturacak şekilde bir araya getirilmiştir.

Araştırmanın örneklemini, gönüllülük esasına göre çalışmaya katılan önlisans ve lisans öğrencileri oluşturmaktadır. Katılımcılar uygun örnekleme yöntemiyle belirlenmiş olup, veri toplama süreci boyunca etik ilkelere uyulmuştur. Çalışma kapsamında tüm katılımcılardan alınan bilgiler ve elde edilen veriler anonimleştirilerek yalnızca bilimsel araştırma amaçları doğrultusunda kullanılmıştır. Öğrencilerin kimlik bilgileri hiçbir aşamada analiz sürecine dâhil edilmemiş ve veriler gizlilik ilkeleri çerçevesinde korunmuştur.

3.1.2 Veri Seti Özellikleri

Bu çalışmada kullanılan veri seti, öğrencilerin akademik performanslarını etkileyebilecek çok sayıda faktörün birlikte ele alınabilmesini sağlamak amacıyla çok boyutlu ve bütünleşik bir yapıdadır. Veri seti; farklı kaynaklardan elde edilen demografik, akademik, davranışsal ve dijital etkileşim temelli değişkenleri içermekte olup, makine öğrenmesi ve açıklanabilir yapay zekâ modellerinin eğitimi için uygun bir zemin sunmaktadır. Bu kapsamda veri seti, bağımlı değişkenler ve bağımsız değişkenler olmak üzere iki ana grupta ele alınmıştır.

Bağımlı değişken, öğrencinin akademik performansıdır. Akademik performans iki başlık altında incelenmiştir; öğrencilerin dönem sonu geçme notu ortalamaları ve ders başarı durumu (başarılı/başarısız). Analiz sürecinde bu değişkenler, araştırma probleminin yapısına bağlı olarak hem regresyon hem de sınıflandırma yaklaşımları kapsamında ele alınmıştır. Bu ikili yapı sayesinde, öğrencilerin hem genel başarı profilleri hem de potansiyel risk durumları bütüncül bir bakış açısıyla değerlendirilebilmiştir.

Bağımsız değişkenler, öğrencilerin akademik performansını (geçme notu ve ders başarı durumu) etkilediği girdilerden oluşmakta ve dört ana kategori altında toplanmaktadır.

- a. **Demografik değişkenler,** öğrencilerin yaş, cinsiyet, anne ve baba eğitim durumu ile aile gelir düzeyi gibi özelliklerini kapsamaktadır. Bu değişkenler, öğrencilerin öğrenme ortamlarını, akademik destek düzeylerini ve eğitim sürecine ilişkin fırsat eşitliğini dolaylı olarak etkileyebilecek faktörleri temsil etmektedir. Alanyazında, sosyo-demografik özelliklerin akademik performans

üzerinde anlamlı etkiler oluşturabildiği belirtilmektedir (Sirin, 2005; OECD, 2019).

b. Akademik değişkenler, öğrencilerin eğitim sürecindeki doğrudan başarı göstergelerini içermektedir. Bu kapsamda vize notları, final notu ve öğrencinin geçme notu ortalamaları kullanılmıştır. Bu değişkenler, modelleme sürecinde akademik performansın en güçlü belirleyicileri arasında yer almakta ve tahmin doğruluğunu artıran temel girdiler olarak değerlendirilmektedir (Romero ve Ventura, 2010; Alnasyan vd., 2024).

c. Davranışsal değişkenler, öğrencilerin ders çalışma alışkanlıkları, ders dışı çalışma süreleri ve devam/devamsızlık durumları öğrenme sürecine aktif katılımı yansıtan göstergelerden oluşmaktadır. Bu tür değişkenler, öğrencilerin akademik disiplinlerini düzeylerini ve öğrenme süreçlerine ayırdıkları zamanı temsil etmektedir. Özellikle devamsızlık ve ders çalışma süresi gibi değişkenlerin, akademik performans üzerinde belirleyici bir role sahip olduğu bilinmektedir (Credé, Roch ve Kieszczyńska, 2010).

d. Dijital etkileşim değişkenleri ise e-öğrenme platformlarından elde edilen verileri kapsamaktadır. Ders materyallerine erişim süreleri, çevrimiçi etkileşim sayıları ve dijital öğrenme etkinliklerine katılım düzeyleri bu grupta yer almaktadır. Bu değişkenler, öğrencilerin dijital öğrenme ortamlarındaki davranışlarını ve öğrenme süreçlerine ne ölçüde entegre olduklarını göstermesi açısından önem taşımaktadır (Siemens ve Long, 2011; Romero ve Ventura, 2013). Özellikle hibrit ve çevrimiçi eğitim modellerinin yaygınlaşmasıyla birlikte, dijital etkileşim verilerinin akademik performans tahmininde kritik bir rol oynadığı görülmektedir (Kovanović vd., 2015; Gašević vd., 2015).

Veri setinde yer alan tüm değişkenler, analiz öncesinde sayısal ve kategorik özelliklerine göre sınıflandırılmış ve makine öğrenmesi modellerine uygun hale getirilmek üzere ön işleme adımlarına tabi tutulmuştur.

Bağımlı ve bağımsız değişkenlerin genel dağılımı, türleri ve ait oldukları kategoriler Tablo 3.1’de özetlenmiştir. Bu yapı, izleyen bölümlerde gerçekleştirilecek özellik seçimi, modelleme ve açıklanabilirlik analizlerine metodolojik bir temel oluşturmaktadır.

Tablo 3.1: Veri seti nitelikleri ve açıklamaları

No	Değişken Adı	Değişken Türü	Değişken Kategorisi	Açıklama
1	Öğrenci ID	Kategorik	Kimlik (Anonim)	Öğrenci numarasından anonimleştirilmiş benzersiz tanımlayıcı
2	Yaş	Sayısal	Demografik	Öğrencinin yaşı
3	Cinsiyet	Kategorik	Demografik	Kadın / Erkek
4	Fakülte-Bölüm	Kategorik	Demografik	Öğrencinin kayıtlı olduğu fakülte ve bölüm
5	Ders Adı	Kategorik	Akademik	Anketin doldurulduğu ders
6	Aile Gelir Düzeyi	Kategorik (Ordinal)	Demografik	Ailenin aylık gelir aralığı
7	Anne Eğitim Durumu	Kategorik (Ordinal)	Demografik	Annenin eğitim seviyesi
8	Baba Eğitim Durumu	Kategorik (Ordinal)	Demografik	Babanın eğitim seviyesi
9	Yaşanılan Yer	Kategorik	Demografik	Şehir merkezi / İlçe / Köy
10	Lise Türü	Kategorik	Ak. Geçmiş	Mezun olunan lise türü
11	Barınma Durumu	Kategorik	Demografik	Öğrencinin öğrenim süresince kaldığı yer
12	Kullanılan Ders Kaynakları	Kategorik (Çoklu)	Davranışsal	Kitap, not, e-ders, internet, yapay zekâ
13	Haftalık Ders Çalışma Süresi	Kategorik (Ordinal)	Davranışsal	Haftalık ortalama çalışma süresi
14	Çalışma Ortamı	Kategorik	Davranışsal	Ev, kütüphane, yurt vb.
15	Çalışma Ortamı Motivasyonu	Kategorik (Ordinal)	Davranışsal	Ortamın motive edicilik düzeyi
16	Ders Çalışırken Ara Verme Sıklığı	Kategorik (Ordinal)	Davranışsal	Çalışma sırasında mola sıklığı
17	Sınava Hazırlık Süresi	Kategorik (Ordinal)	Davranışsal	Sınav öncesi hazırlığa başlama zamanı

Tablo 3.1: (devam ediyor)

18	Sınav Kaygısı Düzeyi	Kategorik (Ordinal)	Psikolojik	Öğrencinin sınav kaygısı durumu
19	Derslere Katılım Düzeyi	Kategorik (Ordinal)	Davranışsal	Derslere devam sıklığı
20	Öğrenme Stili	Kategorik	Bilişsel	Görsel, işitsel, kinestetik vb.
21	Günlük İnternet Kullanım Süresi	Kategorik (Ordinal)	Dijital Davranış	Günlük ortalama internet süresi
22	Akademik Amaçlı İnternet Kullanımı	Kategorik (Ordinal)	Dijital Davranış	İnternetin dersler için kullanım sıklığı
23	Kullanılan Çevrimiçi Öğrenme Araçları	Kategorik (Çoklu)	Dijital Etkileşim	UBYS, YouTube, kurslar, yapay zekâ
24	Çevrimiçi Ders Katılım Düzeyi	Kategorik (Ordinal)	Dijital Etkileşim	Online derslere katılım seviyesi
25	Sosyal Medyanın Ders Çalışmaya Etkisi	Kategorik	Dijital Davranış	Sosyal medyanın etkisi
26	Teknolojinin Ders Çalışmaya Katkısı	Kategorik (Ordinal)	Dijital Davranış	Teknoloji kullanım algısı
27	Vize Notu	Sayısal	Akademik	UBYS'den alınan vize notu
28	Final Notu	Sayısal	Akademik	UBYS'den alınan vize notu
29	Genel Not Ortalaması (GPA)	Sayısal	Akademik	Öğrencinin dönemsel başarı göstergesi
30	Akademik performans /dönem sonu geçme notu	Sayısal	Bağımlı Değişken	100 üzerinden not
31	Akademik performans/Başarılı başarısız	Kategorik	Bağımlı Değişken	Başarılı/Başarısız

3.1.3 Makine Öğrenmesi ve Açıklanabilir Yapay Zekâ Açısından Değişken Seçim Gerekçesi

Bu çalışmadaki değişkenler, makine öğrenmesi modellerinin tahmin doğruluğunu artırmak, açıklanabilir yapay zekâ (XAI) yaklaşımlarıyla model kararlarının yorumlanabilirliğini sağlamak amacıyla bilinçli ve çok boyutlu bir çerçevede seçilmiştir. Akademik performansı tek başına notlar üzerinden modellemek yerine, öğrencilerin demografik özellikleri, akademik geçmişleri, davranışsal öğrenme alışkanlıkları ve dijital etkileşim düzeylerini birlikte ele alan bir değişken yapısı tercih edilmiştir. Bu yaklaşım, makine öğrenmesi modellerinin yalnızca sonuç odaklı değil, neden–sonuç ilişkilerini de öğrenebilmesine olanak tanımaktadır (Arrieta vd., 2020).

Demografik değişkenler, öğrencilerin sosyoekonomik ve çevresel koşullarını temsil ederek akademik performans üzerindeki dolaylı etkilerin modellenmesini sağlamaktadır. Akademik değişkenler, başarıyı doğrudan yansıtan güçlü göstergeler olarak tahmin modellerinin temel girdilerini oluştururken; davranışsal değişkenler, öğrencilerin öğrenme sürecine aktif katılım düzeylerini ve çalışma disiplinlerini ortaya koymaktadır (Baker ve Siemens, 2014). Dijital etkileşim değişkenleri ise özellikle çevrimiçi ve hibrit eğitim ortamlarında öğrencilerin öğrenme süreçlerine ne ölçüde entegre olduklarını göstererek, çağdaş eğitim dinamiklerinin modele yansıtılmasını sağlamaktadır (Gašević vd., 2015).

Açıklanabilir yapay zekâ yöntemleri açısından bakıldığında, bu çeşitlilikteki değişken yapısı; SHAP (Lundberg ve Lee, 2017) ve LIME (Ribeiro vd., 2016) gibi teknikler aracılığıyla hangi faktörlerin akademik performansı olumlu ya da olumsuz yönde etkilediğinin açık biçimde gösterilmesine imkân tanımaktadır. Böylece geliştirilen modeller yalnızca yüksek doğruluk sunduğu gibi eğitimciler ve yöneticiler tarafından anlaşılabilir, güvenilir ve eyleme geçirilebilir çıktılar üretmektedir (Adadi ve Berrada, 2018). Bu yönüyle seçilen değişken seti hem makine öğrenmesi performansını hem de model açıklanabilirliğini birlikte güçlendiren bir temel oluşturmaktadır.

3.2 Model Tasarımı

Bu çalışmada öğrencilerin akademik performansları temel alınarak iki farklı modelleme

yaklaşımı üzerinde durulmuştur. İlk aşamada, demografik, akademik, davranışsal ve dijital etkileşim değişkenleri öznitelik olarak, öğrencilerin dönem sonunda elde ettikleri geçme notu ise doğrudan hedef değişken olarak kullanılarak regresyon modeli (Model 1) geliştirilmiştir. Bu modelde amaç, öğrencilerin akademik performans düzeylerini etkileyen faktörler ile dönem sonu geçme notu arasındaki ilişkiyi ortaya koymak ve öğrencilerin beklenen akademik performans düzeylerini tahmin edebilmektir. Bu yaklaşım, akademik performansın sayısal olarak ele alınmasına ve başarı düzeylerindeki artış ya da azalışların modellenmesine olanak sağlamaktadır.

İkinci aşamada ise demografik, akademik, davranışsal ve dijital etkileşim değişkenlerine dönem sonu geçme notu da öznitelik olarak eklenerek ders başarı durumunu başarılı ve başarısız olmak üzere sınıflandıracak bir makine öğrenmesi modeli (Model 2) geliştirilmiştir. Bu modelde amaç, öğrencilerin yalnızca mevcut başarı durumlarını değil, aynı zamanda başarısızlık riski taşıyıp taşımadıklarını da belirleyebilmektir. Bu yaklaşım sayesinde, akademik açıdan risk altında bulunan öğrencilerin erken aşamada belirlenmesi ve bu öğrencilere yönelik önleyici veya destekleyici müdahalelerin planlanması mümkün hâle gelmektedir. Böylece iki model birlikte kullanılarak hem akademik performansın düzeyi hem de başarı durumuna ilişkin karar verici çıktılar elde edilmiştir.

3.3 Veri İşleme Araçları

Bu çalışmada veri toplama, ön işleme, modelleme, değerlendirme ve açıklanabilirlik analizleri; açık kaynaklı, akademik çalışmalarda yaygın olarak kullanılan yazılım araçları ve kütüphaneler aracılığıyla gerçekleştirilmiştir. Kullanılan araçlar, çalışmanın amaç ve kapsamına uygun olarak seçilmiş ve alanyazında geçerliliği kanıtlanmış yöntemlere dayandırılmıştır (Pedregosa vd., 2011; McKinney, 2010).

3.3.1 Python Programlama Dili

Araştırma kapsamında tüm veri analizi ve modelleme süreçleri Python programlama dili kullanılarak yürütülmüştür. Python; okunabilir sözdizimi, geniş kütüphane ekosistemi ve makine öğrenmesi uygulamalarındaki esnekliği nedeniyle bilimsel araştırmalarda yaygın olarak tercih edilmektedir. Özellikle veri bilimi, makine öğrenmesi ve yapay zekâ

alanlarında sunduğu güçlü araçlar sayesinde, karmaşık veri setlerinin etkin biçimde işlenmesine olanak sağlamaktadır (Van Rossum ve Drake, 2009).

3.3.1.1 Veri Ön İşleme ve Düzenleme Araçları (Pandas ve NumPy)

Veri setinin temizlenmesi, düzenlenmesi ve analiz öncesinde hazırlanması aşamalarında Pandas ve NumPy kütüphanelerinden yararlanılmıştır. Pandas, farklı veri kaynaklarından elde edilen yapılandırılmış verilerin birleştirilmesi, eksik verilerin işlenmesi ve değişken dönüşümlerinin gerçekleştirilmesi için etkin bir altyapı sunmaktadır. NumPy ise çok boyutlu diziler ve sayısal hesaplamalar üzerinde yüksek performans sağlayarak veri işleme sürecini desteklemektedir (McKinney, 2010). Anket verileri, UBYS ve e-ders platformlarından elde edilen kayıtlar bu araçlar kullanılarak tekil ve bütünlük bir veri kümesine dönüştürülmüştür.

3.3.1.2 Makine Öğrenmesi ve Modelleme Araçları (Scikit-learn)

Makine öğrenmesi modellerinin oluşturulması, eğitilmesi ve test edilmesi süreçlerinde Scikit-learn kütüphanesi kullanılmıştır. Scikit-learn; sınıflandırma, regresyon, özellik seçimi, model doğrulama ve performans ölçümü gibi temel makine öğrenmesi işlemleri için kapsamlı ve güvenilir bir araç seti sunmaktadır. Bu çalışma kapsamında karar ağaçları, destek vektör makineleri ve rastgele orman gibi algoritmalar Scikit-learn altyapısı kullanılarak uygulanmıştır. Ayrıca veri setinin eğitim ve test kümelerine ayrılması ile k-fold çapraz doğrulama süreçleri bu kütüphane aracılığıyla gerçekleştirilmiştir (Pedregosa vd., 2011).

3.3.1.3 Veri Görselleştirme ve Analiz Ortamı

Veri dağılımlarının incelenmesi, model performanslarının karşılaştırılması ve elde edilen sonuçların görsel olarak sunulması amacıyla Matplotlib kütüphanesi kullanılmıştır. Görselleştirme, keşifsel veri analizi sürecinde veri yapısının daha iyi anlaşılmasına ve model çıktılarının yorumlanmasına katkı sağlamıştır. Tüm analiz ve deneysel çalışmalar, sürecin şeffaf, tekrarlanabilir ve belgelenebilir olması amacıyla Jupyter Notebook ortamında yürütülmüştür. Bu ortam, kod, çıktı ve açıklamaların birlikte sunulmasına olanak tanıyarak

araştırmanın izlenebilirliğini artırmıştır.

3.4 Yöntem

Bu bölümde, çalışmada kullanılan makine öğrenmesi yöntemleri, oluşturulan tahmin modelleri ve modellerin değerlendirilme süreci ayrıntılı olarak açıklanmaktadır. Araştırmada izlenen yönetsel yaklaşım, öğrencilerin akademik performanslarını etkileyen faktörlerin belirlenmesi ve bu faktörlere dayalı olarak güvenilir, açıklanabilir ve yüksek doğruluk sunan tahmin modellerinin geliştirilmesine olanak sağlayacak biçimde yapılandırılmıştır. Bu kapsamda farklı makine öğrenmesi algoritmaları kullanılarak karşılaştırmalı bir analiz gerçekleştirilmiştir.

- Karar ağaçları
- Destek vektör makineleri (SVM)/ Destek Vektör Regresyonu (SVR)
- Rastgele Orman (RF)
- Yapay Sinir Ağları (ANN)
- XG Boost
- K En Yakın Komşu (K-NN)

Çalışma kapsamında oluşturulan veri seti, modelleme sürecinde %60 eğitim, %40 test ve %70 eğitim, %30 test verisi olacak şekilde ayrılmıştır. Eğitim verileri kullanılarak, öğrencilerin akademik performansını “başarılı/başarısız” olacak şekilde sınıflandırmak üzere her bir makine öğrenmesi algoritması eğitilmiştir. Test verileri kullanılarak her bir model ayrı ayrı test edilmiştir. Böylece her algoritmanın veriler üzerindeki performansları karşılaştırmalı olarak değerlendirilerek başarımı en iyi olan algoritma, baz model olarak seçilmiştir. Modellerin başarımı yalnızca doğruluk oranları üzerinden değil, aynı zamanda kesinlik, duyarlılık, F1 skoru ve ROC-AUC gibi diğer performans ölçütleri de dikkate alınarak analiz edilmiştir. Ayrıca her bir algoritma için çapraz doğrulama uygulanarak, modellerin farklı veri alt kümeleri üzerindeki tutarlılığı ve genellenebilirliği test edilmiştir.

Bu çoklu değerlendirme yaklaşımı çerçevesinde, her bir performans ölçütü model başarısının farklı bir yönünü yansıtacak şekilde ayrı ayrı ele alınmıştır.

- Doğruluk, modelin toplam tahminleri içerisindeki doğru sınıflandırma oranını göstermekte ve genel performans hakkında temel bir değerlendirme

sunmaktadır. Bununla birlikte, özellikle sınıf dağılımının dengesiz olduğu veri setlerinde tek başına yeterli bir ölçüt olmayabilmektedir (Bishop, 2006; James vd., 2013).

- Kesinlik, modelin pozitif olarak tahmin ettiği örneklerin ne kadarının gerçekten pozitif olduğunu ortaya koyarken, duyarlılık ise gerçek pozitif örneklerin ne kadarının model tarafından doğru biçimde tespit edildiğini göstermektedir. Bu iki ölçüt, özellikle yanlış pozitif ve yanlış negatif sonuçların farklı derecelerde önem taşıdığı uygulamalarda daha anlamlı değerlendirmeler yapılmasına katkı sağlamaktadır (Manning vd., 2008; Murphy, 2012).
- F1 skoru ise kesinlik ve duyarlılık değerlerini birlikte dikkate alan dengeli bir performans ölçütü olarak kullanılmakta ve özellikle dengesiz veri setlerinde model başarısının daha sağlıklı yorumlanmasına yardımcı olmaktadır (Hastie vd., 2009; Manning vd., 2008).
- ROC-AUC metriği, modelin farklı karar eşiklerinde sınıfları ayırt etme başarısını değerlendirmekte ve eşik değerinden bağımsız bir performans göstergesi sunmaktadır. Bu nedenle ROC-AUC, sınıflandırma modellerinin genel ayırt edicilik düzeyini incelemeye önemli bir ölçüt olarak kabul edilmektedir (Duda vd., 2001; Hastie vd., 2009).

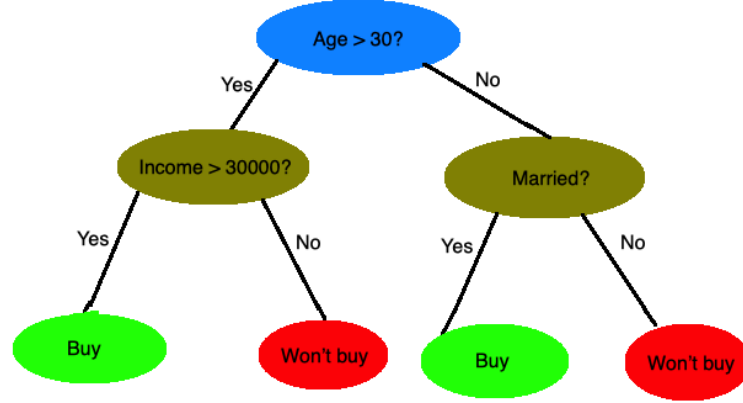
Araştırma süreci; veri toplama, veri ön işleme, özellik seçimi, makine öğrenmesi modellerini eğitimi, model performanslarının değerlendirilmesi ve açıklanabilir yapay zekâ analizleri olmak üzere ardışık aşamalardan oluşmaktadır.

Bu çalışma kapsamında kullanılan makine öğrenmesi algoritmaları aşağıda açıklanmaktadır.

3.4.1 Karar Ağaçları (Decision Trees)

Karar ağaçları, veri setini belirli kurallar çerçevesinde dallara ayırarak karar verme sürecini modelleyen, yorumlanabilirliği yüksek makine öğrenmesi algoritmaları arasında yer almaktadır. Bu algoritma, her bir karar düğümünde veri setini en iyi ayıran değişkeni seçerek hiyerarşik bir yapı oluşturur. Karar ağaçlarının en önemli avantajlarından biri, modelin karar alma sürecinin görsel olarak izlenebilmesi ve sonuçların eğitimciler tarafından kolaylıkla yorumlanabilmesidir. Bu özellik, akademik performans tahmini gibi açıklanabilirliğini

önemli olduğu alanlarda karar ağaçlarını etkili bir yöntem hâline getirmektedir (Quinlan, 1986). Şekil 3.1’de Karar Ağaçları örneği verilmiştir.

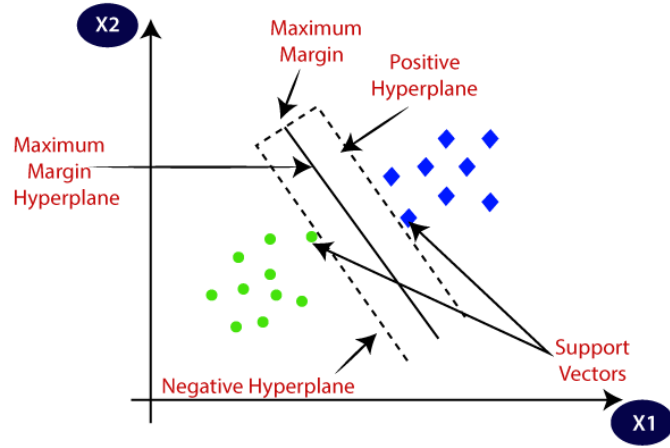


Şekil 3.1: Karar Ağaçları (Kulakowski, 2019)

3.4.2 Destek Vektör Makineleri/ Destek Vektör Regresyonu (Support Vector Machines/ Support Vector Regression)

Destek vektör makineleri, sınıflar arasındaki en uygun ayırıcı sınırı belirleyerek sınıflandırma gerçekleştiren güçlü bir makine öğrenmesi algoritmasıdır. Özellikle yüksek boyutlu veri setlerinde başarılı sonuçlar üreten bu yöntem, genelleme yeteneği yüksek modeller oluşturabilmektedir. Destek vektör makineleri, öğrencilerin akademik performans durumlarını ayırt etmede sınıflar arasındaki karmaşık ilişkileri etkili biçimde öğrenebilmesi nedeniyle bu çalışmada tercih edilmiştir (Cortes ve Vapnik, 1995).

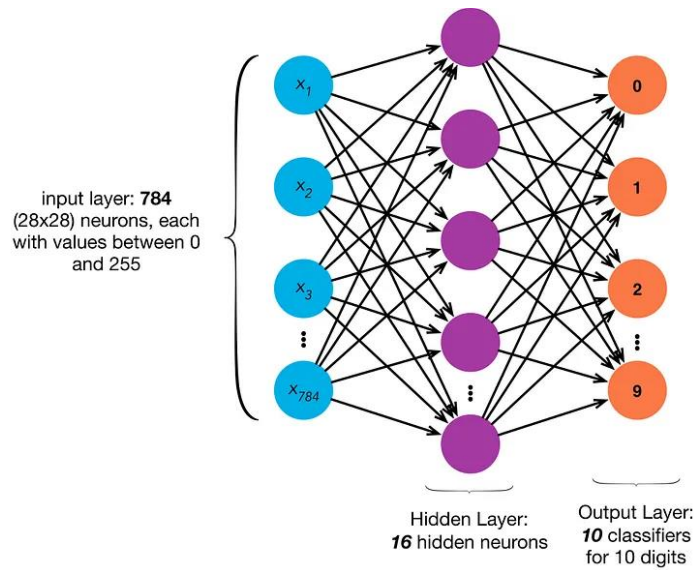
Destek vektör makineleri, sınıflandırma problemlerinin yanı sıra regresyon problemleri için de kullanılabilen esnek bir öğrenme algoritmasıdır. Destek Vektör Regresyonu yaklaşımında amaç, tahmin edilen değerlerin gerçek değerlere belirli bir hata toleransı (ϵ) içerisinde kalmasını sağlayan en uygun fonksiyonun öğrenilmesidir. SVR, marjın kavramını hata toleransı ile birleştirerek aşırı uyumu sınırlandırmakta ve genellenebilirliği yüksek tahminler sunmaktadır (Şekil 3.2.). Bu özellikleri sayesinde SVR, öğrencilerin dönem sonu geçme notu ortalamalarının tahmin edilmesi gibi sürekli değerlerin modellenmesinde etkili bir yöntem olarak alanyazında yaygın biçimde kullanılmaktadır (Vapnik, 1995; Smola ve Schölkopf, 2004).



Şekil 3.2: Destek Vektör Makineleri

3.4.3 Yapay Sinir Ağları (Artificial Neural Networks)

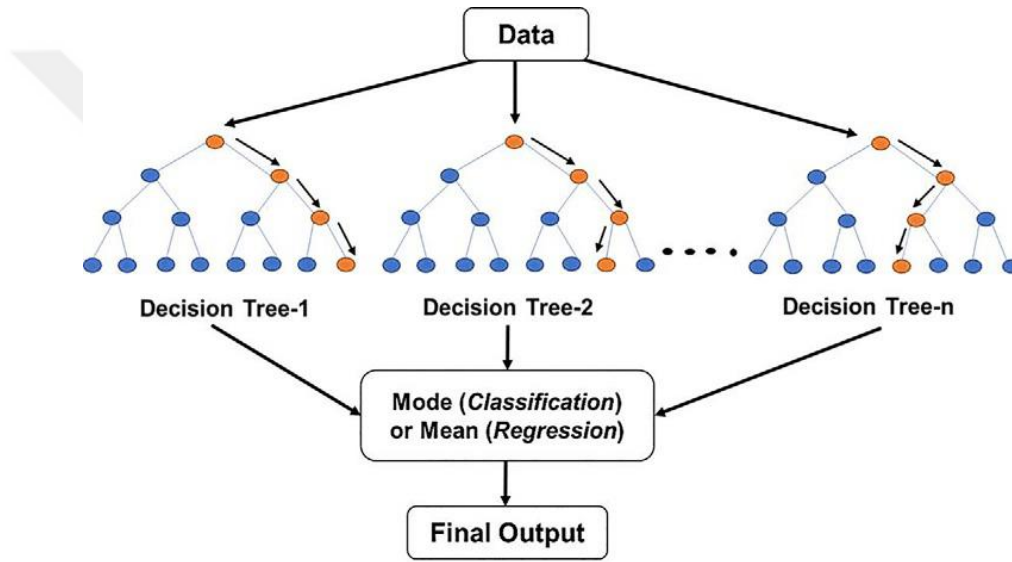
Yapay sinir ağları, insan beyninin çalışma prensibinden esinlenerek geliştirilen ve doğrusal olmayan ilişkileri öğrenebilme yeteneğine sahip modellerdir. Girdi, gizli ve çıktı katmanlarından oluşan bu yapılar, çok boyutlu veri setlerinde karmaşık örüntüleri ortaya çıkarabilmektedir (Şekil 3.3.). Akademik performans tahmini gibi çok sayıda değişkenin etkileşim içinde olduğu problemlerde, yapay sinir ağları yüksek doğruluk sağlayabilen etkili bir yöntem olarak öne çıkmaktadır (Haykin, 2009).



Şekil 3.3: Yapay Sinir Ağları

3.4.4 Rastgele Orman (Random Forest)

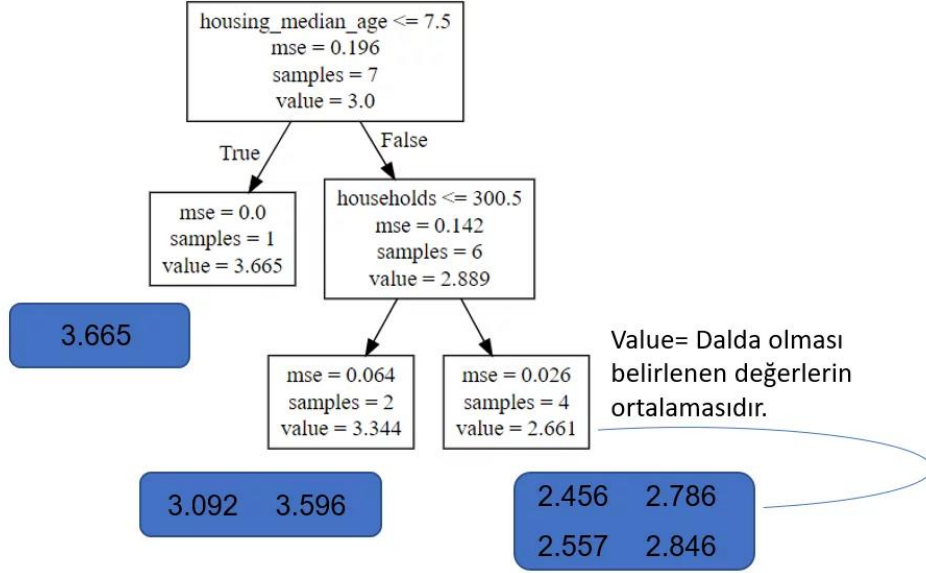
Random Forest algoritması, birden fazla karar ağacının birleşiminden oluşan bir topluluk öğrenme (ensemble) yöntemidir. Bu yaklaşım, tek bir karar ağacına kıyasla aşırı öğrenme riskini azaltmakta ve daha dengeli tahminler sunmaktadır. Ayrıca değişken önem düzeylerini belirleyebilmesi sayesinde, akademik performansı etkileyen temel faktörlerin anlaşılmasına katkı sağlamaktadır. Bu özellikleri nedeniyle Random Forest, çalışmada karşılaştırmalı analiz için önemli bir yöntem olarak kullanılmıştır (Breiman, 2001). Şekil 3.4.'te Rastgele Orman modeli örnek yapısı verilmiştir.



Şekil 3.4: Rastgele Orman (Kumar vd., 2022)

3.4.5 XGBoost (Extreme Gradient Boosting)

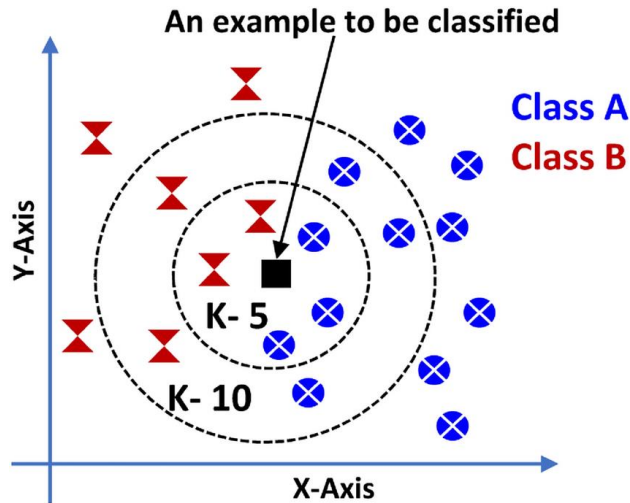
XGBoost, ardışık olarak oluşturulan modeller aracılığıyla hatalı tahminleri düzelten, yüksek doğruluk ve hesaplama verimliliği sunan bir boosting algoritmasıdır. Büyük ve karmaşık veri setlerinde güçlü performans göstermesi ve parametre optimizasyonuna elverişli yapısı sayesinde, akademik performans tahmini problemlerinde sıklıkla tercih edilmektedir. Bu çalışmada XGBoost algoritması, diğer yöntemlerle karşılaştırılarak tahmin doğruluğu ve işlem süresi açısından değerlendirilmiştir (Chen ve Guestrin, 2016). Şekil 3.5'te XGBoost ağaç yapısı örneği verilmiştir.



Şekil 3.5: XGBoost

3.4.6 K-En Yakın Komşu (K-NN)

Bu algoritma, gözlemler arasındaki benzerliklere dayalı olarak çalışan bir makine öğrenmesi yöntemidir. Bir örneğin sınıfı veya değeri, veri uzayında kendisine en yakın k komşunun bilgileri dikkate alınarak belirlenmektedir (Şekil 3.6). K-NN algoritması hem regresyon hem de sınıflandırma problemlerinde kullanılabilmekte olup, sınıflandırmada çoğunluk oyu, regresyonda ise komşu değerlerin ortalaması esas alınmaktadır. Basit yapısı ve veri dağılımına ilişkin varsayım gerektirmemesi nedeniyle, özellikle küçük ve orta ölçekli veri setlerinde etkili sonuçlar üretebilmektedir (Cover ve Hart, 1967; Han vd., 2012).



Şekil 3.6: K-En Yakın Komşu (Kemal, 2025)

3.5 SHAP ve LIME

SHAP (Shapley Additive Explanations), bir makine öğrenmesi modelinin çıktısının oluşumunda her bir özneliliğin katkı payını hesaplayarak tahminin açıklanmasını sağlamaktadır. Böylece her bir değişkenin model tahmini üzerindeki etkisi yönü ve büyüklüğü açısından ayrı ayrı değerlendirilebilmektedir (Lundberg ve Lee, 2017).

SHAP'ta bir gözlem için model çıktısı, temel değer ile öznelilik katkılarının toplamı eşitlik 1'deki gibi ifade edilir.

$$f(x) = \phi_0 + \sum_{i=1}^M (\phi_i) \quad (1)$$

Eşitlik 1'de $f(x)$ model tahminini, ϕ_0 modelin temel (beklenen) çıktısını, ϕ_i ise i . özneliliğin tahmine katkısını göstermektedir. Her bir SHAP değeri aşağıdaki Shapley değeri formülü ile eşitlik 2'deki gibi hesaplanmaktadır.

$$\phi_i = \sum_{S \subseteq F \setminus \{i\}} \frac{|S|! (M - |S| - 1)!}{M!} [f_{S \cup \{i\}}(x_{S \cup \{i\}}) - f_S(x_S)] \quad (2)$$

Eşitlik 2'de S , i . özneliliği içermeyen öznelilik alt kümelerini; M , toplam öznelilik sayısını; $f_{S \cup \{i\}}(x_{S \cup \{i\}}) - f_S(x_S)$ ise ilgili özneliliğin modele eklenmesiyle tahminde oluşan marjinal katkıyı ifade etmektedir. Hesaplanan ortalama marjinal katkı, adil bir özellik katkısı sağlamayı amaçlamaktadır. Bu yaklaşım, aşağıdaki temel özellikleri sağlaması nedeniyle açıklanabilirlik açısından güçlü bir kuramsal temele dayanmaktadır (Givisis vd., 2025).

- **Yerel Doğruluk (Local Accuracy):** Dönüşüm fonksiyonu $h_x(x')$ ifadesi x ile özdeş kabul edildiğinde, açıklama modeli $g(x')$, orijinal model ile aynı çıktıyı üretmektedir. Başka bir ifadeyle, model çıktısı $f(x)$, temel değer ϕ_0 ile tüm özellik katkılarının toplamına eşittir. Bu durum, açıklama modelinin orijinal modeli yerel düzeyde doğru bir şekilde temsil ettiğini göstermektedir.
- **Eksiklik (Missingness):** Eğer $x'_i = 0$ ise, ilgili özelliğin katkısı $\phi_i = 0$ olmaktadır. Bu durum, modele dahil edilmeyen ya da etkisiz olan bir özelliğin tahmin üzerinde herhangi bir katkısının bulunmadığını ifade etmektedir.

Dolayısıyla, eksik ya da kullanılmayan özelliklere sıfır özellik katkı değeri verilmektedir.

- **Tutarlılık (Consistency):** Bir özelliğin model çıktısı üzerindeki etkisi arttığında, bu özelliğe ait Shapley değeri ya artmalı ya da en azından sabit kalmalıdır. Başka bir ifadeyle, bir özelliğin modele katkısı artmasına rağmen atfedilen değerin azalması söz konusu değildir. Bu özellik, model açıklamalarının tutarlı ve güvenilir olmasını sağlamaktadır.

SHAP yönteminin önemli özelliklerinden biri, modelin yalnızca tek bir gözlem için nasıl karar verdiğini değil, aynı zamanda genel olarak hangi özniteliklere daha fazla ağırlık verdiğini de gösterebilmesidir. Bu yönüyle SHAP, karmaşık makine öğrenmesi modellerinin açıklanmasında hem bireysel tahminlerin hem de modelin genel davranışının birlikte değerlendirilmesine olanak tanımaktadır. Ayrıca öznitelik katkılarının pozitif ya da negatif yönde incelenebilmesi, değişkenlerin model çıktısını artırıcı veya azaltıcı etkilerinin daha açık biçimde yorumlanmasını sağlamaktadır (Lundberg vd., 2020).

LIME (Local Interpretable Model-agnostic Explanations) ise karmaşık bir modelin belirli bir gözlem için verdiği tahmini, o gözlemin çevresinde kurulan daha basit ve yorumlanabilir bir yerel model aracılığıyla açıklamaktadır. LIME yaklaşımında, açıklanmak istenen gözlemin çevresinde yeni örnekler üretilmekte ve bu örnekler üzerinden karmaşık modelin ilgili bölgedeki davranışı taklit edilmektedir. Daha sonra bu davranış, doğrusal model gibi daha sade ve yorumlanabilir bir vekil model ile temsil edilmektedir. Böylece modelin tamamı değil, yalnızca ilgili gözleme yakın bölgedeki karar yapısı açıklanmış olmaktadır.

SHAP yöntemi hem yerel hem de küresel düzeyde yorumlanabilirlik sağlaması bakımından açıklanabilir yapay zekâ çalışmalarında yaygın olarak kullanılmaktadır (Lundberg ve Lee, 2017).

LIME'in temel amacı eşitlik 3'te gösterildiği gibi hata (loss) fonksiyonunu minimize etmektir.

$$\xi(x) = \arg \min_{g \in G} L(f, g, \pi_x) + \Omega(g) \quad (3)$$

Eşitlik 3'te f açıklanmak istenen karmaşık modeli, g yerel açıklama için kurulan

yorumlanabilir modeli, $L(f, g, \pi_x)$ yerel bölgede g modelinin f modeline ne kadar iyi yaklaştığını gösteren hata fonksiyonunu, π_x incelenen gözleme yakın örneklerle verilen ağırlıkları, $\Omega(g)$ ise açıklama modelinin karmaşıklığını ifade etmektedir. Bu yapı sayesinde LIME, modelin genel yapısını değil, yalnızca ilgili gözlem çevresindeki davranışını açıklamaya odaklanmaktadır (Ribeiro vd., 2016).

$\pi_x(z)$, z örneği ile x örneği arasındaki yakınlığı (proximity) ölçen bir fonksiyondur ve genellikle eşitlik 4'teki gibi üstel çekirdek (exponential kernel) fonksiyonu ile tanımlanmaktadır (Givisis vd., 2025).

$$\pi_x(z) = \exp\left(-\frac{D(x, z)^2}{\sigma^2}\right) \quad (4)$$

Eşitlik 4'te $D(x, z)$, iki örnek arasındaki uzaklığı ifade etmekte ve σ parametresi ise yerel komşuluk alanının genişliğini belirlemektedir.

Kayıp fonksiyonu $\mathcal{L}$ ise genellikle ağırlıklı kare hata (weighted squared loss) biçiminde tanımlanmaktadır:

$$\mathcal{L}(f, g, \pi_x) = \sum_{z \in Z} \pi_x(z) \cdot (f(z) - g(z))^2 \quad (5)$$

Eşitlik 5'te Z , x örneği etrafında oluşturulan bozunmuş (perturbed) örnekler kümesini ifade etmektedir.

Bu yöntemin genel işleyişi, incelenen bir veri noktasının çevresinde sentetik örnekler oluşturarak bu örnekler üzerinden yerel bir açıklanabilir model elde edilmesine dayanmaktadır. Bu süreç aşağıdaki adımlardan oluşmaktadır:

- **Veri Bozma (Data Perturbation):** İncelenen x örneği etrafında bozunmuş örneklerden oluşan bir Z veri kümesi oluşturulur.
- **Tahmin (Prediction):** Oluşturulan her bir $z \in Z$ örneği için kara kutu modelin $f(z)$ tahminleri elde edilir.
- **Ağırlıklandırma (Weighting):** Her bir örnek için $\pi_x(z)$ yakınlık değeri hesaplanarak, x 'e daha yakın örneklerin modele daha fazla katkı sağlaması sağlanır.

- **Yorumlanabilir Modelin Eğitimi (Interpretable Model Training):** Z veri kümesi ve $\pi_x(z)$ ağırlıkları kullanılarak, $\mathcal{L}(f, g, \pi_x)$ kayıp fonksiyonu en aza indirilecek şekilde ve $\Omega(g)$ karmaşıklık cezası dikkate alınarak yorumlanabilir model g eğitilir.
- **Açıklama (Explanation):** Elde edilen yorumlanabilir model g , $f(x)$ tahmininin açıklaması olarak sunulur. Doğrusal modellerde, model katsayıları her bir özneliğin tahmine yaptığı katkıyı doğrudan göstermektedir.

LIME'in modelden bağımsız bir yöntem olması, farklı makine öğrenmesi algoritmalarına uygulanabilmesini sağlamaktadır. Bu durum, özellikle karar mekanizması doğrudan izlenemeyen sınıflandırma ve regresyon modellerinde tekil tahminlerin anlaşılmasını kolaylaştırmaktadır. Ancak LIME daha çok yerel açıklamalar üretmekte; bu nedenle modelin genel davranışını bütünüyle ortaya koymaktan ziyade belirli bir gözlem özelinde yorum sağlamaktadır (Ribeiro vd., 2016).

SHAP ve LIME birlikte değerlendirildiğinde, bu iki yöntemin birbirini tamamlayan açıklama çerçeveleri sunduğu söylenebilir. SHAP, öznelilik katkılarını kuramsal olarak daha sistematik bir yapıda ortaya koyarken; LIME, tekil tahminleri yerel düzeyde açıklamada işlevsel bir yaklaşım sunmaktadır. Bu nedenle açıklanabilir yapay zekâ çalışmalarında her iki yöntemin birlikte kullanılması, model sonuçlarının hem genel hem de gözlem bazlı olarak daha anlaşılır biçimde değerlendirilmesine katkı sağlamaktadır. Bu durum, özellikle eğitim gibi kararların pedagojik olarak yorumlanmasının önemli olduğu alanlarda açıklanabilirliği daha güçlü bir zemine oturtmaktadır (Lundberg ve Lee, 2017).

3.6 Veri Toplama Süreci

Bu çalışmada veri toplama süreci, öğrencilerin akademik performanslarını etkileyen faktörlerin çok boyutlu biçimde analiz edilebilmesini sağlamak amacıyla birden fazla veri kaynağının bütünleşik kullanımı esas alınarak yürütülmüştür. Veri toplama aşaması; anket uygulamaları, Üniversite Bilgi Yönetim Sistemi (UBYS) kayıtları ve e-öğrenme platformlarından elde edilen verileri kapsamaktadır. Bu yaklaşım, akademik performansı yalnızca notlar üzerinden değil, öğrencilerin öğrenme davranışları, çalışma alışkanlıkları ve dijital etkileşimleriyle birlikte değerlendirmeye olanak tanımaktadır.

Araştırmada kullanılan anket formu, öğrencilerin demografik özellikleri, çalışma alışkanlıkları ve teknoloji kullanımı ile ilgili verileri toplamak amacıyla hazırlanmıştır. Anket uygulaması Microsoft Forms platformu üzerinden gerçekleştirilmiş olup, katılımcı yanıtları sistemin “Yanıtları Topla” bölümü aracılığıyla dijital ortamda elde edilmiştir. Anket soruları; yaş, cinsiyet, aile gelir düzeyi, ebeveyn eğitim durumu gibi demografik değişkenlerin yanı sıra ders çalışma süreleri, derslere katılım düzeyi, sınavlara hazırlık biçimi ve öğrenme stili gibi davranışsal göstergeleri içermektedir. Ayrıca öğrencilerin çevrimiçi öğrenme araçlarını kullanım sıklıkları, internet kullanım süreleri ve teknoloji kullanımının ders çalışmaya etkisine ilişkin algıları da anket kapsamında toplanmıştır.

Akademik performansa ilişkin veriler ise Üniversite Bilgi Yönetim Sistemi (UBYS) üzerinden temin edilmiştir. Öğrencilerin ders sınav notları ile çevrimiçi derslere ait etkileşim bilgileri, sistem üzerinden alınarak Excel formatında düzenlenmiş ve analiz sürecinde kullanılmak üzere veri setine dâhil edilmiştir. Bu veriler, anket yoluyla elde edilen bilgilerle eşleştirilerek her bir öğrenci için bütünlük bir veri yapısı oluşturulmuştur.

3.7 Veri Ön İşleme Süreci

Bu çalışmada veri ön işleme süreci, makine öğrenmesi modellerinin sağlıklı ve güvenilir sonuçlar üretebilmesi amacıyla sistematik bir biçimde yürütülmüştür. Veri işleme ve analiz aşamalarında Python programlama dili (sürüm 3.10.19) kullanılmış, kodlama işlemleri Spyder yazılım geliştirme ortamı üzerinden gerçekleştirilmiştir. Anket, üniversite bilgi yönetim sistemi ve e-öğrenme platformlarından elde edilen veriler, analiz sürecinin düzenli ve izlenebilir şekilde yürütülebilmesi amacıyla ders bazında ayrılmış ve her bir veri dosyası ilgili ders kodları kullanılarak adlandırılmıştır. Bu yaklaşım, veri setlerinin karışmasını önlemiş ve farklı derslere ait verilerin bağımsız olarak analiz edilmesine olanak sağlamıştır. Veri ön işleme sürecinde; eksik ve tutarsız kayıtların belirlenmesi, değişkenlerin uygun biçimde dönüştürülmesi ve modelleme aşamasına hazır hâle getirilmesi hedeflenmiştir.

3.7.1 Verilerin Okunması

Bu çalışmada elde edilen veriler, Python programlama dili içerisinde Pandas kütüphanesi kullanılarak analiz ortamına aktarılmıştır. Anket, UBYS ve e-öğrenme platformlarından

Excel formatında elde edilen veri dosyaları, Pandas aracılığıyla DataFrame yapısı olarak okunmuştur. DataFrame yapısı, verilerin satır–sütun düzeninde temsil edilmesine olanak sağlayarak veri üzerinde filtreleme, düzenleme ve dönüştürme işlemlerinin etkin biçimde gerçekleştirilmesini sağlamıştır.

3.7.2 Eksik Veri Analizi ve Düzenlenmesi

Veri ön işleme sürecinde, veri setinde yer alan eksik gözlemler öncelikle analiz edilmiştir. Veriler Pandas kütüphanesi ile okunmuştur. Ders devam çizelgeleri hazırlanırken imzası bulunan öğrenciler için ilgili haftalara “1” değeri atanmış, imzası bulunmayan öğrenciler ise boş bırakılmıştır. Bu durum, veri setinde eksik değerlerin oluşmasına yol açmış ve söz konusu eksik veriler DataFrame içerisinde “NaN (Not a Number)” olarak yer almıştır. Makine öğrenmesi modellerinin eksik verilerle çalışamaması ve hatalı sonuçlar üretmemesi nedeniyle, bu NaN değerler ön işleme aşamasında anlamlı bir biçimde düzenlenmiştir. Devamsızlık durumunu temsil eden bu eksik değerler, ilgili haftalarda öğrencinin derse katılmadığını ifade ettiğinden, eksik veriler mantıksal olarak “0” değeri ile doldurulmuştur. Böylece veri seti, modelleme sürecine uygun ve tutarlı bir yapıya dönüştürülmüştür.

Çevrimiçi ders videolarını izleyen öğrenciler için de ilgili haftalarda “1” değeri atanmış, izleme gerçekleştirilmeyen öğrenciler ise boş bırakılmıştır. Bu sebepten dolayı DataFrame’de oluşan “NaN” değerlere ilgili videoyu izlemeyen öğrencileri belirtmek amacıyla “0” değeri atanmıştır. Böylece video izleme değişkenleri tutarlı, eksiksiz ve modellemeye uygun bir yapıya dönüştürülmüştür.

Bununla birlikte, sınav notlarına ait veriler içerisinde yer alan ve tamamen “NaN” değerler içeren “t-score” sütunu, analiz ve modelleme açısından anlamlı katkı sağlamayacağından dolayı çıkartılmıştır.

UBYS sistemi sınava girmeyen öğrenciler için Vize ve Final notlarına “GR” değeri atamaktadır. Sınav notlarına ait veriler içerisinde bulunan Vize ve Final notları arasında “GR” değerlerinin bulunması modelde hatalı sonuçlara yol açabileceğinden dolayı “GR” değerleri yerine “0” yazılmıştır.

3.7.3 Sayısal Değişkenlerin Ölçeklendirilmesi

Veri ön işleme sürecinde, sayısal değişkenlerin modelleme aşamasında daha anlamlı ve karşılaştırılabilir hâle getirilebilmesi amacıyla ölçeklendirme işlemleri uygulanmıştır. Bu kapsamda öğrencilerin derslere devam durumları ile çevrimiçi kaynaklar (video) üzerinden gerçekleştirdikleri etkileşimler, farklı ölçüm birimlerinden kaynaklanabilecek tutarsızlıkların önüne geçmek amacıyla yüzdesel oran değerlerine dönüştürülmüştür. Bu dönüşüm sayesinde, öğrencilerin derslere katılım ve çevrimiçi öğrenme süreçlerine olan ilgileri daha standart bir ölçüt üzerinden değerlendirilmiştir.

Öğrencilerin haftalık ders devam durumları tespit edilirken, ilgili haftada imzası bulunan öğrencilere “1” bulunmayan öğrencilere “0” değeri imza kağıtları kontrol edilerek atanmıştır. “Devam Oranı” hesaplanırken “1” olan değerler toplanarak, toplam hafta sayısına bölünmüştür.

$$\text{Devam Oranı} = \frac{\sum_{i=1}^n x_i}{n} \quad x_i \rightarrow \text{Haftalık Devam değerleri (0 veya 1)} \quad (6)$$

Öğrencilerin çevrimiçi videoları izleme oranları tespit edilirken, ilgili videoları izleyen öğrencilere “1”, izlemeyen öğrencilere “0” değeri videoları izleyenler kontrol edilerek atanmıştır. “İzleme Oranı” hesaplanırken “1” olan değerler toplanarak, toplam video sayısına bölünmüştür.

$$\text{İzleme Oranı} = \frac{\sum_{i=1}^n x_i}{n} \quad x_i \rightarrow \text{Videoların İzlenme değerleri (0 veya 1)} \quad (7)$$

Ayrıca sınav notları içerisinde yer alan ve sınava girmeyen öğrencileri ifade eden “GR” kodlu kayıtlar, analiz sürecinde sayısal işlem yapılmasını engellediğinden, bu değerler 0 olarak yeniden düzenlenmiştir. Bu işlem, öğrencinin ilgili değerlendirmeden puan almadığını temsil edecek şekilde ele alınmış ve veri setinin bütünlüğü korunmuştur. Gerçekleştirilen ölçeklendirme ve dönüşüm işlemleri sonucunda, sayısal değişkenler makine öğrenmesi algoritmalarının kullanımına uygun bir yapıya kavuşturulmuştur.

3.7.4 Verilerin Birleştirilmesi

Veri ön işleme sürecinde, farklı kaynaklardan elde edilen verilerin bütüncül bir yapı altında toplanabilmesi amacıyla birleştirme işlemleri gerçekleştirilmiştir. Bu kapsamda, ayrı ayrı dersler hâlinde tutulan devam durumlarına ait imza kayıtları, sınav notları ve çevrimiçi öğrenme kaynaklarına ilişkin etkileşim verileri tek bir veri seti içerisinde bir araya getirilmiştir. Farklı kaynaklardan gelen bu veriler, analiz sürecinde anlamlı ilişkilerin kurulabilmesi amacıyla öğrenci numarası ve ders kodu temel alınarak eşleştirilmiştir.

Gerçekleştirilen bu eşleştirme işlemi sayesinde, her bir öğrenciye ait akademik, davranışsal ve dijital etkileşim bilgileri tekil bir kayıt altında toplanmış ve veri seti bütüncül bir yapıya kavuşturulmuştur. Bu yaklaşım, öğrencilerin akademik performanslarının çok boyutlu biçimde analiz edilmesine olanak sağlamış ve modelleme sürecinde kullanılacak verilerin tutarlılığını artırmıştır.

3.7.5 Öznitelik Seçimi

Veri ön işleme sürecinin bu aşamasında, modelleme sürecinde anlamlı katkı sağlamayacağı değerlendirilen değişkenler veri setinden çıkarılmıştır. Bu kapsamda ad, soyad ve öğrenci numarası gibi yalnızca tanımlayıcı nitelik taşıyan sütunlar, gizlilik ve modelleme açısından herhangi bir öngörü gücü bulunmadığı için analiz dışında bırakılmıştır. Böylece kişisel tanımlayıcı bilgiler veri setinden ayrılarak hem etik ilkelere uygunluk sağlanmış hem de modellerin yalnızca akademik performansla ilişkili değişkenler üzerinden eğitilmesi amaçlanmıştır. Kotsiantis vd. (2006) çalışmalarında; veri ön işleme sürecinde model performansına katkı sağlamayan değişkenlerin çıkarılmasının önemini vurgulakta ve özellikle kimlik bilgileri gibi tahmin gücü olmayan özelliklerin modelden çıkarılması gerektiği belirtmektedir. Ayrıca Aggarwal ve Yu (2008) çalışmalarında; makine öğrenmesinde kişisel tanımlayıcı bilgilerin çıkarılmasının hem etik hem de yasal açıdan gerekli olduğunu belirtmiştir.

Sınıflandırma ve tahmin modellerine girdi olarak, öğrencilerin akademik performanslarını etkileyebileceği öngörülen demografik, davranışsal ve dijital etkileşim temelli öznitelikler seçilmiştir. Bu kapsamda; yaş, cinsiyet, aile gelir düzeyi, anne ve baba eğitim durumu,

yaşanılan şehir, liseden mezun olunan okul türü ve barınma durumu gibi demografik özelliklerin yanı sıra ders çalışırken kullanılan kaynaklar, haftalık ortalama ders çalışma süresi, çalışma ortamının durumu, ders çalışırken verilen ara sıklığı, sınavlara hazırlanma süresi, sınav kaygısı düzeyi, derslere düzenli katılım durumu ve öğrenme stili gibi davranışsal göstergeler modele dâhil edilmiştir. Romero ve Ventura (2010) çalışmalarında; öğrenci performansını tahmin etmek için kullanılan değişkenleri sınıflandırmakta ve demografik (yaş, cinsiyet), akademik geçmiş ve davranışsal faktörlerin (çalışma alışkanlıkları, derslere katılım, öğrenme davranışları) model performansı üzerinde önemli etkisi olduğu vurgulanmaktadır. Ayrıca Cortez ve Silva (2008); aile eğitimi, sosyoekonomik durum gibi demografik özelliklerin öğrenci başarısının tahmininde önemli olduğunu belirtmiştir.

Buna ek olarak öğrencilerin dijital öğrenme alışkanlıklarını yansıtan günlük ortalama internet kullanımı, kullanılan çevrimiçi öğrenme araçları, sosyal medyanın ders çalışmaya etkisi, teknolojinin derslere katkı düzeyi, derslere devam zorunluğu (var/yok), devam oranı ve çevrimiçi içeriklere ait izlenme oranı gibi değişkenler de öznitelik seti içerisinde yer almıştır. Macfadyen vd. (2010) çalışmalarında; öğrencilerin LMS kullanım verileri (içerik görüntüleme, çevrimiçi aktiviteler) analiz edilmekte ve dijital öğrenme alışkanlıklarının akademik başarıyla güçlü şekilde ilişkili olduğu gösterilmektedir. Ayrıca Junco (2012b) çalışmasında; sosyal medya ve internet kullanımının (süre ve kullanım amacı) öğrenci katılımı ve akademik performans üzerinde anlamlı etkileri olduğunu bulmuştur.

Veri sızıntısı, tahmin edilmek istenen hedef değişkeni doğrudan ya da dolaylı biçimde içeren bilgilerin modelin girdileri arasında yer alması durumudur (Kaufman vd., 2012; Kapoor ve Narayanan, 2023). Veri sızıntısına yol açtığından dolayı “Geçme notu, vize ve final notları” öznitelikler arasına dahil edilmemiştir. Hedef değişken olarak kullanılan “Ders başarı durumu (başarılı/başarısız)”, doğrudan geçme notu, final notu ve vize notları temel alınarak oluşturulmuştur. Dolayısıyla model, öğrencinin başarısını tahmin etmeye çalışırken aslında sonucu belirleyen değişkenleri doğrudan görmektedir. Alanyazında bu durum, modelin gerçek anlamda öğrenme gerçekleştirmediği, yalnızca sonucu “ezberlediği” bir yapı olarak tanımlanmaktadır (Aggarwal, 2018).

Bu kapsamlı öznitelik yapısı sayesinde, öğrencilerin akademik performansları yalnızca not

temelli değil; bireysel özellikler, öğrenme davranışları ve dijital etkileşimler birlikte ele alınarak çok boyutlu bir modelleme süreci oluşturulmuştur.

3.7.6 Çoklu Öznitelik Ayırıştırması

Veri ön işleme sürecinde, ankette yer alan çoklu seçimli soruların makine öğrenmesi modellerinde etkin biçimde kullanılabilmesi amacıyla öznitelik ayırıştırma işlemi gerçekleştirilmiştir. Bu kapsamda, öğrencilerin birden fazla seçeneği işaretleyebildiği sorular doğrudan tek bir değişken olarak modele dâhil edilmemiş, her bir seçeneğin ayrı bir öznitelik olarak temsil edilmesi sağlanmıştır.

Ayırıştırma işlemi uygulanan anket soruları; “Ders çalışırken en çok hangi kaynakları kullanırsınız?” ve “Hangi çevrimiçi öğrenme araçlarını kullanıyorsunuz?” şeklinde belirlenmiştir. Bu sorular, her biri en fazla beş farklı seçenek içerecek biçimde yapılandırılmıştır. Çoklu seçimli yanıtların modelleme sürecinde anlamlı hâle getirilebilmesi amacıyla, her bir seçenek ayrı bir sütun olacak şekilde veri setine aktarılmıştır. Bu işlem sonucunda, her çoklu yanıt sorusu için en fazla beş ayrı değişken oluşturulmuş ve öğrencilerin ilgili seçeneği kullanıp kullanmadığı bilgisi ikili (var/yok) biçiminde temsil edilmiştir. Böylece çoklu seçimli anket yanıtları, makine öğrenmesi algoritmalarının doğrudan işleyebileceği bir yapıya dönüştürülmüş ve bilgi kaybı olmaksızın modelleme sürecine dâhil edilmiştir. Uygulanan bu ayırıştırma yaklaşımı, öğrencilerin öğrenme kaynakları ve çevrimiçi araç kullanım tercihlerinin akademik performans üzerindeki etkilerinin daha ayrıntılı biçimde analiz edilmesine olanak sağlamıştır.

3.7.7 Kategorik Verilerin Sayısallaştırılması

Veri ön işleme sürecinin bu aşamasında, modelleme sürecinde kullanılacak olan kategorik yapıya sahip değişkenler sayısal forma dönüştürülmüştür. Bu kapsamda anket verileri, derslere devam durumu ve akademik performans durumu gibi öznitelikler, makine öğrenmesi algoritmalarının işleyebileceği biçimde yeniden düzenlenmiştir. Sayısallaştırma işlemi için ordinal encoding, label encoding ve one-hot encoding gibi yöntemler kullanılmaktadır (Kuhn ve Johnson, 2013). Bu çalışmada; kategorik veriler Ordinal Encoding ve One-Hot Encoding yöntemleri ile kodlanarak makine öğrenmesi modelleri ile test

edilmiş, fakat modellerde daha başarılı olmasından dolayı One-Hot Encoding yöntemi tercih edilmiştir.

3.7.7.1 Ordinal Encoding Kullanılarak Kategorik Verilerin Sayısallaştırılması

Çalışmada kullanılan anket sorularının seçeneklerinin sırası önemlidir. Bu nedenle, kategoriler arasındaki sıralama ve derecelendirme yapısının korunabilmesi amacıyla önce ordinal encoding yöntemi test edilmiştir. Sayısal karşılıklar, anket seçeneklerinde yer alan sıralamaya uygun biçimde atanmış; böylece değişkenlerin anlamı bozulmadan modelleme sürecine aktarılması sağlanmıştır. Bu yaklaşım sayesinde, kategorik veriler hem yapısal bütünlüklerini korumuş hem de algoritmalar tarafından daha doğru biçimde işlenebilir hâle getirilmiştir.

Tablo 3.2: Cinsiyet kategorik verisinin kodlanması

Kategorik Veri	Kodlama
Kadın	0
Erkek	1

Tablo 3.2’de cinsiyet değişkenine ait kategorik verinin sayısallaştırma süreci gösterilmektedir. Bu kapsamda, anket verilerinde yer alan “Kadın” ve “Erkek” kategorileri, makine öğrenmesi algoritmalarının işleyebileceği biçimde sırasıyla 0 ve 1 değerleri ile kodlanmıştır.

Tablo 3.3: ‘Aile gelir düzeyiniz’ kategorik verisinin kodlanması

Kategorik Veri	Kodlama
‘Asgari ücretin altında (20.000 ve altı)',	0
'Asgari ücret düzeyinde (20.000- 30.000)',	1
'Orta gelir düzeyi (30.000- 50.000)',	2
'Orta-Üst gelir (80.000 ve üzeri)']]]	3

Tablo 3.3’te aile gelir düzeyine ilişkin kategorik verinin sayısallaştırma süreci gösterilmektedir. Bu kapsamda gelir düzeyleri, anket seçeneklerinde yer alan sıralamaya uygun biçimde kodlanmıştır. “Asgari ücretin altında (20.000 ve altı)” kategorisi 0, “Asgari ücret düzeyinde (20.000–30.000)” kategorisi 1, “Orta gelir düzeyi (30.000–50.000)” kategorisi 2 ve “Orta-Üst gelir (80.000 ve üzeri)” kategorisi ise 3 olarak sayısal değerlere dönüştürülmüştür. Bu kodlama işlemi sayesinde, gelir düzeyleri arasındaki dereceli ilişki

korunarak ilgili deęiřken makine öğrenmesi modellerinde kullanılabilir hâle getirilmiştir.

Tablo 3.4: 'Anne/Baba eğitim durumunuz' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'İlkokul',	0
'Ortaokul',	1
'Lise',	2
'Üniversite',	3
'Yüksek lisans / Doktora'	4

Anne/baba eğitim durumuna ilişkin kategorik verinin sayısallaştırma sürecinde aynı kodlama sistemi kullanılarak Tablo 3.4'te sunulmuştur. Bu kapsamda, anket sorusunda yer alan eğitim düzeyleri, “İlkokul” düzeyi 0, “Ortaokul” 1, “Lise” 2, “Üniversite” 3 ve “Yüksek lisans/Doktora” düzeyi ise 4 olarak sayısal değerlere dönüřtürülmüřtür. Böylece eğitim düzeyleri arasındaki dereceli yapı korunmuř ve ilgili deęiřken makine öğrenmesi modellerinde kullanılabilir hâle getirilmiştir.

Tablo 3.5: 'Yařadığınız yer' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
řehir merkezi',	0
'İlçe',	1
'Köy/Kasaba',	2
'Mahalle'	3

Tablo 3.5'te öğrencilerin yařadıkları yer bilgisine ait kategorik verinin sayısallaştırma süreci gösterilmektedir. Bu deęiřken, yerleşim yerleri arasındaki yapısal farklılıkları yansıtacak biçimde kodlanmıştır. Buna göre “řehir merkezi” kategorisi 0, “İlçe” 1, “Köy/Kasaba” 2 ve “Mahalle” kategorisi 3 olarak sayısal değerlere dönüřtürülmüřtür. Uygulanan bu kodlama işlemi sayesinde yerleşim yeri bilgisi, makine öğrenmesi modellerinin kullanabileceęi sayısal bir yapıya dönüřtürülmüřtür.

Tablo 3.6: 'Liseden mezun olduęunuz okul türü' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Anadolu lisesi',	0
'İmam hatip lisesi',	1
'Meslek lisesi',	2

Tablo 3.6: (devam ediyor)

'Fen lisesi',	3
'Düz lise'	4

Tablo 3.6’da öğrencilerin mezun oldukları lise türüne ait kategorik verinin sayısallaştırma süreci gösterilmektedir. Bu değişken, anket seçeneklerinde yer alan okul türlerine göre kodlanmıştır. Buna göre “Anadolu lisesi” 0, “İmam hatip lisesi” 1, “Meslek lisesi” 2, “Fen lisesi” 3 ve “Düz lise” kategorisi 4 olarak sayısal değerlere dönüştürülmüştür. Bu kodlama işlemi ile lise türü bilgisi, makine öğrenmesi modellerinde kullanılabilir sayısal bir forma dönüştürülerek veri setine dâhil edilmiştir.

Tablo 3.7: 'Öğrenim süreciniz boyunca barınma durumunuz nedir?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Kendi evimde',	0
'Arkadaşlarımla paylaşımlı evde',	1
'Devlet yurdunda',	2
'Ailemin yanında',	3
'Özel yurttta'	4

Tablo 3.7’de öğrencilerin öğrenim süreci boyunca barınma durumlarına ilişkin kategorik verinin sayısallaştırma işlemi gösterilmektedir. Buna göre “Kendi evimde” kategorisi 0, “Arkadaşlarımla paylaşımlı evde” 1, “Devlet yurdunda” 2, “Ailemin yanında” 3 ve “Özel yurttta” kategorisi 4 olarak sayısal değerlere dönüştürülmüştür. Uygulanan bu kodlama sayesinde barınma durumu değişkeni, makine öğrenmesi algoritmalarında kullanılabilir sayısal bir yapıya dönüştürülmüştür.

Tablo 3.8: 'Ders çalışırken en çok hangi kaynakları kullanırsınız?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Ders kitapları',	Var/Yok
'Ders notları',	Var/Yok
'E-ders Kaynakları',	Var/Yok
'İnternet kaynakları',	Var/Yok
'Yapay zekâ',	Var/Yok

Tablo 3.8’de öğrencilerin ders çalışırken kullandıkları kaynaklara ilişkin kategorik verinin sayısallaştırma işlemi sunulmaktadır. Çoklu seçenek içeren bu değişken, her bir kaynağın

modelleme sürecinde ayrı bir öz nitelik olarak temsil edilebilmesi amacıyla kodlanmıştır. Buna göre “Ders kitapları”, “Ders (çalışma) notları”, “E-ders kaynakları”, “İnternet kaynakları”, “Yapay zekâ” değerleri ayrı ayrı sütunlara bölünerek “Var / Yok” olarak sayısallaştırılmıştır. Bu kodlama işlemiyle, öğrencilerin öğrenme sürecinde tercih ettikleri kaynak türleri sayısal biçimde modele aktarılmıştır.

Tablo 3.9: 'Haftalık ortalama ders çalışma süreniz kaç saattir?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'5 saatten az',	0
'6-10 saat',	1
'11-15 saat',	2
'16-20 saat',	3
'21 saat ve üzeri'	4

Tablo 3.9’da öğrencilerin haftalık ortalama ders çalışma sürelerine ilişkin kategorik verinin sayısallaştırma işlemi gösterilmektedir. Buna göre “5 saatten az” seçeneği 0, “6–10 saat” 1, “11–15 saat” 2, “16–20 saat” 3 ve “21 saat ve üzeri” seçeneği 4 olarak sayısal değerlere dönüştürülmüştür. Bu dönüşüm sayesinde ders çalışma süresi değişkeni, öğrencilerin çalışma yoğunluğunu dereceli biçimde temsil edecek şekilde modele uygun hâle getirilmiştir.

Tablo 3.10: 'Ders çalışmak için en çok kullandığınız ortam/mekân nedir?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Ev',	0
'Kütüphane',	1
'Kafe',	2
'Yurt odası',	3
'Açık alan',	4
'Etüt'	5

Tablo 3.10’da öğrencilerin ders çalışmak için tercih ettikleri ortam ve mekân bilgisine ilişkin verilerin sayısallaştırılması gösterilmektedir. Bu değişken, öğrencilerin çalışma alışkanlıklarını ve öğrenme ortamı tercihlerini yansıtan önemli bir öz nitelik olarak modele dâhil edilmiştir. Anket seçenekleri “Ev”, “Kütüphane”, “Kafe”, “Yurt odası”, “Açık alan” ve “Etüt” olmak üzere altı farklı kategori altında toplanmış ve her bir seçenek 0’dan 5’e kadar artan sayısal değerlerle kodlanmıştır. Bu kodlama işlemi sayesinde öğrencilerin

çalışma ortamı tercihleri, makine öğrenmesi algoritmalarının işleyebileceği standart bir yapıya dönüştürülmüş ve modelleme sürecinde kullanılabilir hâle getirilmiştir.

Tablo 3.11: 'Genel olarak çalışma ortamın seni çalışmaya ne kadar teşvik ediyor?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Çok huzurlu ve motive edici',	0
'Fena değil, genelde odaklanabiliyorum',	1
'Çoğu zaman verimsiz hissediyorum',	2
'Hiç odaklanamıyorum'	3

Tablo 3.11’de öğrencilerin çalışma ortamlarının motivasyon üzerindeki etkisine ilişkin algılarını yansıtan değişkenin kodlama yapısı gösterilmektedir. Bu değişken, öğrencilerin çalışma verimliliğini etkileyen psikolojik ve çevresel faktörleri modelleyebilmek amacıyla sayısal forma dönüştürülmüştür. Anket sorusunda yer alan “Çok huzurlu ve motive edici”, “Fena değil, genelde odaklanabiliyorum” ve “Çoğu zaman verimsiz hissediyorum” seçenekleri, çalışma ortamı motivasyonu düzeyindeki azalışı temsil edecek biçimde sırasıyla 0, 1 ve 2 değerleriyle kodlanmıştır. Bu düzenleme, öğrencilerin çalışma ortamına ilişkin algılarının dereceli olarak modellenmesine ve akademik performans üzerindeki olası etkilerinin analiz edilmesine olanak sağlamıştır.

Tablo 3.12: 'Ders çalışırken ne sıklıkla ara verirsiniz?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Sık sık',	0
'Ara sıra',	1
'Nadiren',	2
'Hiç'	3

Tablo 3.12’de öğrencilerin ders çalışma sırasında ara verme sıklıklarını ifade eden değişkenin kodlama yapısı gösterilmektedir. Bu değişken, öğrencilerin çalışma disiplinlerini ve odaklanma alışkanlıklarını yansıtan önemli bir davranışsal gösterge olarak modele dâhil edilmiştir. Anket seçenekleri, ara verme sıklığındaki azalmayı temsil edecek biçimde “Sık sık” 0, “Ara sıra” 1, “Nadiren” 2 ve “Hiç” 3 olacak şekilde sayısal değerlere dönüştürülmüştür. Bu kodlama sayesinde, öğrencilerin çalışma sürecindeki mola alışkanlıkları dereceli bir yapıda ifade edilerek modelleme aşamasında kullanılabilir hâle getirilmiştir.

Tablo 3.13: 'Sınavlara ne kadar süre önceden hazırlanmaya başlıyorsunuz?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'1 hafta',	0
'2 hafta',	1
'1 ay',	2
'1 aydan fazla',	3
'Sınavdan önceki gün'	4

Tablo 3.13'te öğrencilerin sınavlara hazırlanmaya başlama zamanlarına ilişkin verilerin kodlama yapısı gösterilmektedir. Bu değişken, öğrencilerin sınavlara yönelik planlama ve hazırlık alışkanlıklarını yansıtan önemli bir davranışsal gösterge olarak modele dâhil edilmiştir. Anket seçenekleri, hazırlık süresindeki uzunluğu ifade edecek biçimde derecelendirilmiş ve “1 hafta” 0, “2 hafta” 1, “1 ay” 2, “1 aydan fazla” 3 ve “Sınavdan önceki gün” seçeneği 4 olarak sayısal değerlere dönüştürülmüştür. Bu düzenleme, öğrencilerin sınavlara yaklaşım biçimlerinin karşılaştırılabilir bir yapıda modellenmesine ve akademik performansla ilişkilerinin analiz edilmesine olanak sağlamıştır.

Tablo 3.14: 'Sınav kaygısı yaşıyor musunuz?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Evet, çok fazla',	0
'Evet, biraz',	1
'Hayır'	2

Tablo 3.14'te öğrencilerin sınav kaygısı düzeylerine ilişkin değişkenin kodlama yapısı gösterilmektedir. Bu değişken, öğrencilerin akademik performanslarını etkileyebilecek psikolojik bir faktör olarak modele dâhil edilmiştir. Anket sorusunda yer alan “Evet, çok fazla”, “Evet, biraz” ve “Hayır” seçenekleri, kaygı düzeyindeki azalmayı temsil edecek biçimde sırasıyla 0, 1 ve 2 değerleriyle sayısallaştırılmıştır.

Tablo 3.15: 'Derslere düzenli olarak katılıyor musunuz?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Her zaman',	0
'Çoğunlukla',	1
'Bazen',	2
'Nadiren',	3
'Hiç'	4

Tablo 3.15’te öğrencilerin derslere katılım sıklıklarına ilişkin verilerin kodlama yapısı sunulmaktadır. Bu değişken, öğrencilerin akademik süreçlere aktif katılım düzeylerini ve derslere devam alışkanlıklarını yansıtan önemli bir gösterge olarak değerlendirilmiştir. Anket seçenekleri, katılım sıklığındaki azalışı temsil edecek biçimde derecelendirilmiş ve “Her zaman” 0, “Çoğunlukla” 1, “Bazen” 2, “Nadiren” 3 ve “Hiç” seçeneği 4 olarak sayısal değerlere dönüştürülmüştür. Bu kodlama sayesinde derslere katılım durumu, öğrencilerin akademik performanslarıyla ilişkilendirilebilecek ölçülebilir bir öznel hâline getirilmiştir.

Tablo 3.16: ‘Öğrenme stilinizi nasıl tanımlarsınız?’ kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Görsel',	0
'İşitsel',	1
'Kinestetik (Dokunsal)',	2
'Okuma/yazma',	4
'Karma',	5
'Derste hem görsel hem işitsel öğrenimi tamamladıktan sonra sadece tekrar ederim',	6
'Görsel ve anlatarak',	7
'Görsel İşitsel mantığını anlama kavrama üzerine bir stilim var',	8
'Görsel/İşitsel mantığını öğrenme üzerine bir stilim vardır',	9
'Verileri alıp, örneklerle pekiştirerek öğrenebiliyorum. Tecrübe etmek diyebilirim.'	10

Tablo 3.16’da öğrencilerin kendi öğrenme stillerine ilişkin beyanlarını içeren değişkenin sayısallaştırma yapısı gösterilmektedir. Bu değişken, öğrencilerin bilgi edinme ve öğrenme süreçlerindeki bireysel farklılıklarını modelleyebilmek amacıyla veri setine dâhil edilmiştir. Anket sorusunda yer alan yanıtlar, temel öğrenme stillerini ve öğrenciler tarafından ifade edilen karma öğrenme biçimlerini kapsayacak şekilde geniş bir yelpazede toplanmıştır. “Görsel”, “İşitsel”, “Kinestetik (Dokunsal)” ve “Okuma/Yazma” gibi temel öğrenme stilleri sırasıyla 0–4 aralığında kodlanmış; “Karma” ve diğer birden fazla yöntemi içeren daha ayrıntılı öğrenme biçimleri ise 5–10 aralığında sayısal değerlere dönüştürülmüştür. Bu kodlama yaklaşımı, öğrencilerin öğrenme tercihlerini çeşitlilikleri korunarak temsil etmeyi ve bu tercihlerin akademik performans üzerindeki olası etkilerini analiz etmeyi amaçlamaktadır.

Tablo 3.17: 'Günlük ortalama internet kullanım süreniz kaç saattir?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'1 saatten az',	0
'1-3 saat',	1
'4-6 saat',	2
'7-9 saat',	3
'10 saat ve üzeri'	4

Tablo 3.17’de öğrencilerin günlük ortalama internet kullanım sürelerine ilişkin değişkenin kodlama yapısı sunulmaktadır. Bu değişken, öğrencilerin dijital ortamlarda geçirdikleri zamanın öğrenme süreçleri üzerindeki olası etkilerini analiz edebilmek amacıyla modele dâhil edilmiştir. Anket seçenekleri, internet kullanım süresindeki artışı temsil edecek biçimde derecelendirilmiş ve “1 saatten az” 0, “1–3 saat” 1, “4–6 saat” 2, “7–9 saat” 3 ve “10 saat ve üzeri” seçeneği 4 olarak sayısal değerlere dönüştürülmüştür. Bu düzenleme sayesinde internet kullanım alışkanlıkları, öğrencilerin akademik performanslarıyla ilişkilendirilebilecek ölçülebilir bir öznelilik hâline getirilmiştir.

Tablo 3.18: 'Derslerinizle ilgili olarak interneti ne sıklıkla kullanırsınız?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Günde birkaç kez',	0
'Haftada birkaç kez',	1
'Ayda birkaç kez',	2
'Hiç'	3

Tablo 3.18’de öğrencilerin dersleriyle ilgili olarak interneti kullanma sıklıklarına ilişkin değişkenin kodlama yapısı gösterilmektedir. Bu değişken, öğrencilerin akademik amaçlı dijital kaynaklardan yararlanma düzeylerini değerlendirmek amacıyla modele dâhil edilmiştir. Anket seçenekleri, kullanım sıklığındaki azalmayı yansıtacak biçimde sıralanmış ve “Günde birkaç kez” 0, “Haftada birkaç kez” 1, “Ayda birkaç kez” 2 ve “Hiç” seçeneği 3 olarak sayısal değerlere dönüştürülmüştür. Bu kodlama ile öğrencilerin interneti derslerle ilişkili kullanma alışkanlıkları, karşılaştırılabilir ve analiz edilebilir bir öznelilik hâline getirilmiştir.

Tablo 3.19: 'Hangi çevrimiçi öğrenme araçlarını kullanıyorsunuz?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Öğrenme yönetim sistemi (UBYS, e-ders)',	Var/Yok
'YouTube Ders Videoları',	Var/Yok
'Çevrimiçi Kurslar (Udemy, Coursera vb.)',	Var/Yok
'Yapay Zekâ (Chatgpt, Gemini, Deepseek vb.)',	Var/Yok
'Arama Motorları (Google, Yandex vb.)',	Var/Yok

Tablo 3.19’da öğrencilerin çevrimiçi öğrenme araçlarını kullanım tercihlerine ilişkin verilerin kodlama yapısı sunulmaktadır. Bu değişken, öğrencilerin dijital öğrenme ortamlarıyla etkileşim düzeylerini ve hangi kaynakları daha yoğun kullandıklarını belirleyebilmek amacıyla modele dâhil edilmiştir. Anket kapsamında yer alan seçenekler, her bir öğrenme aracını temsil edecek biçimde ayrı kategoriler hâlinde tanımlanmış ve “Öğrenme yönetim sistemi (UBYS, e-ders)”, “YouTube ders videoları” “Çevrimiçi kurslar”, “Yapay zekâ tabanlı araçlar”, “Arama motorları” değerleri sütunlara ayrılarak “Var / Yok” olarak kodlanmıştır. Bu düzenleme, öğrencilerin dijital öğrenme araçlarına yönelik kullanım alışkanlıklarının sayısal olarak modellenmesine olanak sağlamıştır.

Tablo 3.20: 'Çevrimiçi derslere katılımınız ne düzeyde?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Aktif',	0
'Pasif',	1
'Orta',	2
'Katılmıyorum'	3

Tablo 3.20’de öğrencilerin çevrimiçi derslere katılım düzeylerine ilişkin değişkenin sayısallaştırma yapısı gösterilmektedir. Bu değişken, öğrencilerin dijital öğrenme ortamlarındaki etkinlik ve katılım durumlarını modelleyebilmek amacıyla analiz sürecine dâhil edilmiştir. Anket seçenekleri, katılım düzeyindeki azalmayı yansıtacak biçimde sıralanmış ve “Aktif” 0, “Pasif” 1, “Orta” 2 ve “Katılmıyorum” seçeneği 3 olarak kodlanmıştır. Bu kodlama işlemiyle öğrencilerin çevrimiçi derslere yönelik katılım eğilimleri ölçülebilir bir yapıya dönüştürülmüş ve akademik performansla ilişkilerinin değerlendirilmesine olanak sağlanmıştır.

Tablo 3.21: 'Sosyal medya kullanımı ders çalışmanızı etkiliyor mu?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Evet',	0
'Hayır'	1

Tablo 3.21’de sosyal medya kullanımının ders çalışmaya etkisine ilişkin değişkenin kodlama yapısı gösterilmektedir. Bu değişken, öğrencilerin sosyal medya alışkanlıklarının akademik çalışma süreçleri üzerindeki olası etkilerini değerlendirmek amacıyla modele dâhil edilmiştir. Anket sorusuna verilen “Evet” yanıtı 0, “Hayır” yanıtı ise 1 olarak kodlanmıştır. Bu ikili kodlama sayesinde, sosyal medyanın ders çalışmaya etkisine yönelik algılar sayısal biçimde temsil edilerek analiz sürecinde kullanılabilir hâle getirilmiştir.

Tablo 3.22: 'Teknoloji kullanımının ders çalışmanıza yardımcı olduğunu düşünüyor musunuz?' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Çok Fazla',	0
'Fazla',	1
'Orta Düzeyde',	2
'Az',	3
'Çok Az'	4

Tablo 3.22’de öğrencilerin teknoloji kullanımının ders çalışmaya katkısına ilişkin algılarını içeren değişkenin kodlama yapısı sunulmaktadır. Bu değişken, öğrencilerin teknolojiyi öğrenme süreçlerinde ne ölçüde faydalı gördüklerini değerlendirmek amacıyla modele dâhil edilmiştir. Anket seçenekleri, algılanan katkı düzeyindeki azalışı temsil edecek biçimde sıralı olarak kodlanmış ve “Çok Fazla” 0, “Fazla” 1, “Orta Düzeyde” 2, “Az” 3 ve “Çok Az” seçeneği 4 olarak sayısal değerlere dönüştürülmüştür. Bu kodlama ile öğrencilerin teknolojiye yönelik tutumları dereceli bir öznitelik hâline getirilerek analiz sürecinde kullanılabilir biçimde yapılandırılmıştır.

Tablo 3.23: 'Devam_x (Devam Zorunluluğu)' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Yok',	0
'Var'	1

Tablo 3.23’te öğrencilerin derslere devam durumlarını ifade eden değişkenin kodlama yapısı

gösterilmektedir. Bu değişken, UBYS’de öğrencilerin derse fiilen katılım zorunluluğunun olup olmadığını gösteren temel bir akademik göstergedir. Öğrencinin derse devam zorunluluğunun olmadığını belirten “Yok” seçeneği 0, derse devam zorunluluğunun olduğunu belirten “Var” seçeneği ise 1 olacak şekilde ikili biçimde sayısallaştırılmıştır. Bu kodlama sayesinde derslere devam zorunluluğu, makine öğrenmesi algoritmalarında doğrudan kullanılabilir basit ve anlaşılır bir öznitelik hâline getirilmiştir.

Tablo 3.24: 'Başarı Durumu' kategorik verisinin kodlanması

Kategorik Veri	Kodlama
'Başarısız',	0
'Başarılı'	1

Tablo 3.24’te öğrencilerin akademik performans durumunu ifade eden hedef değişkenin kodlama yapısı gösterilmektedir. Bu değişken, sınıflandırma modellerinde bağımlı değişken olarak kullanılmış ve öğrencilerin derslerden geçme–kalma durumlarını temsil etmiştir. UBYS verilerinden elde edilen sonuçlara göre “Başarısız” durumu 0, “Başarılı” durumu ise 1 olarak kodlanmıştır. Bu ikili kodlama, makine öğrenmesi algoritmalarının öğrencilerin akademik performansını doğru biçimde sınıflandırabilmesi için gerekli olan sayısal hedef yapısını oluşturmuştur.

3.7.7.2 One-Hot Encoding Kullanılarak Kategorik Verilerin Sayısallaştırılması

Makine öğrenmesi algoritmalarının büyük çoğunluğu yalnızca sayısal verilerle çalıştığından, kategorik verilerin uygun biçimde kodlanması model performansı açısından kritik bir öneme sahiptir (Kuhn ve Johnson, 2013). Bu sebeple, kategorik değişkenlerin sayısallaştırılmasında ikinci yöntem olarak One-Hot Encoding yöntemi test edilmiştir. Kodlama işlemi sırasında her kategori için yeni bir sütun oluşturulmuş ve veri setinin boyutu buna bağlı olarak genişlemiştir.

One-Hot Encoding yöntemi, kategorik bir değişkende yer alan her bir sınıfın ayrı bir ikili (binary) sütun olarak temsil edilmesine dayanmaktadır. Bu yöntemde, ilgili kategoriye ait gözlem için “1” değeri atanırken, diğer tüm kategoriler “0” olarak kodlanmaktadır. Böylece kategoriler arasında yapay bir sıralama ilişkisi oluşturulmasının önüne geçilmekte ve modelin kategorik değişkenleri daha doğru şekilde öğrenmesi sağlanmaktadır (James vd.,

2013).

3.7.8 Yöntemsel Sorunlar: Dengesiz Veri, Metrikler ve Veri Sızıntısı

Eğitim veri setlerinde başarısız ya da riskli öğrenci sınıfı çoğu zaman azınlıkta kalabilmektedir. Bu durum dengesiz veri problemine yol açmakta ve doğruluk (accuracy) gibi tekil metriklerle dayalı değerlendirmeleri yanıltıcı hâle getirebilmektedir (He ve Garcia, 2009). Saito ve Rehmsmeier (2015), dengesiz veri koşullarında precision–recall eğrisinin ROC eğrisine kıyasla daha bilgilendirici olabileceğini göstermiştir. Dengesiz veriyle başa çıkmak için yeniden örnekleme teknikleri (ör. SMOTE, ADASYN) ve maliyet duyarlı öğrenme yaklaşımları alanyazında yaygın biçimde ele alınmaktadır (Chawla vd., 2002; Elkan, 2001).

Aghalarova ve Bozkurt Keser (2021), ortaokul öğrencilerinin performansını tahmin etmek amacıyla önerdikleri YSA yaklaşımında SMOTE kullanarak dengesiz sınıf dağılımını dengelemiş; özellik seçimi ve normalizasyon sonrası random search ile hiperparametre optimizasyonu gerçekleştirmiştir. Bu çalışma, eğitim verilerinde önışleme ve model ayarlarının tahmin başarımı üzerinde belirleyici olabileceğini göstermektedir.

He vd. (2008) tarafından önerilen orijinal ADASYN yöntemi, dengesiz veri setlerinde özellikle “öğrenilmesi zor” azınlık sınıfı örneklerine daha fazla sentetik veri üretmeyi hedeflemektedir. Bu yaklaşım, öğrenci performansı sınıflandırmasında (örneğin başarısız öğrenci sayısının az olduğu durumlarda) kritik öneme sahiptir. ADASYN, başarısız öğrencilerin bulunduğu karar sınırına yakın bölgelerde daha fazla veri üretmekle modelin bu öğrencileri daha iyi tanımasını sağlamaktadır. Böylece klasik yöntemlerin göz ardı edebileceği risk altındaki öğrenciler daha doğru şekilde sınıflandırılabilir.

Haixiang vd. (2017) çalışması, dengesiz veri problemlerine yönelik yöntemleri kapsamlı biçimde inceleyerek ADASYN’in eğitim verileri gibi gerçek dünya problemlerindeki etkisini ortaya koymaktadır. Çalışmada, ADASYN’in özellikle sınıflar arası dağılımın karmaşık olduğu durumlarda (örneğin öğrencilerin farklı başarı seviyelerine sahip olduğu veri setleri) daha etkili olduğu belirtilmiştir. Eğitim alanına uygulandığında bu yöntem, düşük başarı gösteren öğrenci grubunun daha iyi temsil edilmesini sağlayarak sınıflandırma

modellerinin genel doğruluğunu ve özellikle azınlık sınıfı üzerindeki duyarlılığını artırmaktadır.

Bu çalışmada; adaptif strateji kullanması, zor öğrenilen örnek sayısını artırması ve SMOTE'dan daha iyi bir performans göstermesinden dolayı veri dengelemede ADASYN tercih edilmiştir. SMOTE ve ADASYN veri artırma yöntemlerinin karşılaştırılması “Bulgular ve Tartışma” bölümünde ayrıntılı olarak verilmiştir.

Veri sızıntısı (data leakage) ise modelin gerçekte öngörmesi gereken bilgiyi, doğrudan ya da dolaylı biçimde özellik seti üzerinden ‘öğrenmesi’ durumudur ve performansı yapay olarak yükseltir. Kaufman vd. (2012), sızıntının formülasyonu, tespiti ve kaçınma yollarını ele alarak eğitim gibi uygulamalı alanlarda deney tasarımının titizlikle yürütülmesi gerektiğini vurgulamaktadır. Kapoor ve Narayanan (2023) ise sızıntı probleminin tekrarlanabilirlik krizine katkısını tartışarak, ML temelli çalışmaların raporlama ve değerlendirme standartlarının güçlendirilmesi gerektiğine işaret etmektedir.

4. BULGULAR VE TARTIŞMA

Bu bölümde; “Seçilen Öznitelikler ve Özellikleri, Öğrenci Performansının Sınıflandırılması, Öğrenci Performansının Sınıflandırılma Sonuçları, Öğrenci Performansı Tahmini, Açıklanabilir Yapay Zeka” başlıkları bulunmaktadır.

4.1 Seçilen Öznitelikler ve Özellikleri

Veriseti 528 kayıttan ve 35 öznitelikten oluşmaktadır. Bu özniteliklerin 3 tanesi sayısal geri kalan 32 tanesi kategorik veriden oluşmaktadır. Kategorik verileri önişleme aşamasında sayısala dönüştürülmüştür. Öznitelik Seçimi, çalışmanın “Materyal ve Yöntem” bölümü içerisinde ayrıntılı olarak, ayrı bir başlık altında sunulmuştur. Şekil 4.1’de önişleme işlemleri sonucunda ortaya çıkan öznitelik tablosu ve özniteliklerin özellikleri verilmiştir. Modeller geliştirilirken bu öznitelikler kullanılmıştır.

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 528 entries, 0 to 527
Data columns (total 35 columns):
#   Column                                                                 Non-Null Count  Dtype
---  -
0   Yaşınız                                                                528 non-null    int64
1   Cinsiyetiniz                                                          528 non-null    object
2   Aile gelir düzeyiniz                                                  528 non-null    object
3   Anne eğitim durumunuz                                                528 non-null    object
4   Baba eğitim durumunuz                                                528 non-null    object
5   Yaşadığınız yer                                                       528 non-null    object
6   Liseden mezun olduğunuz okul türü                                    528 non-null    object
7   Öğrenim süreciniz boyunca barınma durumunuz nedir?                 528 non-null    object
8   Ders kitapları                                                        528 non-null    object
9   Ders notları                                                          528 non-null    object
10  E-ders Kaynakları                                                     528 non-null    object
11  Yapay zeka                                                            528 non-null    object
12  İnternet kaynakları                                                  528 non-null    object
13  Haftalık ortalama ders çalışma süreniz kaç saattir?                 528 non-null    object
14  Ders çalışmak için en çok kullandığınız ortam/mekân nedir?          528 non-null    object
15  Genel olarak çalışma ortamın seni çalışmaya ne kadar teşvik ediyor? 528 non-null    object
16  Ders çalışırken ne sıklıkla ara verirsiniz?                         528 non-null    object
17  Sınavlara ne kadar süre önceden hazırlanmaya başlarsınız?          528 non-null    object
18  Sınav kaygısı yaşıyor musunuz?                                        528 non-null    object
19  Derslere düzenli olarak katılıyor musunuz?                          528 non-null    object
20  Öğrenme stilinizi nasıl tanımlarsınız?                              528 non-null    object
21  Günlük ortalama internet kullanım süreniz kaç saattir?             528 non-null    object
22  Derslerinizle ilgili olarak interneti ne sıklıkla kullanırsınız?     528 non-null    object
23  Arama Motorları (Google, Yandex vb.)                                 528 non-null    object
24  Yapay Zeka (Chatgpt, Gemini, Deepseek vb.)                           528 non-null    object
25  YouTube Ders Videoları                                               528 non-null    object
26  Çevrimiçi Kurslar (Udemy, Coursera vb.)                              528 non-null    object
27  Öğrenme yönetim sistemi (UBYS, e-ders)                              528 non-null    object
28  Çevrimiçi derslere katılımınız ne düzeyde?                          528 non-null    object
29  Sosyal medya kullanımı ders çalışmanızı etkiliyor mu?              528 non-null    object
30  Teknoloji kullanımının ders çalışmanıza yardımcı olduğunu düşünüyor musunuz? 528 non-null    object
31  Devam_x                                                              528 non-null    object
32  Devam Oranı                                                          528 non-null    float64
33  İzleme Oranı                                                         528 non-null    float64
34  Başarı Durumu                                                        528 non-null    object
dtypes: float64(2), int64(1), object(32)
memory usage: 144.5+ KB
```

Şekil 4.1: Seçilen öznitelikler ve özellikleri

4.2 Öğrenci Performansının Sınıflandırılması

Öğrencinin performans sınıflandırması 34 numaralı “Başarı Durumu” özniteliğine göre yapılmıştır. Başarı Durumu özniteliği “Başarılı / Başarısız” olarak değer almaktadır. Kategorik verilerin sayısal değerlere dönüştürülmesi aşamasında “Başarılı” değerine “1”, “Başarısız” değerine “0” atanmıştır.

```
-----  
* Orjinal Veriseti (Satır/Sütun)  
(528, 34)  
-----  
* Orjinal Veriseti Sınıf Sayıları  
Başarı Durumu  
1      382  
0      146  
Name: count, dtype: int64  
-----
```

Şekil 4.2: Veriseti ve sınıf sayıları

Şekil 4.2’de 528 öğrenciden 382’si “Başarılı” ve 146’sı “Başarısız” olarak gözükmektedir. Başarılı öğrenci sayısı, başarısız öğrenci sayısının 2.62 katıdır ve bu durum veri setinde dengesizliğe neden olmaktadır. Bu dengesizlik makine öğrenmesi problemlerinde özellikle gerçek veriler ile çalışırken çok sık karşılaşılan bir durumdur ve modellerin başarısını düşürmektedir (He ve Ma, 2013).

Dengesiz verilerde dengeyi sağlamak amacıyla veri azaltma / undersampling (Batista vd., 2004), veri artırma / oversampling (Chawla vd., 2002), ceza tabanlı yöntemler / cost-sensitive (Elkan, 2001), sınıf ağırlıklı yöntemler / class weighting (Sun vd., 2007) gibi yöntemler bulunmaktadır. Veri azaltma yöntemi veri kaybına, ceza tabanlı ve sınıf ağırlıklı yöntemler model üzerinde müdahaleye ve parametre ayarlamasında karmaşıklığa sebep olmasından dolayı veri artırma (oversampling) tercih edilmiştir.

Veri artırma yöntemleri, azınlıkta olan verilerin arttırılmasına dayanmaktadır (Chawla vd., 2002). Veri artırma yöntemleri Rastgele Veri Arttırma (Random Oversampling) ve Sentetik Veri Arttırma (Synthetic Oversampling) olmak üzere ikiye ayrılmaktadır (Chawla vd., 2002). Rastgele Veri artırma, azınlıkta olan sınıfın örneklerinin kopyalanmasıdır. Bu yöntem aşırı öğrenmeye (overfitting) sebep olabilmektedir. Sentetik Veri artırma yöntemleri arasında SMOTE (Synthetic Minority Over-Sampling Technique) ve ADASYN (Adaptive Synthetic

Sampling) yer almaktadır (Chawla vd., 2002; He vd., 2008). ADASYN yöntemi SMOTE yönteminin gelişmiş bir versiyonudur.

Tablo 4.1: SMOTE ve ADASYN karşılaştırması

Özellik	SMOTE	ADASYN
Örnek üretim mantığı	Her azınlık örneğine eşit sayıda	Zor örneklerle daha fazla
Dengesizlik giderme	Sabit strateji	Adaptif strateji
Sınıra yakın örnekler	Özel ilgi yok	Daha fazla örnek üretir
Performans	İyi	Genelde daha iyi

Tablo 4.1’de SMOTE ve ADASYN yönteminin karşılaştırılması verilmiştir (Fernandez vd., 2018). Tablodan anlaşılacağı gibi adaptif strateji kullanması, zor öğrenilen örnek sayısını artırması ve SMOTE’den daha iyi bir performans göstermesinden dolayı veri dengelemede ADASYN tercih edilmiştir. ADASYN veri artırma yöntemi verilerin eğitim/test olarak bölünmesinden önce yapılmıştır. Alanyazında, ADASYN ve benzeri yeniden örnekleme tekniklerinin eğitim/test ayrımı öncesinde tüm veri kümesine uygulanarak kullanıldığı çalışmalar bulunmaktadır. Örneğin; Dang vd. (2021) ve Reshi vd. (2021) yeniden örnekleme işlemini veri bölme aşamasından önce uygulayan deneysel kurgular raporlamıştır. Bununla birlikte, güncel metodolojik çalışmalar bu yaklaşımın veri sızıntısı ve iyimser performans tahmini riski taşıyabileceğini belirtmektedir. Bu nedenle söz konusu tercih, alanyazında örneği bulunan bir uygulama olmakla birlikte, sonuçlar bu sınırlılık dikkate alınarak yorumlanmıştır.

Öğrenci performansının sınıflandırılması için aşağıda verilen sınıflandırma yöntemleri denenmiştir;

- Karar Ağaçları (Decision Tree)
- Destek Vektör Makinesi (SVM: Support Vector Machine)
- Rastgele Orman (Random Forest)
- K- En Yakın Komşu (K-Nearest Neighbors)
- Çok Katmanlı Algılayıcı (MLP: Multi Layer Perceptron (Yapay Sinir Ağı))
- XGBoost

Bu sınıflandırma yöntemlerinin çeşitli parametreleri bulunmaktadır. Bu parametreler sınıflandırma modelinin performansını doğrudan etkilemektedir. Sınıflandırma yöntemlerinin en iyi performans verdiği değeri bulabilmek amacıyla, Python Scikit-learn kütüphanesi içerisinde yer alan “RandomizerSearchCV” metodu Stratified K-Fold Cross Validation ile kullanılmıştır. Stratified K-Fold Cross Validation belirtilen n değerine göre veri setini n parçaya bölerek ayrı ayrı denemeler yapmaktadır. Bu çalışmada n değeri 10 olarak alınmıştır. RandomizerSearchCV metodu belirtilen değer aralığındaki parametreleri veri setini rastgele 10 parçaya bölerek belirli iterasyon kadar denemektedir. Bu çalışmada iterasyon sayısı 300 olarak belirlenmiştir. İterasyon sayısını artırmak model parametrelerinin daha sağlam belirlenmesini sağlamakta, buna karşın tespit süresini uzatmaktadır.

Modeller, parametreler ve parametre aralıkları aşağıda verildiği gibi seçilmiştir;

- **Decision Tree**
 - **Maximum Dept (Maksimum Derinlik):** 3, 5, 8, 12, 16, 24
 - **Minimum Samples Split:** 2 / 50 aralığı
 - **Minimum Samples Leaf:** 1 / 30 aralığı
 - **Criterion:** Gini, Entropy, Logistic Loss
- **SVM (Support Vector Machine)**
 - **Kernel:** RBF (Radial Basis Function)
 - **C:** 1e-2 / 1e3 aralığı
 - **Gamma:** 1e-4 / 1e0 aralığı
- **Random Forest**
 - **N – Estimators:** 200 / 1200 aralığı
 - **Maximum Depth:** 6, 10, 14, 18, 24, 30
 - **Minimum Sample Split:** 2 / 40 aralığı
 - **Minimum Sample Leaf:** 1 / 20 aralığı
 - **Maximum Features:** Sqrt, Log2
 - **Bootstrap:** True / False
- **KNN (K-Nearest Neighbors)**
 - **N Neighbors:** 3 / 61 aralığı
 - **Weights:** Uniform, Distance
 - **P:** Manhattan, Euclidean
 - **Leaf Size:** 10 / 60 aralığı

- **MLP (Multi Layer Perceptron)**
 - **Maximum İteration :** 1200
 - **Early Stopping:** true (20 iteration hiç değişiklik olmazsa)
 - **Hidden Layer Size:** (64) / (128) / (128,64) / (256,128) / (256, 128, 64)
 - **Activation Function:** Relu, Tanh
 - **Alpha:** 1e-6 / 1e-2 aralığı
 - **Learning Rate:** 1e-4 / 5e-2 aralığı
 - **Solver (Optimizer):** Adam
 - **Batch Size:** 32, 64,128, 256
- **XGBoost**
 - **Eval Metric:** Logistic Loss
 - **N Estimators:** 300 / 1500 aralığı
 - **Maximum Depth (Maksimum Derinlik):** 3 / 10 aralığı
 - **Learning Rate:** 1e-3 / 3e-1 aralığı
 - **Subsampling:** 0.6 / 1.0 aralığı
 - **Col Sampling by Tree:** 0.6 / 1.0 aralığı
 - **Minimum Child Weight:** 1 / 12 aralığı
 - **Gamma:** 0.0 / 5.0 aralığı
 - **Reg Alpha (L1 Regularization):** 0.0 / 1.0 aralığı
 - **Reg Lambda (L2 Regularization):** 0.5 / 2.5 aralığı

4.2.1 Karar Ağacı Modeli Parametreleri

Maximum Depth (Maksimum Derinlik): Ağacın kökten yaprağa kadar izin verilen maksimum seviyesini ifade etmektedir. Derinlik arttıkça model daha karmaşık hale gelmektedir. Çok büyük derinlikler aşırı öğrenmeye (overfitting) yol açabilmektedir (Breiman vd., 1984).

Minimum Samples Split (Minimum Örnek Bölünmesi): Ağacın bir düğümünün bölünebilmesi için gerekli minimum örnek sayısını ifade etmektedir. Değer küçük seçildiğinde ağaç daha fazla derinleşmekte, büyük seçildiğinde ise model daha genelleyici olmaktadır (Hastie vd., 2009).

Minimum Samples Leaf (Minimum Örnek Yaprığı): Ağacın her yaprak düğümünde

bulunması zorunlu minimum örnek sayısını belirlemektedir. Çok küçük seçilmesi, tek örnekli yaprakların oluşmasına ve aşırı uyuma yol açmaktadır. Büyük seçilmesi ise; ağacın daha kararlı ve genelleyici olmasını sağlamaktadır (Breiman vd., 1984).

Criterion (Bölme Kriteri): Bu parametre, düğüm bölme işleminin hangi ölçüte göre yapılacağını belirlemektedir. Gini Index; Sınıf saflığını ölçmektedir. Hesaplama açısından hızlıdır (Breiman, 1984). Entropy; Bilgi kazancı temellidir, daha hassas ayırım yapmaktadır (Quinlan, 1986). Logistic Loss; Olasılıksal tahminlere dayalı modern bir ölçüttür (Friedman, 2001).

4.2.2 Destek Vektör Makinesi Modeli Parametreleri

Kernel: RBF (Radial Basis Function), en yaygın kullanılan çekirdek fonksiyonudur. Doğrusal olmayan karmaşık sınırları modelleyebilmektedir (Schölkopf ve Smola, 2002).

C Parametresi: Ceza parametresidir. Hatalı sınıflandırmalara verilen cezayı kontrol etmektedir. Küçük değerler; daha yumuşak karar sınırları olmakta ve genelleme artmaktadır. Büyük değerler; daha sıkı sınırlar olmakta ve aşırı öğrenme riski bulunmaktadır (Hsu vd., 2003).

Gamma: RBF çekirdeği kullanılırken, karar sınırının esnekliğini kontrol etmektedir. Küçük değerler seçildiğinde daha yumuşak sınırlar oluşmakta ve underfitting riski bulunmaktadır. Büyük değerlerde ise; daha karmaşık sınırlar oluşur ve overfitting riski bulunmaktadır (Hsu vd., 2003).

4.2.3 Rastgele Orman Modeli Parametreleri

n_estimators: Ormandaki ağaç sayısını (tahmin edici) ifade etmektedir. Ağaç sayısı arttıkça performans genelde artmakta, ancak hesaplama maliyeti yükselmektedir (Breiman, 2001).

Maximum Depth: Her bir ağacın maksimum derinliğini sınırlamaktadır. Aşırı öğrenmeyi kontrol etmek için kritik bir parametredir (Liaw ve Wiener, 2002).

Minimum Sample Split (Minimum Örnek Bölünmesi): Karar ağaçlarındaki aynı mantığa sahiptir.

Minimum Samples Leaf (Minimum Örnek Yaprığı): Bu parametre de karar ağaçlarına benzer mantıktadır.

Maximum Features (Maksimum Öznitelik): Her ağaç bölünmesinde kullanılacak maksimum öznitelik sayısını belirlemektedir. Sqrt; öznitelik sayısının karekökü, log2: öznitelik sayısının log2'si olmaktadır (Breiman, 2001).

Bootstrap: Her ağacın eğitimi sırasında örneklerin yeniden örneklenip örneklenmeyeceğini belirlemektedir. Bu değer “true” olarak seçildiğinde bootstrap sampling kullanarak, modelin daha dayanıklı olmasına yardımcı olmaktadır (James vd., 2013).

4.2.4 K- En Yakın Komşu Modeli Parametreleri

N Neighbors: Değer çevresindeki kaç tane komşuya bakılacağını ifade etmektedir. Küçük değerler, daha hassas fakat gürültüye duyarlı, daha büyük değerler; daha genelleyici ve kararlı olmaktadır (Altman, 1992).

Weights: Uniform seçeneğinde tüm komşular eşit ağırlıklı, distance seçeneğinde ise; yakın komşular daha etkili ağırlıklı olarak hesaba katılmaktadır (Altman, 1992).

P Parametresi: Komşu yakınlığını hesaplarken kullanılan mesafe ölçüm yöntemidir. P=1 değeri Manhattan uzaklık ölçümünü, P=2 değeri Euclidean uzaklık ölçümünü ifade etmektedir (Cunningham ve Delany, 2007).

Leaf Size (Yaprak Boyutu): K-NN algoritmasında kullanılan KD-Tree veya Ball-Tree yapılarındaki yaprak düğüm boyutunu ifade etmektedir. Doğrudan sınıflandırma sonucuna etkisi olmamakla birlikte, hesaplama hızını etkilemektedir (Bentley, 1975).

4.2.5 Çok Katmanlı Perceptron (MLP) Modeli Parametreleri

Maximum Iteration: Eğitim döngüsü sayısıdır.

Early Stopping: 20 iterasyon boyunca modelin performansında iyileşme olmazsa eğitimi bitirmektedir. Bu parametre aşırı öğrenmeyi önlemektedir (Prechelt, 1998).

Hidden Layer Size: Ağda bulunan gizli katmanların sayısıdır. Katman ve katmanlarda bulunan nöron sayıları arttıkça ağın öğrenme kapasitesi artmaktadır (Goodfellow vd., 2016).

Activation Function: Ağın eğitimi sırasında hangi nöronların etkinleştirileceğini belirleyen fonksiyondur (Nair ve Hinton, 2010).

Alpha (Regularization): Ağın aşırı öğrenmesini engellemek için L2 ceza terimidir (Bengio, 2012).

Learning Rate: Ağın ağırlıkları güncellenirken güncelleme miktarını temsil etmektedir. Büyük değerler daha hızlı öğrenmeyi sağlamakta, fakat ezberlemeye sebep olabilmektedir. Küçük değerler ise; ağın maksimum performansına daha uzun sürede ulaşmasına sebep olmakta, fakat daha hassas öğrenmeyi sağlamaktadır (Bengio, 2012).

Batch Size: Ağın her öğrenme iterasyonunda kaç tane veri örneğini aynı anda işleyeceğini ifade etmektedir. Modelin kararlılığını etkilemektedir (Bengio, 2012).

4.2.6 XGBoost Modeli Parametreleri

n_estimators: Modeldeki ağaç (tahmin edici) sayısını belirlemektedir (Chen ve Guestrin, 2016).

Maximum Depth: Tahmin edici ağaçların maksimum derinliğini belirlemektedir (Chen ve Guestrin, 2016).

Learnin Rate: Her ağacın modele katkısını ölçeklemektedir (Chen ve Guestrin, 2016).

Subsampling: MLP'deki Batch Size parametresine benzer olarak her eğitim iterasyonunda kullanılacak veri oranını belirtmektedir. Bu değer aşırı öğrenmeyi azaltmaktadır (Chen ve Guestrin, 2016).

Colsample by Tree: Her ağaç içerisinde kullanılarak özellik oranını belirlemektedir (Chen ve Guestrin, 2016).

Min Child Weight: Ağaçların yaprak düğümlerinin bölünme eşiğini belirlemektedir (Chen ve Guestrin, 2016).

Gamma: Ağaçlarda bölme yapabilmek için gerekli minimum kazancı belirlemektedir (Chen ve Guestrin, 2016).

Reg Alpha / Reg Lambda: L1 ve L2 regularization parametreleridir. Aşırı öğrenmeyi önlemek için kullanılmaktadır (Chen ve Guestrin, 2016).

4.3 Öğrenci Performansının Sınıflandırılma Sonuçları

Bu aşamada, ele alınan her bir sınıflandırma modeli için farklı veri ön işleme senaryoları sistematik biçimde uygulanmış ve elde edilen sonuçlar ayrı ayrı raporlanmıştır. İlk olarak, kategorik değişkenler ordinal kodlama yöntemiyle dönüştürülmüş; bu dönüşüm hem ADASYN veri artırma (oversampling) uygulanmadan (Sınıflandırma Sonuçları (Deneme 1)) hem de uygulanarak (Sınıflandırma Sonuçları (Deneme 2)) test edilmiştir. Böylece sınıf dengesizliğinin ve ordinal kodlamanın model performansı üzerindeki etkisi karşılaştırmalı olarak incelenmiştir.

İkinci aşamada, kategorik değişkenler one-hot kodlama yöntemiyle sayısallaştırılmış ve bu yapı altında ADASYN veri artırma tekniği uygulanarak sınıflandırma deneyleri gerçekleştirilmiştir (Sınıflandırma Sonuçları (Deneme 3)). Son olarak, kategorik veriler one-hot kodlama, sayısal veriler de 0–1 aralığında normalize edilmiş ve bu ölçeklendirilmiş veri seti (Sınıflandırma Sonuçları (Deneme 4)) ile modeller yeniden eğitilerek performans çıktıları değerlendirilmiştir.

Bu çoklu deneysel tasarım sayesinde, kodlama yöntemi, veri dengeleme yaklaşımı ve ölçeklendirme işlemlerinin sınıflandırma performansı üzerindeki etkileri karşılaştırmalı ve bütüncül bir çerçevede analiz edilmiştir.

4.3.1 Sınıflandırma Sonuçları (Deneme 1)

Veri setinde bulunan kategorik veriler ordinal kodlama yöntemi ile sayısallaştırılmış, %70 eğitim %30 test verisi (Tablo 4.2) ve %60 eğitim %40 test verisine (Tablo 4.3) ayrılarak, 300 iterasyonlu RandomSearchCV içerisinde Stratified 10-Fold Crossover kullanılarak test edilmiştir ve sınıflandırma modellerinin en iyi parametreler ile elde ettikleri en iyi sonuçlar verilmiştir.

Tablo 4.2: %70 Eğitim %30 Test verileri ve sınıf sayıları (Deneme 1)

Veriseti Türü	Satır Sayısı	Sütun Sayısı	Başarı (1)	Başarısız (0)
Orijinal Veriseti	528	34	382	146
Eğitim Veriseti	369	34	267	102
Test Veriseti	159	34	115	44

Tablo 4.2’de orijinal veri seti ile eğitim ve test veri kümelerine ait temel dağılım bilgileri sunulmaktadır. Orijinal veri seti 528 gözlem ve 34 öznitelikten oluşmakta olup, bu gözlemlerin 382’si başarılı, 146’sı ise başarısız olarak etiketlenmiştir. Veri seti, modelleme sürecinde kullanılmak üzere yaklaşık %70 eğitim ve %30 test olacak şekilde iki alt kümeye ayrılmıştır. Bu doğrultuda eğitim veri seti 369 gözlemden oluşurken, test veri seti 159 gözlem içermektedir.

Sınıf dağılımları incelendiğinde hem eğitim hem de test veri kümelerinde başarılı ve başarısız sınıflarının oransal dağılımlarının korunduğu görülmektedir. Eğitim veri setinde 267 başarılı ve 102 başarısız gözlem bulunurken, test veri setinde 115 başarılı ve 44 başarısız gözlem yer almaktadır. Bu durum, veri bölme işlemi sırasında sınıflar arasındaki dengenin gözetildiğini ve modellerin eğitim ve test aşamalarında benzer dağılımlara sahip verilerle

değerlendirildiğini göstermektedir. Böylece elde edilen performans sonuçlarının daha güvenilir ve genellenebilir olması amaçlanmıştır.

Tablo 4.3: %60 Eğitim %40 Test verileri ve sınıf sayıları (Deneme 1)

Veriseti Türü	Satır Sayısı	Sütun Sayısı	Başarı (1)	Başarısız (0)
Orijinal Veriseti	528	34	382	146
Eğitim Veriseti	316	34	229	87
Test Veriseti	212	34	153	59

Tablo 4.3'te orijinal veri seti ile eğitim ve test veri kümelerine ait dağılım bilgileri sunulmaktadır. Orijinal veri seti 528 gözlem ve 34 öznitelikten oluşmakta olup, bu gözlemlerin 382'si başarılı, 146'sı ise başarısız olarak etiketlenmiştir. Veri seti, modelleme sürecinde kullanılmak üzere yaklaşık %60 eğitim ve %40 test oranında iki alt kümeye ayrılmıştır. Bu doğrultuda eğitim veri seti 316 gözlemden, test veri seti ise 212 gözlemden oluşmaktadır.

Sınıf dağılımları incelendiğinde hem eğitim hem de test veri kümelerinde başarılı ve başarısız sınıflarının oransal yapısının büyük ölçüde korunduğu görülmektedir. Eğitim veri setinde 229 başarılı ve 87 başarısız gözlem yer alırken, test veri setinde 153 başarılı ve 59 başarısız gözlem bulunmaktadır. Bu durum, veri bölme işleminin sınıf dağılımını bozmayacak şekilde gerçekleştirildiğini ve modellerin farklı veri kümeleri üzerinde dengeli biçimde değerlendirilmesine olanak sağladığını göstermektedir. Böylece elde edilen performans ölçütlerinin, modelin genellenebilirliğini daha sağlıklı biçimde yansıtması amaçlanmıştır.

Tablo 4.4: %70 Eğitim %30 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 1)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
KNN	0.764339	0.857805	0.588468	0.681430	0.858396

Tablo 4.4: (devam ediyor)

MLP	0.734384	0.844573	0.521785	0.510254	0.753146
SVM (RBF)	0.753529	0.853141	0.563081	0.681961	0.853319
XGBoost	0.756006	0.851107	0.588157	0.660256	0.843349
Random Forest	0.766967	0.856530	0.608754	0.712444	0.866781
Decision Tree	0.766892	0.854433	0.623372	0.647747	0.817423

Tablo 4.4.'te farklı makine öğrenmesi algoritmalarının çapraz doğrulama sonuçları; doğruluk (CV Accuracy), F1 skoru, dengeli doğruluk (CV Balanced Accuracy), ROC-AUC ve PR-AUC ölçütleri üzerinden karşılaştırmalı olarak sunulmaktadır. Kullanılan performans ölçütleri, özellikle sınıf dağılımının tam dengeli olmadığı problemlerde modellerin gerçek başarılarını daha sağlıklı biçimde değerlendirmek amacıyla tercih edilmiştir

Sonuçlar incelendiğinde, Random Forest ve Decision Tree algoritmalarının genel olarak en yüksek doğruluk değerlerine sahip olduğu görülmektedir. Random Forest modeli, CV Accuracy (0.7669), CV Balanced Accuracy (0.6088) ve özellikle ROC-AUC (0.7124) değerleri ile diğer modellere kıyasla daha dengeli ve ayırt edici bir performans sergilemiştir. Ensemble yapısı sayesinde Random Forest algoritması, tekil karar ağaçlarına kıyasla daha düşük varyans ve daha yüksek genellenebilirlik sunmaktadır.

KNN ve SVM (RBF) modelleri, özellikle F1 skoru ve PR-AUC açısından güçlü sonuçlar üretmiştir. KNN algoritmasının CV F1 Score (0.8578) ve CV PR-AUC (0.8584) değerleri, pozitif sınıfın doğru tahmin edilmesi açısından başarılı bir performansa işaret etmektedir. Benzer şekilde SVM (RBF) modeli de yüksek F1 ve PR-AUC değerleriyle, özellikle sınıflar arasındaki karar sınırlarının doğrusal olmadığı durumlarda etkili sonuçlar sunmuştur.

XGBoost modeli, genel doğruluk ve F1 skorları bakımından rekabetçi bir performans sergilemekle birlikte, ROC-AUC değerinin Random Forest ve KNN'e kıyasla daha düşük kaldığı görülmektedir. Buna rağmen XGBoost'un dengeli doğruluk değerinin görece yüksek olması, modelin her iki sınıfı da dikkate alan bir öğrenme süreci yürüttüğünü göstermektedir. MLP (Yapay Sinir Ağları) modeli ise diğer modellere kıyasla daha düşük Balanced Accuracy ve ROC-AUC değerleri üretmiştir. Bu durum, veri setinin boyutu ve yapısının, derin öğrenme temelli yaklaşımlar için sınırlayıcı olabileceğini göstermektedir. Alanyazında

de küçük ve orta ölçekli veri setlerinde geleneksel makine öğrenmesi algoritmalarının, yapay sinir ağlarına kıyasla daha istikrarlı sonuçlar üretebildiği belirtilmektedir (Goodfellow vd., 2016).

Genel olarak değerlendirildiğinde, Random Forest modeli dengeli doğruluk ve ayırt edicilik ölçütlerinde öne çıkarken; KNN ve SVM (RBF) modelleri pozitif sınıf performansını yansıtan F1 ve PR-AUC ölçütlerinde güçlü sonuçlar sunmuştur. Bu bulgular, tek bir performans ölçütüne bağlı kalınmaksızın çoklu değerlendirme metrikleri kullanmanın model karşılaştırmaları açısından kritik olduğunu göstermektedir.

Tablo 4.5: %70 Eğitim %30 Test verileri ile tüm modellerin test performans sonuçları (Deneme 1)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
KNN	0.748428	0.846154	0.580534	0.695949	0.832425
MLP	0.723270	0.839416	0.500000	0.524111	0.759003
SVM (RBF)	0.729560	0.838951	0.532411	0.669368	0.826589
XGBoost	0.735849	0.838462	0.564822	0.690711	0.853915
Random Forest	0.735849	0.834646	0.585870	0.727273	0.857822
Decision Tree	0.729560	0.832685	0.567490	0.568379	0.763582

Tablo 4.5'te farklı makine öğrenmesi algoritmalarının test veri kümesi üzerindeki performansları doğruluk, F1 skoru, dengeli doğruluk, ROC-AUC ve PR-AUC ölçütleri üzerinden karşılaştırmalı olarak sunulmaktadır. Test veri kümesi sonuçları, modellerin daha önce görülmemiş veriler üzerindeki genellenebilirliğini değerlendirmek açısından önem taşımaktadır.

Sonuçlar incelendiğinde, KNN algoritmasının doğruluk ve F1 skoru bakımından diğer modellere kıyasla daha yüksek değerler elde ettiği görülmektedir. Bu durum, KNN modelinin özellikle başarılı öğrencilerin doğru sınıflandırılmasında etkili bir performans sergilediğini göstermektedir. Random Forest modeli ise dengeli doğruluk ve ROC-AUC değerleri açısından öne çıkmakta olup, her iki sınıfı ayırt etme yeteneğinin daha dengeli olduğunu göstermektedir. XGBoost modelinin PR-AUC değerinin görece yüksek olması,

modelin pozitif sınıfa yönelik tahmin başarısının güçlü olduğuna işaret etmektedir.

SVM (RBF) ve Decision Tree modelleri, genel olarak orta düzeyde performans değerleri üretmiş ve diğer modellerle karşılaştırıldığında daha sınırlı bir ayırt edicilik sergilemiştir. MLP (Yapay Sinir Ağları) modeli ise özellikle dengeli doğruluk ve ROC-AUC ölçütlerinde daha düşük sonuçlar elde etmiş, bu durum veri setinin yapısının bu model için yeterince uygun olmayabileceğini düşündürmektedir.

Genel değerlendirme yapıldığında, test sonuçları doğrultusunda KNN modeli doğruluk ve F1 skoru açısından, Random Forest modeli dengeli doğruluk ve ayırt edicilik açısından, XGBoost modeli ise pozitif sınıf performansını yansıtan PR-AUC ölçütü açısından öne çıkmaktadır. Bu bulgular, akademik performans tahmini gibi problemlerde model başarımının tek bir ölçütle değil, farklı performans göstergeleri birlikte değerlendirilerek ele alınmasının gerekliliğini ortaya koymaktadır.

Tablo 4.6: %60 Eğitim %40 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 1)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
XGBoost	0.772278	0.859785	0.614356	0.680399	0.851962
KNN	0.772480	0.863581	0.594263	0.741502	0.885509
Random Forest	0.765927	0.857607	0.596481	0.731873	0.874207
SVM (RBF)	0.763004	0.857502	0.584360	0.715991	0.867485
Decision Tree	0.756250	0.846263	0.614227	0.665958	0.827304
MLP	0.743952	0.847328	0.551532	0.613516	0.821592

Tablo 4.6’da farklı makine öğrenmesi algoritmalarının çapraz doğrulama (CV) sonuçları doğruluk, F1 skoru, dengeli doğruluk, ROC-AUC ve PR-AUC ölçütleri üzerinden karşılaştırmalı olarak sunulmaktadır. Çapraz doğrulama sonuçları, modellerin farklı veri alt kümeleri üzerindeki tutarlılığını ve genellenebilirliğini değerlendirmek açısından önemli bir gösterge niteliği taşımaktadır.

Sonuçlar incelendiğinde, KNN ve XGBoost modellerinin genel performans açısından öne

çıkacağı görülmektedir. KNN modeli, CV Accuracy (0.7725) ve özellikle CV F1 Score (0.8636) ile başarılı öğrencilerin doğru sınıflandırılmasında güçlü bir performans sergilemiştir. Ayrıca PR-AUC (0.8855) değerinin yüksek olması, pozitif sınıf tahminlerinde istikrarlı sonuçlar ürettiğini göstermektedir. XGBoost modeli ise CV Balanced Accuracy (0.6144) değeriyle sınıflar arası dengeyi gözetten bir performans ortaya koymuştur.

Random Forest ve SVM (RBF) modelleri, doğruluk ve F1 skoru bakımından KNN ve XGBoost modellerine yakın değerler üretmiş olup, özellikle ROC-AUC ve PR-AUC ölçütlerinde rekabetçi sonuçlar sunmuştur. Bu durum, söz konusu modellerin sınıfları ayırt etme yeteneklerinin güçlü olduğunu göstermektedir. Decision Tree modeli, dengeli doğruluk açısından görece yüksek bir değer üretmesine rağmen, diğer ölçütlerde daha sınırlı bir performans sergilemiştir. MLP modeli ise tüm ölçütler dikkate alındığında diğer modellere kıyasla daha düşük sonuçlar elde etmiştir.

Genel değerlendirme yapıldığında, çapraz doğrulama sonuçları KNN modelinin doğruluk ve F1 skoru açısından, XGBoost modelinin ise sınıf dengesi gözetten performans açısından öne çıktığını göstermektedir. Bu bulgular, akademik performans tahmini gibi problemlerde farklı performans ölçütlerinin birlikte değerlendirilmesinin model seçimi açısından önemli olduğunu ortaya koymaktadır.

Tablo 4.7: %60 Eğitim %40 Test verileri ile tüm modellerin test performans sonuçları (Deneme 1)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
XGBoost	0.754717	0.852273	0.574942	0.683062	0.844766
KNN	0.745283	0.845714	0.568406	0.635704	0.804012
Random Forest	0.745283	0.844828	0.573612	0.682065	0.819956
SVM (RBF)	0.735849	0.840909	0.551457	0.615819	0.794162
Decision Tree	0.735849	0.836257	0.577490	0.640744	0.798321
MLP	0.688679	0.788462	0.596876	0.642960	0.828203

Tablo 4.7’de farklı makine öğrenmesi algoritmalarının test veri kümesi üzerindeki performansları doğruluk, F1 skoru, dengeli doğruluk, ROC-AUC ve PR-AUC ölçütleri

açısından karşılaştırmalı olarak sunulmaktadır. Test sonuçları, modellerin eğitim aşamasında görülmeyen veriler üzerindeki performansını ve genellenebilirliğini ortaya koymaktadır. Sonuçlar incelendiğinde, XGBoost modelinin genel olarak en dengeli performansı sergilediği görülmektedir. XGBoost, Test Accuracy (0.7547) ve Test F1 Score (0.8523) açısından tabloda yer alan modeller arasında en yüksek değerlere sahip olup, aynı zamanda ROC-AUC ve PR-AUC ölçütlerinde de rekabetçi sonuçlar üretmiştir. Bu durum, modelin hem genel sınıflandırma başarısı hem de pozitif sınıfı ayırt etme yeteneği açısından güçlü olduğunu göstermektedir.

KNN ve Random Forest modelleri, doğruluk ve F1 skoru bakımından XGBoost modeline yakın performanslar sergilemiştir. Bu iki modelin dengeli doğruluk değerlerinin benzer aralıklarda olması, sınıflar arasındaki ayırt ediciliklerinin karşılaştırılabilir düzeyde olduğunu göstermektedir. Random Forest modelinin ROC-AUC değerinin KNN'e kıyasla daha yüksek olması, sınıf ayırımında daha istikrarlı bir karar mekanizması sunduğunu düşündürmektedir.

Decision Tree ve SVM (RBF) modelleri, genel doğruluk ve F1 skorları bakımından orta düzey performans üretmiş, ancak ROC-AUC ve PR-AUC ölçütlerinde diğer modellere kıyasla daha sınırlı sonuçlar elde etmiştir. Bu durum, söz konusu modellerin test verisi üzerinde ayırt edicilik açısından daha düşük bir performans sergilediğini göstermektedir.

MLP modeli ise doğruluk ve F1 skoru açısından diğer modellere kıyasla daha düşük sonuçlar üretmesine rağmen, Test Balanced Accuracy (0.5969) değerinin görece yüksek olması dikkat çekicidir. Bu bulgu, modelin sınıflar arasındaki dengeyi gözetme konusunda belirli bir avantaj sağladığını, ancak genel sınıflandırma başarısının sınırlı kaldığını göstermektedir.

Genel değerlendirme yapıldığında, test sonuçları doğrultusunda XGBoost modeli genel performans açısından öne çıkarken, KNN ve Random Forest modelleri rekabetçi ve istikrarlı alternatifler olarak değerlendirilebilir. Bu bulgular, model performansının tek bir ölçüt üzerinden değil, farklı değerlendirme metrikleri birlikte ele alınarak yorumlanmasının önemini ortaya koymaktadır.

4.3.1.1 XGBoost Sınıflandırma Modeli Sonuçları

Deneme 1'in sınıflandırma sonuçlarına göre en iyi sonucu %60 eğitim %40 test verileri ile XGBoost sınıflandırma modeli elde etmiştir.

Tablo 4.8: XGBoost en iyi parametreler (Deneme 1)

Parametre	Değer (%60 / %40)
n_estimators	583
max_depth	5
learning_rate	0.25685
gamma	3.53844
min_child_weight	10
subsample	0.89669
colsample_bytree	0.65098
reg_alpha	2.90706

Tablo 4.8'de, %60 eğitim ve %40 test veri bölmesi kullanılarak XGBoost modeli için belirlenen en iyi hiperparametre değerleri sunulmaktadır. Elde edilen parametreler, modelin karmaşıklığı ile genellenebilirliği arasında dengeli bir yapı oluşturacak şekilde optimize edilmiştir. n_estimators değerinin 583 olarak belirlenmesi, modelin yeterli sayıda ağaç ile öğrenme sürecini tamamladığını göstermektedir. max_depth parametresinin 5 olması, aşırı uyumun önüne geçilerek daha kontrollü bir ağaç derinliği sağlandığını ortaya koymaktadır.

Görece yüksek learning_rate değeri, modelin öğrenme sürecinde daha hızlı güncellemeler yaptığını, buna karşın gamma ve min_child_weight parametrelerinin yüksek tutulması, bölünmelerin daha seçici yapılmasını sağlayarak modelin karmaşıklığını sınırlandırmaktadır. Subsample ve colsample_bytree değerlerinin 1'in altında olması, rastgelelik etkisi oluşturarak modelin genellenebilirliğini artırmayı hedeflemektedir. Ayrıca reg_alpha parametresinin sıfırdan büyük bir değer alması, L1 düzenleme etkisiyle gereksiz karmaşıklığın azaltıldığını göstermektedir. Bu parametre kombinasyonu, XGBoost modelinin hem dengeli hem de istikrarlı bir performans sergilemesine katkı sağlamıştır.

Tablo 4.9: XGBoost 10 Fold Cross Validation (CV) sonuçları (Deneme 1)

Metrik	Ortalama $\pm$ Std (%60 / %40)
Accuracy	0.7723 $\pm$ 0.0223
F1 Score	0.8598 $\pm$ 0.0152
Balanced Accuracy	0.6144 $\pm$ 0.0391
ROC-AUC	0.6804 $\pm$ 0.0545
PR-AUC	0.8520 $\pm$ 0.0390

Tablo 4.9’da %60 eğitim ve %40 test veri bölmesi kullanılarak XGBoost modeli için gerçekleştirilen 10 katlı çapraz doğrulama (Cross Validation) sonuçları sunulmaktadır. Elde edilen ortalama performans değerleri ve standart sapmalar, modelin farklı veri alt kümeleri üzerindeki kararlılığını ve genellenebilirliğini değerlendirmek açısından önemli bilgiler sağlamaktadır.

Sonuçlar incelendiğinde, Accuracy (0.7723) ve F1 Score (0.8598) değerlerinin görece yüksek olması, XGBoost modelinin genel sınıflandırma başarısının güçlü olduğunu ve özellikle pozitif sınıfın doğru tahmin edilmesinde etkili bir performans sergilediğini göstermektedir. Balanced Accuracy (0.6144) değerinin doğruluk ölçütüne kıyasla daha düşük olması, veri setinde sınıflar arasında belirli bir dengesizlik bulunduğunu ve modelin her iki sınıfı eşit düzeyde ayırt etme konusunda sınırlı bir performans sergilediğini işaret etmektedir.

ROC-AUC (0.6804) ve PR-AUC (0.8520) ölçütleri birlikte değerlendirildiğinde, modelin sınıflar arasında ayırt edicilik sağladığı ve özellikle pozitif sınıfa yönelik tahminlerde istikrarlı sonuçlar ürettiği görülmektedir. Standart sapma değerlerinin tüm metriklerde görece düşük olması, model performansının farklı katmanlarda büyük dalgalanmalar göstermediğini ve tutarlı bir öğrenme süreci izlediğini ortaya koymaktadır.

Genel olarak bu sonuçlar, XGBoost modelinin çapraz doğrulama sürecinde istikrarlı ve güvenilir bir performans sergilediğini ve test aşamasında elde edilen sonuçları destekler nitelikte olduğunu göstermektedir.

Tablo 4.10: XGBoost Confusion Matrix (Deneme 1)

	%60 / %40	
Gerçek \ Tahmin	0	1
0 (Başarısız)	10	49
1 (Başarılı)	3	150

Tablo 4.10’da %60 eğitim ve %40 test veri bölmesi kullanılarak eğitilen XGBoost modeline ait confusion matrix (karışıklık matrisi) sonuçları sunulmaktadır. Bu tablo, modelin başarılı ve başarısız öğrencileri ne ölçüde doğru ya da hatalı sınıflandırdığını ayrıntılı biçimde göstermektedir.

Sonuçlar incelendiğinde, modelin başarılı (1) sınıfını yüksek doğrulukla tahmin ettiği görülmektedir. Gerçek sınıfı başarılı olan 153 öğrenciden 150’si doğru şekilde başarılı olarak sınıflandırılmış, yalnızca 3 öğrenci başarısız olarak yanlış tahmin edilmiştir. Bu durum, modelin başarılı öğrencileri ayırt etme konusunda oldukça güçlü bir performans sergilediğini göstermektedir.

Buna karşılık başarısız (0) sınıfı için modelin performansının daha sınırlı olduğu görülmektedir. Gerçek sınıfı başarısız olan 59 öğrenciden yalnızca 10’u doğru şekilde başarısız olarak tahmin edilmiş, 49 öğrenci ise başarılı sınıfına yanlış atanmıştır. Bu durum, modelin başarısız öğrencileri tahmin etme konusunda daha temkinli davrandığını ve başarılı sınıfı tercih etme eğilimi gösterdiğini ortaya koymaktadır.

Genel olarak karışıklık matrisi sonuçları, XGBoost modelinin başarılı öğrencilerin tespitinde yüksek duyarlılığa, ancak başarısız öğrencilerin belirlenmesinde daha düşük seçiciliğe sahip olduğunu göstermektedir. Bu bulgu, modelin akademik risk altındaki öğrencileri erken tespit etme amacıyla kullanılacak senaryolarda ek iyileştirme veya sınıf dengesini gözetim yöntemleriyle desteklenebileceğini düşündürmektedir.

Tablo 4.11: XGBoost %60 Eğitim %40 Test verileri ile test sınıflandırma raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7692	0.1695	0.2778	59

Tablo 4.11: (devam ediyor)

1	0.7538	0.9804	0.8523	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.7547	212
Macro Avg	0.7615	0.5749	0.5650	212
Weighted Avg	0.7581	0.7547	0.6924	212

Tablo 4.11’de XGBoost modelinin test veri kümesi üzerindeki precision, recall, F1-score ve support değerleri sınıf bazında sunulmaktadır. Bu ölçütler, modelin her bir sınıfı ne ölçüde doğru tahmin ettiğini ve sınıflar arasındaki performans farklılıklarını ayrıntılı biçimde ortaya koymaktadır.

Sonuçlar incelendiğinde, başarılı (1) sınıfı için modelin oldukça yüksek bir performans sergilediği görülmektedir. Bu sınıfta recall değerinin 0.9804 olması, başarılı öğrencilerin neredeyse tamamının model tarafından doğru şekilde tespit edildiğini göstermektedir. Buna paralel olarak F1-score değerinin 0.8523 olması, modelin bu sınıfta hem duyarlılık hem de kesinlik açısından dengeli bir performans sunduğunu ortaya koymaktadır.

Buna karşılık başarısız (0) sınıfı için model performansının daha sınırlı olduğu görülmektedir. Bu sınıfta precision değeri 0.7692 olmasına rağmen, recall değerinin 0.1695 seviyesinde kalması, başarısız öğrencilerin önemli bir bölümünün model tarafından doğru biçimde yakalanamadığını göstermektedir. Buna bağlı olarak F1-score değeri de 0.2778 ile düşük seviyede gerçekleşmiştir. Bu durum, modelin başarısız sınıfı tahmin etme konusunda temkinli davrandığını ve başarılı sınıfı tercih etme eğilimi gösterdiğini ortaya koymaktadır.

Genel performans değerlendirildiğinde, accuracy değerinin 0.7547 olması modelin test veri kümesi üzerinde kabul edilebilir bir genel doğruluk sunduğunu göstermektedir. Macro average sonuçlarının, weighted average sonuçlarına kıyasla daha düşük olması, sınıflar arasındaki dengesizliğin model performansına yansıdığını ortaya koymaktadır. Weighted average değerlerinin daha yüksek olması ise modelin çoğunluk sınıfı olan başarılı öğrenciler üzerinde daha etkili sonuçlar ürettiğini göstermektedir.

Sonuç olarak sınıflandırma raporu, XGBoost modelinin başarılı öğrencilerin tespitinde

yüksek duyarlılığa sahip olduğunu, ancak başarısız öğrencilerin belirlenmesi konusunda sınırlı bir performans sergilediğini ortaya koymaktadır. Bu bulgu, modelin akademik performansını tahmin etmede güçlü olduğunu, ancak akademik risk altındaki öğrencilerin erken tespiti için ek iyileştirmelere ihtiyaç duyulabileceğini göstermektedir.

Tablo 4.4:12: XGBoost genel test sınıflandırma sonuçları (Deneme 1)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%60 / %40	0.7547	0.8523	0.5749	0.6831	0.8448

Tablo 4.12’de verilen %60 eğitim ve %40 test veri bölmesi kullanılarak elde edilen performans sonuçları, modelin genel sınıflandırma başarımını ve sınıflar arasındaki ayırt ediciliğini ortaya koymaktadır. Accuracy değerinin 0.7547 olması, modelin test veri kümesi üzerinde kabul edilebilir bir genel doğruluk sağladığını göstermektedir. F1 Score’un 0.8523 seviyesinde olması, özellikle pozitif sınıfın doğru tahmin edilmesinde modelin güçlü bir performans sergilediğine işaret etmektedir.

Buna karşılık Balanced Accuracy değerinin 0.5749 olması, sınıflar arasında belirli bir dengesizlik bulunduğunu ve modelin her iki sınıfı eşit düzeyde ayırt etme konusunda sınırlı kaldığını göstermektedir. ROC-AUC (0.6831) değeri, modelin sınıflar arasında ayırt edicilik sağladığını, ancak bu ayırt ediciliğin orta düzeyde olduğunu ortaya koymaktadır. PR-AUC değerinin 0.8448 olması ise modelin pozitif sınıfa yönelik tahminlerde istikrarlı ve güvenilir sonuçlar ürettiğini göstermektedir.

Genel olarak bu sonuçlar, modelin başarılı öğrencilerin tespitinde yüksek performans sunduğunu, ancak başarısız öğrencilerin belirlenmesi açısından ek iyileştirmelere ihtiyaç duyulabileceğini ortaya koymaktadır. Bu durum, performans ölçütlerinin birlikte değerlendirilmesinin model başarımını daha doğru yansıttığını göstermektedir.

4.3.2 Sınıflandırma Sonuçları (Deneme 2)

Veri seti; ADASYN veri arttırımı (Oversampled) yöntemi ile dengelenmiş, kategorik veriler ordinal kodlama yöntemi ile sayısallaştırılmış, %70 eğitim %30 test verisi (Tablo 38) ve %60 eğitim %40 test verisine (Tablo 40) ayrılarak, 300 iterasyonlu RandomSearchCV

içerisinde Stratified 10-Fold Crossover kullanılarak test edilmiştir ve sınıflandırma modellerinin en iyi parametreler ile elde ettikleri en iyi sonuçlar verilmiştir.

Tablo 4.13: Veri Arttırımı yapılmış, %70 Eğitim %30 Test verileri ve sınıf sayıları (Deneme 2)

Veriseti Türü	Satır Sayısı	Sütun Sayısı	Başarı (1)	Başarısız (0)
Orijinal Veriseti	528	34	382	146
Oversampled Toplam	740	34	382	358
Eğitim Veriseti	518	34	267	251
Test Veriseti	222	34	115	107

Tablo 4.13'te, veri setinin oversampling uygulanmadan önceki ve sonraki durumu ile eğitim ve test kümelerine ayrıldıktan sonraki sınıf dağılımları sunulmaktadır. Orijinal veri seti 528 gözlem ve 34 öznitelikten oluşmakta olup, bu gözlemlerin 382'si başarılı, 146'sı ise başarısız olarak etiketlenmiştir. Bu dağılım, veri setinde belirgin bir sınıf dengesizliği bulunduğunu göstermektedir.

Uygulanan oversampling işlemi sonrasında veri seti toplam 740 gözleme çıkarılmış ve başarısız sınıfa ait gözlem sayısı artırılarak sınıflar arasındaki dengesizlik önemli ölçüde azaltılmıştır. Bu aşamada başarılı sınıf gözlem sayısı 382 olarak korunurken, başarısız sınıf gözlem sayısı 358'e yükseltilmiştir. Böylece modelleme sürecinde her iki sınıfın da öğrenme sürecine daha dengeli biçimde katkı sağlaması hedeflenmiştir.

Oversampling uygulanmış veri seti, model eğitimi ve değerlendirmesi amacıyla yaklaşık %70 eğitim ve %30 test olacak şekilde iki alt kümeye ayrılmıştır. Bu doğrultuda eğitim veri seti 518 gözlemden oluşurken, test veri seti 222 gözlem içermektedir. Eğitim ve test kümelerindeki sınıf dağılımlarının birbirine yakın olması, modellerin hem eğitim hem de test aşamalarında benzer sınıf yapılarıyla karşılaşmasını sağlamış ve performans değerlendirmelerinin daha sağlıklı yapılmasına olanak tanımıştır.

Genel olarak tablo, oversampling işleminin sınıf dengesizliğini azaltmada etkili olduğunu ve bu sayede makine öğrenmesi modellerinin özellikle azınlık sınıfı öğrenme kapasitesinin artırılmasının amaçlandığını göstermektedir. Bu yaklaşım, akademik performans tahmini

gibi dengesiz sınıf yapısına sahip problemlerde model başarımını artırmaya yönelik önemli bir ön işleme adımı olarak değerlendirilmektedir.

Tablo 4.14: Veri Arttırımı yapılmış, %60 Eğitim %40 Test verileri ve sınıf sayıları (Deneme 2)

Veriseti Türü	Satır Sayısı	Sütun Sayısı	Başarı (1)	Başarısız (0)
Orijinal Veriseti	528	34	382	146
Oversampled Toplam	740	34	382	358
Eğitim Veriseti	444	34	229	215
Test Veriseti	296	34	153	143

Tablo 4.14’te, orijinal veri seti ile oversampling işlemi sonrası elde edilen veri yapısı ve bu verinin eğitim–test kümelerine ayrılması sonucunda oluşan sınıf dağılımları sunulmaktadır. Orijinal veri seti 528 gözlem ve 34 öznitelikten oluşmakta olup, başarılı (382) ve başarısız (146) sınıfları arasında belirgin bir dengesizlik bulunmaktadır. Bu dengesizliğin giderilmesi amacıyla oversampling yöntemi uygulanmış ve başarısız sınıfa ait gözlem sayısı artırılarak toplam veri seti 740 gözleme çıkarılmıştır. Bu aşamada başarılı sınıf korunurken, başarısız sınıf 358 gözleme yükseltilmiş ve sınıflar arası dağılım büyük ölçüde dengelenmiştir.

Oversampling uygulanmış veri seti, model eğitimi ve değerlendirmesi amacıyla yaklaşık %60 eğitim ve %40 test oranında iki alt kümeye ayrılmıştır. Buna göre eğitim veri seti 444 gözlemden (229 başarılı, 215 başarısız), test veri seti ise 296 gözlemden (153 başarılı, 143 başarısız) oluşmaktadır. Eğitim ve test kümelerinde her iki sınıfın da birbirine oldukça yakın sayılarda temsil edilmesi, modellerin öğrenme sürecinde sınıf yanlılığının azaltılmasına ve performans ölçümlerinin daha güvenilir biçimde değerlendirilmesine olanak sağlamıştır.

Genel olarak tablo, oversampling işleminin sınıf dengesizliğini önemli ölçüde azalttığını ve hem eğitim hem de test aşamalarında dengeli bir veri yapısı oluşturduğunu göstermektedir. Bu yapı, özellikle akademik risk altındaki öğrencilerin daha doğru biçimde tahmin edilmesi hedeflenen sınıflandırma problemleri için model performansını artırmaya yönelik uygun bir zemin sunmaktadır.

Tablo 4.15: %70 Eğitim %30 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 2)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
Random Forest	0.826433	0.828916	0.826630	0.888882	0.899907
SVM (RBF)	0.818590	0.839857	0.815425	0.871336	0.858482
XGBoost	0.822587	0.824210	0.823068	0.878274	0.891688
MLP	0.724057	0.733575	0.723402	0.783312	0.798196
KNN	0.747210	0.708429	0.751858	0.827756	0.842596
Decision Tree	0.724133	0.721431	0.724724	0.729226	0.695120

Tablo 4.15'te oversampling uygulanmış veri seti üzerinde farklı makine öğrenmesi algoritmalarının çapraz doğrulama (CV) performans sonuçları karşılaştırmalı olarak sunulmaktadır. Sonuçlar, doğruluk, F1 skoru, dengeli doğruluk, ROC-AUC ve PR-AUC ölçütleri üzerinden değerlendirilmiş olup, modellerin sınıflar arası dengeyi ne ölçüde sağladığını ve ayırt edicilik düzeylerini ortaya koymaktadır.

Elde edilen bulgular incelendiğinde, Random Forest modelinin tüm performans ölçütleri açısından en dengeli ve güçlü sonuçları ürettiği görülmektedir. Özellikle CV Accuracy (0.8264), CV Balanced Accuracy (0.8266) ve CV ROC-AUC (0.8889) değerlerinin yüksek olması, modelin her iki sınıfı da başarılı biçimde öğrendiğini ve ayırt edicilik gücünün arttığını göstermektedir. PR-AUC (0.8999) değerinin yüksekliği ise pozitif sınıf tahminlerinde modelin istikrarlı bir performans sergilediğine işaret etmektedir.

XGBoost ve SVM (RBF) modelleri, Random Forest'a oldukça yakın performans değerleri üretmiş ve özellikle ROC-AUC ve PR-AUC ölçütlerinde güçlü sonuçlar sunmuştur. Bu durum, her iki modelin de oversampling sonrası dengelenmiş veri yapısından etkin biçimde yararlandığını ve sınıflar arası ayrımı başarılı şekilde gerçekleştirdiğini göstermektedir. SVM (RBF) modelinin F1 skorunun görece yüksek olması, pozitif sınıfın doğru tahmin edilmesi açısından dikkat çekici bir performans sunduğunu ortaya koymaktadır.

Buna karşılık KNN, MLP ve Decision Tree modellerinin performanslarının, ensemble ve kernel tabanlı yöntemlere kıyasla daha düşük seviyelerde kaldığı görülmektedir. Özellikle

MLP ve Decision Tree modellerinin doğruluk ve ayırt edicilik ölçütlerinde sınırlı sonuçlar üretmesi, bu modellerin veri yapısına daha duyarlı olduğunu ve oversampling sonrasında dahi daha istikrarsız bir öğrenme süreci sergileyebileceğini düşündürmektedir.

Genel olarak değerlendirildiğinde, oversampling uygulaması sonrasında modellerin Balanced Accuracy ve ROC-AUC değerlerinde belirgin bir artış gözlenmiştir. Bu durum, sınıf dengesizliğinin giderilmesinin model performansı üzerinde olumlu bir etki yarattığını göstermektedir. Özellikle Random Forest, XGBoost ve SVM (RBF) modelleri, dengeli veri seti üzerinde hem genellenebilirlik hem de sınıflar arası ayırt edicilik açısından öne çıkan yöntemler olarak değerlendirilmektedir.

Tablo 4.16: %70 Eğitim %30 Test verileri ile tüm modellerin test performans sonuçları (Deneme 2)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
Random Forest	0.855856	0.863248	0.855018	0.910646	0.909670
SVM (RBF)	0.842342	0.860558	0.838724	0.913856	0.907397
XGBoost	0.828829	0.840336	0.827306	0.910849	0.919826
MLP	0.801802	0.815126	0.800244	0.857944	0.855212
KNN	0.806306	0.796209	0.809143	0.884112	0.892632
Decision Tree	0.747748	0.754386	0.747745	0.744169	0.697031

Tablo 4.16’da oversampling uygulanmış veri seti kullanılarak eğitilen makine öğrenmesi modellerinin test veri kümesi üzerindeki performansları sunulmaktadır. Sonuçlar doğruluk, F1 skoru, dengeli doğruluk, ROC-AUC ve PR-AUC ölçütleri üzerinden değerlendirilmiş olup, modellerin daha önce görülmemiş veriler üzerindeki genellenebilirlikleri karşılaştırılmıştır.

Elde edilen bulgular incelendiğinde, Random Forest modelinin tüm metrikler açısından en güçlü ve dengeli performansı sergilediği görülmektedir. Model, Test Accuracy (0.8559) ve Test F1 Score (0.8632) değerleriyle genel sınıflandırma başarısında öne çıkarken, Test Balanced Accuracy (0.8550) değeriyle her iki sınıfı da eşit düzeyde ayırt edebildiğini göstermektedir. Ayrıca ROC-AUC (0.9106) ve PR-AUC (0.9097) değerlerinin yüksekliği,

modelin ayırt edicilik gücünün ve pozitif sınıf tahminlerindeki istikrarının oldukça güçlü olduğunu ortaya koymaktadır.

SVM (RBF) ve XGBoost modelleri de Random Forest’a yakın performanslar sergilemiştir. SVM (RBF) modeli özellikle ROC-AUC (0.9139) değeriyle sınıflar arası ayrımı başarılı biçimde gerçekleştirmiştir. XGBoost modeli ise PR-AUC (0.9198) değeriyle pozitif sınıf tahminlerinde dikkat çekici bir performans göstermiştir. Bu sonuçlar, her iki modelin de oversampling sonrası dengelenmiş veri yapısından etkin biçimde faydalandığını göstermektedir.

MLP ve KNN modelleri, doğruluk ve F1 skorları bakımından kabul edilebilir düzeyde sonuçlar üretmiş olmakla birlikte, ensemble ve kernel tabanlı yöntemlerin gerisinde kalmıştır. Decision Tree modeli ise tüm performans ölçütlerinde diğer modellere kıyasla daha düşük sonuçlar elde etmiş ve özellikle ayırt edicilik ölçütlerinde sınırlı bir performans sergilemiştir.

Genel değerlendirme yapıldığında, oversampling sonrası test sonuçları sınıf dengesinin sağlanmasının model performansını belirgin biçimde artırdığını ortaya koymaktadır. Bu bağlamda Random Forest modeli hem doğruluk hem de sınıflar arası denge ve ayırt edicilik ölçütleri açısından en başarılı model olarak öne çıkmaktadır. SVM (RBF) ve XGBoost modelleri ise güçlü alternatifler olarak değerlendirilebilir. Bu bulgular, nihai model seçimi sürecinde oversampling uygulanmış veri seti üzerinde Random Forest modelinin tercih edilmesini destekler niteliktedir.

Tablo 4.17: %60 Eğitim %40 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 2)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.819899	0.832525	0.818784	0.863522	0.854741
Random Forest	0.806616	0.814693	0.806343	0.853808	0.873288
XGBoost	0.808687	0.809838	0.809543	0.864186	0.874667
MLP	0.711970	0.733629	0.710578	0.725608	0.751255
KNN	0.738838	0.694887	0.743888	0.800294	0.823146

Tablo 4.17: (devam ediyor)

Decision Tree	0.702374	0.704307	0.703590	0.730930	0.735500
----------------------	----------	----------	----------	----------	----------

Tablo 4.17’de farklı makine öğrenmesi algoritmalarının çapraz doğrulama (CV) sonuçları sunulmakta ve modellerin farklı veri alt kümeleri üzerindeki performans tutarlılığı değerlendirilmektedir. Elde edilen bulgular, modellerin doğruluk, F1 skoru, dengeli doğruluk ve ayırt edicilik ölçütleri açısından birbirlerinden belirgin biçimde ayrıştığını göstermektedir.

Sonuçlar incelendiğinde, SVM (RBF) modelinin çapraz doğrulama sürecinde genel olarak en istikrarlı performansı sergilediği görülmektedir. Özellikle CV Accuracy (0.8199) ve CV Balanced Accuracy (0.8188) değerlerinin yüksek olması, modelin her iki sınıfı da dengeli biçimde öğrenebildiğini göstermektedir. Buna ek olarak ROC-AUC ve PR-AUC değerleri, SVM modelinin sınıflar arasındaki ayrımı güçlü bir karar sınırıyla gerçekleştirdiğini ortaya koymaktadır.

Random Forest ve XGBoost modelleri, SVM’ye oldukça yakın sonuçlar üretmiş ve özellikle PR-AUC ölçütlerinde dikkat çekici performanslar sergilemiştir. Bu durum, her iki modelin de pozitif sınıf tahminlerinde istikrarlı olduğunu ve veri setindeki örüntüleri etkili biçimde yakalayabildiğini göstermektedir. XGBoost modelinin ROC-AUC değerinin görece yüksek olması, sınıf ayrımında güçlü bir ayırt edicilik sunduğunu düşündürmektedir.

Buna karşılık KNN, MLP ve Decision Tree modellerinin çapraz doğrulama performanslarının daha sınırlı kaldığı görülmektedir. KNN modeli dengeli doğruluk açısından kabul edilebilir bir değer üretmesine rağmen F1 skorunun düşük olması, sınıflardan birine yönelik tahmin başarısının zayıf kaldığını göstermektedir. MLP ve Decision Tree modelleri ise hem doğruluk hem de ayırt edicilik ölçütlerinde diğer yöntemlerin gerisinde kalmıştır.

Genel değerlendirme yapıldığında, çapraz doğrulama sonuçları SVM (RBF), Random Forest ve XGBoost modellerinin veri seti üzerinde daha tutarlı ve dengeli bir öğrenme süreci sergilediğini ortaya koymaktadır. Bu bulgular, model seçimi sürecinde yalnızca tek bir performans ölçütüne değil, çapraz doğrulama kapsamında elde edilen çoklu metriklerin

birlikte değerlendirilmesinin önemini vurgulamaktadır.

Tablo 4.18: %60 Eğitim %40 Test verileri ile tüm modellerin test performans sonuçları (Deneme 2)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
SVM (RBF)	0.837838	0.850932	0.835824	0.895059	0.884479
Random Forest	0.841216	0.847896	0.840692	0.899447	0.898865
XGBoost	0.790541	0.791946	0.791215	0.894282	0.889693
MLP	0.750000	0.767296	0.748343	0.798894	0.810893
KNN	0.773649	0.750929	0.777618	0.850336	0.863673
Decision Tree	0.706081	0.698962	0.707688	0.801385	0.770269

Tablo 4.18’de %60 eğitim ve %40 test veri bölmesi kullanılarak eğitilen tüm makine öğrenmesi modellerinin test performansları karşılaştırmalı olarak sunulmaktadır. Test sonuçları, modellerin eğitim sürecinde öğrenilen örüntüleri daha önce görülmemiş veriler üzerinde ne ölçüde genelleştirebildiğini ortaya koymaktadır.

Sonuçlar incelendiğinde, Random Forest ve SVM (RBF) modellerinin bu veri bölme senaryosunda en güçlü performansı sergilediği görülmektedir. Random Forest modeli, doğruluk ve dengeli doğruluk ölçütlerinde elde ettiği yüksek değerlerle her iki sınıfı da benzer başarı düzeylerinde ayırt edebilmiştir. Buna ek olarak ROC-AUC ve PR-AUC değerlerinin yüksekliği, modelin sınıflar arası ayrımı güvenilir ve istikrarlı biçimde gerçekleştirdiğini göstermektedir. SVM (RBF) modeli ise özellikle dengeli doğruluk ve F1 skoru açısından dikkat çekici sonuçlar üretmiş, bu durum modelin sınıf dağılımını gözeterek bir karar yapısı geliştirdiğine işaret etmiştir.

XGBoost modeli, ayırt edicilik ölçütleri açısından güçlü sonuçlar üretmesine rağmen doğruluk ve F1 skoru bakımından Random Forest ve SVM’nin gerisinde kalmıştır. Bu durum, modelin sınıfları ayırt etme kapasitesinin yüksek olmasına karşın karar eşiği düzeyinde daha temkinli bir yapı sergilediğini düşündürmektedir. KNN modeli ise dengeli doğruluk bakımından kabul edilebilir bir performans ortaya koymuş, ancak F1 skoru değerinin görece düşük olması, sınıflardan biri için tahmin performansının sınırlı kaldığını

göstermektedir.

MLP ve Decision Tree modelleri, genel olarak diğer yöntemlere kıyasla daha düşük performans değerleri üretmiştir. Özellikle Decision Tree modelinin doğruluk ve dengeli doğruluk ölçütlerinde sınırlı sonuçlar vermesi, daha basit karar yapısının karmaşık veri örüntülerini yakalamakta yetersiz kaldığını göstermektedir. MLP modeli ise orta düzeyde bir performans sergilemiş, ancak ensemble ve kernel tabanlı yöntemlerin gerisinde kalmıştır.

Genel olarak bu sonuçlar, %60 eğitim – %40 test veri bölmesi kullanıldığında Random Forest ve SVM (RBF) modellerinin hem sınıf dengesi hem de ayırt edicilik açısından daha güvenilir sonuçlar ürettiğini ortaya koymaktadır. Bu bulgular, nihai model değerlendirmesinde yalnızca doğruluk değerlerinin değil, dengeli doğruluk ve ayırt edicilik ölçütlerinin birlikte ele alınmasının önemini bir kez daha göstermektedir.

4.3.2.1 Rastgele Orman Sınıflandırma Modeli Sonuçları

Deneme 2'in sınıflandırma sonuçlarına göre en iyi sonucu %70 eğitim %30 test verileri ile Rastgele Orman sınıflandırma modeli elde etmiştir.

Tablo 4.19: Rastgele Orman en iyi parametreler (Deneme 2)

Parametre	Değer (%70 / %30)
Bootstrap	False
Max Depth	None
Max Features	sqrt
Min Samples Leaf	2
Min Samples Split	2
n_estimators	503

Tablo 4.19'da %70 eğitim ve %30 test veri bölmesi kullanılarak Random Forest modeli için belirlenen en uygun hiperparametre değerleri sunulmaktadır. n_estimators parametresinin 503 olarak belirlenmesi, modelin yeterli sayıda karar ağacı ile öğrenme sürecini tamamladığını ve sonuçların tekil ağaçlara bağlı kalmadan topluluk etkisiyle üretildiğini

göstermektedir. Bootstrap parametresinin *False* olarak ayarlanması, modelin tüm veri örneklerini kullanarak her bir ağacı oluşturduğunu ve bu sayede veri temsiliinin daha bütüncül biçimde ele alındığını ortaya koymaktadır.

Max Depth parametresinin *None* olması, ağaçların büyüme sürecinde önceden belirlenmiş bir derinlik kısıtı uygulanmadığını, ancak Min Samples Leaf (2) ve Min Samples Split (2) parametreleri aracılığıyla aşırı uyumun kontrol altına alındığını göstermektedir. Max Features parametresinin *sqrt* olarak seçilmesi ise her bölünmede kullanılan öznelik sayısını sınırlandırarak ağaçlar arasında çeşitliliği artırmakta ve modelin genellenebilirliğine katkı sağlamaktadır.

Genel olarak bu parametre kombinasyonu, Random Forest modelinin karmaşıklık ile genellenebilirlik arasında dengeli bir yapı kurmasına olanak tanımış ve %70–%30 veri bölme senaryosunda elde edilen yüksek performans sonuçlarını destekleyen bir hiperparametre yapılandırması sunmuştur.

Tablo 4.20: Rastgele Orman 10 Fold Cross Validation (CV) sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std (%70 / %30)
Accuracy	0.8264 $\pm$ 0.0370
F1 Score	0.8289 $\pm$ 0.0392
Balanced Accuracy	0.8266 $\pm$ 0.0367
ROC-AUC	0.8889 $\pm$ 0.0262
PR-AUC	0.8999 $\pm$ 0.0342

Tablo 4.20’de %70 eğitim ve %30 test veri bölmesi kullanılarak Random Forest modeli için elde edilen çapraz doğrulama sonuçları, modelin hem genel başarımlarını hem de sınıflar arası denge açısından güçlü bir yapı sunduğunu göstermektedir. Accuracy ve Balanced Accuracy değerlerinin birbirine oldukça yakın olması, modelin sınıflar arasında tutarlı bir karar mekanizması geliştirdiğini ve tahmin sürecinde belirgin bir sınıf yanlılığı oluşmadığını ortaya koymaktadır. F1 skorunun yüksekliği, modelin özellikle hatalı sınıflandırmaların maliyetinin önemli olduğu durumlarda dengeli bir performans sergilediğini göstermektedir.

ROC-AUC ve PR-AUC değerlerinin yüksek seviyelerde olması, Random Forest modelinin karar sınırlarını etkin biçimde oluşturabildiğini ve özellikle pozitif sınıfa ilişkin tahminlerde güvenilir sonuçlar ürettiğini işaret etmektedir. Standart sapma değerlerinin tüm metriklerde görece düşük olması ise model performansının farklı çapraz doğrulama katmanları arasında tutarlılığını koruduğunu ve öğrenme sürecinin kararlı bir yapı sergilediğini göstermektedir. Bu sonuçlar, Random Forest modelinin %70–%30 veri bölme senaryosunda genellenebilir ve istikrarlı bir performans sunduğunu ortaya koymaktadır.

Tablo 4.21: Rastgele Orman Confusion Matrix (Deneme 2)

	%70 / %30	
Gerçek \ Tahmin	0	1
0 (Başarısız)	89	18
1 (Başarılı)	14	101

Tablo 4.21’de verilen karışıklık matrisi, %70 eğitim ve %30 test veri bölmesi kullanılarak eğitilen Random Forest modelinin sınıflandırma performansını ayrıntılı biçimde ortaya koymaktadır. Sonuçlara göre, gerçek sınıfı başarısız (0) olan 107 öğrencinin 89’u doğru şekilde başarısız olarak sınıflandırılmış, 18’i ise başarılı sınıfına yanlış atanmıştır. Bu durum, modelin başarısız öğrencileri tespit etme konusunda yüksek bir seçicilik sergilediğini göstermektedir.

Gerçek sınıfı başarılı (1) olan 115 öğrencinin ise 101’i doğru biçimde başarılı olarak tahmin edilmiş, 14 öğrenci başarısız olarak yanlış sınıflandırılmıştır. Bu bulgu, modelin başarılı öğrencileri belirleme konusunda da güçlü bir performans sunduğunu ve her iki sınıf için dengeli bir tahmin yapısı geliştirdiğini ortaya koymaktadır.

Genel olarak karışıklık matrisi incelendiğinde, Random Forest modelinin hem başarılı hem de başarısız sınıfları yüksek doğrulukla ayırt edebildiği, yanlış sınıflandırma sayılarının sınırlı kaldığı görülmektedir. Bu sonuçlar, önceki doğruluk, dengeli doğruluk ve ROC-AUC bulgularını destekler nitelikte olup, modelin %70–%30 veri bölme senaryosunda istikrarlı ve güvenilir bir sınıflandırma performansı sergilediğini göstermektedir.

Tablo 4.22: Rastgele Orman %70 Eğitim %30 Test verileri ile test sınıflandırma raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8641	0.8318	0.8476	107
1	0.8487	0.8783	0.8632	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8559	222
Macro Avg	0.8564	0.8550	0.8554	222
Weighted Avg	0.8561	0.8559	0.8557	222

Tablo 4.22’de, %70 eğitim ve %30 test veri bölmesi kullanılarak eğitilen Random Forest modeline ait sınıflandırma performansı, precision, recall ve F1-score ölçütleri üzerinden sınıf bazında sunulmaktadır. Bu ölçütler, modelin her bir sınıfı ne ölçüde doğru ve tutarlı biçimde tahmin edebildiğini ayrıntılı olarak değerlendirmeye olanak tanımaktadır.

Sonuçlar incelendiğinde, başarısız (0) sınıfı için precision değerinin 0.8641 olması, modelin başarısız olarak etiketlediği öğrencilerin büyük bir bölümünün gerçekten başarısız olduğunu göstermektedir. Recall değerinin 0.8318 seviyesinde olması ise başarısız öğrencilerin önemli bir kısmının doğru biçimde tespit edildiğini ortaya koymaktadır. Bu iki ölçütün dengeli bir şekilde birleşmesiyle elde edilen F1-score (0.8476), modelin bu sınıf için tutarlı bir performans sunduğunu göstermektedir.

Başarılı (1) sınıfı açısından bakıldığında, precision (0.8487) ve recall (0.8783) değerlerinin yüksek olması, modelin başarılı öğrencileri hem doğru tanımlama hem de büyük oranda yakalama konusunda etkili olduğunu ortaya koymaktadır. Bu sınıf için hesaplanan F1-score (0.8632), modelin başarılı öğrenciler üzerindeki güçlü tahmin yeteneğini desteklemektedir. Genel performans değerlendirildiğinde, accuracy değerinin 0.8559 olması, modelin test veri kümesi üzerinde yüksek bir genel sınıflandırma başarısı sergilediğini göstermektedir. Macro average ve weighted average değerlerinin birbirine oldukça yakın olması, model performansının sınıflar arasında dengeli olduğunu ve sınıf dağılımından kaynaklı belirgin bir yanlılık bulunmadığını ortaya koymaktadır. Bu bulgular, Random Forest modelinin %70–%30 veri bölme senaryosunda istikrarlı, dengeli ve güvenilir bir sınıflandırma performansı sunduğunu göstermektedir.

Tablo 4.23: Rastgele Orman genel test sınıflandırma sonuçları (Deneme 2)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8559	0.8632	0.8550	0.9106	0.9097

Tablo 4.23'te %70 eğitim ve %30 test veri bölmesi kullanılarak elde edilen performans sonuçları, modelin genel başarımının yanı sıra sınıflar arası dengeyi ve ayırt edicilik düzeyini ortaya koymaktadır. Accuracy değerinin 0.8559 olması, modelin test veri kümesi üzerinde yüksek bir genel doğruluk sağladığını göstermektedir. F1 Score'un 0.8632 seviyesinde olması, modelin hem kesinlik hem de duyarlılık açısından dengeli bir performans sunduğunu ve yanlış sınıflandırmaların sınırlı kaldığını ortaya koymaktadır.

Balanced Accuracy değerinin 0.8550 olması, modelin her iki sınıfı da benzer başarı düzeylerinde ayırt edebildiğini ve sınıf dağılımından kaynaklanan yanlılıkların büyük ölçüde giderildiğini göstermektedir. ROC-AUC (0.9106) değeri, modelin sınıflar arasında güçlü bir ayırt edicilik sağladığını ortaya koyarken, PR-AUC (0.9097) değeri özellikle pozitif sınıfa yönelik tahminlerin güvenilir ve istikrarlı olduğunu göstermektedir.

Genel olarak bu sonuçlar, modelin %70–%30 veri bölme senaryosunda yüksek doğruluk, dengeli sınıf ayrımı ve güçlü ayırt edicilik sunduğunu ve akademik performans tahmini problemine uygun, güvenilir bir çözüm ürettiğini ortaya koymaktadır.

4.3.3 Sınıflandırma Sonuçları (Deneme 3)

Veri seti; ADASYN veri arttırımı yöntemi ile dengelenmiş, kategorik veriler one-hot encoding yöntemi ile sayısallaştırılmış, %70 eğitim %30 test verisi (Tablo 4.25) ve %60 eğitim %40 test verisine (Tablo 4.26) ayrılarak, 300 iterasyonlu RandomSearchCV içerisinde Stratified 10-Fold Crossover kullanılarak test edilmiştir ve sınıflandırma modellerinin en iyi parametreler ile elde ettikleri en iyi sonuçlar verilmiştir. One-Hot kodlama sonucunda öznitelik sayısı 117'ye çıkmıştır.

Tablo 4.24: Veri Arttırımı yapılmış, %70 Eğitim %30 Test verileri ve sınıf sayıları (Deneme 3)

Veriseti Aşaması	Satır Sayısı	Sütun Sayısı	Başarılı (True)	Başarısız (False)
Orijinal Veriseti	528	117	382	146
Oversampled Veriseti	745	117	382	363
Eğitim (Train) Veriseti	521	117	267	254
Test Veriseti	224	117	115	109

Tablo 4.24’te, veri setinin farklı aşamalardaki yapısı ve sınıf dağılımları ayrıntılı biçimde sunulmaktadır. Orijinal veri seti 528 gözlem ve 117 öznitelikten oluşmakta olup, başarılı (382) ve başarısız (146) sınıfları arasında belirgin bir dengesizlik bulunmaktadır. Bu durum, sınıflandırma modellerinin çoğunluk sınıfı yönelme eğilimi göstermesine neden olabileceğinden, analiz sürecinde dikkate alınması gereken önemli bir yöntemsel kısıt oluşturmaktadır.

Bu dengesizliği gidermek amacıyla uygulanan oversampling işlemi sonrasında veri seti 745 gözleme çıkarılmış, başarısız sınıfa ait gözlem sayısı artırılarak sınıflar arasındaki dağılım büyük ölçüde dengelenmiştir. Oversampling aşamasında başarılı sınıf gözlem sayısı korunurken, başarısız sınıfın 363 gözleme yükseltilmesi, her iki sınıfın öğrenme sürecine daha eşit katkı sağlamasını mümkün kılmıştır.

Dengelenmiş veri seti, modelleme sürecinde kullanılmak üzere eğitim ve test kümelerine ayrılmıştır. Buna göre eğitim veri seti 521 gözlemden (267 başarılı, 254 başarısız), test veri seti ise 224 gözlemden (115 başarılı, 109 başarısız) oluşmaktadır. Eğitim ve test kümelerindeki sınıf dağılımlarının birbirine oldukça yakın olması, modellerin hem öğrenme hem de değerlendirme aşamalarında benzer sınıf yapılarıyla karşılaşmasını sağlamış ve performans ölçümlerinin daha güvenilir biçimde yapılmasına olanak tanımıştır.

Genel olarak tablo, veri setinin modelleme öncesi ve sonrası aşamalarının sistematik biçimde ele alındığını, oversampling yöntemiyle sınıf dengesizliğinin azaltıldığını ve eğitim–test ayrımının dengeli bir yapı oluşturacak şekilde gerçekleştirildiğini göstermektedir. Bu yapı,

geliştirilen makine öğrenmesi modellerinin genellenebilirliğini ve özellikle azınlık sınıfın tahmin başarımını artırmaya yönelik sağlam bir temel sunmaktadır.

Tablo 4.25: Veri Arttırımı yapılmış, %60 Eğitim %40 Test verileri ve sınıf sayıları (Deneme 3)

Veriseti Aşaması	Satır Sayısı	Sütun Sayısı	Başarılı (True)	Başarısız (False)
Orijinal Veriseti	528	117	382	146
Oversampled Veriseti	745	117	382	363
Eğitim (Train) Veriseti	447	117	229	218
Test Veriseti	298	117	153	145

Tablo 4.25'te veri setinin modelleme süreci boyunca geçirdiği aşamalar ve her aşamadaki sınıf dağılımları ayrıntılı biçimde sunulmaktadır. Orijinal veri seti 528 gözlem ve 117 öznitelikten oluşmakta olup, başarılı (382) ve başarısız (146) sınıfları arasında belirgin bir dengesizlik bulunmaktadır. Bu yapı, sınıflandırma modellerinin çoğunluk sınıfa yönelme riskini artırabilecek bir başlangıç durumu oluşturmaktadır.

Bu sorunun giderilmesi amacıyla uygulanan oversampling işlemi sonrasında veri seti 745 gözleme çıkarılmış ve başarısız sınıfa ait gözlem sayısı artırılarak sınıflar arasındaki dağılım büyük ölçüde dengelenmiştir. Bu aşamada başarılı sınıf gözlem sayısı korunurken, başarısız sınıfın 363 gözleme yükseltilmesi, her iki sınıfın öğrenme sürecinde daha dengeli biçimde temsil edilmesini sağlamıştır.

Dengelenmiş veri seti, model eğitimi ve değerlendirmesi için %60 eğitim – %40 test oranında iki alt kümeye ayrılmıştır. Buna göre eğitim veri seti 447 gözlemden (229 başarılı, 218 başarısız), test veri seti ise 298 gözlemden (153 başarılı, 145 başarısız) oluşmaktadır. Eğitim ve test kümelerindeki sınıf dağılımlarının birbirine oldukça yakın olması, modellerin hem öğrenme hem de test aşamalarında benzer sınıf yapılarıyla karşılaşmasını sağlamış ve performans karşılaştırmalarının daha sağlıklı yapılmasına olanak tanımıştır.

Genel olarak bu tablo, veri setinin ön işleme, dengeleme ve bölme adımlarının sistematik ve

kontrollü biçimde yürütüldüğünü göstermektedir. Oluşturulan bu dengeli yapı, geliştirilen makine öğrenmesi modellerinin özellikle her iki sınıf için de genellenebilir ve güvenilir sonuçlar üretmesini destekleyen önemli bir yöntemsel zemin sunmaktadır.

Tablo 4.26: %70 Eğitim %30 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 3)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.857910	0.865051	0.857182	0.895611	0.877010
Random Forest	0.861720	0.867420	0.861396	0.914305	0.908441
KNN	0.846372	0.835304	0.848726	0.886860	0.887555
XGBoost	0.810015	0.814228	0.809943	0.882660	0.871924
MLP	0.800327	0.820821	0.798775	0.849759	0.825285
Decision Tree	0.765856	0.769429	0.765299	0.805802	0.796004

Tablo 4.26’da 117 öznitelikten oluşan dengelenmiş veri seti üzerinde uygulanan makine öğrenmesi algoritmalarının çapraz doğrulama (CV) performans sonuçları yer almaktadır. Bu sonuçlar, modellerin yalnızca tek bir veri bölmesi üzerindeki başarısını değil, farklı alt kümeler üzerinde gösterdikleri kararlılığı ve genellenebilirliği değerlendirmeye imkân tanımaktadır.

Elde edilen bulgular, Random Forest modelinin tüm performans ölçütleri birlikte değerlendirildiğinde en güçlü sonuçları ürettiğini göstermektedir. Özellikle doğruluk, F1 skoru ve dengeli doğruluk değerlerinin birbirine oldukça yakın olması, modelin sınıflar arasında dengeli bir öğrenme gerçekleştirdiğini ortaya koymaktadır. Ayrıca ROC-AUC ve PR-AUC değerlerinin yüksek olması, Random Forest modelinin sınıflar arasındaki ayrımı güvenilir bir biçimde yapabildiğini ve pozitif sınıfa yönelik tahminlerde istikrarlı sonuçlar sunduğunu göstermektedir.

SVM (RBF) modeli, Random Forest’a çok yakın performans değerleri üretmiş ve özellikle dengeli doğruluk ölçütünde dikkat çekici bir başarı sergilemiştir. Bu durum, SVM modelinin sınıf dağılımını etkin biçimde dikkate alan bir karar sınırı oluşturabildiğini göstermektedir. ROC-AUC ve PR-AUC değerleri de modelin ayırt edicilik gücünün yüksek olduğunu

desteklemektedir.

KNN modeli, doğruluk ve dengeli doğruluk açısından rekabetçi sonuçlar üretmiş olmakla birlikte, F1 skorunun görece daha düşük kalması, modelin sınıflardan biri üzerinde daha sınırlı bir tahmin performansı sergilediğine işaret etmektedir. XGBoost ve MLP modelleri ise orta düzeyde performans göstermiş; özellikle doğruluk ve dengeli doğruluk ölçütlerinde üst düzey modellere kıyasla daha düşük sonuçlar elde etmiştir. Decision Tree modeli ise tüm metrikler dikkate alındığında en düşük performansı sergileyerek, tekil ağaç yapısının karmaşık ve yüksek boyutlu veri yapısını yeterince temsil edemediğini ortaya koymuştur.

Genel olarak çapraz doğrulama sonuçları, ensemble tabanlı (Random Forest) ve çekirdek tabanlı (SVM) yöntemlerin bu veri yapısı üzerinde daha dengeli, kararlı ve ayırt edici sonuçlar ürettiğini göstermektedir. Bu bulgular, nihai model seçimi sürecinde bu yöntemlerin öncelikli olarak değerlendirilmesini desteklemektedir.

Tablo 4.27: %70 Eğitim %30 Test verileri ile tüm modellerin test performans sonuçları (Deneme 3)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
SVM (RBF)	0.892857	0.901639	0.891105	0.932509	0.923258
Random Forest	0.883929	0.890756	0.882888	0.934304	0.933861
KNN	0.875000	0.875000	0.875628	0.911647	0.910487
XGBoost	0.852679	0.863071	0.851256	0.917112	0.896369
MLP	0.745536	0.748899	0.745712	0.823055	0.816195
Decision Tree	0.727679	0.738197	0.727124	0.770882	0.752874

Tablo 4.27’de, dengelenmiş ve 117 öznitelikten oluşan veri seti üzerinde eğitilen makine öğrenmesi modellerinin test veri kümesi performansları sunulmaktadır. Test sonuçları, modellerin eğitim sürecinde öğrenilen örüntüleri daha önce görülmemiş veriler üzerinde ne ölçüde doğru ve dengeli biçimde genelleyebildiğini göstermesi açısından kritik öneme sahiptir.

Elde edilen bulgular, SVM (RBF) ve Random Forest modellerinin test aşamasında en yüksek

performansı sergilediğini ortaya koymaktadır. SVM (RBF) modeli, Test Accuracy (0.8929) ve Test F1 Score (0.9016) değerleriyle genel sınıflandırma başarısında öne çıkarken, Balanced Accuracy (0.8911) değeriyle her iki sınıfı da dengeli biçimde ayırt edebildiğini göstermiştir. ROC-AUC ve PR-AUC değerlerinin yüksekliği, modelin sınıflar arasındaki ayrımı güçlü bir karar sınırı ile gerçekleştirdiğini ve pozitif sınıfa yönelik tahminlerde istikrarlı sonuçlar ürettiğini ortaya koymaktadır.

Random Forest modeli ise doğruluk ve F1 skoru bakımından SVM'ye oldukça yakın sonuçlar üretmiş, özellikle ROC-AUC (0.9343) ve PR-AUC (0.9339) değerleriyle ayırt edicilik açısından en güçlü modellerden biri olduğunu göstermiştir. Bu durum, Random Forest'ın topluluk yapısı sayesinde karmaşık öznelik ilişkilerini etkili biçimde yakalayabildiğini düşündürmektedir.

KNN modeli, doğruluk ve dengeli doğruluk açısından yüksek ve tutarlı bir performans sergileyerek dikkat çekmiştir. Ancak ayırt edicilik ölçütleri bakımından SVM ve Random Forest modellerinin gerisinde kalmıştır. XGBoost modeli ise genel olarak güçlü ve dengeli sonuçlar üretmiş olmakla birlikte, test doğruluğu ve F1 skoru bakımından ilk üç modelin altında yer almıştır.

Buna karşılık MLP ve Decision Tree modelleri, tüm performans ölçütleri dikkate alındığında diğer yöntemlere kıyasla daha düşük sonuçlar elde etmiştir. Özellikle Decision Tree modelinin daha sınırlı bir ayırt edicilik sergilemesi, tekil ağaç yapısının yüksek boyutlu ve karmaşık veri yapısını temsil etmede yetersiz kaldığını göstermektedir.

Genel olarak test sonuçları, SVM (RBF) ve Random Forest modellerinin bu veri seti üzerinde yüksek doğruluk, dengeli sınıf ayrımı ve güçlü ayırt edicilik sunduğunu ortaya koymaktadır. Bu bulgular, nihai model seçimi sürecinde bu iki yöntemin öncelikli adaylar olarak değerlendirilmesini desteklemektedir.

Tablo 4.28: %60 Eğitim %40 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 3)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC-AUC	CV PR-AUC
SVM (RBF)	0.859141	0.869671	0.857839	0.890773	0.870842
Random Forest	0.852677	0.862963	0.851534	0.895892	0.889632
KNN	0.801313	0.825532	0.798631	0.852321	0.838621
XGBoost	0.801111	0.807110	0.800616	0.883741	0.877899
MLP	0.780960	0.793498	0.779955	0.830561	0.800135
Decision Tree	0.704899	0.698277	0.706042	0.705251	0.662608

Tablo 4.28’de dengelenmiş veri seti üzerinde uygulanan makine öğrenmesi modellerinin çapraz doğrulama performansları sunulmaktadır. Bu sonuçlar, modellerin farklı veri alt kümeleri üzerinde gösterdikleri kararlılığı, genellenebilirliği ve sınıflar arası ayırt edicilik düzeylerini birlikte değerlendirmeye imkân tanımaktadır.

Elde edilen bulgular incelendiğinde, SVM (RBF) modelinin çapraz doğrulama sürecinde en yüksek ve en istikrarlı performansı sergilediği görülmektedir. Özellikle doğruluk, F1 skoru ve dengeli doğruluk değerlerinin birbirine yakın olması, modelin sınıf dağılımını etkin biçimde gözetten bir öğrenme süreci izlediğini göstermektedir. ROC-AUC ve PR-AUC değerleri, SVM modelinin karar sınırlarını güvenilir biçimde oluşturabildiğini ve pozitif sınıfa yönelik tahminlerde tutarlı sonuçlar ürettiğini desteklemektedir.

Random Forest modeli, SVM’ye oldukça yakın performans değerleri üretmiş ve özellikle ROC-AUC ve PR-AUC ölçütlerinde güçlü sonuçlar sunmuştur. Bu durum, modelin topluluk yapısı sayesinde farklı öznitelik etkileşimlerini etkili biçimde yakalayabildiğini ve sınıflar arasındaki ayrımı başarılı bir şekilde gerçekleştirdiğini göstermektedir.

KNN ve XGBoost modelleri orta düzeyde performans sergilemiş; dengeli doğruluk ve F1 skoru açısından kabul edilebilir sonuçlar üretmişlerdir. Ancak bu modellerin performanslarının, SVM ve Random Forest’a kıyasla daha değişken olduğu görülmektedir. MLP modeli ise genel olarak daha sınırlı bir başarı göstermiş, bu durum ağ yapısının veri dağılımına ve öznitelik yapısına daha duyarlı olabileceğini düşündürmektedir.

Decision Tree modeli, tüm performans ölçütleri dikkate alındığında en düşük sonuçları üretmiş ve özellikle ayırt edicilik ölçütlerinde sınırlı kalmıştır. Bu bulgu, tekil ağaç yapısının yüksek boyutlu ve karmaşık veri örüntülerini temsil etmekte zorlandığını göstermektedir.

Genel olarak bu çapraz doğrulama sonuçları, SVM (RBF) ve Random Forest modellerinin veri seti üzerinde daha kararlı, dengeli ve genellenebilir bir performans sunduğunu ortaya koymaktadır. Bu durum, nihai model seçimi sürecinde bu iki yöntemin öncelikli olarak değerlendirilmesini desteklemektedir.

Tablo 4.29: %60 Eğitim %40 Test verileri ile tüm modellerin test performans sonuçları (Deneme 3)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
SVM (RBF)	0.882550	0.892308	0.880753	0.904710	0.890833
Random Forest	0.855705	0.866044	0.854248	0.918571	0.910025
KNN	0.805369	0.833333	0.801442	0.895695	0.888317
XGBoost	0.815436	0.825397	0.814492	0.916295	0.901108
MLP	0.758389	0.791908	0.754609	0.820825	0.786505
Decision Tree	0.738255	0.741722	0.738427	0.738427	0.687832

Tablo 4.29’da dengelenmiş veri seti üzerinde eğitilen modellerin test veri kümesi performansları karşılaştırmalı olarak sunulmaktadır. Test sonuçları, modellerin öğrenme sürecinde elde ettikleri yapıyı yeni ve daha önce görülmemiş veriler üzerinde ne ölçüde koruyabildiklerini göstermesi açısından önem taşımaktadır.

Sonuçlar incelendiğinde, SVM (RBF) modelinin doğruluk, F1 skoru ve dengeli doğruluk ölçütlerinde en yüksek değerlere ulaştığı görülmektedir. Bu durum, modelin hem genel sınıflandırma başarısı hem de sınıflar arası dengeyi gözeten tahmin yeteneği açısından tutarlı bir performans sunduğunu göstermektedir. Özellikle dengeli doğruluk değerinin doğruluğa oldukça yakın olması, SVM modelinin her iki sınıfı da benzer başarı düzeylerinde ayırt edebildiğini ortaya koymaktadır.

Random Forest ve XGBoost modelleri ise ayırt edicilik ölçütleri bakımından dikkat çekici sonuçlar üretmiştir. ROC-AUC ve PR-AUC değerlerinin yüksek olması, bu modellerin

sınıflar arasındaki ayrımı güçlü bir karar yapısıyla gerçekleştirdiğini ve özellikle pozitif sınıfa ilişkin tahminlerde güvenilir sonuçlar sunduğunu göstermektedir. Buna karşın, genel doğruluk ve F1 skorları bakımından bu modellerin SVM'nin bir miktar gerisinde kaldığı görülmektedir.

KNN modeli, dengeli doğruluk ve ayırt edicilik açısından kabul edilebilir bir performans sergilemiş; ancak doğruluk ve F1 skoru bakımından daha karmaşık modellere kıyasla sınırlı kalmıştır. MLP ve Decision Tree modelleri ise tüm performans ölçütleri birlikte değerlendirildiğinde daha düşük sonuçlar üretmiş, bu durum söz konusu modellerin veri yapısındaki karmaşık ilişkileri yakalamakta zorlandığını düşündürmektedir.

Genel olarak test sonuçları, bu veri yapısı için SVM (RBF) modelinin en dengeli ve güvenilir performansı sunduğunu, Random Forest ve XGBoost modellerinin ise özellikle ayırt edicilik açısından güçlü alternatifler oluşturduğunu göstermektedir. Bu bulgular, nihai model seçimi sürecinde performans ölçütlerinin birlikte ele alınmasının önemini destekler niteliktedir.

4.3.3.1 Destek Vektör Makinesi (SVM) Sınıflandırma Modeli Sonuçları

Deneme 3'ün sınıflandırma sonuçlarına göre en iyi sonucu %70 eğitim %30 test verileri ile Destek Vektör Makinesi sınıflandırma modeli elde etmiştir.

Tablo 4.30: SVM en iyi parametreler (Deneme 3)

Parametre	Değer (%70 / %30)
C	1.9069966
Gamma	0.1382623
Kernel	RBF

Tablo 4.30'da %70 eğitim ve %30 test veri bölmesi kullanılarak Destek Vektör Makineleri (SVM) modeli için belirlenen en uygun hiperparametre değerleri sunulmaktadır. Modelde çekirdek (kernel) fonksiyonu olarak RBF (Radial Basis Function) tercih edilmiştir. RBF çekirdeği, doğrusal olmayan veri yapılarında sınıflar arasındaki karmaşık ilişkileri yakalayabilme yeteneği sayesinde özellikle yüksek boyutlu veri setlerinde etkili sonuçlar

üretmektedir.

C parametresinin 1.9069966 olarak belirlenmesi, modelin hata toleransı ile karar sınırının karmaşıklığı arasında dengeli bir yapı kurduğunu göstermektedir. Bu değer, modelin yanlış sınıflandırmaları tamamen sıfırlamaya çalışmadan, genellenebilirliği koruyacak ölçüde esneklik sağladığını ortaya koymaktadır. Aşırı yüksek C değerlerinin aşırı uyuma, düşük değerlerin ise yetersiz öğrenmeye yol açabildiği dikkate alındığında, elde edilen bu parametre değeri modelin dengeli bir öğrenme süreci izlediğini göstermektedir.

Gamma parametresinin 0.1382623 olması, her bir veri noktasının karar sınırı üzerindeki etki alanının orta düzeyde tutulduğunu ifade etmektedir. Bu durum, modelin karar sınırlarını aşırı karmaşık hâle getirmeden, yerel örüntüleri yeterli düzeyde yakalayabildiğini göstermektedir. Gamma değerinin uygun biçimde ayarlanması, özellikle RBF çekirdeği kullanılan modellerde genellenebilirliği doğrudan etkileyen kritik bir unsurdur.

Genel olarak bu hiperparametre kombinasyonu, SVM (RBF) modelinin karmaşıklık–genellenebilirlik dengesini başarıyla kurduğunu ve %70–%30 veri bölme senaryosunda elde edilen yüksek test performansını destekleyen bir yapı sunduğunu göstermektedir. Bu bulgular, SVM modelinin söz konusu veri yapısı için uygun bir sınıflandırıcı olduğunu ortaya koymaktadır.

Tablo 4.31: SVM 10 Fold Cross Validation (CV) sonuçları (Deneme 3)

Metrik	Ortalama $\pm$ Std (%70 / %30)
Accuracy	0.8579 $\pm$ 0.0348
F1 Score	0.8651 $\pm$ 0.0372
Balanced Accuracy	0.8572 $\pm$ 0.0342
ROC-AUC	0.8956 $\pm$ 0.0408
PR-AUC	0.8770 $\pm$ 0.0543

Tablo 4.31’de, %70 eğitim ve %30 test veri bölmesi kullanılarak SVM (RBF) modeli için elde edilen çapraz doğrulama performans sonuçları sunulmaktadır. Ortalama değerler ve standart sapmalar birlikte değerlendirildiğinde, modelin farklı veri alt kümeleri üzerinde

kararlı ve genellenebilir bir performans sergilediği görülmektedir.

Sonuçlar incelendiğinde, Accuracy (0.8579) ve Balanced Accuracy (0.8572) değerlerinin birbirine oldukça yakın olması, modelin sınıflar arasında tutarlı bir ayırım gerçekleştirdiğini ve sınıf dağılımından kaynaklı belirgin bir yanlılık üretmediğini göstermektedir. F1 Score'un 0.8651 düzeyinde olması, modelin hem yanlış pozitif hem de yanlış negatif tahminleri dengeli biçimde yönettiğini ve sınıflandırma performansının istikrarlı olduğunu ortaya koymaktadır.

ROC-AUC (0.8956) ve PR-AUC (0.8770) değerleri, SVM (RBF) modelinin sınıflar arasındaki ayrımı güçlü bir karar sınırıyla gerçekleştirdiğini ve özellikle pozitif sınıfa yönelik tahminlerde güvenilir sonuçlar ürettiğini göstermektedir. Standart sapma değerlerinin tüm metriklerde görece düşük seviyelerde olması, model performansının çapraz doğrulama katmanları arasında büyük dalgalanmalar göstermediğini ve öğrenme sürecinin kararlı bir yapı izlediğini ortaya koymaktadır.

Genel olarak bu sonuçlar, SVM (RBF) modelinin %70–%30 veri bölme senaryosunda yüksek doğruluk, dengeli sınıf ayrımı ve güçlü ayırt edicilik sunduğunu ve nihai model değerlendirmesi için güvenilir bir aday olduğunu göstermektedir.

Tablo 4.32: SVM Confusion Matrix (Deneme 3)

Gerçek \ Tahmin	%70 / %30	
	0	1
0 (Başarısız)	90	19
1 (Başarılı)	5	110

Tablo 4.32’de verilen karışıklık matrisi, %70 eğitim ve %30 test veri bölmesi kullanılarak eğitilen SVM (RBF) modelinin sınıflandırma performansını ayrıntılı biçimde göstermektedir. Sonuçlara göre, gerçek sınıfı başarısız (0) olan 109 öğrencinin 90’ı doğru biçimde başarısız olarak sınıflandırılmış, 19 öğrenci ise başarılı sınıfına yanlış atanmıştır. Bu durum, modelin başarısız öğrencileri tespit etme konusunda yüksek bir duyarlılık sergilediğini göstermektedir.

Gerçek sınıfı başarılı (1) olan 115 öğrencinin ise 110’u doğru şekilde başarılı olarak tahmin edilmiş, yalnızca 5 öğrenci başarısız olarak yanlış sınıflandırılmıştır. Bu bulgu, modelin başarılı öğrencileri ayırt etme yeteneğinin oldukça güçlü olduğunu ve yanlış negatif oranının düşük seviyede kaldığını ortaya koymaktadır.

Genel olarak değerlendirildiğinde, SVM (RBF) modelinin her iki sınıf için de dengeli ve tutarlı bir sınıflandırma performansı sunduğu görülmektedir. Yanlış sınıflandırma sayılarının sınırlı olması, modelin karar sınırlarını etkin biçimde oluşturduğunu ve önceki doğruluk, dengeli doğruluk ve F1 skoru sonuçlarını destekler nitelikte bir performans sergilediğini göstermektedir. Bu sonuçlar, SVM (RBF) modelinin %70–%30 veri bölme senaryosunda güvenilir ve genellenebilir bir sınıflandırıcı olduğunu ortaya koymaktadır.

Tablo 4.33: SVM %70 Eğitim %30 Test verileri ile test sınıflandırma raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9474	0.8257	0.8824	109
1	0.8527	0.9565	0.9016	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8929	224
Macro Avg	0.9000	0.8911	0.8920	224
Weighted Avg	0.8988	0.8929	0.8923	224

Tablo 4.33’te, SVM (RBF) modelinin test veri kümesi üzerindeki sınıf bazlı precision, recall ve F1-score değerleri sunulmaktadır. Bu ölçütler, modelin her bir sınıfı ne ölçüde doğru tanımladığını ve hata türlerini nasıl yönettiğini ayrıntılı biçimde ortaya koymaktadır.

Sonuçlar incelendiğinde, başarısız (0) sınıfı için precision değerinin 0.9474 olması, modelin başarısız olarak etiketlediği öğrencilerin çok büyük bir bölümünün gerçekten başarısız olduğunu göstermektedir. Bu durum, yanlış pozitiflerin oldukça sınırlı olduğunu ve modelin bu sınıf için seçici bir karar yapısı geliştirdiğini ortaya koymaktadır. Recall değerinin 0.8257 seviyesinde olması ise başarısız öğrencilerin büyük kısmının doğru biçimde tespit edildiğini, ancak sınırlı sayıda örneğin gözden kaçtığını göstermektedir. Bu iki ölçütün birleşimiyle elde edilen F1-score (0.8824), modelin bu sınıf için dengeli ve güçlü bir performans sunduğunu göstermektedir.

Başarılı (1) sınıfı açısından değerlendirildiğinde, recall değerinin 0.9565 olması, başarılı öğrencilerin neredeyse tamamının model tarafından doğru biçimde tanımlandığını ortaya koymaktadır. Precision değerinin 0.8527 olması, başarılı olarak tahmin edilen öğrencilerin büyük bir bölümünün gerçekten başarılı olduğunu göstermektedir. Bu sınıf için hesaplanan F1-score (0.9016), modelin başarılı öğrencileri ayırt etme konusunda oldukça yüksek bir doğruluk sergilediğini göstermektedir.

Genel performans değerlendirildiğinde, accuracy değerinin 0.8929 olması, modelin test veri kümesi üzerinde yüksek bir genel sınıflandırma başarısı sunduğunu göstermektedir. Macro average ve weighted average değerlerinin birbirine oldukça yakın olması, model performansının sınıflar arasında dengeli olduğunu ve sınıf dağılımından kaynaklı belirgin bir yanlılık bulunmadığını ortaya koymaktadır. Bu bulgular, SVM (RBF) modelinin %70–%30 veri bölme senaryosunda hem çoğunluk hem de azınlık sınıfı için güvenilir, dengeli ve genellenebilir bir performans sergilediğini göstermektedir.

Tablo 4.34: SVM genel test sınıflandırma sonuçları (Deneme 3)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8929	0.9016	0.8911	0.9325	0.9233

Tablo 4.34’te %70 eğitim ve %30 test veri bölmesi kullanılarak elde edilen performans sonuçları, SVM (RBF) modelinin sınıflandırma başarımının yüksek ve dengeli olduğunu ortaya koymaktadır. Accuracy değerinin 0.8929 olması, modelin test veri kümesi üzerinde yüksek bir genel doğruluk sağladığını göstermektedir. F1 Score’un 0.9016 seviyesinde gerçekleşmesi, modelin hem yanlış pozitif hem de yanlış negatif hataları etkin biçimde dengelediğini ve sınıflandırma sürecinde tutarlı bir performans sunduğunu ortaya koymaktadır.

Balanced Accuracy değerinin 0.8911 olması, modelin her iki sınıfı da benzer başarı düzeylerinde ayırt edebildiğini ve sınıf dağılımından kaynaklanan yanlılıkların büyük ölçüde giderildiğini göstermektedir. ROC-AUC (0.9325) değeri, SVM modelinin sınıflar arasında güçlü bir ayırt edicilik sağladığını ortaya koyarken, PR-AUC (0.9233) değeri özellikle pozitif sınıfa yönelik tahminlerin yüksek güvenilirlik düzeyine sahip olduğunu göstermektedir.

Genel olarak bu performans değerleri, SVM (RBF) modelinin söz konusu veri yapısı için yüksek doğruluk, dengeli sınıf ayrımı ve güçlü ayırt edicilik sunduğunu ve akademik performans tahmini problemi kapsamında nihai model olarak tercih edilebilecek düzeyde bir performans sergilediğini göstermektedir.

4.3.4 Sınıflandırma Sonuçları (Deneme 4)

Veri seti; ADASYN veri arttırımı yöntemi ile dengelenmiş, kategorik veriler one-hot encoding yöntemi ile sayısallaştırılmış, sayısal veriler 0-1 aralığına normalize edilmiş, %70 eğitim %30 test verisi (Tablo 4.36) ve %60 eğitim %40 test verisine (Tablo 4.37) ayrılarak, 300 iterasyonlu RandomSearchCV içerisinde Stratified 10-Fold Crossover kullanılarak test edilmiştir ve sınıflandırma modellerinin en iyi parametreler ile elde ettikleri en iyi sonuçlar verilmiştir. One-Hot kodlama sonucunda öznitelik sayısı 117'ye çıkmıştır.

Tablo 4.35: Veri Arttırımı yapılmış, %70 Eğitim %30 Test verileri ve sınıf sayıları (Deneme 4)

Veriseti Aşaması	Satır Sayısı	Sütun Sayısı	Başarılı (True)	Başarısız (False)
Orijinal Veriseti	528	117	382	146
Oversampled Veriseti	752	117	382	370
Eğitim (Train) Veriseti	526	117	267	259
Test Veriseti	226	117	115	111

Tablo 4.35'te, veri setinin modelleme süreci boyunca geçirdiği aşamalar ve her aşamadaki sınıf dağılımları ayrıntılı olarak sunulmaktadır. Orijinal veri seti 528 gözlem ve 117 öznitelikten oluşmakta olup, başarılı (382) ve başarısız (146) sınıfları arasında belirgin bir dengesizlik bulunmaktadır. Bu durum, sınıflandırma modellerinin çoğunluk sınıfına yönelme riskini artırabilecek bir başlangıç yapısı ortaya koymaktadır.

Bu dengesizliği azaltmak amacıyla uygulanan oversampling işlemi sonrasında veri seti 752 gözleme çıkarılmış ve başarısız sınıfa ait gözlem sayısı artırılarak sınıflar arasındaki dağılım büyük ölçüde dengelenmiştir. Bu aşamada başarılı sınıf gözlem sayısı korunurken, başarısız

sınıfın 370 gözleme yükseltilmesi, her iki sınıfın öğrenme sürecine daha eşit katkı sağlamasını mümkün kılmıştır.

Dengelenmiş veri seti, model eğitimi ve değerlendirmesi amacıyla eğitim ve test kümelerine ayrılmıştır. Buna göre eğitim veri seti 526 gözlemden (267 başarılı, 259 başarısız), test veri seti ise 226 gözlemden (115 başarılı, 111 başarısız) oluşmaktadır. Eğitim ve test kümelerinde sınıfların neredeyse eşit sayıda temsil edilmesi, modellerin hem öğrenme hem de test aşamalarında sınıf yanlılığından arındırılmış bir veri yapısı ile çalışmasını sağlamıştır.

Genel olarak bu tablo, veri setinin ön işleme, dengeleme ve bölme adımlarının sistematik ve kontrollü biçimde yürütüldüğünü göstermektedir. Oluşturulan bu dengeli yapı, geliştirilen makine öğrenmesi modellerinin özellikle her iki sınıf için de genellenebilir, güvenilir ve karşılaştırılabilir performanslar üretmesini destekleyen sağlam bir yöntemsel zemin sunmaktadır.

Tablo 4.36: Veri Arttırımı yapılmış, %60 Eğitim %40 Test verileri ve sınıf sayıları (Deneme 4)

Veriseti Aşaması	Satır Sayısı	Sütun Sayısı	Başarılı (True)	Başarısız (False)
Orijinal Veriseti	528	117	382	146
Oversampled Veriseti	752	117	382	370
Eğitim (Train) Veriseti	451	117	229	222
Test Veriseti	301	117	153	148

Tablo 4.36’da veri setinin dengeleme ve bölme işlemleri sonrasında aldığı son yapı ve sınıf dağılımları sunulmaktadır. Orijinal veri seti 528 gözlem ve 117 öznitelikten oluşmakta olup, başarılı (382) ve başarısız (146) sınıfları arasında belirgin bir dengesizlik bulunmaktadır. Bu dengesizlik, sınıflandırma modellerinin çoğunluk sınıfına yönelmesine neden olabilecek bir başlangıç durumu oluşturmaktadır.

Bu sorunu gidermek amacıyla uygulanan oversampling işlemi sonucunda veri seti 752 gözleme çıkarılmış ve başarısız sınıfa ait gözlem sayısı artırılarak sınıflar arasındaki dağılım

büyük ölçüde dengelenmiştir. Bu aşamada başarılı sınıf korunurken, başarısız sınıfın 370 gözleme yükseltilmesi, her iki sınıfın öğrenme sürecine daha eşit katkı sağlamasını mümkün kılmıştır.

Dengelenmiş veri seti, model eğitimi ve değerlendirme süreçlerinde kullanılmak üzere yaklaşık %60 eğitim ve %40 test olacak şekilde iki alt kümeye ayrılmıştır. Buna göre eğitim veri seti 451 gözlem (229 başarılı, 222 başarısız), test veri seti ise 301 gözlem (153 başarılı, 148 başarısız) oluşmaktadır. Eğitim ve test kümelerindeki sınıf dağılımlarının birbirine oldukça yakın olması, modellerin hem öğrenme hem de değerlendirme aşamalarında sınıf yanlılığından arındırılmış bir veri yapısı ile çalışmasını sağlamıştır.

Genel olarak bu tablo, veri setinin ön işleme, dengeleme ve eğitim–test bölme adımlarının kontrollü ve sistematik biçimde yürütüldüğünü göstermektedir. Elde edilen bu dengeli yapı, geliştirilen makine öğrenmesi modellerinin her iki sınıf için de genellenebilir ve güvenilir performanslar üretmesine uygun bir yöntemsel zemin sunmaktadır.

Tablo 4.37: %70 Eğitim %30 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 4)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.857910	0.865051	0.857182	0.895611	0.877010
Random Forest	0.861720	0.867420	0.861396	0.914305	0.908441
KNN	0.846372	0.835304	0.848726	0.886860	0.887555
XGBoost	0.810015	0.814228	0.809943	0.882660	0.871924
MLP	0.800327	0.820821	0.798775	0.849759	0.825285
Decision Tree	0.765856	0.769429	0.765299	0.805802	0.796004

Tablo 4.37’de, dengelenmiş veri seti üzerinde uygulanan farklı makine öğrenmesi algoritmalarının çapraz doğrulama (CV) performansları karşılaştırmalı olarak sunulmaktadır. CV sonuçları, modellerin yalnızca tek bir veri bölmesi üzerindeki başarısını değil, farklı alt örneklemeler üzerindeki tutarlılık ve genellenebilirlik düzeylerini ortaya koymaktadır.

Sonuçlar incelendiğinde, Random Forest modelinin tüm performans ölçütleri birlikte değerlendirildiğinde en yüksek ve en dengeli sonuçları ürettiği görülmektedir. Özellikle doğruluk, F1 skoru ve dengeli doğruluk değerlerinin birbirine oldukça yakın olması, modelin her iki sınıfı da benzer başarı düzeylerinde öğrendiğini göstermektedir. Buna ek olarak ROC-AUC ve PR-AUC değerlerinin yüksekliği, Random Forest modelinin sınıflar arasındaki ayrımı güçlü ve istikrarlı bir biçimde gerçekleştirdiğini ortaya koymaktadır.

SVM (RBF) modeli, Random Forest’a çok yakın performans değerleri üretmiş ve özellikle dengeli doğruluk ölçütünde dikkat çekici bir başarı sergilemiştir. Bu durum, SVM modelinin sınıf dağılımını etkin biçimde dikkate alan bir karar sınırı oluşturabildiğini göstermektedir. ROC-AUC ve PR-AUC değerleri de modelin ayırt edicilik açısından güçlü bir performans sunduğunu desteklemektedir.

KNN modeli, doğruluk ve dengeli doğruluk açısından rekabetçi sonuçlar üretmiş olmakla birlikte, F1 skorunun görece daha düşük olması, sınıflardan biri için tahmin performansının sınırlı kaldığını düşündürmektedir. XGBoost ve MLP modelleri orta düzeyde performans sergilemiş; özellikle doğruluk ve dengeli doğruluk ölçütlerinde üst düzey modellere kıyasla daha düşük sonuçlar elde etmiştir. Decision Tree modeli ise tüm metrikler dikkate alındığında en düşük performansı sergilemiş ve tekil ağaç yapısının karmaşık ve yüksek boyutlu veri yapısını temsil etmekte yetersiz kaldığını göstermiştir.

Genel olarak çapraz doğrulama sonuçları, ensemble tabanlı Random Forest ve çekirdek tabanlı SVM (RBF) modellerinin bu veri yapısı üzerinde daha kararlı, dengeli ve genellenebilir performanslar ürettiğini ortaya koymaktadır. Bu bulgular, nihai model değerlendirme sürecinde bu iki yöntemin öncelikli olarak ele alınmasını desteklemektedir.

Tablo 4.38: %70 Eğitim %30 Test verileri ile tüm modellerin test performans sonuçları (Deneme 4)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
SVM (RBF)	0.907080	0.912863	0.906189	0.938895	0.929390
Random Forest	0.898230	0.902954	0.897650	0.943831	0.942801
KNN	0.884956	0.880734	0.885860	0.921778	0.911925

Tablo 4.38: (devam ediyor)

XGBoost	0.867257	0.873950	0.866588	0.938347	0.922940
MLP	0.831858	0.834783	0.831806	0.914297	0.888042
Decision Tree	0.743363	0.738739	0.743909	0.786800	0.796896

Tablo 4.38’de dengelenmiş veri seti kullanılarak eğitilen makine öğrenmesi modellerinin test veri kümesi üzerindeki performansları sunulmaktadır. Elde edilen sonuçlar, modellerin çapraz doğrulama sürecinde öğrenilen yapıyı daha önce görülmemiş veriler üzerinde ne ölçüde başarıyla genelleyebildiğini ortaya koymaktadır.

Sonuçlar incelendiğinde, SVM (RBF) ve Random Forest modellerinin test aşamasında belirgin biçimde öne çıktığı görülmektedir. SVM (RBF) modeli, en yüksek doğruluk (0.9071) ve F1 skoru (0.9129) değerleriyle genel sınıflandırma başarısı açısından ilk sırada yer almıştır. Dengeli doğruluk değerinin doğruluğa oldukça yakın olması, modelin her iki sınıfı da benzer başarı düzeylerinde ayırt edebildiğini göstermektedir. Buna ek olarak ROC-AUC ve PR-AUC değerleri, SVM modelinin güçlü bir ayırt edicilik sunduğunu ortaya koymaktadır.

Random Forest modeli ise özellikle ROC-AUC (0.9438) ve PR-AUC (0.9428) değerleriyle sınıflar arasındaki ayrımı en başarılı biçimde gerçekleştiren yöntem olarak dikkat çekmektedir. Bu durum, modelin karmaşık öznelilik etkileşimlerini etkili biçimde yakalayabildiğini ve karar sınırlarını esnek şekilde oluşturabildiğini göstermektedir. Doğruluk ve F1 skoru bakımından SVM’ye oldukça yakın sonuçlar üretmesi, Random Forest’i güçlü bir alternatif hâline getirmektedir.

KNN ve XGBoost modelleri, genel olarak yüksek ve dengeli performanslar sergilemiş olmakla birlikte, üst düzey iki modelin gerisinde kalmıştır. Bu modellerin ROC-AUC ve PR-AUC değerlerinin yüksek olması, ayırt edicilik açısından güçlü olduklarını; ancak doğruluk ve F1 skorlarının nispeten daha düşük olması, karar eşikleri düzeyinde daha sınırlı bir performans sunduklarını göstermektedir. MLP modeli ise kabul edilebilir düzeyde sonuçlar üretmiş, ancak yüksek boyutlu ve karmaşık veri yapısında ensemble ve kernel tabanlı yöntemlerin gerisinde kalmıştır. Decision Tree modeli ise tüm ölçütler dikkate alındığında en düşük performansı sergilemiş ve tekil ağaç yapısının bu veri yapısı için yeterli olmadığını

ortaya koymuştur.

Genel değerlendirme sonucunda, test performansları SVM (RBF) modelinin genel sınıflandırma başarısı ve denge açısından, Random Forest modelinin ise ayırt edicilik gücü açısından öne çıktığını göstermektedir. Bu bulgular, nihai model seçiminde performans ölçütlerinin birlikte ele alınmasının önemini vurgulamakta ve her iki modelin de problem bağlamına göre güçlü adaylar olduğunu ortaya koymaktadır.

Tablo 4.39: %60 Eğitim %40 Test verileri ile tüm modellerin Çapraz Doğrulama (CV) performans sonuçları (Deneme 4)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.869324	0.876746	0.868577	0.888271	0.862026
Random Forest	0.849324	0.855113	0.849012	0.899699	0.891302
KNN	0.838164	0.835722	0.838636	0.896559	0.902576
XGBoost	0.811498	0.816601	0.811166	0.887455	0.882690
MLP	0.789469	0.798580	0.789130	0.857020	0.837960
Decision Tree	0.733816	0.739823	0.734190	0.766034	0.741469

Tablo 4.39’da dengelenmiş veri seti üzerinde uygulanan makine öğrenmesi modellerinin çapraz doğrulama (CV) sonuçları sunulmaktadır. Bu sonuçlar, modellerin farklı veri alt kümeleri üzerinde ne ölçüde istikrarlı ve genellenebilir performans sergilediğini ortaya koymaktadır.

Elde edilen bulgulara göre SVM (RBF) modeli, doğruluk, F1 skoru ve dengeli doğruluk ölçütlerinde en yüksek değerlere ulaşarak genel performans açısından öne çıkmıştır. Bu durum, modelin sınıflar arasındaki karar sınırlarını tutarlı biçimde oluşturabildiğini ve veri alt kümeleri arasında performans kaybı yaşamadığını göstermektedir. ROC-AUC ve PR-AUC değerlerinin yüksek olması da SVM modelinin ayırt edicilik gücünün güçlü olduğunu desteklemektedir.

Random Forest modeli, özellikle ROC-AUC ve PR-AUC ölçütlerinde yüksek değerler üretmiş ve sınıflar arasındaki ayrımı başarılı biçimde gerçekleştirdiğini göstermiştir.

Bununla birlikte doğruluk ve dengeli doğruluk değerlerinin SVM'ye kıyasla daha düşük olması, modelin bazı alt örneklerde sınıf dengesini tam olarak koruyamadığını düşündürmektedir. Buna rağmen Random Forest, genel olarak kararlı ve güçlü bir ensemble yaklaşımı sunduğunu ortaya koymaktadır.

KNN modeli, özellikle PR-AUC değerinde yüksek bir başarı sergileyerek pozitif sınıf tahminlerinde güvenilir sonuçlar üretmiştir. Ancak doğruluk ve F1 skorunun görece daha düşük olması, komşuluk temelli yaklaşımın yüksek boyutlu veri yapısında sınırlı kaldığını göstermektedir. XGBoost ve MLP modelleri orta düzey performanslar sunmuş; karar ağaçları temelli güçlendirme ve sinir ağı yaklaşımlarının bu veri yapısında üst düzey modellere kıyasla daha sınırlı bir genellenebilirlik sağladığı görülmüştür. Decision Tree modeli ise tüm ölçütler birlikte değerlendirildiğinde en düşük CV performansını sergilemiştir.

Genel olarak bu çapraz doğrulama sonuçları, SVM (RBF) modelinin bu veri seti üzerinde en dengeli ve istikrarlı genelleme performansını sunduğunu göstermektedir. Elde edilen bulgular, test aşamasında da benzer bir performans beklentisini desteklemekte ve SVM modelini nihai değerlendirme sürecinde güçlü bir aday hâline getirmektedir.

Tablo 4.40: %60 Eğitim %40 Test verileri ile tüm modellerin test performans sonuçları (Deneme 4)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
SVM (RBF)	0.873754	0.884146	0.872505	0.924307	0.911139
Random Forest	0.877076	0.882540	0.876546	0.931814	0.916401
KNN	0.870432	0.874598	0.870120	0.928524	0.931111
XGBoost	0.853821	0.858065	0.853559	0.929341	0.903804
MLP	0.823920	0.831746	0.823375	0.887741	0.858776
Decision Tree	0.774086	0.770270	0.774576	0.793941	0.761865

Tablo 4.40'ta, dengelenmiş veri seti kullanılarak eğitilen makine öğrenmesi modellerinin test veri kümesi üzerindeki performansları karşılaştırmalı olarak sunulmaktadır. Elde edilen sonuçlar, modellerin eğitim aşamasında öğrendikleri örüntüleri daha önce görülmemiş veriler üzerinde büyük ölçüde başarıyla genellebildiğini göstermektedir.

Sonuçlar incelendiğinde, Random Forest, SVM (RBF) ve KNN modellerinin test aşamasında birbirine oldukça yakın ve yüksek performans değerleri ürettiği görülmektedir. Random Forest modeli doğruluk, dengeli doğruluk ve ROC-AUC ölçütlerinde hafif bir üstünlük sağlarken, SVM (RBF) modeli F1 skoru bakımından güçlü bir denge sunmuştur. KNN modelinin özellikle PR-AUC değerinin yüksek olması, pozitif sınıfa yönelik tahminlerde güvenilir sonuçlar üretebildiğini göstermektedir.

XGBoost modeli, ayırt edicilik ölçütleri açısından güçlü sonuçlar üretmiş olsa da doğruluk ve F1 skorlarında üst düzey modellere kıyasla daha sınırlı bir performans sergilemiştir. MLP modeli kabul edilebilir düzeyde sonuçlar sunmakla birlikte, ensemble ve kernel tabanlı yöntemlerin gerisinde kalmıştır. Decision Tree modeli ise tüm performans ölçütleri birlikte değerlendirildiğinde en düşük test başarımını sergilemiş ve tekil ağaç yapısının bu veri yapısı için yeterli olmadığını göstermiştir.

Genel olarak test sonuçları, ensemble tabanlı Random Forest ile çekirdek tabanlı SVM (RBF) modellerinin bu problem kapsamında yüksek doğruluk, dengeli sınıf ayrımı ve güçlü ayırt edicilik sunduğunu ortaya koymaktadır. Bu bulgular, nihai model seçiminde bu iki yöntemin öncelikli olarak değerlendirilmesini desteklemektedir.

4.3.4.1 Destek Vektör Makinesi (SVM) Sınıflandırma Modeli Sonuçları

Deneme 4'ün sınıflandırma sonuçlarına göre en iyi sonucu %70 eğitim %30 test verileri ile Destek Vektör Makinesi sınıflandırma modeli elde etmiştir.

Tablo 4.41: SVM en iyi parametreler (Deneme 4)

Parametre	Değer (%70 / %30)
C	56.6335
Gamma	0.1687869
Kernel	RBF

Tablo 4.41'de, %70 eğitim ve %30 test veri bölmesi kullanılarak eğitilen Destek Vektör Makineleri (SVM) modeline ait en iyi hiperparametreler sunulmaktadır. Modelde RBF

(Radial Basis Function) çekirdeğinin tercih edilmesi, veri yapısındaki doğrusal olmayan ilişkilerin etkin biçimde modellenmesini sağlamıştır. C parametresinin 56.6335 olarak belirlenmesi, modelin eğitim hatalarını minimize etme yönünde daha katı bir karar sınırı oluşturduğunu göstermektedir. Bu durum, modelin marj genişliği ile sınıflandırma hataları arasındaki dengeyi yüksek doğruluk lehine ayarladığını ortaya koymaktadır.

Gamma değerinin 0.1687869 olması ise her bir veri noktasının karar sınırı üzerindeki etki alanının kontrollü biçimde belirlendiğini göstermektedir. Bu parametre, modelin aşırı uyum göstermeden karmaşık örüntüleri yakalayabilmesine olanak tanımaktadır. Elde edilen hiperparametre kombinasyonu, SVM (RBF) modelinin test ve çapraz doğrulama aşamalarında gözlenen yüksek doğruluk, dengeli doğruluk ve ayırt edicilik performanslarını destekleyen bir yapı sunduğunu ortaya koymaktadır.

Tablo 4.42: SVM 10 Fold Cross Validation (CV) sonuçları (Deneme 4)

Metrik	Ortalama $\pm$ Std (%70 / %30)
Accuracy	0.8668 $\pm$ 0.0421
F1 Score	0.8743 $\pm$ 0.0382
Balanced Accuracy	0.8661 $\pm$ 0.0423
ROC-AUC	0.8854 $\pm$ 0.0499
PR-AUC	0.8663 $\pm$ 0.0787

Tablo 4.42’de sunulan çapraz doğrulama sonuçları, SVM (RBF) modelinin farklı veri alt kümeleri üzerinde istikrarlı ve dengeli bir performans sergilediğini göstermektedir. Accuracy (0.8668 $\pm$ 0.0421) ve Balanced Accuracy (0.8661 $\pm$ 0.0423) değerlerinin birbirine oldukça yakın olması, modelin her iki sınıfı da benzer başarı düzeylerinde ayırt edebildiğini ve sınıf dengesizliğinden kaynaklı yanlılık üretmediğini ortaya koymaktadır.

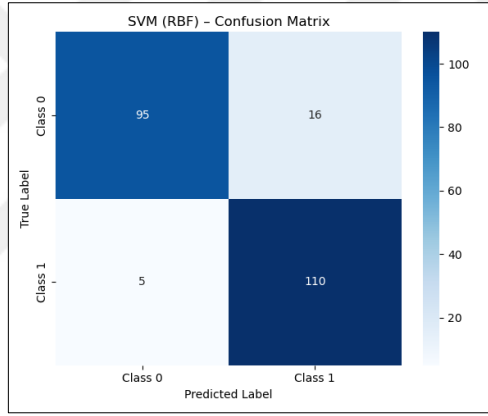
F1 Score’un 0.8743 $\pm$ 0.0382 seviyesinde olması, modelin yanlış pozitif ve yanlış negatif hataları dengeli biçimde yönettiğini göstermektedir. ROC-AUC (0.8854 $\pm$ 0.0499) değeri, SVM modelinin sınıflar arasında güçlü bir ayırt edicilik sunduğunu; PR-AUC (0.8663 $\pm$ 0.0787) değeri ise özellikle pozitif sınıfa yönelik tahminlerin güvenilir olduğunu ortaya koymaktadır. Standart sapma değerlerinin görece düşük olması, model performansının katlar

arasında tutarlı olduğunu ve aşırı uyum eğilimi göstermediğini desteklemektedir.

Genel olarak bu bulgular, SVM (RBF) modelinin %70–%30 veri bölme senaryosunda kararlı, genellenebilir ve dengeli bir sınıflandırma performansı sunduğunu ve test aşamasında elde edilen yüksek başarıyı yöntemsel olarak desteklediğini göstermektedir.

Tablo 4.43: SVM Confusion Matrix (Deneme 4)

%70 / %30		
Gerçek \ Tahmin	0	1
0 (Başarısız)	95	16
1 (Başarılı)	5	110



Şekil 4.3: SVM Confusion Matrix (Deneme 4)

Tablo 4.43 ve Şekil 4.3'te verilen karışıklık matrisi, SVM (RBF) modelinin test veri kümesi üzerindeki sınıflandırma davranışını ayrıntılı biçimde ortaya koymaktadır. Sonuçlara göre, gerçek sınıfı başarısız (0) olan 111 öğrencinin 95'i doğru biçimde başarısız olarak tahmin edilirken, 16 öğrenci başarılı sınıfına yanlış atanmıştır. Bu durum, modelin başarısız öğrencileri tespit etme konusunda yüksek bir doğruluk sergilediğini göstermektedir.

Gerçek sınıfı başarılı (1) olan 115 öğrencinin ise 110'u doğru biçimde başarılı olarak sınıflandırılmış, yalnızca 5 öğrenci başarısız olarak hatalı tahmin edilmiştir. Bu sonuç, modelin başarılı öğrencileri ayırt etme yeteneğinin oldukça güçlü olduğunu ve yanlış negatif oranının düşük seviyede kaldığını ortaya koymaktadır.

Genel olarak değerlendirildiğinde, yanlış sınıflandırma sayılarının sınırlı olması ve her iki sınıf için de doğru tahmin oranlarının yüksekliği, SVM (RBF) modelinin test veri kümesi üzerinde dengeli, güvenilir ve genellenebilir bir performans sergilediğini göstermektedir. Bu bulgular, modelin önceki doğruluk, F1 skoru ve dengeli doğruluk sonuçlarını destekler niteliktedir ve nihai model değerlendirmesinde güçlü bir aday olduğunu ortaya koymaktadır.

Tablo 4.44: SVM %70 Eğitim %30 Test verileri ile test sınıflandırma raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9500	0.8559	0.9005	111
1	0.8730	0.9565	0.9129	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.9071	226
Macro Avg	0.9115	0.9062	0.9067	226
Weighted Avg	0.9108	0.9071	0.9068	226

Tablo 4.44'te verilen sınıflandırma raporu, SVM (RBF) modelinin her iki sınıf için de yüksek ve dengeli bir performans sergilediğini göstermektedir. Başarısız (0) sınıfı için precision değerinin 0.9500 olması, başarısız olarak etiketlenen öğrencilerin büyük çoğunluğunun gerçekten bu sınıfa ait olduğunu ve yanlış pozitif oranının düşük kaldığını ortaya koymaktadır. Recall değerinin 0.8559 olması ise başarısız öğrencilerin önemli bir bölümünün doğru biçimde tespit edildiğini, sınırlı sayıda örneğin gözden kaçtığını göstermektedir. Bu iki ölçütün birleşimiyle elde edilen F1-score (0.9005), modelin bu sınıf için güçlü ve dengeli bir tahmin yeteneğine sahip olduğunu göstermektedir.

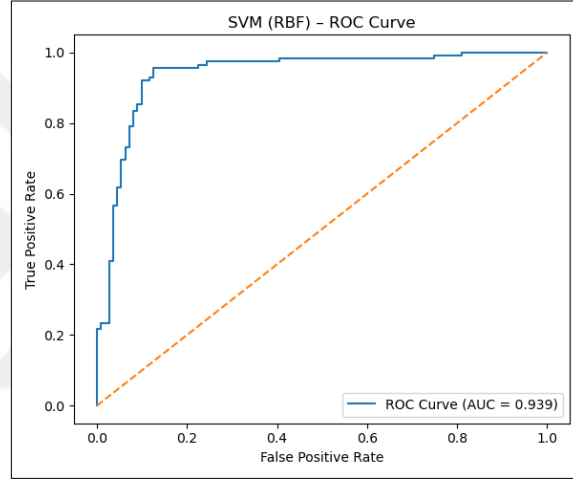
Başarılı (1) sınıfı açısından değerlendirildiğinde, recall değerinin 0.9565 olması, başarılı öğrencilerin neredeyse tamamının model tarafından doğru biçimde sınıflandırıldığını göstermektedir. Precision değerinin 0.8730 olması, başarılı olarak tahmin edilen öğrencilerin büyük bölümünün gerçekten başarılı olduğunu ortaya koymaktadır. Bu sınıf için hesaplanan F1-score (0.9129), modelin pozitif sınıfı ayırt etme konusunda oldukça yüksek bir başarı sergilediğini göstermektedir.

Genel performans ölçütleri incelendiğinde, accuracy değerinin 0.9071 olması modelin test veri kümesi üzerinde yüksek bir genel doğruluk sağladığını ortaya koymaktadır. Macro

average ve weighted average değerlerinin birbirine çok yakın olması, performansın sınıflar arasında dengeli olduğunu ve sınıf dağılımından kaynaklanan belirgin bir yanlılık bulunmadığını göstermektedir. Bu sonuçlar, SVM (RBF) modelinin bu veri yapısı üzerinde güvenilir, dengeli ve genellenebilir bir sınıflandırma performansı sunduğunu ve nihai model olarak tercih edilmesini destekleyen güçlü kanıtlar sunduğunu göstermektedir.

Tablo 4.45: SVM Genel Test sınıflandırma sonuçları (Deneme 4)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.9071	0.9129	0.9062	0.9389	0.9294



Şekil 4.4: SVM ROC Eğrisi (Deneme 4)

Tablo 4.45’te %70 eğitim ve %30 test veri bölmesi kullanılarak elde edilen performans sonuçları, SVM (RBF) modelinin bu çalışma kapsamında yüksek doğruluk, dengeli sınıf ayırımı ve güçlü ayırt edicilik sunduğunu göstermektedir. Accuracy değerinin 0.9071 olması, modelin test veri kümesi üzerinde yüksek bir genel sınıflandırma başarısı sağladığını ortaya koymaktadır. F1 Score’un 0.9129 seviyesinde gerçekleşmesi, modelin yanlış pozitif ve yanlış negatif hataları etkin biçimde dengelediğini ve sınıflandırma sürecinde tutarlı bir performans sergilediğini göstermektedir.

Balanced Accuracy değerinin 0.9062 olması, modelin her iki sınıfı da benzer başarı düzeylerinde ayırt edebildiğini ve sınıf dağılımından kaynaklanan olası yanlılıkların büyük ölçüde giderildiğini göstermektedir. Ayrıca ROC-AUC (0.9389) ve PR-AUC (0.9294) değerleri, SVM (RBF) modelinin sınıflar arasında güçlü bir ayırt edicilik sağladığını ve

özellikle pozitif sınıfa yönelik tahminlerde yüksek güvenilirlik sunduğunu ortaya koymaktadır (Şekil 4.4).

Genel olarak bu sonuçlar, SVM (RBF) modelinin akademik performans tahmini probleminde güvenilir, dengeli ve genellenebilir bir yapı sunduğunu ve çalışmanın amaçları doğrultusunda nihai model olarak tercih edilmesini güçlü biçimde desteklediğini göstermektedir.

4.4 Öğrenci Performansının Tahmini

Bu aşamada öznitelikler kullanılarak **“Başarı Notu”** tahmin edilmiştir. Başarı Notu; 0 ile 100 arasında sayısal bir değer olarak vize ve final notlarından hesaplanmaktadır ve öğrencinin performansını gösteren yegâne değerlerden bir tanesidir.

Başarı notunun tahmin edilmesi regresyon problemidir. Bu sebeple regresyon modellerinden Karar Ağaçları, Destek Vektör, Rastgele Orman, K-En Yakın Komşu, Çok Katmanlı Perceptron ve XGBoost Regresyon modelleri test edilmiştir.

Sınıflandırma probleminde karşılaşılan veri dengesizliği sorunu regresyonda da karşımıza çıkmaktadır. Sınıflandırma problemleri için geliştirilen SMOTE, ADASYN benzeri yöntemler, sürekli hedef değişken varsayımı nedeniyle regresyon problemlerine doğrudan uygulanamamaktadır. Bu sebeple regresyon problemlerine olarak geliştirilen SMOGN (Synthetic Minority Over-sampling Technique for Regression with Gaussian Noise) yöntemi, regresyon problemlerinde dengesiz hedef dağılımını gidermeyi amaçlayan özel bir oversampling yaklaşımıdır (Branco vd., 2017).

SMOGN yöntemi, öncelikle hedef değişkenin dağılımını analiz ederek nadir (rare) değerleri belirlemektedir. Bu değerler genellikle hedef değişkenin uç bölgelerinde yer almakta olup, öğrenme sürecinde yeterince temsil edilmemektedir. Daha sonra bu nadir gözlemler için, özellik uzayında en yakın komşular dikkate alınarak sentetik örnekler üretilir. SMOTE yaklaşımından farklı olarak SMOGN, üretilen sentetik örnekler Gauss gürültüsü (Gaussian noise) ekleyerek hem veri çeşitliliğini artırmakta hem de sürekli hedef değişkenin doğasına daha uygun örnekler oluşturmaktadır (Branco vd., 2017).

Veriseti regresyon modellerine giriş olarak verilmeden önce SMOGN veri artırımı (oversampling) yöntemine tabi tutularak dengesizlik giderilmiştir.

Öğrenci performans sınıflandırmasına benzer olarak modellerin en iyi parametreleri, 10-Fold Cross Validation ve Randomized Search CV ile belirlenmiştir. Parametrelerin çoğunluğu sınıflandırma modelindeki parametreler ile aynı olmaktadır fakat, karar ağaçları parametrelerinden “criterion” parametresinin değerleri “Friedman_mse, absolute_error, poisson, squared_error” biçimindedir. Ayrıca, karar ağaçlarında sınıflandırmadan farklı olarak ccp_alpha parametresi de 0.0 ile 0.02 aralığında değerler kullanılmıştır. Bunun dışında farklı olarak, Destek Vektör Regresyonunda diğer parametrelere ek olarak “epsilon” parametresi 1-4 ile 1-1 aralığında denenmiştir.

4.4.1 Regresyon Sonuçları (Deneme 1)

İlk denemede veriseti, SMOGN veri artırma yöntemi ile dengelenmiş, kategorik verilere ordinal kodlama uygulanmış, sayısal veriler 0-1 aralığına normalize edilmiş ve daha sonra %70 eğitim ve %30 test verilerine ayrılarak modeller test edilmiştir.

Tablo 4.46: Çapraz Doğrulama (CV) performans sonuçları (Deneme 1)

Model	CV	CV	CV
	R ² Score	RMSE Score	MAE Score
XGBoost Reg	0.5137 ± 0.1991	14.4760 ± 1.5975	14.4760 ± 1.5975
Random Forest Reg	0.5094 ± 0.1793	14.7251 ± 1.9264	11.5029 ± 1.5673
SVR (RBF)	0.5549 ± 0.1391	14.1454 ± 2.1242	10.7192 ± 1.8453
KNN Reg	0.5162 ± 0.1470	14.7397 ± 1.8646	10.9197 ± 1.5497
MLP Reg	0.4268 ± 0.1407	16.1944 ± 2.3160	12.9884 ± 1.7028
Decision Tree Reg	0.2305 ± 0.3391	18.1704 ± 2.3882	13.5153 ± 2.1665

Tablo 4.46’da dönem sonu başarı ortalamalarının tahminine yönelik olarak geliştirilen farklı regresyon modellerinin çapraz doğrulama performansları karşılaştırmalı olarak sunulmaktadır. Modellerin başarımı, R², RMSE ve MAE ölçütleri üzerinden değerlendirilmiştir. Bu ölçütler sırasıyla modelin açıklayıcılık düzeyini, ortalama tahmin hatasının büyüklüğünü ve mutlak hata düzeyini yansıtmaktadır.

Sonuçlar incelendiğinde, SVR (RBF) modelinin en yüksek R^2 değeri (0.5549 ± 0.1391) ile bağımlı değişkendeki varyansın en büyük bölümünü açıklayabildiği görülmektedir. Buna ek olarak SVR modelinin RMSE (14.1454) ve özellikle MAE (10.7192) değerlerinin görece düşük olması, tahmin hatalarının hem ortalama hem de mutlak düzeyde daha sınırlı olduğunu göstermektedir. Bu durum, SVR (RBF) modelinin doğrusal olmayan ilişkileri başarılı biçimde yakalayabildiğini ve daha isabetli tahminler üretebildiğini ortaya koymaktadır.

Random Forest Regresyon ve KNN Regresyon modelleri, R^2 değerleri bakımından SVR'ye yakın performanslar sergilemiş, ancak RMSE ve MAE ölçütlerinde görece daha yüksek hata değerleri üretmiştir. Bu durum, bu modellerin genel eğilimleri yakalayabildiğini ancak bireysel tahminlerde daha fazla sapma gösterebildiğini düşündürmektedir. XGBoost Regresyon modeli ise orta düzeyde bir açıklayıcılık sunmuş; ancak hata ölçütleri açısından SVR modelinin gerisinde kalmıştır.

MLP Regresyon ve özellikle Decision Tree Regresyon modellerinin daha düşük R^2 değerleri ve daha yüksek hata oranları sergilemesi, bu yöntemlerin mevcut veri yapısında akademik performans ortalamasını tahmin etmede sınırlı kaldığını göstermektedir. Özellikle Decision Tree regresyon modelinde gözlenen yüksek standart sapma, model performansının katlar arasında tutarsızlık gösterdiğini ortaya koymaktadır.

Genel olarak regresyon sonuçları değerlendirildiğinde, SVR (RBF) modelinin hem açıklayıcılık düzeyi hem de hata ölçütleri açısından en dengeli ve güvenilir performansı sunduğu görülmektedir. Bu bulgular, akademik performans ortalamalarının tahmini için çekirdek tabanlı regresyon yaklaşımlarının bu veri yapısında daha uygun olduğunu ve SVR modelinin regresyon problemleri için nihai aday model olarak öne çıktığını göstermektedir.

Tablo 4.47: Test performans sonuçları (Deneme 1)

Model	TEST R^2 Score	TEST RMSE Score	TEST MAE Score
XGBoost Reg	0.6963	12.4614	8.9904
Random Forest Reg	0.6661	13.0661	9.8173
SVR (RBF)	0.6442	13.4871	10.0474
KNN Reg	0.6287	13.7771	10.1877

Tablo 4.47: (devam ediyor)

MLP Reg	0.5769	14.7069	11.5422
Decision Tree Reg	0.4803	16.3001	11.9639

Tablo 4.47’de, dönem sonu başarı ortalamalarının tahminine yönelik geliştirilen regresyon modellerinin test veri kümesi üzerindeki performansları karşılaştırmalı olarak sunulmaktadır. Test sonuçları, modellerin eğitim aşamasında öğrenilen ilişkileri daha önce görülmemiş veriler üzerinde ne ölçüde başarıyla genelleştirebildiğini ortaya koymaktadır.

Elde edilen bulgulara göre, XGBoost Regresyon modeli test aşamasında en yüksek R^2 değeri (0.6963) ile bağımlı değişkendeki varyansın en büyük bölümünü açıklamayı başarmıştır. Aynı zamanda en düşük RMSE (12.4614) ve en düşük MAE (8.9904) değerlerini üretmesi, bu modelin hem ortalama hata büyüklüğü hem de mutlak hata açısından en isabetli tahminleri sunduğunu göstermektedir. Bu durum, XGBoost’un karmaşık ve doğrusal olmayan ilişkileri test verisi üzerinde de etkili biçimde yakalayabildiğini ortaya koymaktadır.

Random Forest Regresyon modeli, XGBoost’a yakın bir açıklayıcılık düzeyi sergilemiş; ancak hata ölçütlerinin görece daha yüksek olması, bireysel tahminlerde daha fazla sapma oluştuğunu göstermektedir. SVR (RBF) ve KNN Regresyon modelleri orta düzeyde performans üretmiş; bu modellerin test aşamasında kabul edilebilir sonuçlar sunduğu ancak ensemble tabanlı yaklaşımların gerisinde kaldığı görülmüştür.

MLP Regresyon ve özellikle Decision Tree Regresyon modellerinin daha düşük R^2 değerleri ve daha yüksek hata oranları sergilemesi, bu yöntemlerin test verisi üzerinde genellenebilirlik açısından sınırlı kaldığını göstermektedir. Özellikle Decision Tree modelinin performansı, tekil ağaç yapısının akademik performans gibi çok boyutlu bir çıktıyı tahmin etmede yetersiz olduğunu ortaya koymaktadır.

Genel değerlendirme sonucunda, test performansları XGBoost Regresyon modelinin dönem sonu başarı ortalamalarının tahmini açısından en yüksek doğruluk ve en düşük hata düzeylerini sunduğunu göstermektedir. Bu bulgular, regresyon problemi kapsamında XGBoost modelinin nihai tahmin modeli olarak tercih edilmesini güçlü biçimde

desteklemektedir.

Tablo 4.48: XGBoost en iyi parametreler (Deneme 1)

Parametre	Değer
colsample_bytree	0.6071
gamma	3.0128
learning_rate	0.00801
max_depth	9
min_child_weight	3
n_estimators	1840
reg_alpha	0.7808
reg_lambda	1.6495
subsample	0.6233

Tablo 4.48’de sunulan hiperparametreler, XGBoost regresyon modelinin test aşamasında elde edilen yüksek açıklayıcılık ve düşük hata değerlerini destekleyen yapılandırmayı yansıtmaktadır. $n_estimators = 1840$, modelin çok sayıda ağaç kullanarak karmaşık örüntüleri kademeli biçimde öğrenmesine olanak tanımaktadır. Buna karşılık $learning_rate = 0.00801$ gibi düşük bir öğrenme oranı tercih edilmesi, modelin öğrenme sürecini yavaşlatarak aşırı uyum riskini azaltmış ve genellenebilirliği artırmıştır.

$Max_depth = 9$ değeri, ağaçların yeterli derinliğe ulaşmasını sağlayarak doğrusal olmayan ilişkilerin yakalanmasına katkı sunarken, $min_child_weight = 3$ parametresi dallanma kararlarının daha kontrollü alınmasını sağlamıştır. $Gamma = 3.0128$, yeni dallanmaların yalnızca anlamlı kazanım sağladığı durumlarda oluşmasına izin vererek model karmaşıklığını sınırlandırmıştır.

Örnekleme parametreleri olan $subsample = 0.6233$ ve $colsample_bytree = 0.6071$, her bir ağacın veri ve özniteliklerin yalnızca bir kısmı ile eğitilmesini sağlayarak model çeşitliliğini artırmış ve aşırı uyumu önlemiştir. Düzenleştirme açısından $reg_alpha = 0.7808$ (L1) ve $reg_lambda = 1.6495$ (L2) değerlerinin birlikte kullanılması, gereksiz karmaşıklığın bastırılmasına ve daha sade bir model yapısının elde edilmesine katkı sağlamıştır.

Genel olarak bu hiperparametre kombinasyonu, XGBoost regresyon modelinin dönem sonu başarı ortalamalarının tahmininde yüksek doğruluk, düşük hata ve güçlü genellenebilirlik sunmasını sağlayan dengeli bir yapı ortaya koymaktadır.

Tablo 4.49: XGBoost Çapraz Doğrulama performans sonuçları (Deneme 1)

Metrik	Ortalama $\pm$ Std
RMSE	14.4760 $\pm$ 1.5975
MAE	10.7652 $\pm$ 1.4044
R ²	0.5137 $\pm$ 0.1991

Tablo 4.49’da sunulan çapraz doğrulama sonuçları, regresyon modelinin farklı veri alt kümeleri üzerindeki tahmin doğruluğu ve tutarlılığını ortaya koymaktadır. RMSE değerinin 14.4760 $\pm$ 1.5975 olması, modelin tahmin hatalarının ortalama büyüklüğünün makul bir seviyede olduğunu göstermektedir. MAE değerinin 10.7652 $\pm$ 1.4044 olması ise mutlak hata düzeyinin RMSE’ye kıyasla daha düşük olduğunu ve büyük sapmaların sınırlı sayıda gerçekleştiğini ortaya koymaktadır.

R² değerinin 0.5137 $\pm$ 0.1991 olması, modelin bağımlı değişkendeki varyansın yaklaşık yarısını açıklayabildiğini göstermektedir. Standart sapmanın görece yüksek olması, bazı katlarda model performansının değişkenlik gösterebildiğine işaret etmekle birlikte, genel olarak modelin verideki temel eğilimleri yakalayabildiğini ortaya koymaktadır.

Bu sonuçlar, ilgili regresyon modelinin akademik performans ortalamalarının tahmininde orta–yüksek düzeyde açıklayıcılık sunduğunu ve test aşamasında elde edilen daha yüksek performans değerleriyle birlikte değerlendirildiğinde, modelin genellenebilirliğini destekleyen bir yapı ortaya koyduğunu göstermektedir.

Tablo 4.50: XGBoost Test performans sonuçları (Deneme 1)

Metrik	Değer
RMSE	12.4614
MAE	8.9904
R ²	0.6963

Tablo 4.50’de sunulan test sonuçları, regresyon modelinin daha önce görülmemiş veriler üzerindeki tahmin doğruluğunu ve genellenebilirliğini ortaya koymaktadır. RMSE değerinin 12.4614 olması, modelin tahmin hatalarının ortalama büyüklüğünün düşük seviyede kaldığını ve büyük sapmaların sınırlı olduğunu göstermektedir. MAE değerinin 8.9904 olması ise tahminlerin mutlak hata açısından daha isabetli olduğunu ve modelin bireysel gözlemler üzerindeki performansının güçlü olduğunu ortaya koymaktadır.

R^2 değerinin 0.6963 seviyesinde gerçekleşmesi, modelin bağımlı değişkendeki varyansın yaklaşık %70’ini açıklayabildiğini göstermektedir. Bu sonuç, modelin akademik performans ortalamalarını tahmin etmede yüksek bir açıklayıcılık düzeyine ulaştığını ve öğrenilen ilişkilerin test veri kümesine başarıyla genellendiğini ortaya koymaktadır.

Genel olarak bu performans değerleri, regresyon modelinin dönem sonu başarı ortalamalarının tahmininde yüksek doğruluk, düşük hata ve güçlü genellenebilirlik sunduğunu göstermekte ve ilgili modelin çalışmanın regresyon bileşeni için nihai tahmin modeli olarak tercih edilmesini desteklemektedir.

4.4.2 Regresyon Sonuçları (Deneme 2)

İkinci denemede veriseti, SMOGN veri artırma yöntemi ile dengelenmiş, kategorik verilere one-hot kodlama uygulanmış, sayısal veriler de 0-1 aralığına normalize edilmiş ve daha sonra %70 eğitim ve %30 test verilerine ayrılarak modeller test edilmiştir.

Tablo 4.51: Çapraz Doğrulama (CV) performans sonuçları (Deneme 2)

Model	CV	CV	CV
	R^2 Score	RMSE Score	MAE Score
SVR (RBF)	0.5490 ± 0.1507	13.9124 ± 1.5695	10.4385 ± 1.1340
Random Forest Reg	0.5682 ± 0.1042	13.8846 ± 2.1836	10.9533 ± 1.6893
XGBoost Reg	0.5451 ± 0.1943	13.7816 ± 1.4249	10.2795 ± 0.9182
MLP Reg	0.5024 ± 0.1379	14.7755 ± 2.1598	12.0948 ± 1.3527
KNN Reg	0.3988 ± 0.2419	16.0581 ± 1.8273	11.3526 ± 1.1382
Decision Tree Reg	0.2706 ± 0.2840	17.5511 ± 2.2131	13.3408 ± 1.9283

Tablo 4.51’de, dönem sonu başarı ortalamalarının tahminine yönelik geliştirilen regresyon modellerinin çapraz doğrulama performansları karşılaştırmalı olarak sunulmaktadır. Modellerin başarımı, açıklayıcılığı temsil eden R^2 ile tahmin hatalarının büyüklüğünü yansıtan RMSE ve MAE ölçütleri üzerinden değerlendirilmiştir.

Sonuçlar incelendiğinde, Random Forest Regresyon modeli en yüksek ortalama R^2 değeri (0.5682 ± 0.1042) ile bağımlı değişkendeki varyansın en büyük bölümünü açıklamayı başarmıştır. Bununla birlikte RMSE ve MAE değerlerinin görece yüksek olması, bireysel tahminlerde hata düzeyinin tamamen minimize edilemediğini göstermektedir. XGBoost Regresyon ve SVR (RBF) modelleri ise R^2 değerleri bakımından Random Forest’a oldukça yakın performanslar sergilemiş; özellikle XGBoost Regresyon modelinin daha düşük RMSE (13.7816) ve MAE (10.2795) değerleri üretmesi, hata büyüklüğü açısından daha dengeli sonuçlar sunduğunu ortaya koymuştur.

SVR (RBF) modeli, açıklayıcılık ve hata ölçütleri arasında görece dengeli bir yapı sunmuş; çekirdek tabanlı yaklaşımın doğrusal olmayan ilişkileri başarılı biçimde yakalayabildiğini göstermiştir. MLP Regresyon modeli orta düzey bir performans sergilemiş, ancak hata ölçütlerinin yüksekliği nedeniyle üst düzey modellerin gerisinde kalmıştır. KNN Regresyon ve özellikle Decision Tree Regresyon modellerinin düşük R^2 değerleri ve yüksek hata oranları, bu yöntemlerin akademik performans gibi çok boyutlu bir çıktıyı tahmin etmede sınırlı kaldığını göstermektedir. Ayrıca Decision Tree modelinde gözlenen yüksek standart sapma, performansın katlar arasında tutarsızlık gösterdiğine işaret etmektedir.

Genel olarak CV sonuçları, Random Forest, XGBoost ve SVR (RBF) modellerinin regresyon problemi için daha uygun olduğunu; özellikle hata ölçütleri dikkate alındığında XGBoost Regresyon ve SVR (RBF) modellerinin daha dengeli ve kararlı tahminler ürettiğini ortaya koymaktadır. Bu bulgular, test aşamasında elde edilen sonuçlarla birlikte değerlendirildiğinde, nihai regresyon modelinin belirlenmesine yönelik güçlü bir yöntemsel temel sunmaktadır.

Tablo 4.52: Test performans sonuçları (Deneme 2)

Model	TEST R ² Score	TEST RMSE Score	TEST MAE Score
SVR (RBF)	0.6806	12.7653	9.2919
Random Forest Reg	0.6681	13.0115	9.8536
XGBoost Reg	0.6582	13.2046	9.4327
MLP Reg	0.6088	14.1266	11.5385
KNN Reg	0.4612	16.5790	10.6001
Decision Tree Reg	0.4527	16.7092	11.8291

Tablo 4.52’de, dönem sonu başarı ortalamalarının tahminine yönelik geliştirilen regresyon modellerinin test veri kümesi üzerindeki performansları karşılaştırmalı olarak sunulmaktadır. Test sonuçları, modellerin eğitim sürecinde öğrenilen örüntüleri daha önce görülmemiş veriler üzerinde ne ölçüde başarıyla genelledebildiğini ortaya koymaktadır.

Elde edilen bulgulara göre, SVR (RBF) modeli en yüksek R² değeri (0.6806) ile bağımlı değişkendeki varyansın en büyük bölümünü açıklamayı başarmıştır. Aynı zamanda RMSE (12.7653) ve MAE (9.2919) değerlerinin görece düşük olması, modelin hem ortalama hata büyüklüğü hem de mutlak hata açısından isabetli tahminler ürettiğini göstermektedir. Bu durum, çekirdek tabanlı SVR yaklaşımının doğrusal olmayan ilişkileri test verisi üzerinde de etkin biçimde yakalayabildiğini ortaya koymaktadır.

Random Forest Regresyon ve XGBoost Regresyon modelleri, SVR’ye yakın açıklayıcılık düzeyleri sergilemiş olmakla birlikte, hata ölçütlerinin biraz daha yüksek olması nedeniyle tahmin doğruluğu açısından ikinci planda kalmıştır. Bununla birlikte bu iki modelin de test aşamasında istikrarlı performans sunduğu ve genel eğilimleri başarılı biçimde yakalayabildiği görülmektedir.

MLP Regresyon modeli orta düzey bir performans sergilemiş; ancak RMSE ve MAE değerlerinin yüksekliği, bireysel tahminlerde daha fazla sapma oluştuğunu göstermektedir. KNN Regresyon ve özellikle Decision Tree Regresyon modellerinin düşük R² değerleri ve yüksek hata oranları ise bu yöntemlerin test veri kümesi üzerinde genellenebilirlik açısından sınırlı kaldığını ortaya koymaktadır. Özellikle Decision Tree modelinin performansı, tekil

ağaç yapısının akademik performans gibi çok boyutlu bir çıktıyı tahmin etmede yetersiz olduğunu göstermektedir.

Genel değerlendirme sonucunda, test performansları SVR (RBF) modelinin dönem sonu başarı ortalamalarının tahmininde en yüksek açıklayıcılığı ve en düşük hata düzeylerini sunduğunu göstermektedir. Bu bulgular, regresyon problemi kapsamında SVR (RBF) modelinin nihai tahmin modeli olarak tercih edilmesini güçlü biçimde desteklemektedir.

Tablo 4.53: SVR en iyi parametreler (Deneme 2)

Parametre	Değer
Kernel	RBF
C	129.6242
Epsilon	0.07070
Gamma	0.07996

Tablo 4.53'te sunulan hiperparametreler, SVR (Support Vector Regression) modelinin test aşamasında elde edilen yüksek açıklayıcılık ve düşük hata değerlerini destekleyen yapılandırmayı göstermektedir. Modelde RBF (Radial Basis Function) çekirdeğinin kullanılması, veri setindeki doğrusal olmayan ilişkilerin etkin biçimde modellenmesine olanak tanımıştır. C parametresinin 129.6242 olarak belirlenmesi, modelin hata toleransını düşürerek tahmin doğruluğunu artırmaya odaklanan daha katı bir öğrenme yapısı benimsediğini göstermektedir.

Epsilon değerinin 0.07070 olması, regresyon hattı etrafında dar bir hata tolerans bölgesi tanımlandığını ve modelin küçük sapmalara karşı daha duyarlı hâle getirildiğini ortaya koymaktadır. Bu durum, özellikle bireysel tahminlerde MAE ve RMSE değerlerinin düşmesine katkı sağlamıştır. Gamma değerinin 0.07996 olarak seçilmesi ise her bir veri noktasının karar fonksiyonu üzerindeki etki alanını dengeli biçimde sınırlandırarak, modelin aşırı uyum göstermeden karmaşık örüntüleri yakalayabilmesini mümkün kılmıştır.

Genel olarak bu hiper parametre kombinasyonu, SVR (RBF) modelinin dönem sonu akademik performans ortalamalarının tahmininde yüksek doğruluk, düşük hata ve güçlü genellenebilirlik sunmasını sağlayan dengeli bir yapı ortaya koymaktadır.

Tablo 4.54: SVR Çapraz Doğrulama performans sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std
RMSE	13.9124 $\pm$ 1.5695
MAE	10.4385 $\pm$ 1.1340
R ²	0.5490 $\pm$ 0.1507

Tablo 4.54'te sunulan çapraz doğrulama sonuçları, regresyon modelinin farklı veri alt kümeleri üzerindeki tahmin doğruluğu ve kararlılığını göstermektedir. RMSE değerinin 13.9124 $\pm$ 1.5695 olması, tahmin hatalarının ortalama büyüklüğünün kabul edilebilir bir düzeyde seyrettiğini; MAE değerinin 10.4385 $\pm$ 1.1340 olması ise mutlak hata düzeyinin görece daha düşük olduğunu ve büyük sapmaların sınırlı kaldığını ortaya koymaktadır.

R² değerinin 0.5490 $\pm$ 0.1507 seviyesinde gerçekleşmesi, modelin bağımlı değişkendeki varyansın yaklaşık yarısını açıklayabildiğini göstermektedir. Standart sapmanın orta düzeyde olması, performansın katlar arasında belirli bir değişkenlik sergilediğine işaret etse de genel olarak modelin verideki temel eğilimleri yakalayabildiğini ortaya koymaktadır.

Bu bulgular, ilgili regresyon modelinin orta–yüksek düzeyde açıklayıcılık ve istikrarlı hata profili sunduğunu; test aşamasında elde edilen daha yüksek performansla birlikte değerlendirildiğinde, modelin genellenebilirliğini destekleyen bir yapı ortaya koyduğunu göstermektedir.

Tablo 4.55: SVR Test performans sonuçları (Deneme 2)

Metrik	Değer
RMSE	12.7653
MAE	9.2919
R ²	0.6806

Tablo 4.55'te sunulan test sonuçları, regresyon modelinin daha önce görülmemiş veriler üzerindeki tahmin doğruluğunu ve genellenebilirliğini ortaya koymaktadır. RMSE değerinin 12.7653 olması, modelin tahmin hatalarının ortalama büyüklüğünün düşük seviyede kaldığını ve büyük sapmaların sınırlı olduğunu göstermektedir. MAE değerinin 9.2919 olması ise tahminlerin mutlak hata açısından daha isabetli olduğunu ve modelin bireysel

gözlemler üzerindeki performansının güçlü olduğunu ortaya koymaktadır.

R^2 değerinin 0.6806 seviyesinde gerçekleşmesi, modelin bağımlı değişkendeki varyansın yaklaşık %68'ini açıklayabildiğini göstermektedir. Bu sonuç, modelin akademik performans ortalamalarını tahmin etmede yüksek bir açıklayıcılık düzeyine ulaştığını ve eğitim aşamasında öğrenilen ilişkilerin test veri kümesine başarıyla genellendiğini ortaya koymaktadır.

Genel olarak bu performans değerleri, regresyon modelinin dönem sonu başarı ortalamalarının tahmininde yüksek doğruluk, düşük hata ve güçlü genellenebilirlik sunduğunu göstermekte ve ilgili modelin çalışmanın regresyon bileşeni için nihai tahmin modeli olarak tercih edilmesini desteklemektedir.

4.5 Açıklanabilir Yapay Zekâ

Çalışmada elde edilen en iyi sınıflandırma modeli Destek Vektör Makinesi sonuçlarının şeffaf ve açık bir hale getirilebilmesi amacıyla, modele açıklanabilir yapay zekâ yöntemlerinden SHAP ve LIME uygulanmıştır. SHAP ve LIME, makine öğrenmesi modellerinin karar mekanizmalarını açıklamak için yaygın biçimde kullanılan iki yöntemdir. LIME, belirli bir gözlem çevresinde basitleştirilmiş ve yorumlanabilir bir yerel model kurarak tahmini açıklar; bu nedenle daha çok tekil tahminlerin yerel düzeyde yorumlanmasında kullanılmaktadır (Ribeiro vd., 2016). SHAP ise her bir değişkenin tahmine katkısını Shapley değerleri aracılığıyla hesaplamakta ve bu katkıları daha sistematik bir çerçevede sunmaktadır. Lundberg ve Lee'ye (2017) göre SHAP, yerel doğruluk, eksiklik ve tutarlılık gibi önemli kuramsal özellikleri sağlaması bakımından güçlü bir açıklama yaklaşımıdır. Bu nedenle LIME daha pratik ve yerel açıklamalar için uygun bir yöntem olarak değerlendirilirken, SHAP daha tutarlı ve kapsamlı katkı yorumları sunabilmektedir. Bununla birlikte, her iki yöntemin de model yapısından ve değişkenler arasındaki ilişkilerden etkilenebildiği unutulmamalıdır (Salih vd., 2024).

4.5.1 SHAP

Açıklanabilir yapay zekâ yöntemlerinden SHAP, sınıflandırma modelinin sonuçlarını en çok etkileyen öznitelikleri ve etkileme derecelerini ortaya koymaktadır.

Tablo 4.56: SHAP ortalama

Sıra	Değişken	Ortalama SHAP
1	Baba eğitim durumunuz	0.004437
2	Ders notları	0.004286
3	Çevrimiçi Kurslar (Udemy, Coursera vb.)	0.003928
4	Sosyal medya kullanımı ders çalışmanızı etkiliyor mu?	0.003866
5	E-ders Kaynakları	0.003695
6	Ders kitapları	0.003686
7	Ders çalışırken ne sıklıkla ara verirsiniz?	0.003607
8	Çevrimiçi derslere katılımınız ne düzeyde?	0.003522
9	Aile gelir düzeyiniz	0.003264
10	Yaşadığınız yer	0.003246
11	Sınav kaygısı yaşıyor musunuz?	0.003202
12	Derslere düzenli olarak katılıyor musunuz?	0.003169
13	Yapay Zekâ (ChatGPT, Gemini, DeepSeek vb.)	0.002892
14	Sınavlara ne kadar süre önceden hazırlanmaya başlarsınız?	0.002854
15	Genel olarak çalışma ortamınız seni çalışmaya ne kadar teşvik ediyor?	0.002725
16	Haftalık ortalama ders çalışma süreniz kaç saattir?	0.002720
17	YouTube Ders Videoları	0.002544
18	Derslerinizle ilgili olarak interneti ne sıklıkla kullanırsınız?	0.002436
19	Anne eğitim durumunuz	0.002394
20	Teknoloji kullanımının ders çalışmanıza yardımcı olduğunu düşünüyor musunuz?	0.002354
21	Arama Motorları (Google, Yandex vb.)	0.002109
22	Ders çalışmak için en çok kullandığınız ortam/mekân nedir?	0.002099
23	Cinsiyetiniz	0.002084
24	Yapay zekâ	0.002058
25	Öğrenim süreciniz boyunca barınma durumunuz nedir?	0.002034
26	Günlük ortalama internet kullanım süreniz kaç saattir?	0.002033
27	Liseden mezun olduğunuz okul türü	0.001516
28	Öğrenme stilinizi nasıl tanımlarsınız?	0.001468
29	İzleme Oranı	0.001442
30	İnternet kaynakları	0.001356
31	Öğrenme yönetim sistemi (UBYS, e-ders)	0.001234
32	Devam	0.000840
33	Yaşınız	0.000609
34	Devam Oranı	0.000439

Tablo 4.56’da yer alan SHAP analizi ortalama deęerleri incelendięinde, ğrencinin akademik performansını tahmin eden modelde hangi deęiřkenlerin daha nemli olduęu grlmektedir. Ortalama deęer, bir deęiřkenin model tahminlerini ortalama olarak ne kadar etkiledięini ifade eder; deęer bydke deęiřkenin model zerindeki etkisi de artmaktadır (Lundberg ve Lee, 2017). Bu analiz sonuları, ğrencinin akademik performansının tahmin edilmesinde en nemli deęiřkenlerin baba eęitim durumu, ders (alıřma) notu ve evrimii kurslara katılım olduęunu ortaya koymuřtur. Bu durum, ğrencinin bařarısının yalnızca bireysel akademik gstergelerle deęil, aynı zamanda aile eęitim dzeyi ve ęrenmeyi destekleyen ek eęitim faaliyetleriyle de iliřkili olabileceęine iřaret etmektedir.

Elde edilen sonulara gre en etkili deęiřkenin “baba eęitim durumu” olduęu grlmektedir. Bu durum, ğrencinin aile yapısı ve zellikle ebeveynlerin eęitim dzeyinin akademik performans zerinde nemli bir belirleyici olduęunu gstermektedir. Alanyazında da ebeveyn eęitim dzeyinin ğrencilerin akademik performansı ile gl řekilde iliřkili olduęu vurgulanmaktadır (Sirin, 2005).

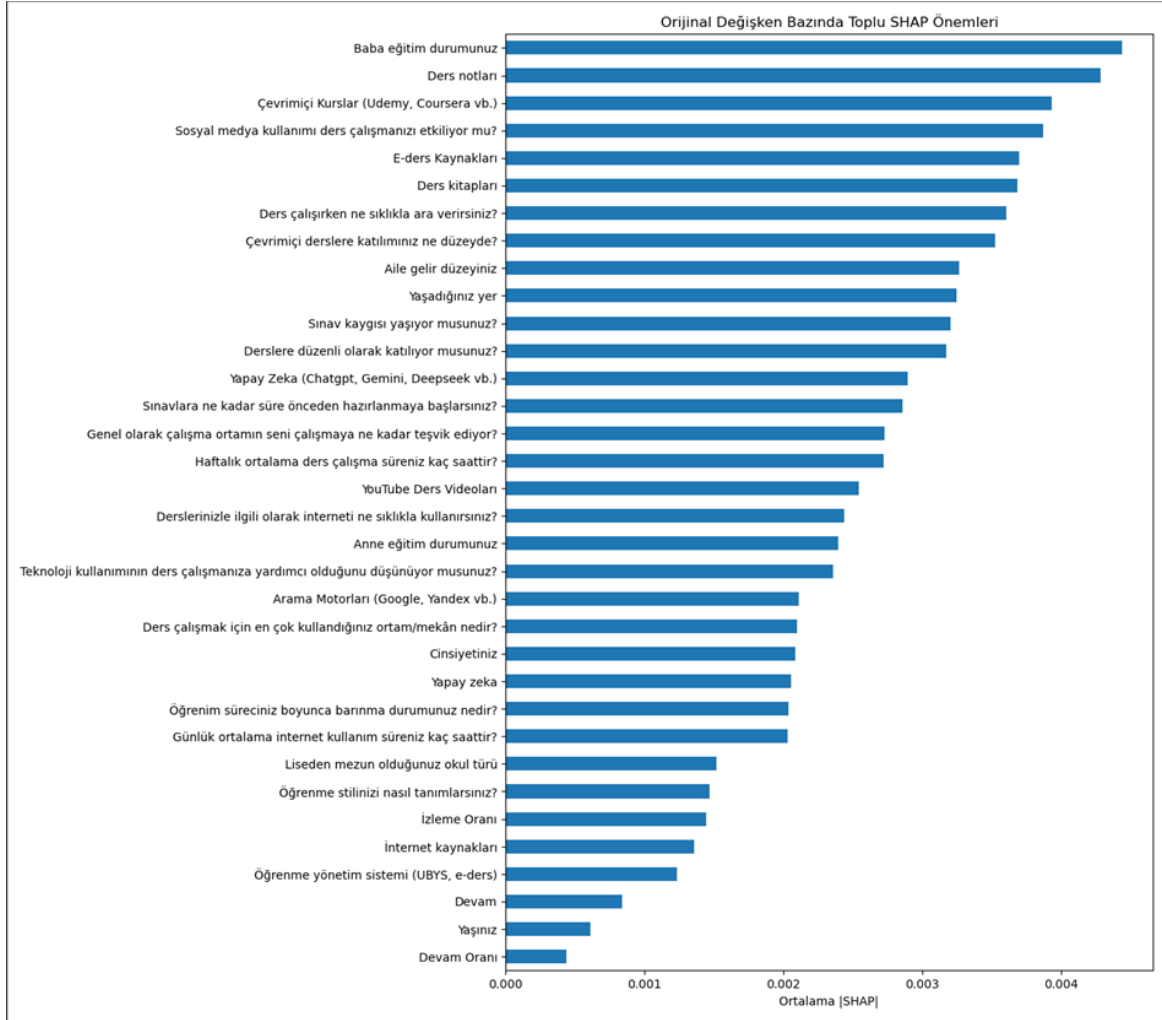
Elde edilen bulgulara gre en etkili ikinci deęiřkenin “ders (alıřma) notları” olduęu grlmektedir. Bu durum, ğrencilerin ders kapsamında oluřturulan ya da kullanılan notlardan yararlanma dzeyinin akademik performans zerinde nemli bir belirleyici olduęunu gstermektedir. Ders notları, bilginin dzenlenmesini, tekrar edilmesini ve sınav odaklı alıřmayı kolaylařtırarak ęrenmenin kalıcılıęına katkı saęlayabilmektedir. Alanyazında da not alma ve ders notlarını etkili kullanmanın ęrenme ıktıları ile akademik bařarı zerinde anlamlı etkileri olduęu vurgulanmaktadır (Kiewra, 1985; Kobayashi, 2005). ne ıkan bir dięer deęiřkenin “evrimii kurslar” olduęu anlařılmaktadır. Bu bulgu, ğrencilerin formal ders ieriklerinin yanında ek dięital ęrenme kaynaklarından yararlanmasının akademik performans aısından nemli bir rol oynadıęını dřndrmektedir. evrimii kurslar; ğrencilerin kendi hızlarında ęrenmelerine, eksik oldukları konuları tamamlamalarına ve farklı anlatım biimleriyle kavramsal anlayıřlarını glendirmelerine katkı saęlayabilmektedir. Alanyazında da evrimii ve harmanlanmış ęrenme ortamlarının uygun řekilde kullanıldıęında ęrenme bařarısını destekledięi ve zellikle z dzenlemeli ęrenme becerileriyle birlikte akademik performansı olumlu ynde glendirebildięi belirtilmektedir (Means ve ark., 2013; Broadbent ve Poon, 2015).

Sosyal medya kullanımının ders çalışmaya etkisi, e-ders kaynakları, ders kitapları ve ders çalışırken ara verme sıklığı gibi değişkenlerin üst sıralarda yer alması, modelin öğrenme sürecinin niteliğine ve çalışma alışkanlıklarına duyarlı olduğunu ortaya koymaktadır. Aile gelir düzeyi, yaşanılan yer, sınav kaygısı, derslere düzenli katılım ve yapay zekâ araçlarının kullanımı ise bu çalışmada “orta düzey etkili” değişkenler olarak değerlendirilmiştir.

Buradaki “orta düzey” ifadesi, mutlak ya da evrensel bir eşik değerine değil, değişkenlerin aynı model içindeki ortalama SHAP dağılımına göre yapılan göreceli bir sınıflamaya dayanmaktadır. Başka bir ifadeyle, bu değişkenler en yüksek katkı sunan özellikler kadar belirleyici değildir; ancak alt sıralardaki değişkenlere kıyasla model tahminine daha anlamlı katkı sağlamaktadır. SHAP temelli yorumlamada değişken öneminin bağlama, veri setine ve model yapısına bağlı olarak göreceli biçimde ele alınması gerektiği vurgulanmaktadır (Lundberg ve Lee, 2017; Molnar, 2022).

Şekil 4.5.’te verilen grafik, modelin öğrenci akademik performans tahmininde kullanılan değişkenlerin etki büyüklüklerine göre önem sıralamasını göstermekte olup, X ekseninde yer alan ortalama mutlak SHAP değerleri bir değişkenin model kararlarını ne ölçüde etkilediğini, Y eksenindeki değerler ise one-hot encoding öncesi birleştirilmiş orijinal değişkenleri ifade etmektedir; üst sıralarda yer alan değişkenler model açısından daha kritik belirleyiciler olup bu grafik etki yönünü değil yalnızca önem düzeyini ortaya koymaktadır. Bulgulara göre baba eğitim durumu en yüksek ortalama SHAP değerine sahip değişken olarak öne çıkmakta ve aile sosyo-kültürel yapısının akademik performans üzerindeki belirleyici rolünü göstermektedir; bunu ders (çalışma) notları ve çevrimiçi kurslara katılım izlemekte olup, süreç değerlendirmelerinin ve öz-yönelimli öğrenme faaliyetlerinin başarı tahmininde kritik öneme sahip olduğu anlaşılmaktadır. Sosyal medya kullanımının ders çalışmaya etkisi, e-ders kaynakları, ders kitapları, ders çalışırken ara verme sıklığı ve çevrimiçi derslere katılım düzeyi gibi değişkenlerin üst grupta yer alması, modelin yalnızca ne kadar çalışıldığına değil, nasıl çalışıldığına ve hangi öğrenme kaynaklarının kullanıldığına duyarlı olduğunu ortaya koymaktadır. Aile gelir düzeyi, yaşanılan yer, sınav kaygısı, derslere düzenli katılım, yapay zekâ araçlarının kullanımı, sınavlara hazırlanmaya başlama süresi ve çalışma ortamının teşvik ediciliği orta düzey etkili değişkenler olarak anlam kazanırken, yaş, devam oranı, öğrenme yönetim sistemi kullanımı, internet kaynakları, öğrenme stili ve izleme oranı gibi değişkenlerin tek başına ayırt edici gücünün

görece düşük olduğu görülmektedir. Şekil 4.6’da görülen SHAP arı sürüsü (beeswarm) grafiği ile değerlendirildiğinde, bu sonuçlar öğrenci başarısının temel olarak ailesel faktörler, çalışma alışkanlıkları ve dijital öğrenme olanaklarının etkin kullanımı ile şekillendiğini ve başarının çok boyutlu bir yapı sergilediğini göstermektedir.



Şekil 4.5: Orijinal değişken bazlı toplu SHAP ortalamaları grafiği

Şekil 4.6’da verilen SHAP arı sürüsü grafiğinde x ekseninde yer alan SHAP değerleri değişkenlerin model çıktısını hangi yönde etkilediğini göstermektedir. Grafikte pozitif SHAP değerleri (sağ taraf) akademik performans tahminini artıran etkileri, negatif SHAP değerleri (sol taraf) ise akademik performans tahminini düşüren etkileri ifade etmektedir. Noktaların renkleri ise sayısal değişkenlerde değişken değerinin maviden kırmızıya doğru büyüdüğünü, kategorik değişkenlerde düşük değerli kategorilerden yüksek değerli kategorilere doğru değiştiğini göstermektedir. Grafikte dar dağılım, bir değişkene ait SHAP değerlerinin X eksenini boyunca geniş bir aralığa yayılmaması ve noktaların çoğunlukla sıfır

etrafında dar bir bantta kümelenmesi durumunu ifade eder. Bu durum, ilgili değişkenin model tahminleri üzerindeki etkisinin sınırlı olduğunu veya gözlemler arasında benzer bir etki gösterdiğini göstermektedir. Başka bir ifadeyle, değişkenin farklı değerleri model çıktısını büyük ölçüde değiştirmemekte ve model bu değişkene karşı yüksek duyarlılık göstermemektedir. Buna karşılık, SHAP değerlerinin geniş bir aralığa yayılması değişkenin model tahmini üzerinde daha güçlü ve değişken bir etkiye sahip olduğunu göstermektedir. Bu nedenle arı sürüsü grafiğinde dar dağılım, genellikle değişkenin model açısından görece düşük veya daha stabil bir katkıya sahip olduğunun bir göstergesi olarak yorumlanmaktadır (Lundberg ve Lee, 2017; Ponce-Bobadilla ve Schmitt, 2024; Bifarin, 2023).



Şekil 4.6: SHAP Arı Sürüsü (Beeswarm) grafiği

*0 noktasında yoğunlaşan noktalar nötr değer alır.

* Veri setinde yer alan “ders notları, çevrimiçi kurslar, ders kitapları, e-ders kaynakları” değişkenleri doğrudan sayısal biçimde bulunmadığından, makine öğrenmesi modellerine uygun hale getirilmesi amacıyla veri dönüştürme işlemleri uygulanmıştır. Bu kapsamda “Ders çalışırken en çok hangi kaynağı kullanıyorsunuz?” sorusu çoktan seçmeli bir yapıya sahip olduğu için modelde kullanılabilmesi amacıyla one-hot encoding yöntemi ile kodlanmıştır. Bu işlem sonucunda her bir kaynak seçeneği (örneğin ders notları, e-ders kaynakları, ders kitapları vb.) ayrı bir sütun olarak temsil edilmiş ve ilgili kaynağı seçen öğrenciler için değer 1, seçmeyenler için ise 0 olacak şekilde veri setine dahil edilmiştir. Bu sayede kategorik bir değişken olan çalışma kaynakları, makine öğrenmesi algoritmalarının işleyebileceği sayısal bir forma dönüştürülmüş ve SHAP analizinde ilgili değişkenlerin model üzerindeki etkisi bağımsız bir öznitelik olarak

değerlendirilebilmiştir. Örneğin ders notu değişkenin “yok” olmasının sebebi çoktan seçmeli yapıda bazı öğrenciler tarafından seçilmediğinin göstergesidir.

SHAP grafiğine bakıldığında model kararlarında en etkili değişkenin “Baba eğitim durumu: Ortaokul” olduğu görülmektedir. Grafikte, baba eğitim durumunun ‘ortaokul’ kategorisine ait değerler negatif SHAP tarafında yoğunlaşmaktadır. Renkler değişken değerini göstermektedir; kırmızı noktalar baba eğitim düzeyinin ortaokul olduğu öğrencileri temsil etmekte ve bu noktalar çoğunlukla negatif SHAP değerlerinde yer almaktadır. Bu da babası ortaokul mezunu olan öğrencilerde modelin akademik performans tahmininin daha düşük olma eğiliminde olduğunu gösterir.

SHAP özet grafiğinde, ‘Derslere düzenli olarak katılıyor musunuz?’ özelliğinin ‘Çoğunlukla’ kategorisine ait değerler pozitif SHAP tarafında yoğunlaşmaktadır. Renkler derse katılım sıklığını göstermektedir; ‘Çoğunlukla’ kategorisindeki öğrenciler (mavi noktalar) genellikle pozitif SHAP değerlerine sahiptir. Bu da derse çoğunlukla katılan öğrencilerde modelin akademik performans tahmininin daha yüksek olma eğiliminde olduğunu gösterir.

SHAP özet grafiğinde, “Ders çalışırken ne sıklıkla ara verirsiniz? – Sık sık” kategorisine ait kırmızı noktaların büyük ölçüde negatif SHAP değerleri tarafında yoğunlaştığı görülmektedir. Bu da ders çalışırken sık sık ara veren öğrencilerde modelin akademik performans tahmininin daha düşük olma eğiliminde olduğunu göstermektedir.

“Yaşadığınız yer – İlçe” değişkeninde mavi noktaların pozitif SHAP değerleri tarafında daha yoğun olduğu görülmektedir. Bu durum, ilçe kategorisi dışında kalan yerleşim yerlerinde bulunan öğrenciler için modelin akademik performans tahmininin daha yüksek olma eğiliminde olduğunu göstermektedir.

“Haftalık ortalama ders çalışma süresi – 6–10 saat” kategorisine ait kırmızı noktaların da ağırlıklı olarak negatif SHAP tarafında yer aldığı görülmekte olup, bu durum haftalık ders çalışma süresi 6–10 saat olan öğrencilerde modelin akademik performans tahmininin daha düşük olma eğiliminde olduğunu göstermektedir. “Çevrimiçi derslere katılım düzeyi – Orta” kategorisine ait değerlerin de pozitif SHAP tarafında yoğunlaştığı görülmekte ve bu durum çevrimiçi derslere orta düzeyde katılan öğrencilerde modelin akademik performans

tahmininin daha yüksek olma eğiliminde olduğunu göstermektedir. 1 Bunun yanında “Derslerle ilgili olarak interneti kullanma sıklığı – Haftada birkaç kez” kategorisine ait değerlerin de pozitif SHAP tarafında yoğunlaştığı görülmektedir. “Çevrimiçi kurslar (Udemy, Coursera vb.) – Var” kategorisine ait değerler de benzer şekilde pozitif SHAP tarafında yer almakta olup, bu durum çevrimiçi kurslara katılımın bulunduğu öğrencilerde modelin akademik performans tahmininin daha yüksek olma eğiliminde olduğunu göstermektedir. Ayrıca “Teknoloji kullanımının ders çalışmaya katkısı – Orta düzeyde”, “Öğrenme stili – Karma” ve “Ders çalışmak için kullanılan ortam – Ev” kategorilerine ait değerlerin de ağırlıklı olarak negatif SHAP tarafında yoğunlaştığı görülmektedir. Bununla birlikte “Sosyal medya kullanımının ders çalışmayı etkilemesi – Evet” ve “Sınav kaygısı – Evet, biraz” kategorilerine ait değerlerin de negatif SHAP tarafında yer aldığı görülmektedir. Benzer biçimde “Aile gelir düzeyi – Orta gelir (30.000–50.000)”, “Sınavlara hazırlanmaya başlama süresi – 2 hafta”, “E-ders kaynakları – Var”, “Ders kitapları – Var”, “Yapay zeka kullanımı – Yok” ve “Çalışma ortamının motive edici olması – Çok huzurlu ve motive edici” kategorilerine ait değerlerin de büyük ölçüde negatif SHAP tarafında yoğunlaştığı görülmekte olup, bu kategorilerde yer alan öğrencilerde modelin akademik performans tahmininin daha düşük olma eğiliminde olduğu söylenebilir.

Genel olarak değerlendirildiğinde modelin tahminlerinde aileye ilişkin özellikler, derslere katılım durumu ve öğrencilerin ders çalışma alışkanlıkları gibi faktörlerin öğrencinin akademik performansını tahmin etmede daha belirleyici oldukları görülmektedir. Buna karşın sosyal medya kullanımı, teknoloji kullanımı (yapay zekâ), çalışma ortamı ve öğrenme tercihlerine ilişkin değişkenlerin modelin tahminine katkılarının görece daha sınırlı olduğu görülmektedir.

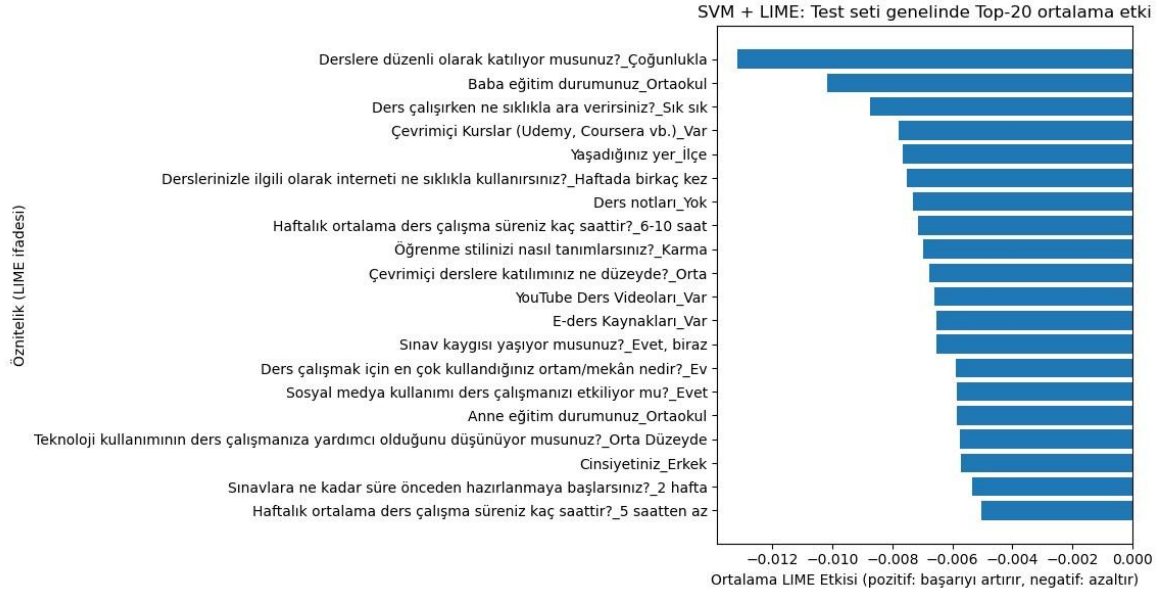
4.5.2 LIME

LIME analizi yorumlanırken, grafikte sunulan değerlerin test kümesindeki her bir değişken için elde edilen yerel LIME katsayılarının ortalamasına dayandığı kabul edilmiştir. LIME yöntemi, her bir değişken için modelin yalnızca o değişkenin yakın çevresindeki davranışını açıklayan yerel ve yorumlanabilir bir vekil model kurmaktadır; dolayısıyla elde edilen katsayılar, modelin genel yapısından çok, ilgili değişkene özgü yerel katkıları yansıtmaktadır (Ribeiro vd., 2016; Molnar, 2024). Bu nedenle test kümesindeki değişken ait yerel LIME

katsayılarının bir araya getirilmesi, değişkenlerin model karar sürecinde hangi ölçüde öne çıktığına ilişkin özet bir değerlendirme sunabilmektedir. Bununla birlikte, bu tür bir toplulaştırma doğrudan tam bir küresel açıklama olarak değil, değişkenlerin göreceli önem düzeyini gösteren bir yaklaşım olarak yorumlanmıştır; çünkü yerel açıklamaların nasıl birleştirildiği, elde edilen genel sonucun yapısını etkileyebilmektedir (Ribeiro vd., 2016; Ahern vd., 2019; van der Linden vd., 2019). (Ribeiro vd., 2016; Molnar, 2022). Bu nedenle Şekil 4.7’deki grafikte üst sıralarda yer alan değişkenler, Destek Vektör Makineleri (SVM) modelinin tahminlerini açıklamada test seti genelinde daha etkili olan değişkenler olarak yorumlanmıştır. Buna göre, derslere düzenli katılım, baba eğitim durumu ve ders çalışırken ara verme sıklığı, model kararında en yüksek ortalama etkiye sahip değişkenler arasında yer almaktadır. Özellikle “derslere düzenli olarak katılım” değişkeninin üst sıralarda bulunması, öğrencinin akademik katılımı ve disiplininin model tarafından ayırt edici bir özellik olarak kullanıldığını göstermektedir. Bu yorum, derslere katılımın akademik başarı ile anlamlı ve güçlü ilişkiler gösterdiğini ortaya koyan çalışmalarla da uyumludur (Credé vd., 2010). Benzer biçimde, baba eğitim durumunun öne çıkması, ebeveyn eğitim düzeyinin öğrencinin akademik çıktıları üzerindeki etkisini vurgulayan alanyazınla örtüşmektedir (Sirin, 2005). Ders çalışırken ara verme sıklığının öne çıkması da çalışma davranışları, öz düzenleme ve öğrenme stratejilerinin akademik performans üzerindeki belirleyici rolünü destekleyen bulgularla paralellik göstermektedir (Richardson vd., 2012; Broadbent ve Poon, 2015).

Şekil 4.7.’de verilen LIME grafiğinde dikkat çeken bir diğer değişken grubu, teknoloji kullanımı ve çevrimiçi kaynak kullanımına ilişkindir. Çevrimiçi kurslara katılım (UDEMY, YouTube) ders videoları kullanımı ve e-ders kaynaklarına erişim gibi faktörlerin üst sıralarda yer alması, modelin öğrencilerin dijital öğrenme davranışlarını da karar sürecine dâhil ettiğini göstermektedir. Bu bulgu, çevrimiçi öğrenme ortamlarının uygun kullanımının öğrenme çıktıları üzerinde etkili olabileceğini gösteren çalışmalarla uyumludur (Means vd., 2013). Bunun yanı sıra sınav kaygısı, sosyal medya kullanımının ders çalışmaya etkisi ve ders (çalışma) notlarının bulunmaması gibi değişkenlerin de dikkate değer bir ortalama etkiye sahip olduğu görülmektedir. Sınav kaygısının öne çıkması, kaygı düzeyinin bilişsel performans ve başarı üzerinde sınırlayıcı etkiler oluşturabileceğini gösteren araştırmalarla desteklenmektedir (Cassady ve Johnson, 2002). Sosyal medya kullanımının ders çalışmaya etkisine ilişkin bulgu da sosyal medya kullanımının dikkat dağınıklığı ve akademik görevlerden sapmanın başarı ile ilişkili olabileceğini ortaya koyan çalışmalarla uyumludur

(Junco, 2012a). Benzer şekilde, ders notlarının bulunmaması ya da not tutma süreçlerinin zayıf olması, öğrenme materyalinin yapılandırılması ve tekrar edilmesi bakımından dezavantaj yaratabileceğinden, akademik performansla ilişkili bir unsur olarak değerlendirilebilir (Kiewra, 1985; Kobayashi, 2005).



Şekil 4.7: Modele en çok etki eden 20 özneliğin ortalama LIME değerleri grafiği

Tablo 4.57 ve Şekil 4.8’de sunulan pozitif ortalama LIME grafiği, SVM modelinin “başarılı” çıktısına ilişkin tahminini artıran başlıca kategori ve değişkenleri göstermektedir. Grafikte en yüksek ortalama pozitif etkinin “derslere düzenli olarak katılıyor musunuz? Her zaman” kategorisinde görülmesi, akademik devamlılığın model tarafından başarıyı ayırt etmede temel bir gösterge olarak kullanıldığını ortaya koymaktadır. Bunu devam oranı, ders çalışırken nadiren ara verme, sınav kaygısının olmaması, yapay zekâ araçlarının kullanılması, kütüphanede çalışma, haftalık 16–20 saat ders çalışma ve sınavlara uzun süre önceden hazırlanmaya başlama gibi değişkenler izlemektedir. Bu bulgular, derslere devamın akademik başarıyla güçlü ilişkisini vurgulayan; öz-düzenleme, zaman yönetimi ve etkili öğrenme stratejilerinin başarı üzerindeki rolünü destekleyen çalışmalarla uyumludur (Credé vd., 2010; Richardson vd., 2012; Broadbent ve Poon, 2015). Ayrıca “Sınav kaygısı yaşıyor musunuz? Hayır” kategorisinin üst sıralarda yer alması, düşük sınav kaygısının akademik performansı desteklediğini gösteren araştırmalarla örtüşmektedir (Cassady ve Johnson, 2002). Yapay zekâ ve dijital öğrenme araçlarının pozitif etkisi ise teknoloji destekli öğrenmenin uygun kullanımının akademik çıktılara katkı sağlayabileceğini

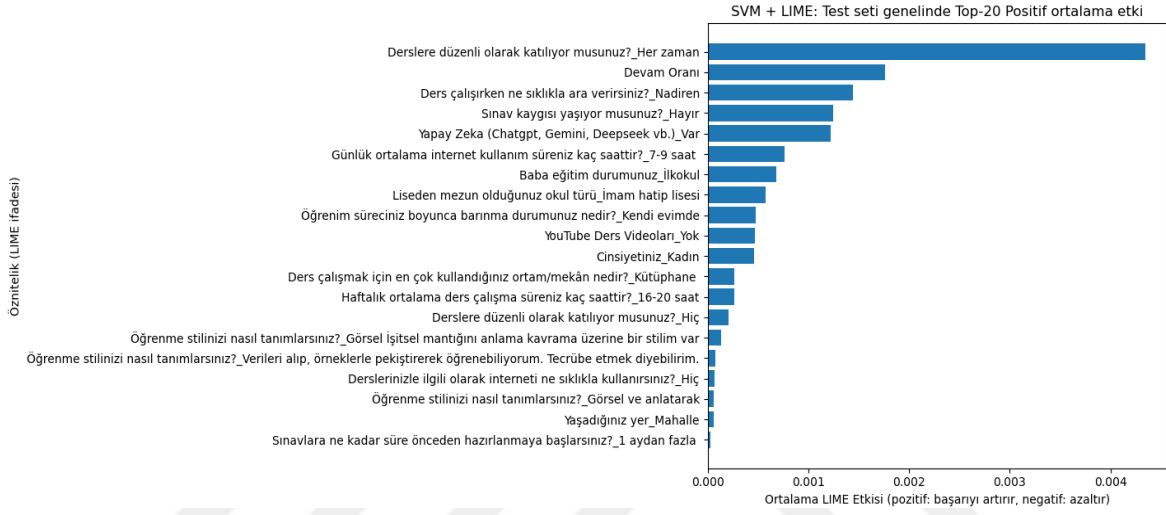
düşündürmektedir (Means vd., 2014). Bununla birlikte, bu sonuçların kategori düzeyinde yorumlanması gerekmektedir. Dolayısıyla grafikte pozitif etki gösteren sosyo-demografik kategoriler, örneğin baba eğitim durumu, mezun olunan okul türü, cinsiyet ya da barınma durumu, doğrudan nedensel bir üstünlük olarak değil; yalnızca veri seti içindeki örüntüler çerçevesinde modelin “başarılı” çıktısına yönelik tahminini artıran bağlamsal göstergeler olarak değerlendirilmektedir (Ribeiro vd., 2016; Molnar, 2020; Sirin, 2005). Sonuç olarak LIME bulguları, SVM modelinin öğrenci başarısını akademik katılım, çalışma disiplini, teknoloji kullanımı ve sosyo-demografik bağlamı birlikte dikkate alan çok boyutlu bir karar mekanizmasıyla açıkladığını göstermektedir.

Tablo 4.57: Modele pozitif etki eden 20 öz niteliğin ortalama LIME değerleri tablosu

	Öznitelik (LIME)	Ortalama LIME Etkisi	Std LIME Etkisi	Mutlak Ortalama Etkisi	Sayı
0	Derslere düzenli olarak katılıyor musunuz? Her zaman	0,0043444	0,0006732	0,0043444	226
1	Devam Oranı	0,0017646	0,0007346	0,0017677	226
2	Ders çalışırken ne sıklıkla ara verirsiniz? Nadiren	0,0014396	0,0006918	0,0014477	226
3	Sınav kaygısı yaşıyor musunuz? Hayır	0,0012481	0,0007356	0,0012994	226
4	Yapay Zekâ (Chatgpt, Gemini, Deepseek vb.)_Var	0,0012239	0,0007275	0,0012504	226
5	Günlük ortalama internet kullanım süreniz kaç saattir?_7-9 saat	0,0007626	0,0007127	0,0008801	226
6	Baba eğitim durumunuz_İlkokul	0,0006816	0,0007876	0,0008437	226
7	Liseden mezun olduğunuz okul türü_İmam hatip lisesi	0,0005749	0,0006681	0,0007484	226
8	Öğrenim süreciniz boyunca barınma durumunuz nedir?_Kendi evimde	0,0004750	0,0007947	0,0007288	226
9	YouTube Ders Videoları_Yok	0,0004668	0,0007316	0,0006934	226
10	Cinsiyetiniz_Kadın	0,0004612	0,0007186	0,0006899	226
11	Ders çalışmak için en çok kullandığınız ortam/mekân nedir?_Kütüphane	0,0002673	0,0008036	0,0007025	226
12	Haftalık ortalama ders çalışma süreniz kaç saattir?_16-20 saat	0,0002623	0,0007412	0,0005868	226
13	Derslere düzenli olarak katılıyor musunuz?_Hiç	0,0002049	0,0007048	0,0005995	226
14	Öğrenme stilinizi nasıl tanımlarsınız?_Görsel	0,0001372	0,0010925	0,0006412	226
15	İşitsel mantığını anlama kavrama üzerine bir stilim var				
	Öğrenme stilinizi nasıl tanımlarsınız?_Verileri alıp, örneklerle pekiştirerek öğrenebiliyorum. Tecrübe etmek diyebilirim.	0,0000737	0,0007363	0,0006004	226

Tablo 4.57: (devam ediyor)

16	Derslerinle ilgili olarak interneti ne sıklıkla kullanırsınız? _Hiç	0,0000717	0,0007083	0,0005722	226
17	Öğrenme stilinizi nasıl tanımlarsınız? _Görsel ve anlatarak	0,0000619	0,0007158	0,0005722	226
18	Yaşadığınız yer _Mahalle	0,0000614	0,0007357	0,0005878	226
19	Sınavlara ne kadar süre önceden hazırlanmaya başlarsınız? _1 aydan fazla	0,0000273	0,0007002	0,0005478	226



Şekil 4.8: Modele pozitif etki eden 20 özneliliğin ortalama LIME değerleri grafiği

Tablo 4.58 ve Şekil 4.9’da sunulan negatif ortalama LIME grafiği, işaretli ortalama LIME katsayılarına dayanması koşuluyla, SVM modelinin “başarılı” çıktısına ilişkin tahminini azaltan başlıca kategori örüntülerini göstermektedir (Ribeiro vd., 2016; Molnar, 2020). Grafikte en yüksek negatif etkiye sahip kategorinin “derslere düzenli olarak katılıyor musunuz? Çoğunlukla” olması, modelin başarıyı açıklarken yalnızca derse katılımı değil, tam ve sürekli katılım düzeyini de ayırt edici bir özellik olarak kullandığını göstermektedir. Bu bulgu, pozitif LIME grafiğinde öne çıkan “Her zaman” katılım kategorisi ile değerlendirildiğinde, akademik devamlılığın başarı tahmininde temel belirleyicilerden biri olduğunu ortaya koymaktadır (Credé vd., 2010). Benzer biçimde, “ders çalışırken sık sık ara verme”, “haftalık 6–10 saat ders çalışma” ve “sınavlara 2 hafta kala hazırlanmaya başlama” gibi kategoriler, model tarafından daha düşük başarı olasılığıyla ilişkilendirilen çalışma alışkanlıkları olarak öne çıkmaktadır. Bu durum, öz-düzenleme, zaman yönetimi ve çalışma stratejilerinin akademik başarı üzerindeki rolünü vurgulayan araştırmalarla uyumludur (Richardson vd., 2012; Broadbent ve Poon, 2015). Ayrıca “sınav kaygısı yaşıyor musunuz? Evet, biraz” kategorisinin negatif etkisi, sınav kaygısının akademik performansı sınırlayıcı

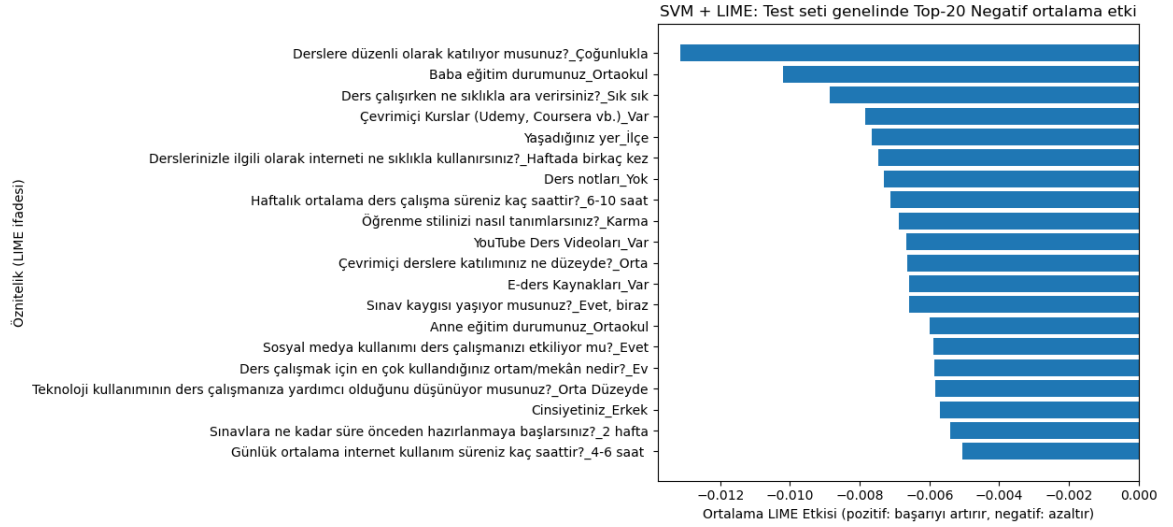
etkileri olabileceğini gösteren bulgularla örtüşmektedir (Cassady ve Johnson, 2002). “Ders notları_yok” kategorisinin üst sıralarda yer alması da not alma ve not gözden geçirme süreçlerinin öğrenme başarısı bakımından önemini desteklemektedir (Kiewra, 1985; Kobayashi, 2006). Dijital öğrenme ile ilişkili olarak çevrimiçi kurslara katılım, YouTube ders videoları kullanımı, e-ders kaynaklarına erişim ve günlük internet kullanımına ilişkin bazı kategorilerin negatif yönde görünmesi ise teknoloji kullanımının kendisinden çok, kullanımın niteliği ve bağlamının belirleyici olduğunu düşündürmektedir (Means vd., 2014; Junco, 2012a). Son olarak ebeveyn eğitim düzeyi, yaşanan yer ve cinsiyet gibi sosyo-demografik kategorilerin de model kararında yer alması, öğrenci başarısının çok boyutlu bir yapı sergilediğini göstermektedir; ancak bu kategorilerin yorumunun nedensel değil, veri setine özgü model temelli örüntüler olarak yapılması gerekmektedir (Sirin, 2005). Bu çerçevede negatif LIME bulguları, SVM modelinin düşük başarı olasılığını daha çok kısmi katılım, sınırlı öz-düzenleme, yetersiz planlama, kaygı ve bazı bağlamsal sosyo-demografik özelliklerle ilişkilendirdiğini ortaya koymaktadır.

Tablo 4.58: Modele negatif etki eden 20 öz niteliğin ortalama LIME değerleri tablosu

Öznitelik (LIME)	Ortalama LIME Etkisi	Std LIME Etkisi	Mutlak Ortalama Etkisi	Sayı
0 Derslere düzenli olarak katılıyor musunuz?_Çoğunlukla	-0,0131421	0,0007667	0,0131421	226
1 Baba eğitim durumunuz_Ortaokul	-0,0101977	0,0007387	0,0101977	226
2 Ders çalışırken ne sıklıkla ara verirsiniz?_Sık sık	-0,0088581	0,0007428	0,0088581	226
3 Çevrimiçi Kurslar (Udemy, Coursera vb.)_Var	-0,0078332	0,0006961	0,0078332	226
4 Yaşadığınız yer_İlçe	-0,0076476	0,0007528	0,0076476	226
5 Derslerinizle ilgili olarak interneti ne sıklıkla kullanırsınız?_Haftada birkaç kez	-0,0074709	0,0006382	0,0074709	226
6 Ders notları_Yok	-0,0073133	0,0007359	0,0073133	226
7 Haftalık ortalama ders çalışma süreniz kaç saattir?_6-10 saat	-0,0071260	0,0007255	0,0071260	226
8 Öğrenme stilinizi nasıl tanımlarsınız?_Karma	-0,0068812	0,0007702	0,0068812	226
9 YouTube Ders Videoları_Var	-0,0066681	0,0007656	0,0066681	226
10 Çevrimiçi derslere katılımınız ne düzeyde?_Orta	-0,0066333	0,0007362	0,0066333	226
11 E-ders Kaynakları_Var	-0,0065861	0,0007052	0,0065861	226
12 Sınav kaygısı yaşıyor musunuz?_Evet, biraz	-0,0065818	0,0007574	0,0065818	226

Tablo 4.58: (devam ediyor)

13	Anne eğitim durumunuz_Ortaokul	-0,0060066	0,0007292	0,0060066	226
14	Sosyal medya kullanımı ders çalışmanızı etkiliyor mu?_Evet	-0,0058868	0,0007396	0,0058868	226
15	Ders çalışmak için en çok kullandığınız ortam/mekân nedir?_Ev	-0,0058578	0,0007054	0,0058578	226
16	Teknoloji kullanımının ders çalışmanıza yardımcı olduğunu düşünüyor musunuz?_Orta Düzeyde	-0,0058410	0,0007703	0,0058410	226
17	Cinsiyetiniz_Erkek	-0,0056992	0,0007658	0,0056992	226
18	Sınavlara ne kadar süre önceden hazırlanmaya başlarsınız?_2 hafta	-0,0053984	0,0007071	0,0053984	226
19	Günlük ortalama internet kullanım süreniz kaç saattir?_4-6 saat	-0,0050543	0,0007882	0,0050543	226



Şekil 4.9: Modele negatif etki eden 20 özneliliğin ortalama LIME değerleri grafiği

5. SONUÇ, TARTIŞMA VE ÖNERİLER

Bu bölümde; “Sınıflandırma Sonuçlarının Değerlendirilmesi, Regresyon Sonuçlarının Değerlendirilmesi, Model Karşılaştırması, Açıklanabilir Yapay Zeka Bulgularının Değerlendirilmesi” başlıkları yer almaktadır.

5.1 Sınıflandırma Sonuçlarının Değerlendirilmesi

Bu çalışmada öğrencilerin akademik performanslarını tahmin etmek amacıyla farklı makine öğrenmesi yaklaşımları kullanılmış ve sınıflandırma modellerinin performansları çeşitli değerlendirme ölçütleri aracılığıyla analiz edilmiştir. Sınıflandırma modellerinin değerlendirilmesinde doğruluk oranı (Accuracy), F1 skoru, dengeli doğruluk (Balanced Accuracy), ROC-AUC ve PR-AUC gibi ölçütler kullanılmıştır. Bu ölçütlerin birlikte değerlendirilmesi model performansının yalnızca doğru tahmin oranı açısından değil, aynı zamanda sınıflar arasındaki ayırım gücü açısından da analiz edilmesini mümkün kılmıştır.

Yapılan analizler sonucunda Destek Vektör Makine (Support Vector Machine-SVM) algoritmasının sınıflandırma performansı açısından en başarılı model olduğu belirlenmiştir. Modelin doğruluk oranı 0.9071 olarak hesaplanmıştır. Bununla birlikte F1 skoru 0.9129, dengeli doğruluk değeri 0.9062 olarak elde edilmiştir. Bu değerler modelin yalnızca doğru sınıflandırma oranının yüksek olmadığını, aynı zamanda modelin farklı sınıfları dengeli bir şekilde ayırt edebildiğini göstermektedir.

SVM modelinin ROC-AUC değeri 0.9389 olarak hesaplanmıştır. ROC-AUC değeri modelin sınıfları ayırt etme yeteneğini ölçen önemli bir performans göstergesidir. Bu değer 0.90’ın üzerinde olması modelin güçlü bir ayırt ediciliğe sahip olduğunu göstermektedir. Benzer şekilde PR-AUC değerinin 0.9294 olarak hesaplanması modelin özellikle dengesiz veri setlerinde bile güvenilir tahminler üretebildiğini göstermektedir.

Bu çalışmada elde edilen sınıflandırma sonuçları eğitim analitiği alanyazını ile de uyum göstermektedir. Eğitim veri madenciliği çalışmalarında makine öğrenmesi algoritmalarının öğrenci başarısını tahmin etmede etkili araçlar sunduğu belirtilmektedir (Romero ve Ventura, 2010). Örneğin; Alsaman ve diğerleri (2019) tarafından gerçekleştirilen bir

çalışmada karar ağaçları ve yapay sinir ağları kullanılarak öğrencilerin akademik performansları tahmin edilmiş ve makine öğrenmesi algoritmalarının eğitim verileri üzerinde yüksek doğruluk oranlarına ulaşabildiği belirtilmiştir. Benzer şekilde Alnasyan vd. (2024) tarafından yapılan araştırmada farklı makine öğrenmesi algoritmaları kullanılarak öğrencilerin akademik performansları tahmin edilmiş ve özellikle destek vektör makineleri gibi güçlü sınıflandırma algoritmalarının eğitim verilerinde başarılı sonuçlar üretebildiği vurgulanmıştır. Ayrıca Su vd. (2023) tarafından gerçekleştirilen çalışmada büyük veri ve makine öğrenmesi algoritmaları kullanılarak öğrencilerin akademik performansları tahmin edilmiş ve sınıflandırma modellerinin öğrencilerin akademik performanslarını yüksek doğruluk oranları ile tahmin edebildiği belirtilmiştir. Bu çalışmalar makine öğrenmesi yöntemlerinin eğitim analitiği alanında güçlü araçlar sunduğunu göstermektedir.

Bu çalışmada elde edilen yaklaşık %91 doğruluk oranı, alanyazında rapor edilen çalışmalarla karşılaştırıldığında benzer ya da bazı durumlarda daha yüksek bir model performansına işaret etmektedir. Bu sonuç, kullanılan veri setinin farklı veri kaynaklarını bir araya getirmesinin model performansını artıran önemli bir unsur olabileceğini göstermektedir. Özellikle Üniversite Bilgi Yönetim Sistemi (UBYS) üzerinden elde edilen akademik veriler ile öğrencilerin öğrenme alışkanlıklarını ve davranışlarını yansıtan anket verilerinin birlikte kullanılması, makine öğrenmesi modelinin öğrenci akademik performansını daha doğru ve bütüncül bir biçimde tahmin etmesine katkı sağlamış olabilir. Bu bağlamda çoklu veri kaynaklarının entegrasyonu, öğrenci akademik performansının tahminine yönelik geliştirilen makine öğrenmesi modellerinin açıklayıcılığını ve tahmin gücünü artıran önemli bir yaklaşım olarak değerlendirilebilir.

Bu bağlamda elde edilen sonuçlar, makine öğrenmesi tabanlı modellerin eğitim kurumlarında erken uyarı sistemlerinin geliştirilmesinde etkili biçimde kullanılabileceğini göstermektedir. Nitekim alanyazında, öğrenme yönetim sistemi ve akademik süreç verilerine dayalı tahmin modellerinin öğrencilerin başarı olasılığını erken dönemde öngörebildiği ve bu sayede öğretim elemanlarına risk altındaki öğrencileri zamanında belirleme olanağı sunduğu belirtilmektedir (Macfadyen ve Dawson, 2010; Akçapınar vd., 2019). Benzer biçimde, sistematik derleme bulguları makine öğrenmesi tekniklerinin öğrenci başarısızlığı ve okul terk riski taşıyan bireylerin belirlenmesinde önemli bir rol oynadığını ve bu modellerin öğrenci performansını iyileştirmeye dönük kurumsal stratejilere

katkı sağlayabildiğini ortaya koymaktadır (Albreiki vd., 2021). Ayrıca erken tahmin yaklaşımlarının ders sürecinin henüz erken aşamalarında dahi uygulanabildiği, böylece akademik destek, yönlendirme ve kaynak tahsisinin daha zamanında gerçekleştirilebildiği ifade edilmektedir (Santos ve Henriques, 2023). Bunun yanında, makine öğrenmesi temelli müdahale çerçevelerinin yalnızca riskli öğrencilerin belirlenmesine değil, aynı zamanda öğrencilerin bilgi eksiklerinin saptanmasına, kişiselleştirilmiş çalışma planlarının oluşturulmasına ve eğitim çıktılarının anlamlı biçimde iyileştirilmesine de katkı sunduğu bildirilmektedir (Alalawi vd., 2024). Bu doğrultuda, öğrencilerin akademik performanslarının erken aşamada tahmin edilmesi, başarısızlık riski taşıyan öğrencilerin belirlenmesine ve gerekli akademik destek mekanizmalarının zamanında uygulanmasına katkı sağlayabilecek önemli bir karar destek aracı olarak değerlendirilebilir.

5.2 Regresyon Sonuçlarının Değerlendirilmesi

Araştırmanın ikinci aşamasında öğrencilerin akademik performansları sürekli bir değişken olarak ele alınmış ve regresyon yaklaşımı kullanılarak tahmin edilmiştir. Regresyon analizlerinde model performanslarının değerlendirilmesinde determinasyon katsayısı (R^2), kök ortalama kare hata (RMSE) ve ortalama mutlak hata (MAE) gibi ölçütler kullanılmıştır.

Yapılan analizler sonucunda XGBoost regresyon modelinin diğer modeller arasında en başarılı performansı gösterdiği belirlenmiştir. Test veri seti üzerinde elde edilen sonuçlara göre modelin R^2 değeri 0.6963 olarak hesaplanmıştır. Bu değer modelin akademik performans tahmini için kullanılan Başarı Durumu bağımlı değişkenindeki varyansın yaklaşık %69'unu açıklayabildiğini göstermektedir. Eğitim verilerinde bu düzeyde bir açıklayıcılık oranı modelin güçlü bir tahmin performansı sunduğunu göstermektedir (Lou ve Colvin, 2025).

Model performansının hata ölçütleri açısından değerlendirilmesi de önemli sonuçlar ortaya koymaktadır. XGBoost modeli için hesaplanan RMSE değeri 12.4614 olarak belirlenmiştir. RMSE değeri modelin tahmin hatalarının ortalama büyüklüğünü göstermektedir. Bu değer düşük olması modelin gerçek değerlere yakın tahminler üretebildiğini göstermektedir (Hudson, 2022). Ayrıca ortalama mutlak hata (MAE) değerinin 8.9904 olarak hesaplanması modelin tahmin hatalarının genel olarak düşük seviyede olduğunu ortaya koymaktadır.

Bu çalışmada elde edilen sonuçlar XGBoost algoritmasının akademik performans tahmini gibi çok değişkenli veri yapılarında güçlü bir model olduğunu göstermektedir. Özellikle öğrencilerin akademik verileri ile davranışsal verilerinin birlikte kullanılması modelin tahmin gücünü artıran önemli bir faktör olarak ortaya çıkmıştır. Bu bulgular alanyazında yer alan benzer çalışmalarla da paralellik göstermektedir. Örneğin; Escamilla ve Hosseini (2023) tarafından gerçekleştirilen çalışmada makine öğrenmesi algoritmalarının öğrencilerin akademik performanslarını tahmin etmede etkili sonuçlar ürettiği belirtilmiş ve özellikle topluluk öğrenmesi algoritmalarının yüksek tahmin doğruluğu sağlayabildiği vurgulanmıştır. Benzer şekilde; Kumar ve Sharma (2024) tarafından yapılan çalışmada farklı makine öğrenmesi algoritmaları karşılaştırılmış ve boosting tabanlı algoritmaların eğitim verilerinde güçlü tahmin performansı sunduğu ifade edilmiştir. Ayrıca Bressane ve diğerleri (2024) tarafından gerçekleştirilen araştırmada öğrencilerin akademik performanslarının tahmin edilmesinde öğrenme stratejileri, öğrenme güçlükleri ve öğrencilerin öğrenme davranışlarına ilişkin değişkenlerin modele dahil edilmesinin tahmin doğruluğunu artırdığı belirtilmiştir. Bu çalışmada da öğrencilerin akademik verilerinin yanı sıra davranışsal verilerin kullanılması model performansını olumlu yönde etkilemiş olabilir. Benzer şekilde Ali ve Khan (2024) tarafından yapılan çalışmada hibrit makine öğrenmesi modellerinin öğrencilerin akademik performanslarını tahmin etmede güçlü sonuçlar üretebildiği ve özellikle topluluk (ensemble) tabanlı modellerin daha yüksek performans sağlayabildiği belirtilmektedir. Bu bulgular XGBoost algoritmasının bu çalışmada en iyi regresyon performansını göstermesini destekleyen alanyazın bulguları ile uyumludur.

Bu bağlamda elde edilen sonuçlar akademik performans tahmini çalışmalarında topluluk öğrenmesi (ensemble) algoritmalarının önemli bir potansiyele sahip olduğunu göstermektedir. Özellikle öğrencilerin akademik performanslarını etkileyen farklı veri kaynaklarının bir araya getirilmesi makine öğrenmesi modellerinin tahmin doğruluğunu artırabilmektedir. Eğitim kurumlarında bu tür veri temelli yaklaşımların kullanılması öğrencilerin akademik gelişimlerinin daha iyi anlaşılmasına ve eğitim süreçlerinin iyileştirilmesine katkı sağlayabilir.

5.3 Model Karşılaştırması

Sınıflandırma ve regresyon modellerinin birlikte değerlendirilmesi akademik performans

tahmini açısından farklı yaklaşımların avantajlarını ortaya koymaktadır (Strecht vd., 2015). Sınıflandırma modelleri öğrencilerin belirli başarı kategorilerine ayrılmasını sağlarken regresyon modelleri öğrencilerin akademik performanslarını sayısal bir değer olarak tahmin etmeye olanak tanımaktadır (Sayed vd., 2025).

Sınıflandırma modelleri özellikle erken uyarı sistemleri açısından önemli avantajlar sunmaktadır. Öğrencilerin başarılı veya başarısız olma risklerinin önceden belirlenmesi eğitim kurumlarının gerekli müdahale stratejilerini geliştirmesine olanak sağlayabilmektedir (Macfadyen ve Dawson, 2010). Regresyon modelleri ise öğrencilerin akademik performanslarını daha ayrıntılı bir şekilde analiz etmeye yardımcı olmaktadır (Sayed vd., 2025).

Bu çalışmada elde edilen sonuçlar sınıflandırma ve regresyon modellerinin birbirini tamamlayan yaklaşımlar olduğunu göstermektedir. Sınıflandırma modelleri risk analizi için kullanılabilirken regresyon modelleri öğrencilerin akademik performanslarının daha ayrıntılı tahmin edilmesine katkı sağlayabilmektedir.

5.4 Açıklanabilir Yapay Zekâ Bulgularının Değerlendirilmesi

Bu çalışmanın önemli katkılarından biri geliştirilen makine öğrenmesi modellerinin açıklanabilir yapay zekâ yöntemleri kullanılarak yorumlanmasıdır. Makine öğrenmesi modelleri yüksek tahmin doğruluğu sağlayabilmelerine rağmen çoğu zaman karar süreçlerinin anlaşılması zor olan “kara kutu” yapılar olarak değerlendirilmektedir (Adadi ve Berrada, 2018). Bu durum özellikle eğitim gibi insan merkezli alanlarda model sonuçlarının yorumlanmasını zorlaştırabilmektedir. Bu nedenle model çıktılarının yorumlanabilirliğini artırmak amacıyla çalışmada SHAP ve LIME açıklanabilir yapay zekâ yöntemleri kullanılmıştır. Açıklanabilir yapay zekâ yaklaşımları, modelin hangi değişkenlere dayanarak tahmin ürettiğini ortaya koyarak model kararlarının daha anlaşılır ve güvenilir hale gelmesini sağlamaktadır (Adadi ve Berrada, 2018; Arrieta vd., 2020).

5.4.1 SHAP Analizi Sonuçlarının Değerlendirilmesi

SHAP analizi sonuçlarına göre öğrencilerin akademik performansını etkileyen en önemli

değişkenlerin başında baba eğitim durumu gelmektedir. Bu değişkenin SHAP değeri 0.004437 olarak hesaplanmıştır. Baba eğitim durumunu sırasıyla ders (çalışma) notları, çevrimiçi kurs kullanımı, sosyal medya kullanımının ders çalışmaya etkisi ve e-ders kaynaklarının kullanımı değişkenleri takip etmektedir.

Bu bulgular öğrencilerin akademik performanslarının yalnızca akademik faktörlerden değil aynı zamanda sosyo-ekonomik ve davranışsal faktörlerden de etkilendiğini göstermektedir. Özellikle ebeveyn eğitim düzeyinin öğrencilerin akademik performansları üzerindeki etkisi eğitim alanyazınında sıklıkla vurgulanan bir konudur. Sirin (2005) tarafından gerçekleştirilen meta-analiz çalışmasında ailelerin sosyo-ekonomik durumlarının öğrencilerin akademik performansları üzerinde anlamlı bir etkisi olduğu belirtilmektedir. Benzer şekilde Gašević vd. (2016), çalışmalarında öğrencilerin akademik performanslarının yalnızca akademik verilerle değil, aynı zamanda demografik ve çevresel faktörlerle de ilişkili olduğunu ifade etmektedir. Bu çalışmada elde edilen SHAP bulguları söz konusu literatür sonuçları ile paralellik göstermektedir. Ayrıca SHAP bulgularına göre ders notlarını çalışma kaynağı olarak kullanmayan öğrencilerde modelin akademik performans tahmini daha düşük olma eğilimindedir.

Bununla birlikte çevrimiçi kurs kullanımı ve e-ders kaynaklarının kullanımı gibi değişkenlerin önemli faktörler arasında yer alması, dijital öğrenme ortamlarının öğrencilerin öğrenme süreçlerinde giderek daha belirleyici hale geldiğini göstermektedir. Günümüzde çevrimiçi öğrenme platformlarının yaygınlaşması öğrencilerin öğrenme süreçlerini destekleyen önemli öğrenme araçları oluşturmuştur. Bressane vd. (2024) tarafından gerçekleştirilen çalışmada öğrenme stratejileri ile dijital öğrenme kaynaklarının öğrencilerin akademik performansı üzerinde önemli etkileri olduğu belirtilmektedir. Bu bulgular çalışmada elde edilen sonuçlarla uyumludur.

5.4.2 LIME Analizi Sonuçlarının Değerlendirilmesi

LIME analizi sonuçları, öğrencilerin akademik performanslarını etkileyen değişkenlerin bireysel tahminler düzeyinde nasıl rol oynadığını ortaya koymaktadır. Elde edilen bulgular, modelin özellikle derslere düzenli katılım, devam oranı, ders çalışma sırasında ara verme sıklığı, sınav kaygısı, ders (çalışma) notlarının varlığı, çevrimiçi kurslara katılım ve

öğrencilerin yaşadığı yer gibi değişkenleri tahmin sürecinde dikkate aldığını göstermektedir. Bu yönüyle LIME analizi, öğrencilerin akademik performanslarının yalnızca not ya da başarı puanı gibi doğrudan akademik göstergelerle değil; aynı zamanda çalışma disiplini, öğrenme alışkanlıkları ve öğrenme çevresiyle birlikte şekillendiğini ortaya koymaktadır.

Özellikle derslere düzenli katılım değişkeninin hem pozitif hem de negatif LIME örüntülerinde öne çıkması dikkat çekicidir. “Her zaman” katılım kategorisinin başarı tahminini artıran değişkenler arasında yer alması, buna karşılık “çoğunlukla” katılım kategorisinin negatif etki göstermesi, modelin yalnızca derse katılmayı değil, katılımın süreklilik ve istikrar düzeyini de ayırt edici bir unsur olarak değerlendirdiğini göstermektedir. Bu bulgu, derslere devamın akademik başarı üzerindeki belirleyici etkisini vurgulayan alanyazınla uyumludur (kaynak). Aynı şekilde ders çalışırken nadiren ara verme, düzenli çalışma, sınav kaygısının düşük olması ve ders notlarının bulunması gibi unsurların daha olumlu tahminlerle ilişkilendirilmesi; öz düzenleme, zaman yönetimi ve etkili öğrenme stratejilerinin öğrencilerin başarı düzeyleri üzerinde önemli rol oynadığını desteklemektedir.

Bununla birlikte çevrimiçi kurslar, YouTube ders videoları, e-ders kaynakları ve internet kullanımı gibi dijital öğrenme göstergelerinin LIME sonuçlarında öne çıkması, teknoloji kullanımının akademik performans açısından tek başına değil, kullanımın niteliği ve amacı doğrultusunda değerlendirilmesi gerektiğini göstermektedir. Benzer şekilde baba eğitim durumu, anne eğitim durumu, cinsiyet ve yaşanan yer gibi sosyo-demografik değişkenlerin model tahminlerinde etkili olması, öğrenci başarısının çok boyutlu bir yapıya sahip olduğunu ortaya koymaktadır (Costa vd., 2024; Tadesse vd., 2022; Huang ve Mok, 2025). Ancak bu değişkenlerin yorumlanmasında nedensel bir çıkarım yapmak yerine, bunların veri seti içindeki örüntülere dayalı bağlamsal göstergeler olarak ele alınması daha uygun olacaktır; çünkü tahminsel modellerde ortaya çıkan ilişkiler açıklama ve nedensellik ile özdeş değildir ve gözlemsel verilerden elde edilen örüntüler ek kuramsal ve yöntemsel dayanak olmadan doğrudan nedensel etki olarak yorumlanamaz (Shmueli, 2010; Weidlich vd., 2022). Ayrıca, eğitimde kullanılan tahmin modellerinde demografik değişkenler çoğu zaman diğer akademik ve kurumsal değişkenlerle güçlü biçimde ilişkili olduğundan, bu değişkenler çoğu kez doğrudan nedenlerden çok daha geniş bağlamsal eşitsizliklerin ve örüntülerin göstergeleri olarak değerlendirilmektedir (Bird vd., 2021).

Genel olarak deęerlendirildięinde LIME analizi, SVM modelinin öęrenci akademik performansını açıklarken akademik katılım, alıřma davranıřları, dijital öęrenme pratikleri ve sosyo-demografik baęlamı birlikte dikkate alan ok boyutlu bir karar yapısı geliřtirdięini gstermektedir. Bu sonular, açıklanabilir yapay zekâ yöntemlerinin yalnızca model performansını yorumlamada deęil, aynı zamanda öęrencilerin başarılarını etkileyen temel faktrleri daha ayrıntılı ve anlaşılır biimde ortaya koymada da önemli katkılar sunduęunu gstermektedir. Bu aıdan LIME, eęitim alanında geliřtirilen makine öęrenmesi modellerinin daha řeffaf, yorumlanabilir ve uygulamaya dnük biimde kullanılmasına katkı saęlayan işlevsel bir yöntem olarak deęerlendirilebilir.

5.5 Eęitim Sistemine Ynelik ıkarımlar

Arařtırma bulguları eęitim kurumlarının veri temelli karar verme srelerinin geliřtirilmesi aısından önemli sonular ortaya koymaktadır. Eęitim kurumlarında öęrenci verilerinin analiz edilmesi öęrencilerin akademik geliřimlerinin daha yakından takip edilmesine olanak saęlayabilir. Gnmzde eęitim kurumlarında retilen byk veri miktarının artması, öęrencilerin öęrenme srelerinin daha ayrıntılı bir řekilde analiz edilmesini mmkn kılmaktadır. Bu durum eęitsel veri madencilięi yaklařımlarının eęitim alanında giderek daha fazla kullanılmasına yol amıřtır (Romero ve Ventura, 2010).

Makine öęrenmesi tabanlı modellerin eęitim sistemine entegre edilmesi öęrencilerin akademik performanslarının erken ařamada tahmin edilmesine ve gerekli akademik destek mekanizmalarının geliřtirilmesine katkı saęlayabilir. Bu tr sistemler zellikle akademik performanssızlık riski tařıyan öęrencilerin erken tespit edilmesine yardımcı olabilir. Öęrencilerin akademik performanslarının erken ařamada tahmin edilmesi, öęretim elemanlarının ve eęitim yneticilerinin gerekli mdahaleleri zamanında gerekleřtirebilmesine olanak tanımaktadır (Lopez vd., 2021).

Bu baęlamda makine öęrenmesi tabanlı akademik performans tahmin modelleri eęitim kurumları iin önemli bir karar destek mekanizması olarak deęerlendirilebilir. Öęrencilerin akademik performanslarının nceden tahmin edilmesi, dřk performans gsterme riski tařıyan öęrencilerin belirlenmesine ve bu öęrenciler iin rehberlik, akademik danıřmanlık veya destekleyici öęrenme programlarının uygulanmasına katkı saęlayabilir (Fahd vd.,

2022).

Ayrıca bu tür veri temelli yaklaşımlar eğitim kurumlarının eğitim politikalarını daha etkili bir şekilde planlamasına da yardımcı olabilir. Öğrenci başarısını etkileyen faktörlerin belirlenmesi, eğitim yöneticilerinin kaynakları daha verimli kullanmasına ve öğrencilerin öğrenme süreçlerini destekleyen stratejiler geliştirmesine olanak sağlayabilir. Gašević vd. (2015) makine öğrenmesinin eğitim sistemlerinde veri temelli karar alma kültürünün gelişmesine katkı sağladığını ve eğitim süreçlerinin daha sistematik bir şekilde değerlendirilmesine olanak tanıdığını belirtmektedir.

Bu çalışmada elde edilen bulgular makine öğrenmesi ve açıklanabilir yapay zekâ yaklaşımlarının eğitim alanında yalnızca akademik performans tahmini yapmakla kalmayıp aynı zamanda öğrencilerin öğrenme süreçlerini daha iyi anlamaya da katkı sağlayabileceğini göstermektedir. Bu nedenle eğitim kurumlarında veri analitiği temelli yaklaşımların yaygınlaştırılması öğrencilerin akademik performanslarının artırılması açısından önemli bir potansiyel taşımaktadır.

5.6 Öneriler

Bu çalışma sonucunda elde edilen bulgular doğrultusunda çeşitli öneriler geliştirilebilir.

- Üniversitelerde makine öğrenmesi tabanlı akademik performans tahmin sistemlerinin geliştirilmesi önerilebilir. Bu tür sistemler öğrencilerin akademik gelişimlerinin düzenli olarak izlenmesini sağlayabilir.
- Üniversite bilgi sistemlerinde öğrencilerin öğrenme davranışlarına ilişkin daha fazla veri toplanması önerilebilir. Öğrencilerin dijital öğrenme materyalleri ile etkileşimleri, ders çalışma alışkanlıkları ve çevrimiçi öğrenme platformlarını kullanım düzeyleri gibi veriler akademik performans tahmin modellerinin doğruluğunu artırabilir.
- Eğitim kurumlarında veri temelli karar verme süreçlerinin güçlendirilmesi ve eğitim verilerinin düzenli olarak analiz edilerek öğretmenler ve yöneticiler tarafından kullanılması önerilebilir.
- Gelecekte yapılacak çalışmalarda daha büyük veri setlerinin kullanılması ve farklı makine öğrenmesi algoritmalarının karşılaştırılması önerilebilir.

- Derin öğrenme yöntemleri ve hibrit modeller kullanılarak akademik performans tahmini üzerine daha kapsamlı araştırmalar gerçekleştirilebilir.
- Eğitim kurumlarında, öğrencilerin akademik verilerini düzenli ve bütüncül biçimde depolayabilecek veri yönetim sistemlerinin oluşturulması önerilebilir. Bu sistemler aracılığıyla öğrencilerin akademik performanslarının anlık olarak analiz edilmesi, düşük başarı riski taşıyan öğrencilerin erken dönemde belirlenmesi ve öğrencilerin başarılı/başarısız ya da riskli/risksiz gibi kategorilerde izlenebilmesi mümkün olabilir. Ayrıca bu tür sistemlere entegre edilecek geri bildirim mekanizmaları sayesinde öğretim elemanları ve eğitim yöneticileri, öğrencilerin mevcut durumlarını zamanında değerlendirebilir ve gerekli akademik destek, yönlendirme ve müdahale süreçlerini daha etkili biçimde planlayabilir.

KAYNAKLAR

- Adadi, A., ve Berrada, M. (2018). Peeking inside the black-box: a survey on explainable artificial intelligence (XAI). *IEEE Access*, 6, 52138-52160.
- Aggarwal, C. C. (2018). *Neural Networks and Deep Learning*. Springer International Publishing.
- Aghalarova, S. ve Bozkurt, S. (2021). Öğrencilerin Akademik Performanslarının Tahmin Edilmesi için AutoML Tekniğinin Uygulanması. *El-Cezeri Fen ve Mühendislik Dergisi*, 9(2), 394-412
- Ahern, I., Noack, A., Guzman-Nateras, L., Dou, D., Li, B.A., & Huan, J. (2019). NormLime: A New Feature Importance Metric for Explaining Deep Neural Networks. *ArXiv*.
- Akçapınar, G., Altun, A., ve Aşkar, P. (2019). Using learning analytics to develop early-warning system for at-risk students. *International Journal of Educational Technology in Higher Education*, 16, 40.
- Alalawi, K., Athauda, R., Chiong, R., ve Renner, I. (2024). Evaluating the student performance prediction and action framework through a learning analytics intervention study. *Education and Information Technologies*, 30, 2887-2916.
- Albreiki, B., Zaki, N., ve Alashwal, H. (2021). A systematic literature review of student' performance prediction using machine learning techniques. *Education Sciences*, 11(9), 552.
- Almalawi, A., Soh, B., Li, A., ve Samra, H. (2024). Predictive models for educational purposes: A systematic review. *Big Data and Cognitive Computing*, 8(12), 187.
- Alnasyan, B., Basher, M., & Alassafi, M. (2024). The power of Deep Learning techniques for predicting student performance in virtual learning environments: A systematic literature review. *Computers and Education: Artificial Intelligence*, 6, 100231.
- Alnasyan, E., Obeidat, A., Al-Tamimi, A., ve Hammad, M. (2024). Predicting student academic performance using machine learning techniques. *Education and Information Technologies*, 29, 1–25.
- Als Salman, Y. S., Abu Halemah, N. K., AlNagi, E. S. ve Salameh, W. (2019). Using decision tree and artificial neural network to predict students' academic performance. *2019 10th International Conference on Information and Communication Systems (ICICS)* içinde (104–109). IEEE.
- Altman, N.S. (1992) An Introduction to Kernel and Nearest Neighbor Nonparametric Regression. *The American Statistician*, 46, 175-185.
- Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., vd. (2020). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–

- Bakan, Z. ve Kanbay, F. (2024). Makine öğrenmesi yöntemleri ile eğitim başarısına etki eden faktörlerin modellenmesi. *İstanbul Ticaret Üniversitesi Fen Bilimleri Dergisi*, 23(45), 27–41.
- Baker, R. S. ve Siemens, G. (2014). Educational data mining and learning analytics. R. K. Sawyer (Ed.), *The Cambridge handbook of the learning sciences* (2. Baskı) içinde (253–274). Cambridge University Press.
- Basu, P., ve Rao, K. (2023). The role of explainable AI in educational data mining for predicting student outcomes. *Journal of Artificial Intelligence in Education*, 19(4), 456-478.
- Batista, G. E. A. P. A., Prati, R. C. ve Monard, M. C. (2004). A study of the behavior of several methods for balancing machine learning training data. *ACM SIGKDD Explorations Newsletter*, 6(1), 20–29.
- Bentley, J. L. (1975). Multidimensional binary search trees used for associative searching. *Communications of the ACM*, 18(9), 509–517.
- Bergstra, J., ve Bengio, Y. (2012). *Random search for hyper-parameter optimization*. *Journal of Machine Learning Research*, 13, 281–305.
- Bifarin, O. O. (2023). Interpretable machine learning with tree-based SHAP: Application to metabolomics datasets for binary classification. *PLOS ONE*, 18(5), e0284315.
- Bird, K. A., Castleman, B. L., Mabel, Z. ve Song, Y. (2021). Bringing transparency to predictive analytics: A systematic comparison of predictive modeling methods in higher education. *AERA Open*, 7, 1–19.
- Bishop, C. M. (2006). *Pattern recognition and machine learning*. Springer.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
- Breiman, L., Friedman, J., Olshen, R.A., ve Stone, C.J. (1984). Classification and Regression Trees (2. baskı). *Chapman and Hall/CRC*.
- Bressane, A., Zwirn, D., Essiptchouk, A., Saraiva, A. C. V., Carvalho, F. L. C., Formiga, J. K. S., Medeiros, L. C. C. ve Negri, R. G. (2024). Understanding the role of study strategies and learning disabilities on student academic performance to enhance educational approaches: A proposal using artificial intelligence. *Computers and Education: Artificial Intelligence*, 6, 100196.
- Broadbent, J., ve Poon, W. L. (2015). Self-regulated learning strategies ve academic achievement in online higher education learning environments: A systematic review. *The Internet and Higher Education*, 27, 1–13.
- Canale, L., Cagliero, L., Farinetti, L. ve Torchiano, M. (2024). On predicting exam performance using version control systems' features. *Computers*, 13(6), 150.

- Cassady, J. C. ve Johnson, R. E. (2002). Cognitive test anxiety and academic performance. *Contemporary Educational Psychology*, 27(2), 270–295.
- Chawla, N. V., Bowyer, K. W., Hall, L. O. ve Kegelmeyer, W. P. (2002). SMOTE: Synthetic minority over-sampling technique. *Journal of Artificial Intelligence Research*, 16, 321–357.
- Chen, T. ve Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* içinde (785–794). ACM.
- Cortes, C. ve Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20(3), 273–297.
- Cortez, P. ve Silva, A. M. G. (2008). Using data mining to predict secondary school student performance. A. Brito ve J. Teixeira (Ed.), *Proceedings of the 5th Future Business Technology Conference (FUBUTEC 2008)* içinde (5–12).
- Costa, A., Moreira, D., Casanova, J., Azevedo, Â., Gonçalves, A., Oliveira, Í., Azevedo, R. ve Dias, P. C. (2024). Determinants of academic achievement from the middle to secondary school education: A systematic review. *Social Psychology of Education*, 27(6), 3533–3572.
- Credé, M., Roch, S. G. ve Kieszczyńska, U. M. (2010). Class attendance in college: A meta-analytic review of the relationship of class attendance with grades and student characteristics. *Review of Educational Research*, 80(2), 272–295.
- Cunningham, P. ve Delany, S. J. (2007). *k-nearest neighbour classifiers* (UCD CSI Technical Report No. UCD-CSI-2007-3). University College Dublin, School of Computer Science and Informatics.
- Çoban Kapuoğlu, E. (2024). Eğitim felsefesi perspektifinden makine öğrenmesi. *Felsefe Dünyası*, 79, 245–264.
- Dang, T. K., Tran, T. C., Tuan, L. M. ve Tiep, M. V. (2021). Machine learning based on resampling approaches and deep reinforcement learning for credit card fraud detection systems. *Applied Sciences*, 11(21), 10004.
- Dilbaz Sayın, S. ve Arslan, H. (2017). Öğretmen ve okul yöneticilerinin öğretmen performans değerlendirme sürecindeki çoklu veri kaynakları ile ilgili görüşleri ve öz değerlendirmeleri. *Uluslararası Türkçe Edebiyat Kültür Eğitim (TEKE) Dergisi*, 6(2), 1222–1241.
- Duda, R. O., Hart, P. E. ve Stork, D. G. (2000). *Pattern classification* (2. baskı). Wiley.
- Elkan, C. (2001). The foundations of cost-sensitive learning. *Proceedings of the 17th International Joint Conference on Artificial Intelligence (IJCAI 2001)* içinde (973–978). Morgan Kaufmann.
- Escamilla, J., Hosseini, S. ve González, J. M. (2023). Analyzing and predicting students'

- performance by means of machine learning: A review. *Education and Information Technologies*, 28(2), 1345–1370.
- Fahd, K., Venkatraman, S., Miah, S. J. ve Ahmed, K. (2022). Application of machine learning in higher education to assess student academic performance, at-risk, and attrition: A meta-analysis of literature. *Education and Information Technologies*, 27(3), 3743–3775.
- Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *Annals of Statistics*, 29(5), 1189–1232.
- Gašević, D., Dawson, S. ve Rogers, T. (2016). Learning analytics should not promote one size fits all: The effects of instructional design on the predictive power of demographic variables and learning analytics metrics. *Internet and Higher Education*, 28, 68–84.
- Gašević, D., Dawson, S. ve Siemens, G. (2015). Let's not forget: Learning analytics are about learning. *TechTrends*, 59(1), 64–71.
- Géron, A. (2019). *Hands-on machine learning with Scikit-Learn, Keras, and TensorFlow: Concepts, tools, and techniques to build intelligent systems* (2. baskı). O'Reilly Media.
- Givisis, I., Kalatzis, D., Christakis, C. ve Kiouvrekis, Y. (2025). Comparing explainable AI models: SHAP, LIME, and their role in electric field strength prediction over urban areas. *Electronics*, 14(23), 4766.
- Goodfellow, I., Bengio, Y. ve Courville, A. (2016). *Deep learning*. MIT Press.
- Gottfried, M. A. (2010). Evaluating the relationship between student attendance and achievement in urban elementary and middle schools: An instrumental variables approach. *American Educational Research Journal*, 47(2), 434–465.
- Gunasekara, S. ve Saarela, M. (2024). Explainability in educational data mining and learning analytics: An umbrella review. *Computers and Education: Artificial Intelligence*, 6, 100229.
- Gunasekara, S. ve Saarela, M. (2025). Systematic literature review on explainable learning analytics and educational data mining. *IEEE Access*, 13, 33560–33580.
- Guo, H., Li, Y., Shang, J., Gu, M., Huang, Y. ve Gong, B. (2017). Learning from class-imbalanced data: Review of methods and applications. *Expert Systems with Applications*, 73, 220–239.
- Han, J., Kamber, M. ve Pei, J. (2012). *Data mining: Concepts and techniques* (3. baskı). Morgan Kaufmann.
- Hastie, T., Tibshirani, R. ve Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction* (2. baskı). Springer.

- Haykin, S. (2009). *Neural networks and learning machines* (3. baskı). Pearson.
- He, H., Bai, Y. ve Garcia, E. A. (2008). ADASYN: Adaptive synthetic sampling approach for imbalanced learning. *IEEE International Joint Conference on Neural Networks (IJCNN 2008)* içinde (ss. 1322–1328). IEEE.
- He, H. ve Garcia, E. A. (2009). Learning from imbalanced data. *IEEE Transactions on Knowledge and Data Engineering*, 21(9), 1263–1284.
- He, H. ve Ma, Y. (Ed.). (2013). *Imbalanced learning: Foundations, algorithms, and applications*. Wiley-IEEE Press.
- Holmes, W., Porayska-Pomsta, K., Holstein, K., Sutherland, E., Baker, T., Buckingham Shum, S., Santos, O. C., Rodrigo, M. T., Cukurova, M., Bittencourt, I. I. ve Koedinger, K. R. (2022). *Ethics of AI in education: Towards a community-wide framework*. *International Journal of Artificial Intelligence in Education*, 32(3), 504–526.
- Hsu, C.-W., Chang, C.-C. ve Lin, C.-J. (2003). *A practical guide to support vector classification* [Teknik rapor]. Department of Computer Science, National Taiwan University.
- Huang, J. ve Mok, K. H. (2025). Contributions of tertiary students' background factors, learning strategies, behavioural engagement, and cognitive skills to problem solving in technology-rich environments: Evidence from 24 countries. *Higher Education*, 89(2), 873–894.
- Ilić, M., Keković, G., Mikić, V., Mangaroska, K., Kopanja, L. ve Vesin, B. (2021). Predicting student performance in a programming tutoring system using AI and filtering techniques. *IEEE Transactions on Learning Technologies*, 14(4), 468–479.
- James, G., Witten, D., Hastie, T. ve Tibshirani, R. (2013). *An introduction to statistical learning: With applications in R*. Springer.
- Jang, Y., Choi, S., Jung, H. ve Kim, H. (2022). Practical early prediction of students' performance using machine learning and eXplainable AI. *Education and Information Technologies*, 27(9), 12855–12889.
- Junco, R. (2012a). Too much face and not enough books: The relationship between multiple indices of Facebook use and academic performance. *Computers in Human Behavior*, 28(1), 187–198.
- Junco, R. (2012b). The relationship between frequency of Facebook use, participation in Facebook activities, and student engagement. *Computers & Education*, 58(1), 162–171.
- Kapoor, S. ve Narayanan, A. (2023). Leakage and the reproducibility crisis in machine-learning-based science. *Patterns*, 4(9), 100804.
- Kaufman, S., Rosset, S., Perlich, C. ve Stitelman, O. (2012). Leakage in data mining: Formulation, detection, and avoidance. *ACM Transactions on Knowledge*

Discovery from Data, 6(4), 1–21.

- Kiewra, K. A. (1985). Investigating note-taking and review: A depth of processing alternative. *Educational Psychologist*, 20(1), 23–32.
- Kim, S., Cho, H., & Kim, L. Y. (2019). Socioeconomic Status and Academic Outcomes in Developing Countries: A Meta-Analysis. *Review of Educational Research*, 89, 875–916.
- Kobayashi, K. (2005). What limits the encoding effect of note-taking? A meta-analytic examination. *Contemporary Educational Psychology*, 30(2), 242–262.
- Kobayashi, K. (2006). Combined effects of note-taking/reviewing on learning and the enhancement through interventions: A meta-analytic review. *Educational Psychology*, 26(3), 459–477.
- Kotsiantis, S., Kanellopoulos, D. ve Pintelas, P. (2006). Handling imbalanced datasets: A review. *GESTS International Transactions on Computer Science and Engineering*, 30(1), 25–36.
- Kovanović, V., Gašević, D., Dawson, S., Joksimović, S. ve Siemens, G. (2015). Does time-on-task estimation matter? Implications on validity of learning analytics findings. *Journal of Learning Analytics*, 2(3), 81–110.
- Kuhn, M. ve Johnson, K. (2013). *Applied predictive modeling*. Springer.
- Kulakowski, Mikolaj. (2019). *Estimating Football Match Outcomes with Machine Learning*.
- Kumar, D., ve Sharma, S. (2024). Machine Learning in Education: A Comparative Study of Techniques for Academic Performance Prediction. *International Journal of Educational Technology and Innovation*, 9(1), 56-72.
- Kumar, J. ve Rani, A. ve Kumar, N. ve Sinha, N. (2022). Machine learning for soil moisture assessment. *Deep Learning for Sustainable Agriculture*.
- Lexcellent, C. (2019). Artificial intelligence. *Artificial intelligence versus human intelligence* içinde. Springer.
- Liaw, A. ve Wiener, M. (2002). Classification and regression by randomForest. *R News*, 2(3), 18–22.
- Liu, Q. ve Khalil, M. (2024). Explainable AI in learning analytics: Improving predictive models and advancing transparency trust. *2024 IEEE Global Engineering Education Conference (EDUCON)* içinde (1–7). IEEE.
- Lou, Y. ve Colvin, K. F. (2025). Performance prediction using educational data mining techniques: A comparative study. *Discover Education*, 4, 112.
- Lundberg, S. M., ve Lee, S.-I. (2017). *A unified approach to interpreting model predictions*. *Advances in Neural Information Processing Systems*, 30.

- Lundberg, S. M., Erion, G., Chen, H., DeGrave, A. J., Prutkin, J. M., Nair, B., Katz, N. H., Himmelfarb, J., Bansal, N. ve Lee, S.-I. (2020). From local explanations to global understanding with explainable AI for trees. *Nature Machine Intelligence*, 2(1), 56–67.
- Macfadyen, L. P. ve Dawson, S. (2010). Mining LMS data to develop an “early warning system” for educators: A proof of concept. *Computers & Education*, 54(2), 588–599.
- Manning, C. D., Raghavan, P. ve Schütze, H. (2008). *Introduction to information retrieval*. Cambridge University Press.
- McKinney, W. (2010). Data structures for statistical computing in Python. *Proceedings of the 9th Python in Science Conference*, 51–56.
- Means, B., Bakia, M. ve Murphy, R. (2014). *Learning online: What research tells us about whether, when and how*. Routledge.
- Means, B., Toyama, Y., Murphy, R., Bakia, M. ve Jones, K. (2013). The effectiveness of online and blended learning: A meta-analysis of the empirical literature. *Teachers College Record*, 115(3), 1–47.
- Molnar, C. (2022). *Interpretable Machine Learning (2. baskı)*.
- Morandín-Ahuerma, F. (2022). What is artificial intelligence? *International Journal of Research Publication and Reviews*, 3(12), 1947–1951.
- Murphy, K. P. (2012). *Machine learning: A probabilistic perspective*. MIT Press.
- Nagendraswamy, C. ve Salis, A. (2021). A review article on artificial intelligence. *Annals of Biomedical Science and Engineering*, 5(1), 013–014.
- Nair, V. ve Hinton, G. E. (2010). Rectified linear units improve restricted Boltzmann machines. *Proceedings of the 27th International Conference on Machine Learning (ICML-10)* içinde (ss. 807–814).
- OECD. (2019). *PISA 2018 results (Cilt 2): Where all students can succeed*. PISA. OECD Publishing.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., VanderPlas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M. ve Duchesnay, É. (2011). Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 12, 2825–2830.
- Ponce-Bobadilla, A. V., Schmitt, V., Maier, C. S., Mensing, S. ve Stodtmann, S. (2024). Practical guide to SHAP analysis: Explaining supervised machine learning model predictions in drug development. *Clinical and Translational Science*, 17, e70056.

- Prechelt, L. (2000). Early Stopping- But When? *Lecture Notes in Computer Science*.
- Quinlan, J. R. (1986). Induction of decision trees. *Machine Learning*, 1(1), 81–106.
- Rahman, N. F. A., Wang, S. L., Ng, T. F. ve Ghoneim, A. S. (2024). Artificial intelligence in education: A systematic review of machine learning for predicting student performance. *Journal of Advanced Research in Applied Sciences and Engineering Technology*, 54(1), 198–221.
- Reshi, A. A., Ashraf, I., Rustam, F., Shahzad, H. F., Mehmood, A. ve Choi, G. S. (2021). Diagnosis of vertebral column pathologies using concatenated resampling with machine learning algorithms. *PeerJ Computer Science*, 7, e547.
- Ribeiro, M. T., Singh, S. ve Guestrin, C. (2016). “Why should I trust you?”: Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (1135–1144). ACM.
- Richardson, M., Abraham, C. ve Bond, R. (2012). Psychological correlates of university students’ academic performance: A systematic review and meta-analysis. *Psychological Bulletin*, 138(2), 353–387.
- Romero, C. ve Ventura, S. (2010). Educational data mining: A review of the state of the art. *IEEE Transactions on Systems, Man, and Cybernetics, Part C (Applications and Reviews)*, 40(6), 601–618.
- Romero, C. ve Ventura, S. (2013). Data mining in education. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 3(1), 12–27.
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1(5), 206–215.
- Saeys, Y., Inza, I. ve Larrañaga, P. (2007). A review of feature selection techniques in bioinformatics. *Bioinformatics*, 23(19), 2507–2517.
- Saito, T. ve Rehmsmeier, M. (2015). The precision–recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets. *PLOS ONE*, 10(3), e0118432.
- Salih, A. M., Raisi-Estabragh, Z., Boscolo Galazzo, I., Radeva, P., Petersen, S. E., Menegaz, G. ve Lekadir, K. (2024). A perspective on explainable artificial intelligence methods: SHAP and LIME. *Advanced Intelligent Systems*, 6(3), 2300252.
- Santos, R. M. ve Henriques, R. (2023). Accurate, timely, and portable: Course-agnostic early prediction of student performance from LMS logs. *Computers and Education: Artificial Intelligence*, 5, 100175.
- Sayed, A. F. A., Arafa, M. A., El-Nimr, N. A., Banawan, K. A. S. ve Abdou, M. S. (2025). Comparison of machine learning classification and regression models for prediction of academic performance among postgraduate public health students. *Scientific*

Reports, 15, 44056.

Shmueli, G. (2010). To explain or to predict? *Statistical Science, 25*(3), 289–310.

Siemens, G., ve Long, P. (2011). Penetrating the fog: Analytics in learning and education. *EDUCAUSE Review, 46*(5), 31–40.

Şirin, S. R. (2005). Socioeconomic status and academic achievement: A meta-analytic review of research. *Review of Educational Research, 75*(3), 417–453.

Smola, A. J. ve Schölkopf, B. (2004). A tutorial on support vector regression. *Statistics and Computing, 14*(3), 199–222.

Strecht, P., Cruz, L., Soares, C., Mendes-Moreira, J. ve Abreu, R. (2015). A comparative study of classification and regression algorithms for modelling students' academic performance. *Proceedings of the 8th International Conference on Educational Data Mining* içinde (392–395).

Sun, Y., Wong, A. K. C. ve Kamel, M. S. (2009). Classification of imbalanced data: A review. *International Journal of Pattern Recognition and Artificial Intelligence, 23*(4), 687–719.

Suresh, S., Newton, D. T., Everett, T. H., IV, Lin, G. ve Duerstock, B. S. (2022). Feature selection techniques for a machine learning model to detect autonomic dysreflexia. *Frontiers in Neuroinformatics, 16*, 851148.

Tadese, M., Yeshaneh, A. ve Baye, G. B. (2022). Determinants of good academic performance among university students in Ethiopia: A cross-sectional study. *BMC Medical Education, 22*, 395.

Türkmen, G. (2025). The review of studies on explainable artificial intelligence in educational research. *Journal of Educational Computing Research, 63*(2), 277–310.

Van der Linden, I., Haned, H. ve Kanoulas, E. (2019). Global aggregations of local explanations for black box models. *Proceedings of the 28th ACM International Conference on Information and Knowledge Management (CIKM)* içinde (1783–1792). ACM.

Van Rossum, G. ve Drake, F. L. (2009). Python 3 reference manual. CreateSpace.

Vapnik, V. N. (1995). *The nature of statistical learning theory*. Springer.

Weidlich, J., Gašević, D. ve Drachsler, H. (2022). Causal inference and bias in learning analytics: A primer on pitfalls using directed acyclic graphs. *Journal of Learning Analytics, 9*(3), 183–199.

EKLER

Ek 1: Yapay zekâ kullanım tablosu

No	Yapay Zekâ Aracı	Kullanım Amacı	Kullanıldığı Bölüm	Kullanım Şekli
1	ChatGpt	Cümle düzeltme, makale özetleme	Materyal, yöntem	Yazılan cümleleri ve eklenen dosyaları düzeltme
2	Google Gemini	Metin düzeltme	Tezin geneli	Cümle düzenleme
3	Elicit	Makale araştırması	Özet okuma	Konu ile ilgili detaylı araştırma yapma



Ek 2: Sosyal ve Beşerî Bilimler Etik Kurulu Onay Belgesi



T.C.
BARTIN ÜNİVERSİTESİ REKTÖRLÜĞÜ
Sosyal ve Beşerî Bilimler Etik Kurulu
Onay Belgesi

TOPLANTI SAYISI
11

KARAR SAYISI
1

TOPLANTI TARİHİ
18.06.2025

Protokol No:	2025-SBB-0530
Araştırmanın Başlığı:	Öğrencinin Akademik Performansını Tahmin Etmede Açıklanabilir Yapay Zeka ve Makine Öğrenimi Kullanımı
Proje Yürütücüsü:	Müjde Dilek SOBUTAY
Başvuru Formunun Geliş Tarihi:	11.05.2025

Başvuru dosyasında etik sorun oluşturabilecek sorular/maddeler, süreçler ya da unsurlar bulunmadığından 18.06.2025 tarihli ve 11 numaralı toplantıda 2025-SBB-0530 numaralı başvuruya araştırma için ETİK KURUL ONAY belgesinin verilmesine karar verilmiştir.

Prof. Dr. Ayla ÇETİN DİNDAR
Başkan

Doç. Dr. Hilal UYSAL
Başkan yardımcısı

Doç. Dr. Mehmet
ALTUNMERAL
Üye

Doç. Dr. Sevdâ KAMAN
Üye

Dr. Öğr. Üyesi Nergiz TEKE
Üye

Dr. Öğr. Üyesi Sevim
Handan YILMAZ
Üye

Dr. Öğr. Üyesi Sevinç
AKKAYA TAHTA
Üye

Belge Doğrulama Kodu: FHDA7DH

Bu belge, güvenli elektronik imza ile imzalanmıştır.

Belge Takip Adresi: <https://www.turkiye.gov.tr/bartın-universitesi-ebys>

Adres: Ağdacı Mahallesi Fakülte Caddesi No:54 Bartın

Telefon No: (0 378) 2235500

e-Posta:

Kep Adresi: bartinuniversitesi@hs01.kep.tr

Faks No: (0 378) 2235042

İnternet Adresi: <http://www.bartın.edu.tr/>

Bilgi için :

Hasan Hüseyn Gürcan
Sekreter

Telefon No:

Direkt Hat:



Ek 3: Sınıflandırma Sonuçları Tabloları

Tablo Ek3.1: Kara Ağaçları En İyi Parametreler (Deneme 1)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Criterion	gini	log_loss
Max Depth	3	8
Min Samples Leaf	29	29
Min Samples Split	20	16

Tablo Ek3.2: Kara Ağaçları 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7669 ± 0.0473	0.7563 ± 0.0377
F1 Score	0.8544 ± 0.0283	0.8463 ± 0.0269
Balanced Acc	0.6234 ± 0.0768	0.6142 ± 0.0422
ROC-AUC	0.6477 ± 0.0793	0.6660 ± 0.0615
PR-AUC	0.8174 ± 0.0381	0.8273 ± 0.0417

Tablo Ek3.3: Kara Ağaçları Confusion Matrix (Deneme 1)

	%70 / %30		%60 / %30	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	9	35	13	46
1 (Başarılı)	8	107	10	143

Tablo Ek3.4: Kara Ağaçları %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.5294	0.2045	0.2951	44
1	0.7535	0.9304	0.8327	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7296	159
Macro Avg	0.6415	0.5675	0.5639	159

Weighted Avg	0.6915	0.7296	0.6839	159
---------------------	--------	--------	--------	-----

Tablo Ek3.5: Kara Ağaçları %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.5652	0.2203	0.3171	59
1	0.7566	0.9346	0.8363	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7358	212
Macro Avg	0.6609	0.5775	0.5767	212
Weighted Avg	0.7033	0.7358	0.6918	212

Tablo Ek3.6: Kara Ağaçları Genel Test Sınıflandırma Sonuçları (Deneme 1)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7296	0.8327	0.5675	0.5684	0.7636
%60 / %40	0.7358	0.8363	0.5775	0.6407	0.7983

Tablo Ek3.7: Destek Vektör Makinesi En İyi Parametreler (Deneme 1)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
C	0.7459343	503.8403
Gamma	0.6351221	0.6542057
Kernel	RBF	RBF

Tablo Ek3.8: Destek Vektör Makinesi 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7535 ± 0.0244	0.7630 ± 0.0328
F1 Score	0.8531 ± 0.0136	0.8575 ± 0.0181
Balanced Accuracy	0.5631 ± 0.0408	0.5844 ± 0.0554
ROC-AUC	0.6820 ± 0.0657	0.7160 ± 0.0608
PR-AUC	0.8533 ± 0.0411	0.8675 ± 0.0380

Tablo Ek3.9: Destek Vektör Makinesi Confusion Matrix (Deneme 1)

	%70 / %30		%60 / %30	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	4	40	8	51
1 (Başarılı)	3	112	5	148

Tablo Ek3.10: Destek Vektör Makinesi %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.5714	0.0909	0.1569	44
1	0.7368	0.9739	0.8390	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7296	159
Macro Avg	0.6541	0.5324	0.4979	159
Weighted Avg	0.6911	0.7296	0.6502	159

Tablo Ek3.11: Destek Vektör Makinesi %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.6154	0.1356	0.2222	59
1	0.7437	0.9673	0.8409	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7358	212
Macro Avg	0.6796	0.5515	0.5316	212
Weighted Avg	0.7080	0.7358	0.6687	212

Tablo Ek3.12: Destek Vektör Makinesi Genel Test Sınıflandırma Sonuçları (Deneme 1)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7296	0.8390	0.5324	0.6694	0.8266
%60 / %40	0.7358	0.8409	0.5515	0.6158	0.7942

Tablo Ek3.13: Rastgele Orman En İyi Parametreler (Deneme 1)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Bootstrap	False	False
Max Depth	10	None
Max Features	log2	log2
Min Samples Leaf	1	1
Min Samples Split	15	19
n_estimators	570	423

Tablo Ek3.14: Rastgele Orman 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7670 ± 0.0322	0.7659 ± 0.0281
F1 Score	0.8565 ± 0.0208	0.8576 ± 0.0179
Balanced Accuracy	0.6088 ± 0.0496	0.5965 ± 0.0395
ROC-AUC	0.7124 ± 0.0697	0.7319 ± 0.0495
PR-AUC	0.8668 ± 0.0395	0.8742 ± 0.0345

Tablo Ek3.15: Rastgele Orman Confusion Matrix (Deneme 1)

	%70 / %30		%60 / %30	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	11	33	11	48
1 (Başarılı)	9	106	6	147

Tablo Ek3.16: Rastgele Orman %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.5500	0.2500	0.3438	44
1	0.7626	0.9217	0.8346	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7358	159
Macro Avg	0.6563	0.5859	0.5892	159
Weighted Avg	0.7038	0.7358	0.6988	159

Tablo Ek3.17: Rastgele Orman %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.6471	0.1864	0.2895	59
1	0.7538	0.9608	0.8448	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7453	212
Macro Avg	0.7005	0.5736	0.5672	212
Weighted Avg	0.7241	0.7453	0.6903	212

Tablo Ek3.18: Rastgele Orman Genel Test Sınıflandırma Sonuçları (Deneme 1)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7358	0.8346	0.5859	0.7273	0.8578
%60 / %40	0.7453	0.8448	0.5736	0.6821	0.8200

Tablo Ek3.19: K-En Yakın Komşu En İyi Parametreler (Deneme 1)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
n_neighbors	32	35
leaf_size	30	46
p	1 (Manhattan Distance)	1 (Manhattan Distance)
weights	distance	distance

Tablo Ek3.20: K-En Yakın Komşu 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7643 ± 0.0263	0.7725 ± 0.0345
F1 Score	0.8578 ± 0.0144	0.8636 ± 0.0184
Balanced Accuracy	0.5885 ± 0.0460	0.5943 ± 0.0609
ROC-AUC	0.6814 ± 0.0593	0.7415 ± 0.0456
PR-AUC	0.8584 ± 0.0319	0.8855 ± 0.0220

Tablo Ek3.21: K-En Yakın Komşu Confusion Matrix (Deneme 1)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	9	35	10	49
1 (Başarılı)	5	110	5	148

Tablo Ek3.22: K-En Yakın Komşu %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.6429	0.2045	0.3103	44
1	0.7586	0.9565	0.8462	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7484	159
Macro Avg	0.7007	0.5805	0.5782	159
Weighted Avg	0.7266	0.7484	0.6979	159

Tablo Ek3.23: K-En Yakın Komşu %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.6667	0.1695	0.2703	59
1	0.7513	0.9673	0.8457	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7453	212
Macro Avg	0.7090	0.5684	0.5580	212
Weighted Avg	0.7277	0.7453	0.6856	212

Tablo Ek3.24: K-En Yakın Komşu Genel Test Sınıflandırma Sonuçları (Deneme 1)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7484	0.8462	0.5805	0.6959	0.8324
%60 / %40	0.7453	0.8457	0.5684	0.6357	0.8040

Tablo Ek3.25: Çok Katmanlı Perceptron En İyi Parametreler (Deneme 1)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Activation	relu	relu
Alpha	0.0036485	0.0005876
Batch Size	256	64
Hidden Layer Sizes	(64,)	(64,)
Learning Rate Init	0.0068058	0.0274253
Solver	adam	adam

Tablo Ek3.26: Çok Katmanlı Perceptron 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7344 ± 0.0205	0.7440 ± 0.0337
F1 Score	0.8446 ± 0.0098	0.8473 ± 0.0171
Balanced Accuracy	0.5218 ± 0.0453	0.5515 ± 0.0720
ROC-AUC	0.5103 ± 0.1500	0.6135 ± 0.1099
PR-AUC	0.7531 ± 0.0811	0.8216 ± 0.0638

Tablo Ek3.27: Çok Katmanlı Perceptron Confusion Matrix (Deneme 1)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	0	44	23	36
1 (Başarılı)	0	115	30	123

Tablo Ek3.28: Çok Katmanlı Perceptron %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.0000	0.0000	0.0000	44
1	0.7233	1.0000	0.8394	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7233	159
Macro Avg	0.3616	0.5000	0.4197	159
Weighted Avg	0.5231	0.7233	0.6071	159

Tablo Ek3.29: Çok Katmanlı Perceptron %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.4340	0.3898	0.4107	59
1	0.7736	0.8039	0.7885	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.6887	212
Macro Avg	0.6038	0.5969	0.5996	212
Weighted Avg	0.6791	0.6887	0.6833	212

Tablo Ek3. 30: Çok Katmanlı Perceptron Genel Test Sınıflandırma Sonuçları (Deneme 1)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7233	0.8394	0.5000	0.5241	0.7590
%60 / %40	0.6887	0.7885	0.5969	0.6430	0.8282

Tablo Ek3.31: XGBoost En İyi Parametreler (Deneme 1)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
n_estimators	1398	583
max_depth	4	5
learning_rate	0.01693	0.25685
Gamma	1.30347	3.53844
min_child_weight	5	10
Subsample	0.92139	0.89669
colsample_bytree	0.73625	0.65098
reg_alpha	2.33768	2.90706

Tablo Ek3.32: XGBoost 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7560 ± 0.0457	0.7723 ± 0.0223
F1 Score	0.8511 ± 0.0271	0.8598 ± 0.0152
Balanced Accuracy	0.5882 ± 0.0707	0.6144 ± 0.0391
ROC-AUC	0.6603 ± 0.0759	0.6804 ± 0.0545

PR-AUC	0.8433 ± 0.0370	0.8520 ± 0.0390
---------------	---------------------	---------------------

Tablo Ek3.33: XGBoost Confusion Matrix (Deneme 1)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	8	36	10	49
1 (Başarılı)	6	109	3	150

Tablo Ek3.34: XGBoost %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.5714	0.1818	0.2759	44
1	0.7517	0.9478	0.8385	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7358	159
Macro Avg	0.6616	0.5648	0.5572	159
Weighted Avg	0.7018	0.7358	0.6828	159

Tablo Ek3.35: XGBoost %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 1)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7692	0.1695	0.2778	59
1	0.7538	0.9804	0.8523	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7547	212
Macro Avg	0.7615	0.5749	0.5650	212
Weighted Avg	0.7581	0.7547	0.6924	212

Tablo Ek3.36: XGBoost Genel Test Sınıflandırma Sonuçları (Deneme 1)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7358	0.8385	0.5648	0.6907	0.8539
%60 / %40	0.7547	0.8523	0.5749	0.6831	0.8448

Tablo Ek3.37 Tüm Modeller %70 Eğitim %30 Test Verileri ile Tüm Modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 1)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
KNN	0.764339	0.857805	0.588468	0.681430	0.858396
MLP	0.734384	0.844573	0.521785	0.510254	0.753146
SVM (RBF)	0.753529	0.853141	0.563081	0.681961	0.853319
XGBoost	0.756006	0.851107	0.588157	0.660256	0.843349
Random Forest	0.766967	0.856530	0.608754	0.712444	0.866781
Decision Tree	0.766892	0.854433	0.623372	0.647747	0.817423

Tablo Ek3.38 Tüm Modeller %70 Eğitim %30 Test Verileri ile Tüm Modellerin Test Performans Sonuçları (Deneme 1)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC- AUC	Test PR- AUC
KNN	0.748428	0.846154	0.580534	0.695949	0.832425
MLP	0.723270	0.839416	0.500000	0.524111	0.759003
SVM (RBF)	0.729560	0.838951	0.532411	0.669368	0.826589
XGBoost	0.735849	0.838462	0.564822	0.690711	0.853915
Random Forest	0.735849	0.834646	0.585870	0.727273	0.857822
Decision Tree	0.729560	0.832685	0.567490	0.568379	0.763582

Tablo Ek3.39 Tüm Modeller %60 Eğitim %40 Test Verileri ile Tüm Modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 1)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
XGBoost	0.772278	0.859785	0.614356	0.680399	0.851962
KNN	0.772480	0.863581	0.594263	0.741502	0.885509
Random Forest	0.765927	0.857607	0.596481	0.731873	0.874207
SVM (RBF)	0.763004	0.857502	0.584360	0.715991	0.867485
Decision Tree	0.756250	0.846263	0.614227	0.665958	0.827304
MLP	0.743952	0.847328	0.551532	0.613516	0.821592

Tablo Ek3.40 Tüm Modeller %60 Eğitim %40 Test Verileri ile Tüm Modellerin Test Performans Sonuçları (Deneme 1)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
XGBoost	0.754717	0.852273	0.574942	0.683062	0.844766
KNN	0.745283	0.845714	0.568406	0.635704	0.804012
Random Forest	0.745283	0.844828	0.573612	0.682065	0.819956
SVM (RBF)	0.735849	0.840909	0.551457	0.615819	0.794162
Decision Tree	0.735849	0.836257	0.577490	0.640744	0.798321
MLP	0.688679	0.788462	0.596876	0.642960	0.828203

Tablo Ek3.41: Kara Ağaçları En İyi Parametreler (Deneme 2)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Criterion	log_loss	log_loss
Max Depth	24	8
Min Samples Leaf	1	9
Min Samples Split	4	8

Tablo Ek3.42: Kara Ağaçları 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7241 ± 0.0589	0.7024 ± 0.0814
F1 Score	0.7214 ± 0.0641	0.7043 ± 0.0785
Balanced Accuracy	0.7247 ± 0.0584	0.7036 ± 0.0818
ROC-AUC	0.7292 ± 0.0533	0.7309 ± 0.0906
PR-AUC	0.6951 ± 0.0510	0.7355 ± 0.0968

Tablo Ek3.43: Kara Ağaçları Confusion Matrix (Deneme 2)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	80	27	108	35
1 (Başarılı)	29	86	52	101

Tablo Ek3.44: Kara Ağaçları %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7339	0.7477	0.7407	107
1	0.7611	0.7478	0.7544	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7477	222
Macro Avg	0.7475	0.7477	0.7476	222
Weighted Avg	0.7480	0.7477	0.7478	222

Tablo Ek3.45: Kara Ağaçları %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.6750	0.7552	0.7129	143
1	0.7426	0.6601	0.6990	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7061	296
Macro Avg	0.7088	0.7077	0.7059	296
Weighted Avg	0.7100	0.7061	0.7057	296

Tablo Ek3.46: Kara Ağaçları Genel Test Sınıflandırma Sonuçları (Deneme 2)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7477	0.7544	0.7477	0.7442	0.6970
%60 / %40	0.7061	0.6990	0.7077	0.8014	0.7703

Tablo Ek3.47: Destek Vektör Makinesi En İyi Parametreler (Deneme 2)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
C	17.7078	47.3538
Gamma	0.249398	0.163653
Kernel	RBF	RBF

Tablo Ek3.48: Destek Vektör Makinesi 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std	Ortalama $\pm$ Std
Accuracy	0.8186 $\pm$ 0.0509	0.8199 $\pm$ 0.0627
F1 Score	0.8399 $\pm$ 0.0392	0.8325 $\pm$ 0.0596
Balanced Accuracy	0.8154 $\pm$ 0.0531	0.8188 $\pm$ 0.0624
ROC-AUC	0.8713 $\pm$ 0.0492	0.8635 $\pm$ 0.0762
PR-AUC	0.8585 $\pm$ 0.0616	0.8547 $\pm$ 0.0841

Tablo Ek3.49: Destek Vektör Makinesi Confusion Matrix (Deneme 2)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	79	28	111	32
1 (Başarılı)	7	108	16	137

Tablo Ek3.50: Destek Vektör Makinesi %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9186	0.7383	0.8187	107
1	0.7941	0.9391	0.8606	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8423	222
Macro Avg	0.8564	0.8387	0.8396	222
Weighted Avg	0.8541	0.8423	0.8404	222

Tablo Ek3.51: Destek Vektör Makinesi %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8740	0.7762	0.8222	143
1	0.8107	0.8954	0.8509	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8378	296
Macro Avg	0.8423	0.8358	0.8366	296
Weighted Avg	0.8413	0.8378	0.8371	296

Tablo Ek3.52: Destek Vektör Makinesi Genel Test Sınıflandırma Sonuçları (Deneme 2)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8423	0.8606	0.8387	0.9139	0.9074
%60 / %40	0.8378	0.8509	0.8358	0.8951	0.8845

Tablo Ek3.53: Rastgele Orman En İyi Parametreler (Deneme 2)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Bootstrap	False	False
Max Depth	None	10
Max Features	sqrt	log2
Min Samples Leaf	2	1
Min Samples Split	2	15
n_estimators	503	570

Tablo Ek3.54: Rastgele Orman 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8264 ± 0.0370	0.8066 ± 0.0692
F1 Score	0.8289 ± 0.0392	0.8147 ± 0.0626
Balanced Accuracy	0.8266 ± 0.0367	0.8063 ± 0.0698
ROC-AUC	0.8889 ± 0.0262	0.8538 ± 0.0598
PR-AUC	0.8999 ± 0.0342	0.8733 ± 0.0569

Tablo Ek3.55: Rastgele Orman Confusion Matrix (Deneme 2)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	89	18	118	25
1 (Başarılı)	14	101	22	131

Tablo Ek3.56: Rastgele Orman %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8641	0.8318	0.8476	107
1	0.8487	0.8783	0.8632	115

Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.8559	222
Macro Avg	0.8564	0.8550	0.8554	222
Weighted Avg	0.8561	0.8559	0.8557	222

Tablo Ek3.57: Rastgele Orman %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8429	0.8252	0.8339	143
1	0.8397	0.8562	0.8479	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.8412	296
Macro Avg	0.8413	0.8407	0.8409	296
Weighted Avg	0.8412	0.8412	0.8411	296

Tablo Ek3.58: Rastgele Orman Genel Test Sınıflandırma Sonuçları (Deneme 2)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8559	0.8632	0.8550	0.9106	0.9097
%60 / %40	0.8412	0.8479	0.8407	0.8994	0.8989

Tablo Ek3.59: K-En Yakın Komşu0:1 En İyi Parametreler (Deneme 2)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
n_neighbors	3	3
leaf_size	13	13
p	1 (Manhattan Distance)	1 (Manhattan Distance)
weights	distance	distance

Tablo Ek3.60: K-En Yakın Komşu 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7472 ± 0.0437	0.7388 ± 0.0517
F1 Score	0.7084 ± 0.0562	0.6949 ± 0.0698
Balanced Accuracy	0.7519 ± 0.0421	0.7439 ± 0.0517

ROC-AUC	0.8278 ± 0.0475	0.8003 ± 0.0419
PR-AUC	0.8426 ± 0.0451	0.8231 ± 0.0357

Tablo Ek3.61: K-En Yakın Komşu Confusion Matrix (Deneme 2)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	95	12	128	15
1 (Başarılı)	31	84	52	101

Tablo Ek3.62: K-En Yakın Komşu %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7540	0.8879	0.8155	107
1	0.8750	0.7304	0.7962	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.8063	222
Macro Avg	0.8145	0.8091	0.8058	222
Weighted Avg	0.8167	0.8063	0.8055	222

Tablo Ek3.63: K-En Yakın Komşu %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7111	0.8951	0.7926	143
1	0.8707	0.6601	0.7509	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.7736	296
Macro Avg	0.7909	0.7776	0.7717	296
Weighted Avg	0.7936	0.7736	0.7710	296

Tablo Ek3.64: K-En Yakın Komşu Genel Test Sınıflandırma Sonuçları (Deneme 2)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8063	0.7962	0.8091	0.8841	0.8926
%60 / %40	0.7736	0.7509	0.7776	0.8503	0.8637

Tablo Ek3.65: Çok Katmanlı Perceptron En İyi Parametreler (Deneme 2)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Activation	relu	relu
Alpha	0.0003573	0.00014797
Batch Size	32	128
Hidden Layer Sizes	(256, 128, 64)	(128, 64)
Learning Rate Init	0.0044773	0.0052173
Solver	adam	adam

Tablo Ek3.66: Çok Katmanlı Perceptron 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7241 ± 0.0662	0.7120 ± 0.0701
F1 Score	0.7336 ± 0.0477	0.7336 ± 0.0488
Balanced Accuracy	0.7234 ± 0.0682	0.7106 ± 0.0720
ROC-AUC	0.7833 ± 0.0547	0.7256 ± 0.0770
PR-AUC	0.7982 ± 0.0570	0.7513 ± 0.0816

Tablo Ek3.67: Çok Katmanlı Perceptron Confusion Matrix (Deneme 2)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	81	26	100	43
1 (Başarılı)	18	97	31	122

Tablo Ek3.68: Çok Katmanlı Perceptron %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8182	0.7570	0.7864	107
1	0.7886	0.8435	0.8151	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8018	222
Macro Avg	0.8034	0.8002	0.8008	222
Weighted Avg	0.8029	0.8018	0.8013	222

Tablo Ek3.69: Çok Katmanlı Perceptron %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7634	0.6993	0.7299	143
1	0.7394	0.7974	0.7673	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7500	296
Macro Avg	0.7514	0.7483	0.7486	296
Weighted Avg	0.7510	0.7500	0.7492	296

Tablo Ek3.70: Çok Katmanlı Perceptron Genel Test Sınıflandırma Sonuçları (Deneme 2)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8018	0.8151	0.8002	0.8579	0.8552
%60 / %40	0.7500	0.7673	0.7483	0.7989	0.8109

Tablo Ek3.71: XGBoost En İyi Parametreler (Deneme 2)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
n_estimators	585	713
max_depth	6	8
learning_rate	0.0670015	0.1066518
Gamma	0.1001318	0.0453620
min_child_weight	1	3
Subsample	0.6356498	0.9360486
colsample_bytree	0.9896658	0.7506781
reg_alpha	0.5209578	0.0184277
reg_lambda	2.0589642	1.5806020

Tablo Ek3.72: XGBoost 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8226 ± 0.0425	0.8087 ± 0.0557
F1 Score	0.8242 ± 0.0413	0.8098 ± 0.0542
Balanced Accuracy	0.8231 ± 0.0430	0.8095 ± 0.0564

ROC-AUC	0.8783 $\pm$ 0.0362	0.8642 $\pm$ 0.0684
PR-AUC	0.8917 $\pm$ 0.0397	0.8747 $\pm$ 0.0602

Tablo Ek3.73: XGBoost Confusion Matrix (Deneme 2)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	84	23	116	27
1 (Başarılı)	15	100	35	118

Tablo Ek3.74: XGBoost %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8485	0.7850	0.8155	107
1	0.8130	0.8696	0.8403	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8288	222
Macro Avg	0.8307	0.8273	0.8279	222
Weighted Avg	0.8301	0.8288	0.8284	222

Tablo Ek3.75: XGBoost %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7682	0.8112	0.7891	143
1	0.8138	0.7712	0.7919	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7905	296
Macro Avg	0.7910	0.7912	0.7905	296
Weighted Avg	0.7918	0.7905	0.7906	296

Tablo Ek3.76: XGBoost Genel Test Sınıflandırma Sonuçları (Deneme 2)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8288	0.8403	0.8273	0.9108	0.9198
%60 / %40	0.7905	0.7919	0.7912	0.8943	0.8897

Tablo Ek3.77 Tüm Modeller %70 Eğitim %30 Test Verileri ile Tüm Modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 2)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
Random Forest	0.826433	0.828916	0.826630	0.888882	0.899907
SVM (RBF)	0.818590	0.839857	0.815425	0.871336	0.858482
XGBoost	0.822587	0.824210	0.823068	0.878274	0.891688
MLP	0.724057	0.733575	0.723402	0.783312	0.798196
KNN	0.747210	0.708429	0.751858	0.827756	0.842596
Decision Tree	0.724133	0.721431	0.724724	0.729226	0.695120

Tablo Ek3.78 Tüm Modeller %70 Eğitim %30 Test Verileri ile Tüm Modellerin Test Performans Sonuçları (Deneme 2)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC- AUC	Test PR- AUC
Random Forest	0.855856	0.863248	0.855018	0.910646	0.909670
SVM (RBF)	0.842342	0.860558	0.838724	0.913856	0.907397
XGBoost	0.828829	0.840336	0.827306	0.910849	0.919826
MLP	0.801802	0.815126	0.800244	0.857944	0.855212
KNN	0.806306	0.796209	0.809143	0.884112	0.892632
Decision Tree	0.747748	0.754386	0.747745	0.744169	0.697031

Tablo Ek3.79 Tüm Modeller %60 Eğitim %40 Test Verileri ile Tüm Modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 2)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.819899	0.832525	0.818784	0.863522	0.854741
Random Forest	0.806616	0.814693	0.806343	0.853808	0.873288
XGBoost	0.808687	0.809838	0.809543	0.864186	0.874667
MLP	0.711970	0.733629	0.710578	0.725608	0.751255
KNN	0.738838	0.694887	0.743888	0.800294	0.823146
Decision Tree	0.702374	0.704307	0.703590	0.730930	0.735500

Tablo Ek3.80 Tüm Modeller %60 Eğitim %40 Test Verileri ile Tüm Modellerin Test Performans Sonuçları (Deneme 2)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
SVM (RBF)	0.837838	0.850932	0.835824	0.895059	0.884479
Random Forest	0.841216	0.847896	0.840692	0.899447	0.898865
XGBoost	0.790541	0.791946	0.791215	0.894282	0.889693
MLP	0.750000	0.767296	0.748343	0.798894	0.810893
KNN	0.773649	0.750929	0.777618	0.850336	0.863673
Decision Tree	0.706081	0.698962	0.707688	0.801385	0.770269

Tablo Ek3.81: Kara Ağaçları En İyi Parametreler (Deneme 3)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Criterion	gini	gini
Max Depth	8	24
Min Samples Leaf	3	1
Min Samples Split	2	2

Tablo Ek3.82: Kara Ağaçları 10 Fold Cross Validation (CV) Sonuçları (Deneme 3)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7659 ± 0.0418	0.7049 ± 0.0494
F1 Score	0.7694 ± 0.0392	0.6983 ± 0.0517
Balanced Accuracy	0.7653 ± 0.0430	0.7060 ± 0.0496
ROC-AUC	0.8058 ± 0.0536	0.7053 ± 0.0488
PR-AUC	0.7960 ± 0.0709	0.6626 ± 0.0446

Tablo Ek3.83: Kara Ağaçları Confusion Matrix (Deneme 3)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	77	32	108	37
1 (Başarılı)	29	86	41	112

Tablo Ek3.84: Kara Ağaçları %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7264	0.7064	0.7163	109
1	0.7288	0.7478	0.7382	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7277	224
Macro Avg	0.7276	0.7271	0.7272	224
Weighted Avg	0.7276	0.7277	0.7275	224

Tablo Ek3.85: Kara Ağaçları %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7248	0.7448	0.7347	145
1	0.7517	0.7320	0.7417	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7383	298
Macro Avg	0.7383	0.7384	0.7382	298
Weighted Avg	0.7386	0.7383	0.7383	298

Tablo Ek3.86: Kara Ağaçları Genel Test Sınıflandırma Sonuçları (Deneme 3)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7277	0.7382	0.7271	0.7709	0.7529
%60 / %40	0.7383	0.7417	0.7384	0.7384	0.6878

Tablo Ek3.87: Destek Vektör Makinesi En İyi Parametreler (Deneme 3)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
C	1.9069966	90.7602
Gamma	0.1382623	0.1440365
Kernel	RBF	RBF

Tablo Ek3.88: Destek Vektör Makinesi 10 Fold Cross Validation (CV) Sonuçları (Deneme 3)

Metrik	Ortalama $\pm$ Std (%70 / %30)	Ortalama $\pm$ Std (%60 / %40)
Accuracy	0.8579 $\pm$ 0.0348	0.8591 $\pm$ 0.0313
F1 Score	0.8651 $\pm$ 0.0372	0.8697 $\pm$ 0.0287
Balanced Accuracy	0.8572 $\pm$ 0.0342	0.8578 $\pm$ 0.0312
ROC-AUC	0.8956 $\pm$ 0.0408	0.8908 $\pm$ 0.0415
PR-AUC	0.8770 $\pm$ 0.0543	0.8708 $\pm$ 0.0577

Tablo Ek3.89: Destek Vektör Makinesi Confusion Matrix (Deneme 3)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	90	19	118	27
1 (Başarılı)	5	110	8	145

Tablo Ek3.90: Destek Vektör Makinesi %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9474	0.8257	0.8824	109
1	0.8527	0.9565	0.9016	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.8929	224
Macro Avg	0.9000	0.8911	0.8920	224
Weighted Avg	0.8988	0.8929	0.8923	224

Tablo Ek3.91: Destek Vektör Makinesi %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9365	0.8138	0.8708	145
1	0.8430	0.9477	0.8923	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.8826	298
Macro Avg	0.8898	0.8808	0.8816	298
Weighted Avg	0.8885	0.8826	0.8819	298

Tablo Ek3.92: Destek Vektör Makinesi Genel Test Sınıflandırma Sonuçları (Deneme 3)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8929	0.9016	0.8911	0.9325	0.9233
%60 / %40	0.8826	0.8923	0.8808	0.9047	0.8908

Tablo Ek3.93: Rastgele Orman En İyi Parametreler (Deneme 3)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Max Depth	30	30
Min Samples Leaf	1	1
Min Samples Split	2	2
n_estimators	576	576

Tablo Ek3.94: Rastgele Orman 10 Fold Cross Validation (CV) Sonuçları (Deneme 3)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8617 ± 0.0479	0.8527 ± 0.0514
F1 Score	0.8674 ± 0.0520	0.8630 ± 0.0466
Balanced Accuracy	0.8614 ± 0.0468	0.8515 ± 0.0515
ROC-AUC	0.9143 ± 0.0257	0.8959 ± 0.0416
PR-AUC	0.9084 ± 0.0289	0.8896 ± 0.0480

Tablo Ek3.95: Rastgele Orman Confusion Matrix (Deneme 3)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	92	17	116	29
1 (Başarılı)	9	106	14	139

Tablo Ek3.96: Rastgele Orman %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9109	0.8440	0.8762	109
1	0.8618	0.9217	0.8908	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8839	224

Macro Avg	0.8863	0.8829	0.8835	224
Weighted Avg	0.8857	0.8839	0.8837	224

Tablo Ek3.97: Rastgele Orman %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8923	0.8000	0.8436	145
1	0.8274	0.9085	0.8660	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.8557	298
Macro Avg	0.8598	0.8542	0.8548	298
Weighted Avg	0.8590	0.8557	0.8551	298

Tablo Ek3.98: Rastgele Orman Genel Test Sınıflandırma Sonuçları (Deneme 3)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8839	0.8908	0.8829	0.9343	0.9339
%60 / %40	0.8557	0.8660	0.8542	0.9186	0.9100

Tablo Ek3.99: K-En Yakın Komşu En İyi Parametreler (Deneme 3)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
n_neighbors	3	23
p	1 (Manhattan Distance)	1 (Manhattan Distance)
weights	uniform	distance

Tablo Ek3.100: K-En Yakın Komşu 10 Fold Cross Validation (CV) Sonuçları (Deneme 3)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8464 ± 0.0525	0.8013 ± 0.0607
F1 Score	0.8353 ± 0.0617	0.8255 ± 0.0508
Balanced Accuracy	0.8487 ± 0.0513	0.7986 ± 0.0610
ROC-AUC	0.8869 ± 0.0499	0.8523 ± 0.0685
PR-AUC	0.8876 ± 0.0444	0.8386 ± 0.0843

Tablo Ek3.101: K-En Yakın Komşu Confusion Matrix (Deneme 3)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	98	11	95	50
1 (Başarılı)	17	98	8	145

Tablo Ek3.102: K-En Yakın Komşu %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8522	0.8991	0.8750	109
1	0.8991	0.8522	0.8750	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8750	224
Macro Avg	0.8756	0.8756	0.8750	224
Weighted Avg	0.8763	0.8750	0.8750	224

Tablo Ek3.103: K-En Yakın Komşu %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9223	0.6552	0.7661	145
1	0.7436	0.9477	0.8333	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8054	298
Macro Avg	0.8330	0.8014	0.7997	298
Weighted Avg	0.8306	0.8054	0.8006	298

Tablo Ek3.104: K-En Yakın Komşu Genel Test Sınıflandırma Sonuçları (Deneme 3)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8750	0.8750	0.8756	0.9116	0.9105
%60 / %40	0.8054	0.8333	0.8014	0.8957	0.8883

Tablo Ek3.105: Çok Katmanlı Perceptron En İyi Parametreler (Deneme 3)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Activation	tanh	relu

Alpha	0.0066936	0.00002647
Batch Size	32	32
Hidden Layer Sizes	(256, 128, 64)	(128, 64)
Learning Rate Init	0.0005000	0.00063376
Solver	adam	adam

Tablo Ek3.106: Çok Katmanlı Perceptron 10 Fold Cross Validation (CV) Sonuçları (Deneme 3)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8003 ± 0.0519	0.7810 ± 0.0598
F1 Score	0.8208 ± 0.0357	0.7935 ± 0.0552
Balanced Accuracy	0.7988 ± 0.0551	0.7800 ± 0.0598
ROC-AUC	0.8498 ± 0.0686	0.8306 ± 0.0590
PR-AUC	0.8253 ± 0.0815	0.8001 ± 0.0637

Tablo Ek3.107: Çok Katmanlı Perceptron Confusion Matrix (Deneme 3)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	82	27	89	56
1 (Başarılı)	30	85	16	137

Tablo Ek3.108: Çok Katmanlı Perceptron %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7321	0.7523	0.7421	109
1	0.7589	0.7391	0.7489	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7455	224
Macro Avg	0.7455	0.7457	0.7455	224
Weighted Avg	0.7459	0.7455	0.7456	224

Tablo Ek3.109: Çok Katmanlı Perceptron %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 2)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8476	0.6138	0.7120	145
1	0.7098	0.8954	0.7919	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7584	298
Macro Avg	0.7787	0.7546	0.7520	298
Weighted Avg	0.7769	0.7584	0.7530	298

Tablo Ek3.110: Çok Katmanlı Perceptron Genel Test Sınıflandırma Sonuçları (Deneme 3)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7455	0.7489	0.7457	0.8231	0.8162
%60 / %40	0.7584	0.7919	0.7546	0.8208	0.7865

Tablo Ek3.111: XGBoost En İyi Parametreler (Deneme 3)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
n_estimators	1091	629
max_depth	5	7
learning_rate	0.1734215	0.1871962
gamma	0.2492101	0.4579104
min_child_weight	1	1
subsample	0.6656327	0.6346821
colsample_bytree	0.8524683	0.7443896
reg_alpha	0.7295082	0.6905602
reg_lambda	2.4137822	1.8032517

Tablo Ek3.112: XGBoost 10 Fold Cross Validation (CV) Sonuçları (Deneme 3)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8100 ± 0.0518	0.8011 ± 0.0552
F1 Score	0.8142 ± 0.0599	0.8071 ± 0.0579
Balanced Accuracy	0.8099 ± 0.0512	0.8006 ± 0.0547

ROC-AUC	0.8827 ± 0.0391	0.8837 ± 0.0354
PR-AUC	0.8719 ± 0.0444	0.8779 ± 0.0471

Tablo Ek3.113: XGBoost Confusion Matrix (Deneme 3)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	87	22	113	32
1 (Başarılı)	11	104	23	130

Tablo Ek3.114: XGBoost %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8878	0.7982	0.8406	109
1	0.8254	0.9043	0.8631	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8527	224
Macro Avg	0.8566	0.8513	0.8518	224
Weighted Avg	0.8557	0.8527	0.8521	224

Tablo Ek3.115: XGBoost %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 3)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8309	0.7793	0.8043	145
1	0.8025	0.8497	0.8254	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8154	298
Macro Avg	0.8167	0.8145	0.8148	298
Weighted Avg	0.8163	0.8154	0.8151	298

Tablo Ek3.116: XGBoost Genel Test Sınıflandırma Sonuçları (Deneme 3)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8527	0.8631	0.8513	0.9171	0.8964
%60 / %40	0.8154	0.8254	0.8145	0.9163	0.9011

Tablo Ek3.117 Tüm Modeller %70 Eğitim %30 Test Verileri ile Tüm Modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 3)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.857910	0.865051	0.857182	0.895611	0.877010
Random Forest	0.861720	0.867420	0.861396	0.914305	0.908441
KNN	0.846372	0.835304	0.848726	0.886860	0.887555
XGBoost	0.810015	0.814228	0.809943	0.882660	0.871924
MLP	0.800327	0.820821	0.798775	0.849759	0.825285
Decision Tree	0.765856	0.769429	0.765299	0.805802	0.796004

Tablo Ek3.118 Tüm Modeller %70 Eğitim %30 Test Verileri ile Tüm Modellerin Test Performans Sonuçları (Deneme 3)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC- AUC	Test PR- AUC
SVM (RBF)	0.892857	0.901639	0.891105	0.932509	0.923258
Random Forest	0.883929	0.890756	0.882888	0.934304	0.933861
KNN	0.875000	0.875000	0.875628	0.911647	0.910487
XGBoost	0.852679	0.863071	0.851256	0.917112	0.896369
MLP	0.745536	0.748899	0.745712	0.823055	0.816195
Decision Tree	0.727679	0.738197	0.727124	0.770882	0.752874

Tablo Ek3.119 Tüm Modeller %60 Eğitim %40 Test Verileri ile Tüm Modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 3)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.859141	0.869671	0.857839	0.890773	0.870842
Random Forest	0.852677	0.862963	0.851534	0.895892	0.889632
KNN	0.801313	0.825532	0.798631	0.852321	0.838621
XGBoost	0.801111	0.807110	0.800616	0.883741	0.877899
MLP	0.780960	0.793498	0.779955	0.830561	0.800135
Decision Tree	0.704899	0.698277	0.706042	0.705251	0.662608

Tablo Ek3.120 Tüm Modeller %60 Eğitim %40 Test Verileri ile Tüm Modellerin Test Performans Sonuçları (Deneme 3)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
SVM (RBF)	0.882550	0.892308	0.880753	0.904710	0.890833
Random Forest	0.855705	0.866044	0.854248	0.918571	0.910025
KNN	0.805369	0.833333	0.801442	0.895695	0.888317
XGBoost	0.815436	0.825397	0.814492	0.916295	0.901108
MLP	0.758389	0.791908	0.754609	0.820825	0.786505
Decision Tree	0.738255	0.741722	0.738427	0.738427	0.687832

Tablo Ek3.121: Kara Ağaçları En İyi Parametreler (Deneme 4)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
Criterion	gini	entropy
Max Depth	None	16
Min Samples Leaf	1	2
Min Samples Split	48	15

Tablo Ek3.122: Kara Ağaçları 10 Fold Cross Validation (CV) Sonuçları (Deneme 4)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.7111 ± 0.0540	0.7338 ± 0.0655
F1 Score	0.7183 ± 0.0546	0.7398 ± 0.0643
Balanced Accuracy	0.7108 ± 0.0544	0.7342 ± 0.0655
ROC-AUC	0.7559 ± 0.0454	0.7660 ± 0.0770
PR-AUC	0.7671 ± 0.0371	0.7415 ± 0.0946

Tablo Ek3.123: Kara Ağaçları Confusion Matrix (Deneme 4)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	86	25	119	29
1 (Başarılı)	33	82	39	114

Tablo Ek3.124: Kara Ağaçları %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7227	0.7748	0.7478	111
1	0.7664	0.7130	0.7387	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7434	226
Macro Avg	0.7445	0.7439	0.7433	226
Weighted Avg	0.7449	0.7434	0.7432	226

Tablo Ek3.125: Kara Ağaçları %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.7532	0.8041	0.7778	148
1	0.7972	0.7451	0.7703	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.7741	301
Macro Avg	0.7752	0.7746	0.7740	301
Weighted Avg	0.7755	0.7741	0.7740	301

Tablo Ek3.126: Kara Ağaçları Genel Test Sınıflandırma Sonuçları (Deneme 4)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.7434	0.7387	0.7439	0.7868	0.7969
%60 / %40	0.7741	0.7703	0.7746	0.7939	0.7619

Tablo Ek3.127: Destek Vektör Makinesi En İyi Parametreler (Deneme 4)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
C	56.6335	111.3796
Gamma	0.1687869	0.1739617
Kernel	RBF	RBF

Tablo Ek3.128: Destek Vektör Makinesi 10 Fold Cross Validation (CV) Sonuçları (Deneme 4)

Metrik	Ortalama $\pm$ Std (%70 / %30)	Ortalama $\pm$ Std (%60 / %40)
Accuracy	0.8668 $\pm$ 0.0421	0.8693 $\pm$ 0.0505
F1 Score	0.8743 $\pm$ 0.0382	0.8767 $\pm$ 0.0485
Balanced Accuracy	0.8661 $\pm$ 0.0423	0.8686 $\pm$ 0.0502
ROC-AUC	0.8854 $\pm$ 0.0499	0.8883 $\pm$ 0.0431
PR-AUC	0.8663 $\pm$ 0.0787	0.8620 $\pm$ 0.0734

Tablo Ek3.129: Destek Vektör Makinesi Confusion Matrix (Deneme 4)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	95	16	118	30
1 (Başarılı)	5	110	8	145

Tablo Ek3.130: Destek Vektör Makinesi %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9500	0.8559	0.9005	111
1	0.8730	0.9565	0.9129	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.9071	226
Macro Avg	0.9115	0.9062	0.9067	226
Weighted Avg	0.9108	0.9071	0.9068	226

Tablo Ek3.131: Destek Vektör Makinesi %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9365	0.7973	0.8613	148
1	0.8286	0.9477	0.8841	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.8738	301
Macro Avg	0.8825	0.8725	0.8727	301
Weighted Avg	0.8816	0.8738	0.8729	301

Tablo Ek3.132: Destek Vektör Makinesi Genel Test Sınıflandırma Sonuçları (Deneme 4)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.9071	0.9129	0.9062	0.9389	0.9294
%60 / %40	0.8738	0.8841	0.8725	0.9243	0.9111

Tablo Ek3.133: Rastgele Orman En İyi Parametreler (Deneme 4)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
max_depth	30	30
min_samples_leaf	1	1
min_samples_split	4	2
n_estimators	855	576

Tablo Ek3.134: Rastgele Orman 10 Fold Cross Validation (CV) Sonuçları (Deneme 4)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8496 ± 0.0444	0.8493 ± 0.0583
F1 Score	0.8560 ± 0.0431	0.8551 ± 0.0577
Balanced Accuracy	0.8490 ± 0.0443	0.8490 ± 0.0580
ROC-AUC	0.9067 ± 0.0420	0.8997 ± 0.0378
PR-AUC	0.9035 ± 0.0488	0.8913 ± 0.0531

Tablo Ek3.135: Rastgele Orman Confusion Matrix (Deneme 4)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	96	15	125	23
1 (Başarılı)	8	107	14	139

Tablo Ek3.136: Rastgele Orman %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.9231	0.8649	0.8930	111
1	0.8770	0.9304	0.9030	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8982	226

Macro Avg	0.9001	0.8976	0.8980	226
Weighted Avg	0.8997	0.8982	0.8981	226

Tablo Ek3.137: Rastgele Orman %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8993	0.8446	0.8711	148
1	0.8580	0.9085	0.8825	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	–	–	0.8771	301
Macro Avg	0.8787	0.8765	0.8768	301
Weighted Avg	0.8783	0.8771	0.8769	301

Tablo Ek3.138: Rastgele Orman Genel Test Sınıflandırma Sonuçları (Deneme 4)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8982	0.9030	0.8976	0.9438	0.9428
%60 / %40	0.8771	0.8825	0.8765	0.9318	0.9164

Tablo Ek3.139: K-En Yakın Komşu0:2 En İyi Parametreler (Deneme 4)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
n_neighbors	3	6
p	1	1
weights	uniform	distance

Tablo Ek3.140: K-En Yakın Komşu0:3 10 Fold Cross Validation (CV) Sonuçları (Deneme 4)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8649 ± 0.0442	0.8382 ± 0.0506
F1 Score	0.8516 ± 0.0527	0.8357 ± 0.0531
Balanced Accuracy	0.8662 ± 0.0448	0.8386 ± 0.0511
ROC-AUC	0.8962 ± 0.0362	0.8966 ± 0.0414
PR-AUC	0.8978 ± 0.0342	0.9026 ± 0.0510

Tablo Ek3.141: K-En Yakın Komşu0:4 Confusion Matrix (Deneme 4)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	104	7	126	22
1 (Başarılı)	19	96	17	136

Tablo Ek3.142: K-En Yakın Komşu0:5 %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8455	0.9369	0.8889	111
1	0.9320	0.8348	0.8807	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8850	226
Macro Avg	0.8888	0.8859	0.8848	226
Weighted Avg	0.8895	0.8850	0.8847	226

Tablo Ek3.143: K-En Yakın Komşu0:6 %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8811	0.8514	0.8660	148
1	0.8608	0.8889	0.8746	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8704	301
Macro Avg	0.8709	0.8701	0.8703	301
Weighted Avg	0.8708	0.8704	0.8704	301

Tablo Ek3.144: K-En Yakın Komşu0:7 Genel Test Sınıflandırma Sonuçları (Deneme 4)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8850	0.8807	0.8859	0.9218	0.9119
%60 / %40	0.8704	0.8746	0.8701	0.9285	0.9311

Tablo Ek3.145: Çok Katmanlı Perceptron En İyi Parametreler (Deneme 4)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
activation	relu	relu

alpha	0.0007128	1.5519e-06
batch_size	32	64
hidden_layer_sizes	(256, 128, 64)	(128, 64)
learning_rate_init	0.0032724	0.0111021
solver	adam	adam

Tablo Ek3.146: Çok Katmanlı Perceptron 10 Fold Cross Validation (CV) Sonuçları (Deneme 4)

Metrik	Ortalama ± Std	Ortalama ± Std
Accuracy	0.8097 ± 0.0530	0.7895 ± 0.0522
F1 Score	0.8261 ± 0.0501	0.7986 ± 0.0513
Balanced Accuracy	0.8084 ± 0.0530	0.7891 ± 0.0517
ROC-AUC	0.8613 ± 0.0522	0.8570 ± 0.0429
PR-AUC	0.8386 ± 0.0641	0.8380 ± 0.0527

Tablo Ek3.147: Çok Katmanlı Perceptron Confusion Matrix (Deneme 4)

	%70 / %30		%60 / %40	
Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	92	19	117	31
1 (Başarılı)	19	96	22	131

Tablo Ek3.148: Çok Katmanlı Perceptron %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8288	0.8288	0.8288	111
1	0.8348	0.8348	0.8348	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8319	226
Macro Avg	0.8318	0.8318	0.8318	226
Weighted Avg	0.8319	0.8319	0.8319	226

Tablo Ek3.149: Çok Katmanlı Perceptron %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8417	0.7905	0.8153	148
1	0.8086	0.8562	0.8317	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8239	301
Macro Avg	0.8252	0.8234	0.8235	301
Weighted Avg	0.8249	0.8239	0.8237	301

Tablo Ek3.150: Çok Katmanlı Perceptron Genel Test Sınıflandırma Sonuçları (Deneme 4)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8319	0.8348	0.8318	0.9143	0.8880
%60 / %40	0.8239	0.8317	0.8234	0.8877	0.8588

Tablo Ek3.151: XGBoost En İyi Parametreler (Deneme 4)

Parametre	Değer (%70 / %30)	Değer (%60 / %40)
colsample_bytree	0.7480	0.9897
gamma	0.1371	0.1001
learning_rate	0.00941	0.0670
max_depth	9	6
min_child_weight	2	1
n_estimators	739	585
reg_alpha	0.1656	0.5210
reg_lambda	2.1144	2.0589
subsample	0.6179	0.6356

Tablo Ek3.152: XGBoost 10 Fold Cross Validation (CV) Sonuçları (Deneme 4)

Metrik	Ortalama ± Std (%70 / %30)	Ortalama ± Std (%60 / %40)
Accuracy	0.8230 ± 0.0366	0.8115 ± 0.0599
F1 Score	0.8293 ± 0.0370	0.8166 ± 0.0601
Balanced Accuracy	0.8226 ± 0.0366	0.8112 ± 0.0598

ROC-AUC	0.8977 ± 0.0367	0.8875 ± 0.0479
PR-AUC	0.8954 ± 0.0416	0.8827 ± 0.0491

Tablo Ek3.153: XGBoost Confusion Matrix (Deneme 4)
%70 / %30 %60 / %40

Gerçek \ Tahmin	0	1	0	1
0 (Başarısız)	92	19	124	24
1 (Başarılı)	11	104	20	133

Tablo Ek3.154: XGBoost %70 Eğitim %30 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8932	0.8288	0.8598	111
1	0.8455	0.9043	0.8739	115
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8673	226
Macro Avg	0.8694	0.8666	0.8669	226
Weighted Avg	0.8689	0.8673	0.8670	226

Tablo Ek3.155: XGBoost %60 Eğitim %40 Test Verileri ile Test Sınıflandırma Raporu (Deneme 4)

Sınıf	Precision	Recall	F1-Score	Support
0	0.8611	0.8378	0.8493	148
1	0.8471	0.8693	0.8581	153
Ortalama Türü	Precision	Recall	F1-Score	Support
Accuracy	—	—	0.8538	301
Macro Avg	0.8541	0.8536	0.8537	301
Weighted Avg	0.8540	0.8538	0.8538	301

Tablo Ek3.156: XGBoost Genel Test Sınıflandırma Sonuçları (Deneme 4)

	Accuracy	F1 Score	Balanced Accuracy	ROC-AUC	PR-AUC
%70 / %30	0.8673	0.8739	0.8666	0.9383	0.9229
%60 / %40	0.8538	0.8581	0.8536	0.9293	0.9038

Tablo Ek3.157: Tüm Modeller %70 Eğitim %30 Test Verileri ile Tüm Modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 4)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.857910	0.865051	0.857182	0.895611	0.877010
Random Forest	0.861720	0.867420	0.861396	0.914305	0.908441
KNN	0.846372	0.835304	0.848726	0.886860	0.887555
XGBoost	0.810015	0.814228	0.809943	0.882660	0.871924
MLP	0.800327	0.820821	0.798775	0.849759	0.825285
Decision Tree	0.765856	0.769429	0.765299	0.805802	0.796004

Tablo Ek3.158: Tüm Modeller %70 Eğitim %30 Test Verileri ile Tüm Modellerin Test Performans Sonuçları (Deneme 4)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC- AUC	Test PR- AUC
SVM (RBF)	0.907080	0.912863	0.906189	0.938895	0.929390
Random Forest	0.898230	0.902954	0.897650	0.943831	0.942801
KNN	0.884956	0.880734	0.885860	0.921778	0.911925
XGBoost	0.867257	0.873950	0.866588	0.938347	0.922940
MLP	0.831858	0.834783	0.831806	0.914297	0.888042
Decision Tree	0.743363	0.738739	0.743909	0.786800	0.796896

Tablo Ek3.159: Tüm Modeller %60 Eğitim %40 Test Verileri ile Tüm Modellerin Çapraz Doğrulama (CV) Performans Sonuçları (Deneme 4)

Model	CV Accuracy	CV F1 Score	CV Balanced Accuracy	CV ROC- AUC	CV PR- AUC
SVM (RBF)	0.869324	0.876746	0.868577	0.888271	0.862026
Random Forest	0.849324	0.855113	0.849012	0.899699	0.891302
KNN	0.838164	0.835722	0.838636	0.896559	0.902576
XGBoost	0.811498	0.816601	0.811166	0.887455	0.882690
MLP	0.789469	0.798580	0.789130	0.857020	0.837960
Decision Tree	0.733816	0.739823	0.734190	0.766034	0.741469

Tablo Ek3.160: Tüm Modeller %60 Eğitim %40 Test Verileri ile Tüm Modellerin Test Performans Sonuçları (Deneme 4)

Model	Test Accuracy	Test F1 Score	Test Balanced Accuracy	Test ROC-AUC	Test PR-AUC
SVM (RBF)	0.873754	0.884146	0.872505	0.924307	0.911139
Random Forest	0.877076	0.882540	0.876546	0.931814	0.916401
KNN	0.870432	0.874598	0.870120	0.928524	0.931111
XGBoost	0.853821	0.858065	0.853559	0.929341	0.903804
MLP	0.823920	0.831746	0.823375	0.887741	0.858776
Decision Tree	0.774086	0.770270	0.774576	0.793941	0.761865



Ek 4: Regresyon Sonuçları Tabloları

Tablo Ek4.1: Karar Ağaçları En İyi Parametreler (Deneme 1)

Parametre	Değer
Criterion	squared_error
max_depth	5
min_samples_leaf	3
min_samples_split	2
ccp_alpha	0.01877

Tablo Ek4.2: Karar Ağaçları 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama $\pm$ Std
RMSE	18.1704 $\pm$ 2.3882
MAE	13.5153 $\pm$ 2.1665
R ²	0.2305 $\pm$ 0.3391

Tablo Ek4.3: Karar Ağaçları Test Sonuçları (Deneme 1)

Metrik	Değer
RMSE	16.3001
MAE	11.9639
R ²	0.4803

Tablo Ek4.4: Destek Vektör Makinesi En İyi Parametreler (Deneme 1)

Parametre	Değer
Kernel	RBF
C	668.5510
Epsilon	0.000236
Gamma	0.08384

Tablo Ek4.5: Destek Vektör Makinesi Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama $\pm$ Std
RMSE	14.1454 $\pm$ 2.1242
MAE	10.7192 $\pm$ 1.8453
R ²	0.5549 $\pm$ 0.1391

Tablo Ek4.6: Destek Vektör Makinesi Test Sonuçları (Deneme 1)

Metrik	Değer
RMSE	13.4871
MAE	10.0474
R ²	0.6442

Tablo Ek4.7: Rastgele Orman En İyi Parametreler (Deneme 1)

Parametre	Değer
criterion	poisson
max_depth	30
max_features	log2
min_samples_leaf	2
min_samples_split	4
n_estimators	1090
bootstrap	False

Tablo Ek4.8: Rastgele Orman 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama $\pm$ Std
RMSE	14.7251 $\pm$ 1.9264
MAE	11.5029 $\pm$ 1.5673
R ²	0.5094 $\pm$ 0.1793

Tablo Ek4.9: Rastgele Orman Test Sonuçları (Deneme 1)

Metrik	Değer
RMSE	13.0661
MAE	9.8173
R ²	0.6661

Tablo Ek4.10: K-En Yakın Komşu En İyi Parametreler (Deneme 1)

Parametre	Değer
n_neighbors	8
leaf_size	13
p	1
weights	distance

Tablo Ek4.11: K-En Yakın Komşu 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama $\pm$ Std
RMSE	14.7397 $\pm$ 1.8646
MAE	10.9197 $\pm$ 1.5497
R ²	0.5162 $\pm$ 0.1470

Tablo Ek4.12: K-En Yakın Komşu Test Sonuçları (Deneme 1)

Metrik	Değer
RMSE	13.7771
MAE	10.1877
R ²	0.6287

Tablo Ek4.13: Çok Katmanlı Perceptron En İyi Parametreler (Deneme 1)

Parametre	Değer
activation	tanh
alpha	1.8405e-06
batch_size	32
hidden_layer_sizes	(128)
learning_rate_init	0.0047375

solver

adam

Tablo Ek4.14: Çok Katmanlı Perceptron 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama $\pm$ Std
RMSE	16.1944 $\pm$ 2.3160
MAE	12.9884 $\pm$ 1.7028
R ²	0.4268 $\pm$ 0.1407

Tablo Ek4.15: Çok Katmanlı Perceptron Test Sonuçları (Deneme 1)

Metrik	Değer
RMSE	14.7069
MAE	11.5422
R ²	0.5769

Tablo Ek4.16: XGBoost En İyi Parametreler (Deneme 1)

Parametre	Değer
colsample_bytree	0.6071
gamma	3.0128
learning_rate	0.00801
max_depth	9
min_child_weight	3
n_estimators	1840
reg_alpha	0.7808
reg_lambda	1.6495
subsample	0.6233

Tablo Ek4.17: XGBoost 10 Fold Cross Validation (CV) Sonuçları (Deneme 1)

Metrik	Ortalama $\pm$ Std
RMSE	14.4760 $\pm$ 1.5975
MAE	10.7652 $\pm$ 1.4044
R ²	0.5137 $\pm$ 0.1991

Tablo Ek4.18: XGBoost Test Sonuçları (Deneme 1)

Metrik	Değer
RMSE	12.4614
MAE	8.9904
R ²	0.6963

Tablo Ek4.19: Karar Ağaçları En İyi Parametreler (Deneme 2)

Parametre	Değer
criterion	squared_error
max_depth	24
min_samples_leaf	5
min_samples_split	23
ccp_alpha	0.01034

Tablo Ek4.20: Karar Ağaçları 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std
RMSE	17.5511 $\pm$ 2.2131
MAE	13.3408 $\pm$ 1.9283
R ²	0.2706 $\pm$ 0.2840

Tablo Ek4.21: Karar Ağaçları Test Sonuçları (Deneme 2)

Metrik	Değer
RMSE	16.7092
MAE	11.8291
R ²	0.4527

Tablo Ek4.22: Destek Vektör Makinesi En İyi Parametreler (Deneme 2)

Parametre	Değer
Kernel	RBF
C	129.6242
Epsilon	0.07070
Gamma	0.07996

Tablo Ek4.23: Destek Vektör Makinesi Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std
RMSE	13.9124 $\pm$ 1.5695
MAE	10.4385 $\pm$ 1.1340
R ²	0.5490 $\pm$ 0.1507

Tablo Ek4.24: Destek Vektör Makinesi Test Sonuçları (Deneme 2)

Metrik	Değer
RMSE	12.7653
MAE	9.2919
R ²	0.6806

Tablo Ek4.25: Rastgele Orman En İyi Parametreler (Deneme 2)

Parametre	Değer
bootstrap	True
criterion	friedman_mse
max_depth	18
max_features	sqrt
min_samples_leaf	1
min_samples_split	5
n_estimators	797

Tablo Ek4.26: Rastgele Orman 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std
RMSE	13.8846 $\pm$ 2.1836
MAE	10.9533 $\pm$ 1.6893
R ²	0.5682 $\pm$ 0.1042

Tablo Ek4.27: Rastgele Orman Test Sonuçları (Deneme 2)

Metrik	Değer
RMSE	13.0115
MAE	9.8536
R ²	0.6681

Tablo Ek4.28: K-En Yakın Komşu En İyi Parametreler (Deneme 2)

Parametre	Değer
n_neighbors	3
leaf_size	13
p	1
weights	distance

Tablo Ek4.29: K-En Yakın Komşu 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std
RMSE	16.0581 $\pm$ 1.8273
MAE	11.3526 $\pm$ 1.1382
R ²	0.3988 $\pm$ 0.2419

Tablo Ek4.30: K-En Yakın Komşu Test Sonuçları (Deneme 2)

Metrik	Değer
RMSE	16.5790
MAE	10.6001
R ²	0.4612

Tablo Ek4.31: Çok Katmanlı Perceptron En İyi Parametreler (Deneme 2)

Parametre	Değer
activation	tanh
alpha	0.001216
batch_size	128
hidden_layer_sizes	(128)
learning_rate_init	0.000153
solver	adam

Tablo Ek4.32: Çok Katmanlı Perceptron 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std
RMSE	14.7755 $\pm$ 2.1598
MAE	12.0948 $\pm$ 1.3527
R ²	0.5024 $\pm$ 0.1379

Tablo Ek4.33: Çok Katmanlı Perceptron Test Sonuçları (Deneme 2)

Metrik	Değer
RMSE	14.1266
MAE	11.5385
R ²	0.6088

Ek4.34: XGBoost En İyi Parametreler (Deneme 2)

Parametre	Değer
colsample_bytree	0.6061
gamma	4.6672
learning_rate	0.0174
max_depth	7
min_child_weight	4
n_estimators	1339
reg_alpha	0.0706
reg_lambda	2.1060
subsample	0.6106

Tablo Ek4.35: XGBoost 10 Fold Cross Validation (CV) Sonuçları (Deneme 2)

Metrik	Ortalama $\pm$ Std
RMSE	13.7816 $\pm$ 1.4249
MAE	10.2795 $\pm$ 0.9182
R ²	0.5451 $\pm$ 0.1943

Tablo Ek4.36: XGBoost Test Sonuçları (Deneme 2)

Metrik	Değer
RMSE	13.2046
MAE	9.4327
R ²	0.6582

ÖZGEÇMİŞ

Kişisel Bilgiler

Adı Soyadı : Müjde Dilek SOBUTAY

Eğitim Durumu

Lisans Öğrenimi : Anadolu Üniversitesi Açıköğretim Fakültesi Yönetim Bilişim Sistemleri Bölümü

Yüksek Lisans Öğrenimi :

Bildiği Yabancı Diller : İngilizce

Bilimsel Faaliyet/Yayınlar :

Aldığı Ödüller :

İş Deneyimi

Stajlar :

Projeler ve Kurs Belgeleri :

Çalıştığı Kurumlar :

