

KARADENİZ TEKNİK ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ

BİLGİSAYAR MÜHENDİLİĞİ ANABİLİM DALI

**İSTİFLEME MODELLERİYLE EL YAZISI VE ÇİZİMDEN
DEPRESYON, ANKSİYETE VE STRES SINIFLANDIRMASI**

YÜKSEK LİSANS TEZİ

Semra BAYRAK

MART 2025
TRABZON



KARADENİZ TEKNİK ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ

BİLGİSAYAR MÜHENDİSLİĞİ ANABİLİM DALI

**İSTİFLEME MODELLERİYLE EL YAZISI VE ÇİZİMDEN
DEPRESYON, ANKSİYETE VE STRES SINIFLANDIRMASI**

Semra BAYRAK

Karadeniz Teknik Üniversitesi Fen Bilimleri Enstitüsünde
"BİLGİSAYAR YÜKSEK MÜHENDİSİ"
Unvanı Verilmesi İçin Kabul Edilen Tezdir.

Tezin Enstitüye Verildiği Tarih : 07 / 02 / 2025

Tezin Savunma Tarihi : 11 / 03 / 2025

Tez Danışmanı : Dr. Öğr. Üyesi Murat AYKUT

Trabzon 2025

ÖNSÖZ

“İstifleme Modellerle El Yazısı ve Çizimden Depresyon, Anksiyete ve Stres Sınıflandırması” isimli bu tez Karadeniz Teknik Üniversitesi Fen Bilimleri Enstitüsü Bilgisayar Bilimleri Anabilim Dalı, Yüksek Lisans Programı’nda hazırlanmıştır.

Bu tezin oluşmasında katkıda bulunan tez danışmanım Dr. Öğr. Üyesi Murat AYKUT’a, sabır ve rehberlikleri için teşekkür ederim. Ayrıca, bu çalışmada görev alarak bilgi ve deneyimlerini paylaşan Dr. Öğr. Üyesi Sedat GOLGİYZ’a yardımlarından dolayı teşekkür ederim. Son olarak, bu süreçte beni destekleyen aileme ve arkadaşlarıma teşekkür eder, minnettarlığımı sunarım. Onların desteği ve anlayışı olmadan bu çalışmayı tamamlayamazdım.

Semra BAYRAK
Trabzon 2025

TEZ ETİK BEYANNAMESİ

Yüksek Lisans Tezi olarak sunduğum “İstifleme Modellerle El Yazısı ve Çizimden Depresyon, Anksiyete ve Stres Sınıflandırması” başlıklı bu çalışmayı baştan sona kadar danışmanım Dr. Öğr. Üyesi Murat AYKUT’un sorumluluğunda tamamladığımı, deneyleri/analizleri kendim yaptığımı, başka kaynaklardan aldığım bilgileri metinde ve kaynakçada eksiksiz olarak gösterdiğimi, çalışma sürecinde bilimsel araştırma ve etik kurallara uygun olarak davrandığımı ve aksinin ortaya çıkması durumunda her türlü yasal sonucu kabul ettiğimi beyan ederim. 11/03/2025

Semra BAYRAK

İÇİNDEKİLER

	<u>Sayfa No</u>
ÖNSÖZ	III
TEZ ETİK BEYANNAMESİ.....	IV
İÇİNDEKİLER.....	V
ÖZET	VIII
SUMMARY.....	IX
ŞEKİLLER DİZİNİ	X
TABLolar DİZİNİ.....	XI
SEMBOLLER VE KISALTMALAR DİZİNİ	XII
1. GENEL BİLGİLER	1
1.1. Giriş.....	1
1.2. Literatür Taraması	2
1.3. Duygular	6
1.3.1. Depresyon.....	6
1.3.2. Anksiyete.....	6
1.3.3. Stres.....	7
1.3.4. DAS Ölçeği.....	7
1.4. Veri Ön İşleme	9
1.4.1. EMOTHAW Veri Tabanındaki Ham Özellikler.....	9
1.4.2. Öznitelik Türetme.....	9
1.4.3. Zaman Tabanlı Özellikler.....	10
1.4.4. Frekans ve Kepstral Özellikler.....	11
1.4.5. Kinematik Özellikler.....	16
1.4.6. İstatistiksel Özellikler.....	17
1.5. Boyut İndirgeme.....	19
1.5.1. Temel Bileşen Analizi.....	20
1.6. Öznitelik Seçimi	22
1.7. Öznitelik Seçimi yöntemleri.....	22
1.8. Veri Dengeleme.....	23
1.8.1. Örnek Çoğaltma	24
1.8.2. Örnek Azaltma	24

1.8.3. Hibrit Yaklaşımlar.....	25
1.9. Topluluk Öğrenme	25
1.9.1. Artırma Yöntemi	26
1.9.2. Örneklem Tabanlı Toplulaştırma	28
1.9.3. İstifleme	29
2. YAPILAN ÇALIŞMALAR.....	31
2.1. EMOTHAW Veri Tabanı	33
2.1.1. Görevler Seti	33
2.1.2. Katılımcılar	35
2.2. Özellik Analizi	37
2.3. Öznitelik Seçimi	37
2.3.1. Gradyan Artırma Sınıflandırıcısı ile Özellik Önem Derecesi Değerlendirme.....	38
2.4. Veri Dengeleme Yöntemi:ADASYN+Tomek Links Hibrit Yaklaşımı	40
2.5. Sınıflandırma	42
2.5.1. Lojistik Regresyon	42
2.5.2. Rastgele Orman Algoritması.....	43
2.5.3. Gradyan Artırma Algoritması	45
2.5.4. Aşırı Gradyan Artırma Algoritması.....	46
2.5.5. Hafif Gradyan Artırma Makinesi	47
2.6. İstifleme Modeller	49
2.7. Optuna: Dinamik Hiperparametre Optimizasyonu	50
2.8. Performans Değerlendirme.....	52
2.8.1. Doğruluk	53
2.8.2. Özgüllük.....	53
2.8.3. Hassasiyet.....	53
2.8.4. Kesinlik	54
2.8.5. F1 Skoru.....	54
2.8.6. Eğri Altındaki Alan (AUC).....	54
2.9. k-Katlı Çapraz Doğrulama ile Eğitim-Test-Doğrulama Setlerinin Kullanımı.....	54
3. BULGULAR VE TARTIŞMA	56
3.1. Sınıflandırma Modelleri.....	57
3.2. Meta Model Seçimi Karşılaştırılması.....	59

3.3. Temel Bileşen Analizi ile Boyut İndirgeme Yaklaşımı	60
3.4. Gradyan Artırma ile Özellik Seçimi Yaklaşımı	62
3.5. Özellik Seçimi ve Veri Artırma Sırasının Performansa Etkisi.....	63
3.6. Veri Artırma Olmadan Model Performansı.....	65
3.7. Temel Bileşen Analizi ve Gradyan Artırma Yaklaşımının Birleşik Kullanımı.....	66
3.8. Sadece Veri Artırma ile Model Performansının Değerlendirilmesi	67
4. SONUÇLAR.....	69
5. ÖNERİLER.....	72
6. KAYNAKLAR	73
ÖZGEÇMİŞ	

Yüksek Lisans Tezi

ÖZET

İSTİFLEME MODELLERİYLE EL YAZISI VE ÇİZİMDEN DEPRESYON,
ANKSİYETE VE STRES SINIFLANDIRMASI
Semra BAYRAK

Karadeniz Teknik Üniversitesi
Fen Bilimleri Enstitüsü
Bilgisayar Mühendisliği Anabilim Dalı
Danışman: Dr. Öğr. Üyesi Murat AYKUT
2025, 78 Sayfa

El yazısı ve çizim verileri, bireylerin duygusal durumlarını yansıtan dinamik ve istatistiksel özellikler taşıması nedeniyle, duygudurum analizinde anlamlı bir veri kaynağı olarak değerlendirilmektedir. Bu çalışma, el yazısı ve çizimlerden elde edilen dinamik ve istatistiksel özellikleri kullanarak duygusal durumları tespit etmek için en uygun model kombinasyonunu belirlemeyi amaçlamaktadır. Öznitelik vektörleri, fiziksel (kinematik), istatistiksel ve sinyal işleme (spektral, kepsstral ve frekans alanı) analizlerinden türetilmiştir. Sınıf dengesizliği problemini ele almak amacıyla, aşırı örnekleme ve örnek azaltma yöntemleri birlikte uygulanmıştır. Bu yaklaşımla, azınlık sınıfındaki örneklerin sayısı çoğaltılarak aşırı örnekleme uygulanmış, çoğunluk sınıfı ise az örnekleme yöntemiyle azaltılmıştır. Boyut indirgeme, öznitelik çıkarımı ve öznitelik seçimi yöntemleri kullanılarak gerçekleştirilmiştir. En az ve en etkili öznitelikler belirlendikten sonra, örnekler iki seviyeli topluluk öğrenme yöntemi ile sınıflandırılmıştır. Çalışmada, karar ağaçları temelli popüler dört farklı yöntem temel ve meta model olarak belirlenmiştir. Model ve hiper parametre optimizasyonu Optuna çerçevesinde gerçekleştirilmiştir. Modellerin performansı, kamuya açık EMOTHAW veri tabanı üzerinde yapılan deneylerle değerlendirilmiştir. Bu çalışmada, 129 kişinin yedi farklı çizim/yazı görevinden elde edilen veriler kullanılarak depresyon, anksiyete ve stresin duygusal durumlarını belirlemek hedeflenmiştir. Elde edilen sonuçlar, modelin dikkate değer bir performans sergilediğini göstermektedir.

Anahtar Kelimeler: İstifleme Model, Duygu Durumu Tespiti, Çevrimiçi El Yazısı Analizi

Master's Thesis

SUMMARY

CLASSIFICATION OF DEPRESSION, ANXIETY, AND STRESS FROM HANDWRITING AND DRAWING USING STACKING MODELS

Semra BAYRAK

Karadeniz Technical University
Institute of Science
Department of Computer Science
Supervisor: Asst. Prof. Murat AYKUT
2025, 78 Pages

Handwriting and drawing data are considered a meaningful data source in emotion analysis, as they contain dynamic and statistical features that reflect individuals' emotional states. This study aims to determine the optimal model combination for detecting emotional states using dynamic and statistical features obtained from handwriting and drawings. Feature vectors were derived from physical (kinematic), statistical, and signal processing (spectral, cepstral, and frequency domain) analyses. To address the class imbalance problem, oversampling and undersampling methods were applied together. With this approach, the number of samples in the minority class was increased through oversampling, while the majority class was reduced using undersampling methods. Dimensionality reduction was performed using feature extraction and feature selection methods. After identifying the minimum and most effective features, the samples were classified using a two-level ensemble learning method. In the study, four popular tree-based methods were selected as base and meta models. Model and hyperparameter optimization were carried out within the Optuna framework. The performance of the models was evaluated through experiments conducted on the publicly available EMOTHAW dataset. In this study, the emotional states of depression, anxiety, and stress were aimed to be determined using data obtained from seven different handwriting/drawing tasks of 129 individuals. The results show that the model demonstrates a noteworthy performance.

Key Words: Stacking Model, Emotional State Detection, Online Handwriting Analysis

ŞEKİLLER DİZİNİ

	<u>Sayfa No</u>
Şekil 1. Ham veriler için örnek bir SVC dosyası	9
Şekil 2. AdaBoost yönteminde zayıf öğrencilerin sıralı eğitimi.....	27
Şekil 3. Örnekleme tabanlı toplulaştırma yöntemi ile paralel model eğitim süreci	28
Şekil 4. İstifleme yönteminin temel mantığı ve işleyiş diyagramı.....	29
Şekil 5. Tez kapsamında yapılan çalışmanın akış diyagramı	32
Şekil 6. Her katılımcı için gerçekleştirilen görevler ve tüm görevlerden toplanan yazı ve çizim örnekleri.....	34
Şekil 7. Beşgen çizim göreviyle ilgili bir svc dosyasının bir bölümü	35
Şekil 8. Veri tabanındaki örneklerin ikili sınıflandırma üzerindeki dağılımı.	36
Şekil 9. PCA-GB tabanlı özellik seçiminin akış şeması	39
Şekil 10. Rastgele orman sınıflandırıcısının çalışma mimarisi.....	43
Şekil 11. Aşırı Gradyan Artırma algoritmasının çalışma mantığı	47
Şekil 12. XGBoost ve LightGBM algoritmalarında ağaç büyüme stratejileri.....	48
Şekil 13. 5 -katlı statik tek katmanlı hibrit istifleme mimarisi.....	49
Şekil 14. Karışıklık matrisi	53
Şekil 15. k-Kat çapraz doğrulama ve eğitim-test ayrımının gösterimi	55

TABLolar DİZİNİ

	<u>Sayfa No</u>
Tablo 1. Duygusal duruma göre DAS puan aralıkları	8
Tablo 2. İkili etiketleme.....	8
Tablo 3. Kalem hareketlerinin zamansal dinamiklerini gösteren özellikler	10
Tablo 4. Frekans analizinde çıkarılan özellik notasyonları ve tanımları	12
Tablo 5. Kepstral analizde çıkarılan özellik notasyonları ve tanımları.....	14
Tablo 6. Kinematik özellikler için notasyonlar ve tanımları	17
Tablo 7. İstatistiksel özellikler için notasyonlar ve tanımları	18
Tablo 8. Modellerin bireysel performans analizi.....	57
Tablo 9. Ön işleme uygulanmadan, LGB + XGB temel modelleri ve LR meta model ile elde edilen sonuçlar	59
Tablo 10. Ön işleme uygulanmadan, LGB + XGB temel modelleri ve XGBoost meta model ile elde edilen sonuçlar.....	60
Tablo 11. PCA ile model kombinasyonlarının sınıflandırma performansları.....	61
Tablo 12. GBC ile model kombinasyonlarının sınıflandırma performansları	62
Tablo 13. PCA+GBC özellik seçiminden sonra veri artırma, LGB + XGB + XGB modeli ile elde edilen sonuçlar	63
Tablo 14. PCA+GBC özellik seçiminden önce veri artırma, LGB + XGB + XGB modeli ile elde edilen sonuçlar	64
Tablo 15. PCA + GBC veri artırma olmadan, LGB + XGB + XGB modeli ile elde edilen sonuçlar	65
Tablo 16. PCA + GBC ile istifleme modellerinin sınıflandırma performansı	66
Tablo 17. Özellik seçimi uygulanmadan yalnızca veri artırma ile elde edilen model sonuçları.....	67
Tablo 18. EMOTHAW veri seti kullanılarak yapılan çalışma sonuçlarının karşılaştırılması.....	69

SEMBOLLER VE KISALTMALAR DİZİNİ

WHO	: Dünya Sağlık Örgütü
HOG	: Yönlendirilmiş Gradyan Histogramı
BDI	: Beck Depresyon Envanteri
DASS	: Depresyon, anksiyete, stres ölçeği
FFT	: Hızlı Fourier Dönüşümü
IFT	: Ters Fourier Dönüşümü
DC	: Doğru Akım
(μ)	: Ortalama
(S)	: Çarpıklık
(K)	: Basıklık
(σ)	: Standart sapma

1. GENEL BİLGİLER

1.1. Giriş

Duygular, bireylerin yaşamlarında merkezi bir role sahiptir; belirli bir durum karşısında nasıl tepki vereceklerini yönetir ve diğer insanlarla olan etkileşimlerini şekillendirir. Schlenker & Leary'ye (1982) göre, üzüntü, öfke veya ilgisizlik gibi duygular sürekli bir şekilde bireyde mevcutsa, bu durumlar duygusal durumlar olarak sınıflandırılabilir. Bu tür duygusal durumlar, bireyin insanlarla olan etkileşimlerini, problem çözme becerisini ve yaşam dengesi üzerindeki kontrolünü olumsuz etkileyebilir. Özellikle kronik hastalıklarla ilişkili olarak depresyon, anksiyete ve stres gibi duygusal durumlar, önemli bir problem alanı oluşturmaktadır (Christensen vd., 2020; Turner vd., 2020).

Bu duygusal durumların tanınması ve yönetilmesi, bireylerin sağlığı ve refahı açısından kritik bir rol oynamaktadır. Makine öğrenmesi ve derin öğrenme alanlarındaki gelişmeler, el yazısı ve çizim gibi biyometrik verilerden duygusal durumların yüksek doğrulukla tespitine imkân sağlamıştır. El yazısı analizi, kişinin karakteristik özelliklerini ve duygu durumunu anlamak ve tahmin etmek amacıyla uzun yıllardır kullanılan bir yöntemdir.

Bu çalışmada, Likforman-Sulem vd., (2017) geliştirdiği EMOTHAW (EMOTion recognition from HAndWriting and DraWing) veri tabanı temel alınmıştır. Çalışmanın amacı, mevcut veri tabanı üzerinde farklı yaklaşımlar deneyerek daha iyi yaklaşımlar önerip sınıflandırma performansını artırmaktır. EMOTHAW veri tabanı, depresyon, anksiyete ve stres ölçeğiyle değerlendirilen 129 katılımcının verilerinden oluşmaktadır. Katılımcılar, yedi görevden oluşan bir testi tamamlamış ve bu görevler kapsamında ev çizimi, beşgen çizimi, kelime ve cümle yazımı gibi işlemler gerçekleştirmiştir. Bu testlerden elde edilen öznitelikler arasında kalemin durumu (havada, kâğıt üzerinde), basınç, kalem pozisyonu (x, y), azimut ve yükseklik bilgileri bulunmaktadır. Bu öznitelikler üzerinden türetilen yeni öznitelikler, makine öğrenimi yöntemi olan istifleme (stacking) modelleri kullanılarak analiz edilmiştir. Stacking yöntemi, farklı türdeki sınıflandırıcıların güçlü yönlerini birleştirerek daha genelleyici ve dengeli bir performans sunması nedeniyle tercih edilmiştir. Bu yaklaşım, yalnızca tek bir sınıflandırıcı kullanmak yerine, modeller arası hata örüntülerini azaltarak genel başarıyı artırmayı hedeflemiştir. Yapılan çalışma, Likforman-Sulem vd., (2017)'nin çalışmasına dayandırılmıştır. Bu çalışmada, sınıflandırma modelinde, öznitelik üretimi,

öznitelik seçimi ve boyut indirgeme süreçlerini genişletmek amaçlanmıştır. Bu çalışmanın teorik altyapısı, el yazısı ve çizim analizine dayanırken, veri artırma ve öznitelik işleme yöntemleriyle sonuçların genelleştirilebilirliğini artırmayı hedeflemektedir.

Bu tez çalışması beş bölümden oluşmaktadır. 1. Bölüm 'de Literatür İncelemesi, literatürde yer alan farklı yöntemlerin, algoritmaların ve süreçlerin detaylı bir incelemesini, DASS ölçeğini, veri ön işlemeyi (özellik çıkarımı, özellik seçimi teknikleri), veri Artırma, dengesiz veri tabanlarının ne olduğu ve bu dağılım sorununu nasıl çözülebileceği ve makine öğrenimi modelleme ve doğrulama süreçlerine yönelik açıklamaları içerir. 2.Bölüm olan yapılan çalışmalar kısmında, EMOTHAW veri tabanının nasıl oluşturulduğunu, kimler tarafından ve hangi demografiden oluştuğunu detaylı bir şekilde açıklanır ve bu çalışmada izlenen metodoloji detaylı bir şekilde anlatılır. Bu süreç, ham özniteliklerden öznitelik çıkarma yöntemlerine, öznitelik seçim tekniklerine, dengeleme/veri artırma algoritmalarına ve nihayetinde makine öğrenimi sınıflandırma modeline ve doğrulama testlerine kadar tüm adımları açıklar. Daha sonra yapılan analizlerin tüm sonuçları, 3. Bölüm 'de sunulmakta olup, her biri kendi amaçlarına sahip sonuç özet tablolarını içerir. Son olarak, 4. ve 5. Bölümler sonuçları ve paylaşılan düşünceleri içerir.

1.2. Literatür Taraması

Likforman-Sulem vd. (2017), el yazısı ve çizim örneklerinden duygusal durumların tanınması amacıyla, Depresyon-Anksiyete-Stres Ölçeği (DASS) ile değerlendirilmiş 129 katılımcıya ait örneklerden oluşan halka açık EMOTHAW (EMOTion recognition from HAndWriting and draWing) veri tabanını sunmuştur. Bu çalışmada, zaman ve duktus (kalem hareketi) temelli öznitelikler çıkarılmış ve sınıflandırma süreci Rastgele Orman (Random Forest, RF) algoritması ile gerçekleştirilmiştir. Sonuçlar, anksiyete ve stres tanınmanın, depresyona kıyasla daha yüksek performansla yapıldığını göstermiştir. Çalışma, el yazısının dinamik bileşenlerinin psikolojik durumları yansıttığını vurgulaması açısından alandaki temel çalışmalardan biri olmuştur.

Chitlangia & Malathi (2019), el yazısından kişilik özelliklerini otomatik olarak tanımayı amaçladıkları çalışmalarında, 50 katılımcının dijital el yazısı örneklerinden oluşan bir veri seti kullanmışlardır. Görsel özellik çıkarımı için Histogram of Oriented Gradient (HOG) yöntemi uygulanmış ve bu öznitelikler Destek Vektör Makinesi (Support Vector Machine, SVM) sınıflandırıcısına girdi olarak verilmiştir. Modelin başarımı, hem %90

eğitim - %10 test ayrımıyla hem de Leave-One-Out (LOO) çapraz doğrulama ile test edilmiştir. Beş farklı kişilik özelliğinin tanımlanmasında her iki yaklaşımda da %80 doğruluk elde edilmiştir. Ayrıca, duygusal durum yerine kişilik analizi yapılması açısından da farklı bir hedef taşımaktadır.

Ayzeren vd. (2019), duygusal durumların (mutlu, üzgün, stresli) tespiti amacıyla 134 katılımcının çevrimdışı ve çevrimiçi el yazısı ile imza biyometriklerinden oluşan yeni bir veri seti oluşturmuşlardır. Çalışmada, ham verilerden dinamik öznitelikler çıkarılmış ve bu öznitelikler frekans alanında analiz edilmiştir. Sınıflandırma için K-En Yakın Komşu (k-NN), JRip ve Rastgele Orman (RF) gibi farklı algoritmalar kullanılmıştır. Deneysel sonuçlar, özellikle stresin el yazısından yüksek doğrulukla tespit edilebildiğini göstermiştir. Bu çalışma, hem veri çeşitliliği hem de öznitelik türü açısından önceki çalışmalardan ayrılmaktadır. Likforman-Sulem vd. (2017) yalnızca çevrimiçi yazı verisini kullanırken, Ayzeren vd. çalışmasında çevrimdışı verilerle zenginleştirilmiş bir yaklaşım benimsenmiştir.

Cordasco vd. (2019) tarafından yapılan çalışmada, olumsuz ruh halleri (depresyon, stres ve anksiyete) yaşayan bireylerin el yazısı ve çizim öznitelikleri yaş ve cinsiyet açısından uyumlu bir kontrol grubu ile karşılaştırılmıştır. Karışık ANOVA analizleri sonucunda, gruplar arasında anlamlı farklar olduğu tespit edilmiş; bu farkların ilgili egzersiz türlerine ve öznitelik kategorilerine göre değiştiği ortaya konmuştur. Zaman ve frekans alanındaki özniteliklerin olumsuz duygusal durumların belirlenmesinde etkili olduğu görülmüştür. Bu çalışma, sınıflandırıcı algoritmalarından çok istatistiksel hipotez testleri kullanarak literatürde yöntemsel olarak farklılık göstermektedir. Ayrıca, kontrol grubu içeren yapı sayesinde elde edilen bulguların genelleştirilebilirliğini artırması bakımından önceki çalışmalardan ayrılmaktadır.

Nolazco-Flores vd. (2021) çalışmalarında, tablet üzerinde kaydedilen sinyallerden zaman, spektral ve kepsral alanlara ait öznitelikler birleştirilmiştir. EMOTHAW veri setindeki ham veriler kullanılarak; havada ve kâğıtta geçirilen zaman, görev süresi ve çeşitli dinamik özniteliklerin yanı sıra spektral ve kepsral alanlara ilişkin öznitelikler de çıkartılmıştır. Öznitelik seçimi için Fast Correlation-Based Filtering (FCBF) yöntemi uygulanmış ve azınlık sınıfların temsil gücünü artırmak amacıyla veri artırımı yapılmıştır. Sınıflandırmada Support Vector Machine (SVM) algoritması kullanılmış ve değerlendirme Leave-One-Out (LOO) stratejisiyle gerçekleştirilmiştir. Bu çalışma, önceki araştırmalardan

farklı olarak hem daha fazla öznitelik alanı birleştirmiş hem de veri dengesizliği sorununa doğrudan müdahale etmiştir.

Rahman & Halim (2022), el yazısı örneklerinden elde ettikleri on bir özniteliği graf tabanlı yazı temsili yöntemiyle çıkarmış ve bu öznitelikleri kullanarak kişilik tanıma süreci gerçekleştirmiştir. Sınıflandırma performansını artırmak amacıyla Yarı Denetimli Çekişmeli Üretici Ağ (Semi-supervised Generative Adversarial Network, SGAN) modeli kullanılmıştır. Deneysel sonuçlara göre, 173 katılımcının el yazısı özniteliklerinden %91,3 doğrulukla kişilik özellikleri sınıflandırılabilmiştir. Bu çalışma, klasik sınıflandırıcılardan farklı olarak derin öğrenme temelli bir yaklaşım sunmakta ve etiketli veri gereksinimini azaltan yarı denetimli yapı sayesinde öğrenme sürecini daha esnek hale getirmektedir. Ayrıca, Likforman-Sulem vd. (2017) ve Nolzco-Flores vd. (2021) gibi çalışmalar duygusal durum odaklıyken, bu çalışmanın kişilik özelliklerini hedeflemesi, literatürdeki uygulama alanlarının çeşitliliğini göstermektedir.

Nolzco-Flores vd. (2022) çalışmasında, EMOTHAW veri setindeki ham veriler zaman, kinematik, istatistiksel, spektral ve kepsral alanlardaki özniteliklere dönüştürülmüştür. Öznitelik seçimi sürecinde Temel Bileşenler Analizi (PCA) ve modifiye Fast Correlation-Based Filtering (mFCBF) yöntemleri bir arada kullanılmış, veri artırımı için ise Gaussian gürültüsü uygulanmıştır. Model eğitimi ve test süreci Otomatik Makine Öğrenimi (AutoML) çerçevesinde gerçekleştirilmiştir. İki sınıflı duygu şiddetinin tespitinde %100 doğruluk elde edilirken, üç sınıflı sınıflandırmada (normal,orta,aşırı) depresyon için %82,5, anksiyete için %72,8 ve stres için %74,56 doğruluk oranları elde edilmiştir. AutoML yaklaşımı sayesinde, model ve hiperparametre seçim süreci otomatikleştirilmiş ve insan müdahalesine ihtiyaç duyulmadan gerçekleştirilmiştir. Özellikle, AutoML yaklaşımı sayesinde model ve hiperparametre seçiminin insan müdahalesi olmadan optimize edilmesi, hem de öznitelik mühendisliği sürecinde çok katmanlı bir strateji izlemesi bakımından önceki çalışmalardan ayrılmaktadır.

Bhattacharya vd. (2022) tarafından yapılan çalışmada, el yazısı görüntülerinden elde edilen piksel verileri üzerinde aglomeratif hiyerarşik kümeleme yöntemi uygulanmıştır. Bu yöntem ile ön işleme aşamasından geçirilen görüntü pikselleri, benzerliklerine göre gruplandırılmış ve her bir kümenin belirli bir duygusal durumu temsil etmesi sağlanmıştır. Model, öfke, üzüntü, depresyon, mutluluk ve heyecan olmak üzere beş temel duygu sınıfı üzerinde test edilmiş ve %75'in üzerinde doğruluk oranı elde edilmiştir. Bu çalışma, önceki araştırmalardan farklı olarak doğrudan ham görüntü verileriyle çalışması ve piksel

düzeyinde bir yaklaşımla duygu tanıma gerçekleştirmesi bakımından özgündür. Öznitelik çıkarımı yerine görüntülerin kümeleme ile doğrudan temsil edilmesi, özellikle geleneksel sınıflandırma temelli yaklaşımlardan metodolojik olarak ayrılmaktadır.

Greco vd. (2023) çalışmalarında, sağlıklı katılımcılar (n=28) ve depresyon tanısı konmuş hastalar (n=27) arasında el yazısı ve çizim öznitelikleri kullanılarak el yazısı ve çizim performanslarının dinamik bir değerlendirmesi yapılmıştır. Katılımcılardan elde edilen veriler istatistiksel olarak analiz edilmiştir. Çalışmanın bulguları, kâğıt üzerindeki baskı kuvvetinin gruplar arasında ayırım yapamadığını, ancak diğer tüm özniteliklerin depresif ve depresif olmayan bireyleri ayırt etmede başarılı olduğunu göstermiştir. Bu yönüyle çalışma, Cordasco vd. (2019) gibi kontrol grubu temelli istatistiksel karşılaştırmalar yapan yaklaşımlar arasında yer almakta; ancak odak noktasını sadece depresyon tanısı ve dinamik özniteliklerle sınırlandırması açısından daha derinlemesine bir inceleme sunmaktadır.

Rahman & Halim (2023) çalışmalarında, el yazısı ve çizim örneklerinden elde edilen sinyaller zaman, spektral ve Mel Frekans Kepstral Katsayıları (MFCC) yöntemleriyle analiz edilmiştir. Elde edilen öznitelik vektörleri, çift yönlü Uzun Kısa Süreli Bellek (Bidirectional Long Short-Term Memory, BiLSTM) ağı ile sınıflandırılmıştır. Model, birden fazla halka açık veri kümesi üzerinde test edilmiş ve özniteliklerin birleştirilmesinin tanıma başarısını anlamlı şekilde artırdığı gözlemlenmiştir. Bu çalışma, derin öğrenme tabanlı sıralı veri işleme yeteneği sunan BiLSTM yapısı sayesinde zamana bağlı dinamik değişimlerin modellenmesini sağlamış, bu yönüyle geleneksel makine öğrenimi tabanlı çalışmalardan ayrılmıştır. Ayrıca, modelin birden fazla veri kümesi üzerinde değerlendirilmesi, genelleştirilebilirliğine dair güçlü bir kanıt sunmakta ve bu yönüyle çalışmayı literatürdeki benzerlerinden ayırmaktadır.

Khan vd. (2024), el yazısı ve çizim verilerinden elde edilen öznitelikleri kullanarak dikkat tabanlı bir dönüştürücü (transformer) model uygulamış ve EMOTHAW veri seti üzerinde %92,64 doğruluk oranı elde etmiştir. Bu çalışma, BiLSTM yapısını kullanan Rahman & Halim (2023) çalışmasıyla benzerlik göstermekte; ancak transformer mimarisi, daha güncel ve güçlü bir temsil öğrenme yeteneği sunmaktadır. Elde edilen yüksek doğruluk oranı, modern derin öğrenme yaklaşımlarının el yazısı ve çizim verilerine dayalı duygusal durum analizinde etkili bir alternatif olabileceğini ortaya koymaktadır.

1.3. Duygular

1.3.1. Depresyon

Depresyon ya da majör depresif bozukluk, Dünya Sağlık Örgütü'ne (WHO) göre, derin üzüntü ve ilgi kaybı gibi belirtilerle kendini gösteren yaygın bir ruhsal sağlık sorunudur. Bu durum, sıklıkla yanlış teşhis edilmesi ve diğer rahatsızlıklarla karıştırılması nedeniyle karmaşık bir yapıya sahiptir (Beck & Beamesderfer, 1974). Depresyon, bireylerin daha önce zevk aldıkları etkinliklere karşı ilgilerini yitirmelerine neden olur; aynı zamanda enerji kaybı, isteksizlik ve yoğun üzüntü duygularının zihinsel kontrolü ele geçirmesiyle vücutta ciddi bir yorgunluk hissi yaratır. Ağır vakalarda bu durum intihar ile sonuçlanabilir (Beck & Beamesderfer, 1974). Beck, depresyonun tanı konulmasındaki zorlukları fark etmiş ve bu zorlukları aşmak için Beck Depresyon Envanteri'ni (BDI) geliştirmiştir. Bu araç, bireylerin depresyon düzeylerini değerlendiren 21 sorudan oluşan, öz-yönetimli bir değerlendirme ölçeğidir (Beck & Beamesderfer, 1974). Depresyonun tanımlanması sürecindeki zorluklar, EMOTHAW çalışmasında da öne çıkmıştır. Bu çalışma depresyonun, özellikle düşük motivasyon ve özsaygı eksikliği gibi belirtiler üzerinden değerlendirildiğini göstermiştir. Araştırmada kullanılan DAS Ölçeği, depresyonun bu özelliklerini analiz etmek için geliştirilmiştir (Likforman-Sulem vd., 2017).

1.3.2. Anksiyete

Schlenker ve Leary, anksiyeteyi fizyolojik uyarılma (sinir sisteminin aktivasyonu) ve bireyin olumsuz bir sonucu yakın zamanda gerçekleşecekmiş gibi algılaması nedeniyle hissettiği endişe olarak tanımlamaktadır (Schlenker & Leary, 1982). Anksiyetenin tanımlanmasında temel belirtiler arasında terleme, çarpıntı, hızlı kalp atışı (taşikardi) ve hızlı nefes alıp verme gibi semptomlar yer alır. Bu belirtiler, sağlık uzmanlarının hastalarda anksiyeteyi doğru şekilde teşhis etmesine olanak tanır (Beck & Beamesderfer, 1974; Spielberg, 2010). Spielberg, bu bağlamda Freud'un, anksiyetenin psikiyatrik bozukluklar arasında yaygın olduğuna dair gözlemlerine dikkat çekmektedir (Spielberg, 2010). Anksiyete, bireyin sürekli olarak hem zihinsel hem de fiziksel olarak tetikte kalmasına neden olur. Bu durum, enerji tüketimini artırır ve bireyin yorgun, huzursuz ve tahammülsüz hissetmesine yol açar (Spielberg, 2010). Costello & Comrey (1967) ile Beck vd., (1988), anksiyete ve depresyonu doğru şekilde ölçmek amacıyla ölçekler geliştirmiştir. Costello ve

Comrey'in geliřtirdiđi ölçekler, mevcut ölçeklerle karşılaştırıldığında 0,50 düzeyinde bir korelasyon göstermiştir. Beck'in geliřtirdiđi Beck Anksiyete Envanteri (BAI) ise Beck Depresyon Envanteri (BDI) ile benzer bir korelasyona sahiptir ve her iki ölçek de 21 maddeden oluşmaktadır (Beck vd., 1988). Anksiyete durumunun değerlendirilmesinde DAS ölçeğinde korku ve algılanan panik gibi unsurlar temel alınmıştır (Likforman-Sulem vd., 2017).

1.3.3. Stres

Hans Selye, stres kavramını, vücudun herhangi bir talebe karşı verdiđi belirli olmayan bir tepki olarak tanımlamaktadır (Selye, 1956). Bununla birlikte, stres aynı zamanda öngörülebilir biyokimyasal, fizyolojik ve davranışsal deđişikliklerle birlikte görülen duygusal bir durum olarak da ifade edilmektedir (Baum, 1990). Tıp dünyasında stres, günümüzde birçok hastalığın temel nedenlerinden biri olarak kabul edilmektedir. Diđer iki duygusal durumla kıyaslandığında, stres doğal bir süreçtir ve bireylerin günlük yaşamlarında karmaşık durumlarla karşılaştıklarında hissettikleri bir tepki olarak ortaya çıkar. Selye'ye göre, stres yalnızca insanlarla sınırlı olmayıp hayvanlar, bakteriler ve hatta bitkilerde bile gözlemlenebilir (Selye, 1956). Bununla birlikte, stresin hissedilme sıklığının zamanla deđiřtiđi ve arttıđı belirtilmiştir. Bu artış, vücut tarafından salgılanan stres hormonlarının aşırı düzeylere ulaşarak diyabet ve kardiyovasküler hastalıklar gibi rahatsızlıklara neden olmasına yol açmıştır (Christensen vd., 2020). Bu durum, stresin bir zihinsel bozukluk olarak başlayıp zamanla vücut geneline yayılan sistemik bir hastalık haline gelebileceđini göstermektedir. Bu nedenle, stres, üç temel duygu durumu arasında en zoru olarak değerlendirilmektedir. EMOTHAW çalışmasında ise, stresi tanımlamak için gerginlik ve sinirlilik gibi unsurlar temel alınabilir (Likforman-Sulem vd., 2017).

1.3.4. DAS Ölçeđi

Depresyon, anksiyete ve stres gibi üç temel duygusal durumu ölçmek için geliřtirilmiş olan DAS (Depresyon, Anksiyete, Stres) Ölçeđi, psikometrik değerlendirme alanında önemli bir araçtır. Ölçek, 1995 yılında Syd Lovibond ve Peter Lovibond tarafından New South Wales Üniversitesi'nde tasarlanmıştır (Lovibond & Lovibond, 1995). Bu ölçek, bireylerin semptomlarını ölçmek için 1 ile 4 arasında derecelendirilen 42 maddeden oluşur ve her bir duygusal durum için ayrı ayrı analiz yapılmasına olanak tanır.

Tablo 1. Duygusal duruma göre DAS puan aralıkları (Likforman-Sulem vd., 2017)

Duygusal Durum Düzeyi	Depresyon	Anksiyete	Stres
Normal	0-9	0-7	0-14
Hafif	10-13	8-9	15-18
Orta	14-20	10-14	19-25
Şiddetli	21-27	15-19	26-33
Aşırı Şiddetli	28+	20+	34+

Tablo 1, ölçeğin depresyon, anksiyete ve stres durumlarına göre beş farklı şiddet düzeyini göstermektedir. Ancak, çalışmanın kapsamı gereği bu şiddet düzeylerinin tamamının detaylı bir incelemesi yapılmamıştır. DAS Ölçeği, EMOTHAW metodolojisi bağlamında, her bir görevi bireylerin verdiği anket yanıtlarına göre sınıflandırmak için kullanılmıştır. Bu ölçek, 129 katılımcı üzerinde uygulanmış ve Lovibond'un geliştirdiği aracın, Crawford ve Henry tarafından 2003 yılında İngilizce dilinde normlara uygun olarak değerlendirilmesi sağlanmıştır (Crawford & Henry, 2003). Tablo 1'de yer alan tam kapsamlı aralıklar, çalışmanın süresi ve kapsamı içinde ele alınamayacak bir karmaşıklık oluşturduğu için analiz dışında bırakılmıştır.

Tablo 2. İkili Etiketleme

	Depresyon	Anksiyete	Stres
Normal (0)	0-9	0-7	0-14
Normal Üstü (1)	10-28+	8-20+	15-34+

Nolazco-Flores vd. (2021), DASS puanlarını ruh hali durumlarını ikili olarak sınıflandırmak amacıyla kullanmıştır. İkili sınıflandırma, ölçeğin en basit hali olarak kabul edilmiş ve Tablo 2'de belirtildiği gibi bu yöntemle, puanlar 'normal' ve 'normal üstü' olarak sınıflandırılmıştır. Depresyon ölçeğinde 9 puanın üzerinde alan bir kişi depresyonda kabul edilirken, anksiyete ölçeğinde 7 ve üzeri puan alan kişi anksiyeteli olarak değerlendirilmiştir. Stres ölçeğinde ise 14 ve üzeri puan alan kişi stresli kabul edilmiştir. Bu çalışmada, DASS puanlarının ruh hali durumlarını sınıflandırmak için ikili etiketleme yöntemi benimsenmiştir.

X-koordinatı ve Y-koordinatı, kullanıcının kalem hareketinin kâğıt üzerindeki yatay ve dikey konumunu belirler.

1.4. Veri Ön İşleme

1.4.1. EMOTHAW Veri Tabanındaki Ham Özellikler

EMOTHAW veri tabanı, el yazısı ve çizim verilerini analiz etmek için yedi başlangıç ham sensör verileri içermektedir. Bu veriler, kullanıcıların el hareketlerini detaylı bir şekilde inceleyebilmek için çok boyutlu bilgiler sunar. X-koordinatı ve Y-koordinatı, kullanıcının kalem hareketinin kağıt üzerindeki yatay ve dikey eksenlerdeki konumunu belirler. Zaman Damgası, her veri noktasının kaydedildiği zamanı gösterir, böylece hareketin zamansal ilerleyişi takip edilebilir. Kalem pozisyonu, kalemin havada mı yoksa kâğıt üzerinde mi olduğunu belirten önemli bir özelliktir. Azimut açısı, kalemin yatay düzlemdeki yönünü ifade ederken yükseklik ise, kalemin yüzeye olan yüksekliğini gösterir. Son olarak basınç yüzeye olan basınç miktarını temsil eder. Bu ham nitelikler, zaman, kinematik, istatistiksel, spektral ve kepsral özellikleri üretmek için kullanılmıştır.

1796								
49076	34584	17606448	1	1870	560	45		
49025	34608	17606456	1	1870	560	81		
49009	34613	17606463	1	1870	560	157		
48995	34614	17606478	1	1870	560	193		
48993	34614	17606486	1	1870	560	219		
48993	34614	17606493	1	1860	560	246		
48993	34614	17606501	1	1860	550	284		
...	...	...	...	...	...	...		
50786	33795	17606756	1	1900	550	305		
50727	33808	17606764	1	1900	540	130		
50727	33808	17606771	0	1900	540	0		
50640	33840	17606779	0	1900	540	0		
50621	33860	17606786	0	1900	540	0		
50619	33878	17606794	0	1900	540	0		
...	...	...	...	...	...	...		
51032	33781	17607320	0	1940	510	0		
51032	33781	17607328	1	1940	510	84		
51056	33773	17607336	1	1940	510	118		
...	...	...	...	...	...	...		

Şekil 1. Ham veriler için örnek bir SVC dosyası (Likforman-Sulem vd., 2017)

1.4.2. Öznitelik Türetme

Özellik Çıkartımı (Feature Extraction), bir veri setindeki ham bilgilerden yeni ve anlamlı nitelikler elde etmeye odaklanan önemli bir makine öğrenimi tekniğidir (Khalid vd., 2014). Bu teknik, büyük veri setlerini minimum bilgi kaybıyla açıklama yöntemi olarak kullanıldığında, genellikle özellik seçimi ve boyut indirgeme yaklaşımlarıyla ilişkilendirilir.

Özellik çıkartımı sürecinde, mevcut niteliklere istatistiksel ölçümlerin uygulanması yoluyla yeni ve anlamlı bir nitelik seti oluşturulması yaygın bir yöntemdir (Guyon vd., 2008). Bu çalışmada elde edilecek özellikler: Zaman özellikleri (Likforman-Sulem vd., 2017), kinematik özellikler (Impedovo, 2019), istatistiksel özellikler (Drotár vd., 2014), spektral ve kepsral özellikler (Nolazco-Flores vd., 2021). Bu özelliklerin her biri, sonuçta elde edilecek özellik vektörüne anlamlı bilgiler ekler.

1.4.3. Zaman Tabanlı Özellikler

Zaman tabanlı özellik vektörü, Likforman vd., (2017) tarafından önerilmiştir. Zamansal öznitelikler, kalemin hava ve kâğıt üzerindeki hareketleri ile ilgilidir. Bu analizde, kalemin havada kaldığı sürenin toplamı, kâğıt üzerinde hareket ettiği sürenin toplamı, görevin başlangıç ve bitiş zaman damgaları arasında geçen toplam süre ve kalemin kâğıt ile hava arasındaki toplam geçiş sayısı zamansal öznitelikler olarak hesaplanmıştır.

Tablo 3. Kalem hareketlerinin zamansal dinamiklerini gösteren özellikler.

Notasyon	Tanım
T_{havada}	Kalemin havada kaldığı toplam süre
$T_{kâğıtta}$	Kalemin kâğıt üzerinde kaldığı toplam süre
T_{Toplam}	Görev boyunca geçen toplam süre
N_{st}	Kalemin kâğıt ve hava arasında yaptığı geçişlerin toplam sayısı

Tablo 3'te, kalem hareketlerine ilişkin zamansal özellikler detaylandırılmaktadır. Zamansal özellikler, kullanıcıların yazma dinamiklerini anlamada kritik öneme sahiptir, çünkü bu özellikler hem motor becerileri hem de yazma sürecindeki duraklamalar hakkında önemli bilgiler sağlayabilir.

$$T_{air} = \sum_{i=1}^{N-1} (t_{i+1} - t_i) \quad \text{öyle ki } S_i = 0 \quad (1)$$

Burada t_i , i numaralı zaman damgasını, S_i ise kalemin durumunu (0: havada, 1: kâğıtta) temsil eder. Yani, kalemin havada olduğu anlar arasında zaman farklarının toplamı alınır.

$$T_{paper} = \sum_{i=1}^{M-1} (t_{i+1} - t_i) \quad \text{burada } S_i = 1 \quad (2)$$

Benzer şekilde, t_i zaman damgalarının toplamı, kalemin kâğıtta olduğu ($S_i = 1$) durumlar arasında hesaplanır.

$$T_{total} = t_N - t_1 \quad (3)$$

Buradaki formül, ilk ve son zaman damgası arasındaki toplam süreyi hesaplar. t_N Son zaman damgasını, t_1 ise ilk zaman damgasını temsil eder.

$$N_{st} = \sum_{i=1}^{N-1} \delta(S_i, S_{i+1}) \quad \text{burada } \delta(S_i, S_{i+1}) = \begin{cases} 1, & \text{eğer } S_i = 1 \text{ ve } S_{i+1} = 0 \\ 0, & \text{değilse} \end{cases} \quad (4)$$

Formül, $\delta(S_i, S_{i+1})$ ifadesi, kalemin kâğıttan havaya (veya tersi) geçtiği durumların sayısını hesaplar.

1.4.4. Frekans ve Kepstral Özellikler

Frekans ve kepsral özelliklerin çıkarımı, sinyalin zaman alanından frekans ve spektral alana dönüştürülmesiyle başlar. Bu süreç, sinyal analizi ve özellik mühendisliği açısından önemlidir. Bu çalışmada, sinyallerin frekans alanındaki temel özniteliklerini incelemek ve anlamlı veriler elde etmek amacıyla frekans ve kepsral özniteliklerin temel hesaplamaları yapılmıştır.

Frekans Özellikleri: Frekans ile ilgili özniteliklerin çıkarımı için, kalemin X ve Y eksenindeki konumu ile yüzeye uyguladığı basıncı temsil eden ham sinyal verileri işlenmiştir. Her bir sinyalin Hızlı Fourier Dönüşümü (FFT) alınarak frekans spektrumu elde edilmiştir. Bu dönüşüm, sinyallerin frekans bileşenlerini analiz etmeye olanak tanır. Daha sonra, frekans spektrumunun genlik değerleri hesaplanmış ve analiz sonuçlarını

etkileyebilecek sabit ofset bileşeni olan DC bileşeni (ilk eleman) çıkarılmıştır. DC bileşeninin çıkarılmasının temel sebebi, sinyaldeki sıfır frekanslı sabit bileşenlerin dinamik analizi gölgeleyebilmesidir. Frekans spektrumunun en baskın frekansı, genlik değerlerinin en yüksek olduğu nokta olarak tespit edilmiş ve sinyal enerjisi hesaplanmıştır. Enerji, sinyalin genel frekans yoğunluğunu temsil eder ve sinyalin genel frekans dağılımının bir ölçüsü olarak kullanılmıştır. Bu adımlar sonucunda, her bir sütun için en baskın frekans ve enerji değerleri belirlenmiştir.

Kepstral Özellikler: Sinyallerin daha derinlemesine analiz edilebilmesi amacıyla, her bir sinyalin logaritmik spektrumu ($\log|FFT|$) hesaplanmış ve bu spektrum üzerinden kepsral özellikler çıkarılmıştır. Elde edilen logaritmik spektruma Ters Fourier dönüşümü (IFT) uygulanarak kepsrum elde edilmiştir. kepsrum, sinyalin spektral yapılarını analiz etmek için kullanılır ve bu süreç, frekans alanındaki düzenliliklerin daha ayrıntılı şekilde incelenmesine olanak tanır. Kepstral özelliklerin belirlenmesinde, kepsrumun ortalama, standart sapma ve maksimum değerleri hesaplanmıştır.

Bu süreç hem frekans hem de kepsral özelliklerin detaylı bir şekilde çıkarılmasını sağlayarak sinyallerin karakteristik yapısını anlamaya katkıda bulunur. Frekans analizi, sinyalin güçlü bileşenlerini tespit ederken, kepsral analiz sinyalin daha ince yapısal düzenliliklerini ortaya çıkararak ayırt edici özellikler sunmaktadır.

Tablo 4. Frekans analizinde kullanılan notasyonlar ve tanımları

Notasyon	Tanım
f_{dom}^X	X eksenini için baskın frekans
f_{dom}^Y	Y eksenini için baskın frekans
f_{dom}^P	Basınç verisi için baskın frekans
E_{freq}^X	X eksenini için frekans enerjisi
E_{freq}^Y	Y eksenini için frekans enerjisi
E_{freq}^P	Basınç verisi için frekans enerjisi

Tablo 4'te, çıkarılan frekans özelliklerinin notasyonları ve tanımları verilmiştir. Bu notasyonlar, X ve Y eksenleri ile kalem basıncına ait baskın frekans (f_{dom}) ve frekans enerjisi (E_{freq}) değerlerini ifade etmektedir. Her bir notasyon, ilgili özelliğin hangi sinyal bileşenine ait olduğunu netleştirmek amacıyla sistematik olarak tanımlanmıştır. Bu değerler,

sinyallerin karakteristik yapılarını anlamak ve analiz etmek için temel metrikler olarak kullanılmıştır.

$$X(k) = \sum_{n=0}^{N-1} x(n) \cdot e^{-j\frac{2\pi}{N}kn} \quad (5)$$

Burada sinyal $x(n)$ 'in Fourier dönüşümü $X(k)$ olarak tanımlanır; bu dönüşüm, sinyali zaman domaininden frekans domainine taşıyarak sinyalin farklı frekanslardaki faz ve genlik bileşenlerini belirlemeyi sağlar.

- $X(k)$: Frekans alanındaki değerler.
- N Sinyalin uzunluğunu temsil eder.
- n : Zaman örneği
- k Frekans bileşeninin indisidir.
- j Sanal birimdir.

$$|X(k)| = \sqrt{R(X(k))^2 + I(X(k))^2} \quad (6)$$

Fourier dönüşümünden elde edilen frekans bileşenlerinin genlik ($|X(k)|$) değerleri hesaplanıyor. Burada $R(X(k))$, Fourier dönüşümünün reel (gerçek) kısmını; $I(X(k))$ ise imajiner (sanal) kısmını temsil etmektedir.

$$\text{Frekans_Genliği} = |X(k)| \quad \text{for } k = 1, 2, \dots, N - 1 \quad (7)$$

İlk eleman ($k = 0$, yani DC bileşeni) çıkarılıyor çünkü bu, sinyalin statik ofsetini temsil eder ve analiz için genellikle gereksizdir.

$$\text{Baskın_Frekans} = \arg \max_k (|X(k)|) \quad (8)$$

En baskın frekans, genlik spektrumunda en büyük değere sahip olan k frekans bileşeninin indisidir. Bu, sinyaldeki en baskın frekansı temsil eder.

$$\text{Frekans_Enerjisi} = \frac{1}{N-1} \sum_{k=1}^{N-1} |X(k)|^2 \quad (9)$$

Frekans enerjisi, sinyalin frekans spektrumundaki enerji yoğunluğunu veya dağılımını hesaplamak için kullanılır. Bu metrik, sinyalin toplam gücünün bir ölçüsü olarak, farklı frekans bileşenlerinin sinyale katkılarını değerlendirme imkânı sağlar. Böylece, sinyalin frekans spektrumunda hangi bileşenlerin baskın olduğunu ve enerjinin nasıl dağıldığını anlamaya yardımcı olur.

Tablo 5. Kepstral analizde kullanılan notasyonlar ve tanımları

Notasyon	Tanım
C_{ort}^X	X eksenini için kepsral ortalama.
C_{std}^X	X eksenini için kepsral standart sapma.
C_{max}^X	X eksenini için kepsral maksimum değeri.
C_{ort}^Y	Y eksenini için kepsral ortalama.
C_{std}^Y	Y eksenini için kepsral standart sapma.
C_{max}^Y	Y eksenini için kepsral maksimum değeri.
C_{ort}^P	Basınç verisi için kepsral ortalama.
C_{std}^P	Basınç verisi için kepsral standart sapma.
C_{max}^P	Basınç verisi için kepsral maksimum değeri.

Tablo 5'te, çıkarılan kepsral özelliklerin notasyonları ve tanımları verilmiştir. Her bir notasyon, X ve Y eksenleri ile basınç verisi üzerinde hesaplanan kepsral ortalama (C_{mean}), standart sapma (C_{std}) ve maksimum değeri (C_{max}) gibi temel özellikleri ifade etmektedir. Bu özellikler, sinyalin spektral yapısındaki düzenliliklerin daha detaylı incelenmesine olanak sağlayarak, duygu durumu tanımlama süreçlerinde ayırt edici bir temel oluşturur. Tablo 5'te kepsral özelliklerin adım adım nasıl çıkarıldığını gösteren matematiksel formülleri:

$$\text{Log_Spektrum} = \log(|X(k)|) \quad (10)$$

Frekans spektrumunun genlik değerlerinin logaritması alınır. Bu işlem, büyük değerleri sıkıştırarak küçük değerleri daha görünür hâle getirir.

$$c(n) = \mathcal{F}^{-1}(\log(|X(k)|)) \quad (11)$$

Logaritmik spektrumun ters fourier dönüşümü uygulanarak kepsrum ($c(n)$) elde edilir. Kepsrum, sinyalin spektral düzenliliklerini analiz etmek için kullanılır ve $c(n)$, kepsral katsayıları temsil eder. Bu katsayılar, zaman alanındaki bileşenlerin düzenliliklerini inceleme imkânı sağlar. Kepsral analiz sonucunda, sinyalin karakteristik özelliklerini tanımlamak için ortalama, standart sapma ve maksimum değer gibi metrikler çıkarılmıştır. Ortalama, sinyalin genel düzenini; standart sapma, kepsrumun yayılımını ve varyasyonunu; maksimum değer ise sinyalin en güçlü düzenlilik bileşenini ifade eder.

Ortalama $\mu_{kepsrum}$

$$\mu_{kepsrum} = \frac{1}{N} \sum_{n=0}^{N-1} c(n) \quad (12)$$

Standart Sapma $\sigma_{kepsrum}$

$$\sigma_{kepsrum} = \sqrt{\frac{1}{N} \sum_{n=0}^{N-1} (c(n) - \mu_{kepsrum})^2} \quad (13)$$

Maksimum Değer $\max(c(n))$

$$\text{Max_kepsrum} = \max(c(n)) \quad (14)$$

1.4.5. Kinematik Özellikler

Kinematik özellikler, hareketin temel dinamiklerini anlamak ve analiz etmek için kullanılan önemli metriklerdir. Bu çalışma kapsamında kinematik analiz, çizim ve yazım verilerinden hareketin mekânsal (X, Y pozisyonları) ve zamansal (zaman damgaları) ham verileri ayrıştırılmış ve bu verilerden yer değiştirme, hız, ivme ve sarsıntı (jerk) olmak üzere dört kinematik özellik çıkarılmıştır. İlk olarak, iki boyutlu pozisyon verileri (X, Y) kullanılarak, ardışık konumlar arasındaki Öklid mesafeye dayalı yer değiştirme hesaplanmıştır. Yer değiştirme, hareketin uzayda ne kadar yol aldığını ifade eder ve şu şekilde formüle edilir:

$$\Delta d_i = \sqrt{(x_{i+1} - x_i)^2 + (y_{i+1} - y_i)^2} \quad (15)$$

Burada, yer değiştirme (Δd_i), iki ardışık konumun (x_i, y_i), oluşturduğu pozisyon çiftleri arasındaki mesafeyi ifade eder. X_i ve y_i , i . adımda X ve Y pozisyonlarını temsil eder. Elde edilen yer değiştirme değerleri ve zaman farkları temel alınarak hız bilgisi türetilmiştir:

$$v_i = \frac{\Delta d_i}{t_{i+1} - t_i} \quad (16)$$

Hız (v_i), veri noktaları arası yer değiştirme (Δd_i) ile bu yer değiştirmenin gerçekleştiği zaman farkı ($\Delta t = t_{i+1} - t_i$) oranıdır. Bu, hareketin ne kadar hızlı gerçekleştiğini gösterir.

$$a_i = \frac{v_{i+1} - v_i}{t_{i+2} - t_{i+1}} \quad (17)$$

Bir sonraki aşamada, hız değerlerinin türevi alınarak ivme hesaplanmıştır. İvme (a_i), hızdaki değişim oranıdır. Başka bir deyişle, hızın zamanla değişim hızını ifade eder.

$$j_i = \frac{a_{i+1} - a_i}{t_{i+3} - t_{i+2}} \quad (18)$$

Son olarak, ivme değerlerinin ikinci türevi alınarak sarsıntı bilgisi elde edilmiştir. Sarsıntı (j_i), ivmedeki değişim oranıdır. Hareketin ne kadar ani ve düzensiz olduğunu gösterir. Özellikle, yüksek sarsıntı değerleri hareketin düzensizliğini ifade eder.

Tablo 6. Kinematik özellikler için notasyonlar ve tanımları

Notasyon	Tanım
d_{ort}	Yer değiştirme için ortalama değer.
d_{max}	Yer değiştirme için maksimum değer.
v_{ort}	Hız için ortalama değer.
v_{max}	Hız için maksimum değer.
a_{ort}	İvme için ortalama değer.
a_{max}	İvme için maksimum değer
j_{ort}	Sarsıntı için ortalama değer
j_{std}	Sarsıntı için standart sapma.

Tablo 6, yer değiştirme, hız, ivme ve sarsıntı gibi kinematik özelliklere uygulanan istatistiksel işlemleri göstermektedir. İlk olarak, ham verilerden sırasıyla yer değiştirme, hız, ivme ve sarsıntı hesaplanmıştır. Daha sonra bu temel kinematik özellikler üzerine ortalama, maksimum ve standart sapma gibi istatistiksel ölçümler uygulanmıştır. Bu istatistiksel özet, bireylerin motor davranışlarının nicel bir şekilde ifade edilmesine ve modelleme süreçlerinde girdi olarak kullanılmasına olanak sağlamıştır.

Bu süreç, kinematik özelliklerin zamana bağlı değişimlerinin incelenmesi açısından önemli bir adım olarak değerlendirilmektedir. Hız, ivme ve sarsıntı gibi kinematik parametreler, bireylerin çizim ve yazım sırasında gerçekleştirdiği motor hareketlerin ve davranışların detaylı bir şekilde analiz edilmesine olanak tanımaktadır.

1.4.6. İstatistiksel Özellikler

İstatistiksel özellikler, veri kümesinin temel yapısını anlamak ve analiz etmek için kullanılan önemli metriklerdir. Bu çalışma kapsamında, ham verilerin merkezi eğilimini, yayılımını, asimetrisini ve uç değerlerini değerlendirmek amacıyla ortalama (μ), standart

sapma, (σ) medyan (Med), çeyrekler (Q_{75}, Q_{25}), çarpıklık (S) ve basıklık (K) gibi temel istatistiksel özellikler çıkarılmıştır. Aşağıdaki tablo, bu metriklerin notasyon ve tanımlarını özetlemektedir:

Tablo 7. İstatistiksel özellikler için notasyonlar ve tanımları

Notasyon	Tanım
μ	Verilerin ortalama değeri
σ	Verilerin standart sapması
Med	Verilerin ortanca değeri
K	Verilerin basıklığı (uç değerlerinin varlığı)
S	Verilerin Çarpıklığı (asimetri derecesi)
Q_{25}	Verilerin 25. yüzdalık dilimi
Q_{75}	Verilerin 75. yüzdalık dilimi
x_{max}	Verilerin maksimum değeri

Ortalama, tüm veri noktalarının aritmetik ortalaması olarak tanımlanır ve aşağıdaki formülle hesaplanır:

$$\mu = \frac{1}{N} \sum_{i=1}^N x_i \quad (19)$$

Ortalama, veri kümesinin merkezi eğilimini ölçmek için kullanılır ve dağılımın genel merkezini temsil eder.

$$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2} \quad (20)$$

Standart sapma, veri noktalarının ortalama etrafındaki yayılımını ölçer ve verilerin ne kadar dağınık olduğunu anlamak için kullanılır.

Medyan, veri kümesinin sıralandıktan sonraki ortanca değeridir. Medyan, uç değerlerden etkilenmediği için özellikle asimetrik dağılımlar için önemlidir.

$$S = \frac{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^3}{\sigma^3} \quad (21)$$

Çarpıklık, veri dağılımının simetrik olup olmadığını gösterir. Burada, $(x_i - \mu)^3$, verinin simetri yapısını ölçen küp sapmasıdır. Pozitif çarpıklık ($S > 0$) verilerin sağa eğimli, negatif çarpıklık ($S < 0$) ise sola eğimli olduğunu gösterir.

$$K = \frac{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^4}{\sigma^4} \quad (22)$$

Basıklık, veri dağılımındaki uç değerlerin yoğunluğunu ölçer. Burada, Yüksek basıklık ($K > 3$), uç değerlerin yoğun olduğunu, düşük basıklık ($K < 3$) ise uç değerlerin az olduğunu gösterir.

Çeyrekler, sıralı veri kümesinin yüzde 25'lik Q_{25} ve yüzde 75'lik Q_{75} noktalarını ifade eder ve veri dağılımının belirli yüzdeliklerini ölçerek daha detaylı bir yapısal analiz sunar. Maksimum değer ise veri kümesindeki en büyük değeri ifade eder ve uç değerleri analiz etmek ve dağılımın sınırlarını belirlemek için kullanılır.

Bu istatistiksel özellikler, bireyin yazma veya çizim hareketlerindeki düzenlilikleri ve dalgalanmaları analiz ederek, duygusal durumların dolaylı ipuçlarını sağlamaktadır. Örneğin, standart sapma ve çarpıklık gibi metrikler hareketlerin ne kadar tutarlı olduğunu gösterirken, basıklık ve maksimum değer gibi ölçütler uç davranışları değerlendirmemize yardımcı olur. Bu metrikler bir arada kullanılarak, bireyin stres, kaygı, huzursuzluk veya sakinlik gibi duygu durumlarını anlamak için güçlü bir temel oluşturulmaktadır.

1.5. Boyut İndirgeme

Günümüzde özellikle yapılandırılmamış veriler, artan karmaşıklıkları nedeniyle analiz süreçlerini zorlaştırmaktadır (Van der Maaten vd., 2009). Bu tür veriler genellikle modelleme veya tahmin süreçlerini olumsuz etkileyebilecek gürültülü ya da gereksiz öznitelikler içerebilir. Bu nedenle, boyut indirgeme teknikleri kullanılarak anlamlı bilgi

korunurken, gereksiz özelliklerin ortadan kaldırılması hedeflenir (Van der Maaten vd., 2009; Shaw & Jebara, 2009). Bu yöntemler sayesinde hesaplama maliyeti azaltılır, sınıflandırma ve model eğitim süreleri kısaltılır ve ayrıca veri desenlerinin ve yapılarının görselleştirilmesi ve kümeleme analizi gibi süreçlerde avantaj sağlanır (Van der Maaten vd., 2009). Boyut indirgeme yöntemleri genel olarak doğrusal ve doğrusal olmayan yöntemler olarak iki kategoriye ayrılır. Kullanılacak yöntemin seçimi, verinin türü ve çalışmanın amacına göre belirlenir. Geleneksel veri türleri doğrusal yöntemlerle daha iyi analiz edilirken, insan davranışı veya el yazısı gibi karmaşık örüntüler içeren verilerde doğrusal olmayan yöntemler daha başarılı sonuçlar verir (Van der Maaten vd., 2009). Bununla birlikte, Öznitelik çıkarımı ve öznitelik seçimi gibi yöntemler, boyut indirgeme süreçlerinin önemli unsurlarıdır (Shaw & Jebara, 2009). Öznitelik çıkarımı, veri içindeki gürültüyü azaltarak daha anlamlı öznitelikler üretirken, öznitelik seçimi ise modelleme başarısına doğrudan katkı sağlayan öznitelikleri belirlemeyi amaçlar (Shaw & Jebara, 2009). Boyut indirgeme teknikleri arasında yaygın olarak kullanılan Temel Bileşen Analizi (PCA), verileri düşük boyutlu ve anlamlı forma dönüştürerek analiz süreçlerini kolaylaştırmaktadır. PCA, verideki önemli bilgiyi koruyarak modelleme süreçlerini optimize eder ve veri yapısının anlaşılmasını kolaylaştıran kritik bir araç olarak kabul edilir (Cao vd., 2003). Bu tür yaklaşımlar, boyut indirgemeyi sadece veri işlemede bir araç değil, aynı zamanda veri yapısının anlaşılmasını kolaylaştıran kritik bir süreç haline getirir.

1.5.1. Temel Bileşen Analizi

Büyük veri setlerini anlamlandırmak için, içerdikleri bilgilerin büyük bir kısmını koruyarak boyutlarını önemli ölçüde ve anlaşılır bir şekilde azaltan yöntemlere ihtiyaç vardır. Bu amaçla birçok teknik geliştirilmiş olsa da en eski ve en yaygın kullanılan yöntemlerden biri Temel Bileşen Analizidir (PCA). PCA, veri boyutunu azaltarak veri setinin karmaşıklığını azaltmayı ve önemli bilgileri korumayı amaçlayan sıkça kullanılan bir boyut indirgeme tekniğidir (Jolliffe, 2002). Temel prensibi, bir veri setinin boyutunu azaltırken mümkün olduğunca fazla varyansı (istatistiksel bilgiyi) muhafaza etmektir. PCA, orijinal veri setinin varyansını korurken, veriyi ortogonal bileşenlere dönüştürerek yeni bir özellik uzayı oluşturur. Bu süreçte, verinin yüksek boyutlu yapısı, daha düşük boyutlu ve anlamlı bir forma indirgenir.

PCA'nın matematiksel işleyişinde ilk adım, veri setindeki her sütunun ortalamasının çıkarılmasıyla veriyi merkezlemektir. Bu işlem, aşağıdaki gibi tanımlanır:

$$X_{\text{merkezlenmiş}} = X - \bar{X} \quad (23)$$

Bu merkezlenmiş veri üzerinden, değişkenler arası ilişkiyi temsil eden kovaryans matrisi hesaplanır. Kovaryans matrisi, aşağıdaki gibi tanımlanır:

$$C = \frac{1}{n} \sum_{i=1}^n X X^T \quad (24)$$

Burada n , örnek sayısını; X , merkezlenmiş veri; X^T , merkezlenmiş verinin transpozunu ifade eder.

Kovaryans matrisi kullanılarak, verinin temel bileşenlerini belirlemek için özdeğerler (λ) ve özvektörler (v) hesaplanır:

$$Cv = \lambda v \quad (25)$$

Burada, her bir özdeğer (λ) varyansı temsil eder ve bu özdeğerlere karşılık gelen özvektörler (v) temel bileşenleri oluşturur. PCA, en yüksek varyansı açıklayan ilk temel bileşeni belirler ve ardından kalan bileşenleri, orijinal veri setine kıyasla varyansın azaldığı sıralamada hesaplar. Verinin yeni uzaya dönüştürülmesi şu şekilde gerçekleştirilir:

$$Z = X_{\text{merkezlenmiş}} W \quad (26)$$

Burada W , en büyük özdeğerlerle ilişkilendirilmiş özvektörlerden oluşan bir dönüşüm matrisidir. Maksimum bileşen sayısı, genellikle veri setindeki gözlem sayısı (N) ve özellik sayısı (M) arasındaki minimum değere eşittir ($\min(N, M)$). Ancak, tüm bileşenlerin modele dahil edilmesi verimli olmayabilir. Bunun yerine, yalnızca önemli varyansı açıklayan bileşenler seçilir. Bileşenlerin seçilmesinde açıklanan varyans oranı aşağıdaki formül ile hesaplanır:

$$\text{Açıklanan Varyans Oranı} = \frac{\lambda_i}{\sum_{j=1}^p \lambda_j} \quad (27)$$

Bu süreç, verinin yüksek boyutlu yapısını azaltarak, daha düşük boyutlu ve daha anlamlı bir forma indirgenmesini sağlar (Abdi & Williams, 2010; Partridge & Calvo, 1997).

1.6. Öznitelik Seçimi

Yeni nesil veri tabanlarında artan boyutlar, Özellik Seçimi tekniklerinin uygulanmasında zorluklar oluşturmuştur. Yüksek boyutlu veri setlerinin karmaşıklığı nedeniyle, geleneksel algoritmalar günümüz uygulamalarında gereken hesaplama verimliliğini sağlayamamaktadır (Yu & Liu, 2003). Bununla birlikte, bu teknikler hâlâ boyut indirgeme, gürültü azaltma ve model genelleştirilebilirliğini artırma gibi önemli faydalar sunmaktadır. Özellik Seçimi, eğitim sürelerini kısaltır, veri boyutunu küçültür ve makine öğrenimi modellerinin doğruluğunu artırır (Sammur & Webb, 2011). Özellik seçimi süreci, belirli kriterlere göre öznitelikleri sıralar ve bu kriterleri karşılamayanları elimine eder. Örneğin, korelasyona dayalı bir özellik seçimi tekniği, yalnızca 0.4 veya daha düşük korelasyon katsayısına sahip öznitelikleri seçebilir (Dash & Liu, 1997). Özellik Seçimi yöntemleri üç ana kategoriye ayrılmaktadır: filtre yöntemleri, sarma yöntemleri ve gömülü yöntemler (Guyon & Elisseeff, 2003). Bu tekniklerin performansı, bir algoritmanın ne kadar etkili olduğunu değerlendirmek için kullanılan metriklerden büyük ölçüde etkilenir.

1.7. Öznitelik Seçimi Yöntemleri

Öznitelik seçim yöntemleri; filtre yöntemleri, sarma yöntemleri ve gömülü yöntemler olarak üç ana gruba ayrılır.

Sarma yöntemleri, model performansına dayalıdır ve belirli bir öğrenme algoritmasını kullanarak özniteliklerin alt kümelerini iteratif olarak seçer. Bu yöntemler, seçilen her bir öznitelik alt kümesini doğrudan model üzerinde deneyerek en iyi performansı veren alt kümeyi bulmaya çalışır ve bu nedenle filtre yöntemlerine göre daha fazla hesaplama yükü gerektirir, ancak daha yüksek doğruluk sağlayabilirler (Guyon & Elisseeff, 2003).

Filtre yöntemleri, hesaplama maliyeti açısından sarma algoritmalarına göre daha avantajlıdır, çünkü bu yöntemler belirli bir kullanım durumuna göre özelleştirilmemiştir.

Buna rağmen, filtre algoritmaları hâlâ önemli bilgiler taşır ve verimli öznitelik seçim süreçleri sunar (Guyon & Elisseeff, 2003). Bu algoritmalar, öznitelikleri sıralamak ve önemli olanları seçmek için farklı ölçütler kullanır. Korelasyon tabanlı analiz, karşılıklı bilgi ve mesafe tabanlı ölçümler, filtre algoritmalarında kullanılan yaygın yöntemler arasındadır. Bu ölçütler, nispeten etkili bir öznitelik alt kümesine ulaşmayı ve makine öğrenimi modelleri için daha kısa eğitim süreleri sağlamayı mümkün kılar (Guyon & Elisseeff, 2003).

Son olarak gömülü yöntemler, bağımsız bir algoritma olmaktan ziyade, modelleme sürecinin ayrılmaz bir parçası olarak çalışır. Bu yöntemler, özellikle verimlidir; çünkü sürekli olarak veriyi alt gruplara ayırma işlemi yapmazlar. Bunun yerine, modelleme sürecinin içinden yararlanarak bir özniteliğin önemli olup olmadığını belirler ve bu değerlendirmeyi metriklerle destekler (Guyon & Elisseeff, 2003). Gömülü yöntemlerin etkili bir örneği, L1 metriği ile katsayıları cezalandırarak gereksiz öznitelikleri ortadan kaldıran LASSO (Least Absolute Shrinkage and Selection Operator) algoritmasıdır (Tibshirani, 1996). LASSO algoritmasının başarısı, birçok başka algoritmanın geliştirilmesinde temel olmuştur. Bu çalışmada kullanılan gömülü yöntemlerin başka bir örneği olan Gradyan Artırma (GB) algoritmasının detayları, bölüm 2 kısım 2.3.1’de ele alınmıştır ve öznitelik seçim sürecinin genel performansını artırmaya yönelik katkıları bu bölümde açıklanmıştır.

1.8. Veri Dengeleme

Veri tabanları, her sınıfın sahip olduğu örnek sayısına göre değerlendirildiğinde dengeli ve dengesiz veri tabanları diye iki ana yapıya ayrılır. Eğer bir veri tabanındaki tüm sınıflar aynı sayıda örneğe sahipse ya da sınıflar arasındaki fark önemsiz düzeydeyse, bu veri tabanı dengeli olarak adlandırılır (Barandela vd., 2003). Dengesiz veri tabanları ise bir veya daha fazla sınıfta belirli bir popülasyona ait az sayıda örnek veya aykırı değerler içeren veri yapılarıdır. Dengesiz veri tabanlarında genellikle iki ana az örnekleme türü mevcuttur: Tahmin değişkenlerinin (özniteliklerin) yeterince temsil edilmemesi ve bir veri kümesindeki bağımlı değişkenlerin (sınıfların) dengesiz bir şekilde dağılması (Esfahani & Dougherty, 2014). Özellikle ikinci durum, tıp ve psikoloji gibi alanlarda sıklıkla karşılaşılan bir problemdir. Bunun nedeni, belirli bir hastalığa veya bozukluğa sahip olmayan bireylerin sayısının, hasta bireylere kıyasla genellikle daha fazla olmasıdır (Seera & Lim, 2014). Bu bağlamda, sağlıklı bireylerin sayısının hasta bireylere kıyasla yüksek olması yaygın bir durumdur. Her ne kadar dengesiz sınıflar genelleme açısından zorluklar oluştursa da bu

problemi ele almak için çeşitli yöntemler geliştirilmiştir. Dengesiz sınıf sorunlarını çözmek için örneklem azaltma (undersampling) ve örneklem artırma (oversampling) gibi yöntemler kullanılır. Örneklem azaltma yöntemleri, çoğunluk sınıfından örneklerin azaltılmasıyla veri dengesini sağlarken, örneklem artırma yöntemleri ise azınlık sınıfına yeni örnekler ekleyerek veri setini dengeler. Bu yöntemler arasında sentetik veri oluşturma yaklaşımları (örneğin SMOTE), hibrit yaklaşımlar ve veri setini benzer popülasyonlara genişletme gibi farklı stratejiler bulunmaktadır (He & Ma, 2013). Bu çözümler, özellikle dengesiz sınıf probleminin eğitim ve tahmin süreçlerinde meydana getirdiği olumsuz etkileri azaltmak için kullanılmaktadır.

1.8.1. Örnek Çoğaltma

Örnek çoğaltma, dengesiz veri setlerinde azınlık sınıfının örnek sayısını artırmak için kullanılan temel veri artırma yöntemlerinden biridir. Bu teknik, az temsil edilen sınıfın örneklerini çoğaltarak sınıflar arasındaki dengesizliği giderir (Chawla vd.,2002). Örnek çoğaltma yöntemleri arasında en popüler olanlardan biri, Sentetik Azınlık Fazla Örneklem Tekniği (SMOTE)'dir. SMOTE, azınlık sınıfındaki örnekler arasında sentetik veri noktaları oluşturarak veri setini dengeler. Bu yöntem, azınlık sınıfındaki her bir veri noktasını en yakın komşularına dayalı olarak değerlendirir ve belirli bir ağırlıkla yeni veri noktaları üretir. Bu sürecin temel metodolojisi, bir veri noktasından en yakın komşusuna kadar olan vektörü hesaplamak, bu vektörü $[0,1]$ arasında rastgele bir sayı ile çarpmak ve sonuç değerini mevcut veri setine ekleyerek yeni bir sentetik gözlem oluşturmaktır (Chawla vd., 2002). SMOTE'un en önemli avantajı, veri setine yeni bilgi ekleyerek aşırı öğrenme riskini azaltmasıdır. Bununla birlikte, bu yöntem, sınıf sınırında yer alan gürültülü verilerin çoğaltmasına neden olabilir ve bu durum model performansını olumsuz etkileyebilir (Chawla vd., 2002).

1.8.2. Örnek Azaltma

Örnek azaltma, dengesiz veri setinde fazla temsil edilen sınıfın örnek sayısını azaltarak veri setini dengelemeyi amaçlayan bir tekniktir (Mohammed vd., 2020). Bu yöntem, eğitim sürecinde verimliliği artırırken, çoğunluk sınıfındaki fazla örneklerin model üzerindeki baskısını hafifletir (He & Ma, 2013). Örnek azaltma tekniklerinden biri Tomek Link yöntemidir. Tomek Link, dengesiz veri setlerinde sınıf sınırlarını daha net hale getirmek için kullanılan bir veri temizleme ve azaltma tekniğidir. Bu yöntem, azınlık ve çoğunluk

sınıflarına ait ve birbirine en yakın olan örnek çiftlerini tespit eder ve bu çiftlerden çoğunluk sınıfı örneğini veri setinden çıkartır (Elhassan vd.,2016). Böylece, sınıf sınırları daha netleşir ve gürültü azaltılır, bu da modelin azınlık sınıfını daha iyi öğrenmesine olanak tanır ve yanlış sınıflandırma oranını düşürür (Pereira vd.,2020).

Bir veri setinde (x_i, x_j) çifti Tomek Link olarak tanımlanırsa:

- x_i ve x_j , birbirine en yakın komşularsa,
- x_i ve x_j farklı sınıflara aitse,
- Bu çiftlerin arasındaki en yakın mesafe, her iki örneğin başka herhangi bir veri örneğine olan mesafesinden daha kısa ise, bu çift bir Tomek Link'tir.

Tomek Link, aşırı örnekleme yöntemleri ile kullanıldığında, veri setinde hem azınlık sınıfını güçlendirir hem de gürültüyü ve yanlış sınıflandırmalara neden olabilecek örnekleri temizler (Devi vd.,2017).

1.8.3. Hibrit Yaklaşımlar

Hibrit yaklaşımlar, dengesiz veri setlerini dengelemek için örnek çoğaltma ve örnek azaltma yöntemlerinin bir kombinasyonu olarak kullanılır. Bu teknik, her iki yöntemin avantajlarından faydalanarak daha dengeli ve etkili bir veri seti oluşturmayı hedefler. Hibrit yaklaşımlar, özellikle sınıf sınırlarının belirgin olmadığı ve gürültülü verilerin olduğu durumlarda, veri setinin genel kalitesini artırmada kritik bir rol oynar (He & Ma, 2013). Örneğin, SMOTE + Tomek Link yöntemi, önce azınlık sınıfı için SMOTE kullanılarak sentetik örnekler oluşturur, ardından Tomek Link algoritması ile sınıf sınırındaki gürültülü veya fazla noktalar kaldırılır (Chawla vd., 2002). Hibrit yaklaşımlar hem sentetik örneklerin doğruluğunu artırır hem de veri setindeki genel dengesizliği azaltır (Batista vd., 2004). Bu çalışmada kullanılan ADASYN + Tomek Link hibrit yaklaşımın detayları, bölüm 2 kısım 2.4.'te ele alınmıştır ve veri dengeleme sürecinin genel performansını artırmaya yönelik katkıları bu bölümde açıklanmıştır.

1.9. Topluluk Öğrenme

Topluluk öğrenimi, makine öğrenimi alanında en temel bileşenlerden biri haline gelmiştir ve son yıllarda yapay zekâ, örüntü tanıma, sinir ağları ve veri madenciliği gibi

birçok disiplinde önemli bir ilgi görmüştür (Zhou, 2012). Bu yaklaşım, birden fazla sınıflandırıcı veya temel öğrenici setini bir araya getirerek, bireysel modellerin çıktılarının birleştirilmesi yoluyla genel varyansı azaltmayı hedefler. Bu sayede, tek bir modelin performansından daha yüksek doğruluk ve genelleme kapasitesi elde edilir (Dietterich, 2000).

Topluluk öğrenme teknikleri, özellikle karmaşık veri setleriyle çalışırken ve belirsizliğin yüksek olduğu durumlarda büyük avantajlar sunar. Örneğin, birden fazla modelin çıktılarının birleştirilmesi, hem sistemin hatalarını dengelemeye yardımcı olur hem de bireysel modellerin zayıf yönlerini telafi eder. Bu tekniklerin etkili olduğu, farklı problem alanlarında ve gerçek dünya uygulamalarında kanıtlanmıştır (Polikar, 2006). Topluluk öğrenimi hem sınıflandırma hem de regresyon problemlerinde başarılı sonuçlar sunarak, makine öğrenimi sistemlerinin doğruluğunu ve güvenilirliğini artırır (Zhou, 2012).

1.9.1. Artırma Yöntemi

Artırma (Boosting), makine öğreniminde zayıf öğrenicileri bir araya getirerek güçlü bir model oluşturmayı amaçlayan bir topluluk yöntemidir. Bu teknik, sıralı bir şekilde çalışır; her bir öğrenici, önceki öğrenicilerin hatalarını dikkate alarak eğitilir ve bu süreçte verilerin ağırlıkları güncellenir. (Schapire, 1990). Artırma yaklaşımının en bilinen türlerinden biri olan AdaBoost algoritması, zayıf öğrenicileri sıralı bir şekilde eğiterek güçlü bir sınıflandırıcı üretmeyi amaçlar. Bu süreçte, her yeni model önceki modelin hatalarına daha fazla ağırlık vererek eğitilir. AdaBoost'un temel çalışma adımları şu şekildedir:

Algoritma:

- Başlangıç Ağırlıklarının Belirlenmesi: Tüm veri noktalarına eşit ağırlık ($w_i = 1/n$) atanır.
- Zayıf Öğrenicinin (Temel Algoritma) Eğitilmesi: Verilen ağırlıklarla, ilk zayıf öğrenici eğitilir. Temel algoritma, her veri örneği için tahminlerde bulunur.
- Hata Hesabı: Eğitilen modelin tahmin performansı, ağırlıklı hata oranı (ϵ) ile ölçülür:

$$\epsilon = \sum_{i=1}^N w_i \cdot I(y_i \neq \hat{y}_i) \quad (28)$$

- Burada $I(y_i \neq \hat{y}_i)$, yanlış tahmin edilen örneklerin bir göstergesidir.

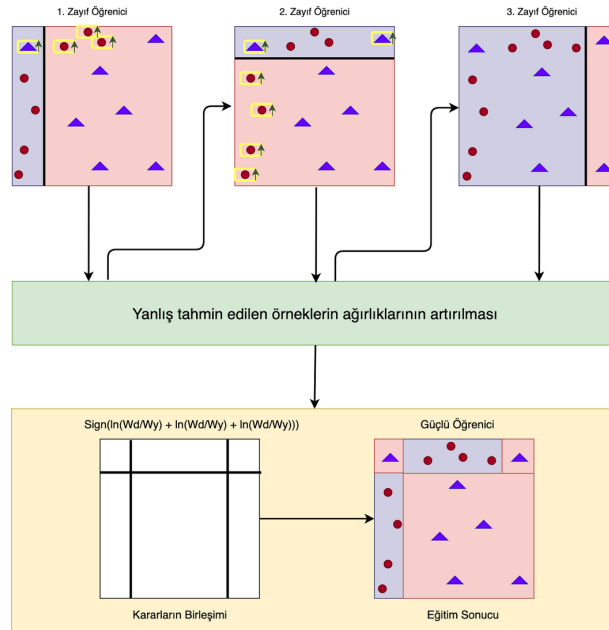
- Model Ağırlığının Belirlenmesi: Modelin önemi, hata oranına göre belirlenir:

$$\alpha = \ln\left(\frac{1 - \epsilon}{\epsilon}\right) \quad (29)$$

- Ağırlıkların Güncellenmesi: Yanlış tahmin edilen örneklerin ağırlığı artırılır:

$$w_i \leftarrow w_i \cdot e^{\alpha \cdot I(y_i \neq \hat{y}_i)} \quad (30)$$

- Güncellenen ağırlıklar, bir sonraki modelin eğitimi sırasında kullanılarak, yanlış sınıflandırılan örneklerin daha fazla dikkate alınmasını sağlar. Bu adımları, belirlenen iterasyon sayısına veya hata oranı istenilen seviyeye gelene kadar tekrarlanır.
- Sonuçların Birleştirilmesi: Tüm zayıf öğrencilerin çıktıları birleştirilerek güçlü bir öğrenici oluşturulur. Şekil 2, bu süreçte zayıf öğrencilerin sıralı eğitimini görselleştirmektedir.



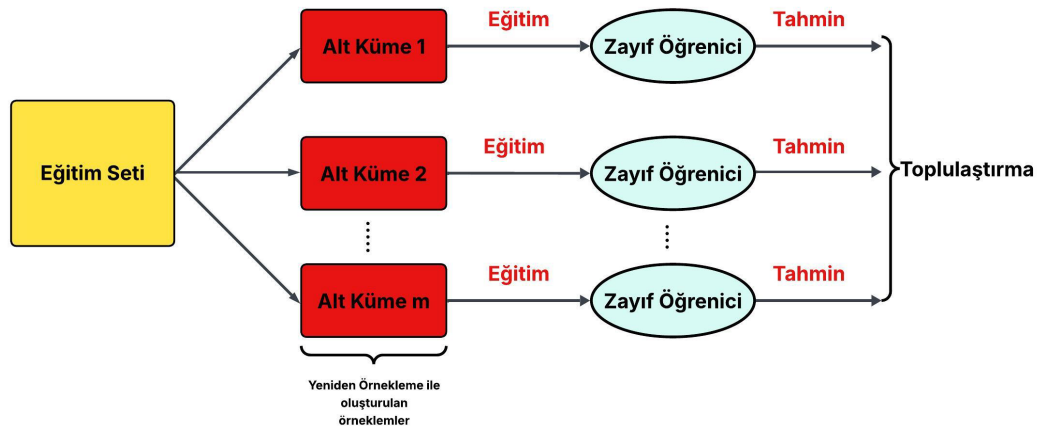
Şekil 2. AdaBoost yönteminde zayıf öğrencilerin sıralı eğitimi (Güzel, 2020)

1.9.2. Örneklem Tabanlı Toplulaştırma

Örneklem Tabanlı Toplulaştırma (Bagging - Bootstrap Aggregating) algoritması, tekrar örneklem (bootstrap) ve toplulaştırma yöntemlerini birleştirerek paralel topluluk yöntemleri oluşturur ve temel öğrenicilerin doğruluğunu artırmayı hedefler (Efron & Tibshirani, 1994; Breiman, 1996). Bu yöntemde, eğitim verisinden bootstrap yöntemi ile birçok alt örneklem (S') oluşturulur ve bu örneklemeler üzerinde temel öğrenme algoritmaları (B) bağımsız ve paralel olarak eğitilir. Ortaya çıkan temel sınıflandırıcıların ($C_1, C_2, \dots, C_T$) çıktıları bir araya getirilerek sınıflandırma problemleri için çoğunluk oylaması, regresyon problemleri için ise çıktıların ortalaması alınır. Bagging algoritmasının adımları şu şekilde tanımlanır:

- Eğitim kümesinden (S) T adet bootstrap örneği oluşturulur (S' , tekrarlı örneklemle ile alınır).
- Her bir bootstrap örneği üzerinde temel öğrenme algoritması (B) çalıştırılarak yeni bir sınıflandırıcı (C_i) üretilir.
- Test aşamasında, sınıflandırma için her sınıf (y) için toplam oylama ($\sum_i 1_{C_i(x)=y}$) hesaplanır ve en yüksek oyu alan sınıf ($C^*(x)$) tahmin olarak atanır.

Bagging algoritması hem ikili sınıflandırma hem de çok sınıflı sınıflandırma problemlerinde etkili bir yöntem olarak kullanılabilir (Breiman, 1996).



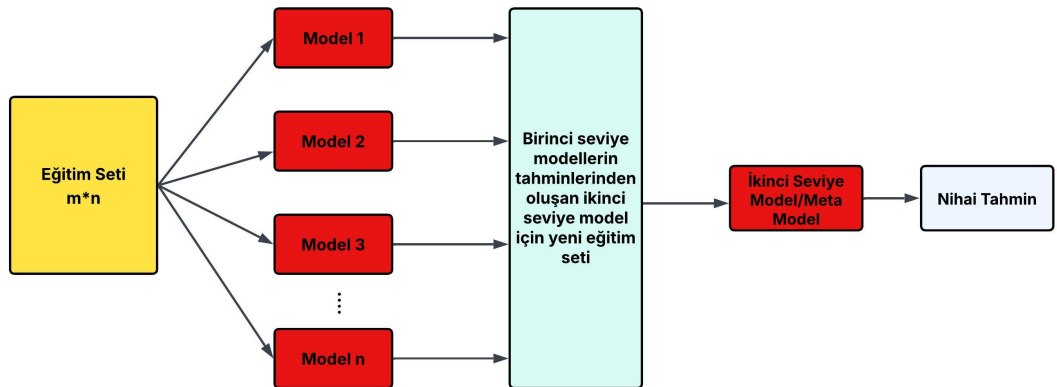
Şekil 3. Örneklem tabanlı toplulaştırma yöntemi ile paralel model eğitim süreci

1.9.3. İstifleme

Makine öğreniminde topluluk yöntemleri arasında öne çıkan tekniklerden biri olan yığılma/istifleme (stacking), farklı modellerin güçlü yanlarını bir araya getirerek hataları en aza indirme amacını taşır (Wolpert, 1992). Stacking, özellikle karmaşık problemler üzerinde yüksek doğruluk sağlayabilmesi nedeniyle yaygın olarak kullanılmaktadır (Zhou, 2012).

Stacking, bir topluluk öğrenme yöntemi olarak, çeşitli makine öğrenimi veya derin öğrenme algoritmalarının bir kombinasyonunu kullanır. Stacking'in temel mantığı, birden fazla temel modelin her birinin tahmin sonuçlarını bir meta-öğrenici aracılığıyla birleştirmektir (Ting & Witten, 1997). Temel modeller genellikle farklı özelliklere veya hiper parametrelere sahip olabilir, böylece veri üzerinde daha geniş bir bakış açısı sağlarlar. Stacking şu şekilde işler:

- Temel Modellerin Eğitimi: Farklı modeller, eğitim verisi üzerinde bağımsız olarak eğitilir. Genellikle k-kat çapraz-doğrulama gibi bir teknikle tahmin çıktıları elde edilir.
- Meta Modelin Eğitimi: Temel modellerin doğrulama setindeki tahmin çıktıları, meta modelin giriş verisi olarak kullanılır ve bu model, nihai tahmini yapmak için eğitilir.
- Son Tahmin: Test verisi yalnızca meta model aracılığıyla değerlendirilir ve stacking modelin performansı ölçülür.
- Bu yapı, stacking'in bireysel modellerin tahmin hatalarını dengelemek için etkili bir yöntem olmasını sağlar. Ayrıca, meta modelin doğrusal olmayan ilişkileri öğrenebilme kapasitesi, stacking'in karmaşık veri setlerinde yüksek performans sunmasına olanak tanır (Sagi & Rokach, 2018).



Şekil 4. İstifleme yönteminin temel mantığı ve işleyiş diyagramı

Stacking, topluluk yöntemleri arasında sıklıkla bagging ve boosting ile karşılaştırılır. Bagging, özellikle model varyansını azaltmayı hedeflerken (örneğin, Random Forest), boosting model hatasını iteratif olarak düzeltmeye odaklanır (örneğin, AdaBoost). Stacking'in bu yöntemlerden farkı, şu avantajları ile ortaya çıkar:

- Bagging ile Karşılaştırma: Bagging, genellikle aynı türde modellerin (homojen modeller) birleştirilmesiyle çalışırken, stacking farklı türde modelleri (heterojen modeller) birleştirerek daha geniş bir genelleme yeteneği sunar (Breiman, 1996).
- Boosting ile Karşılaştırma: Boosting'de modeller sıralı bir şekilde eğitilir ve her bir model bir önceki modelin hatalarını düzeltir. Buna karşın stacking, modelleri paralel olarak eğitir ve meta model üzerinden birleştirir. Bu durum, stacking'in daha esnek bir yapı sunmasını sağlar ve hata yayılımını önler (Dietterich, 2000).

Stacking yönteminin uygulama alanlarına göre çeşitli türleri bulunmaktadır:

- Single-Layer Stacking: Yalnızca bir meta model katmanı içerir. Genelde daha az karmaşıklık gerektirir ve birçok uygulamada yeterli doğruluk sağlar.
- Multi-Layer Stacking: Birden fazla meta model katmanından oluşur ve daha karmaşık problemleri çözmek için kullanılır.
- Weighted Stacking: Temel modellerin performansına göre farklı ağırlıklar atar ve meta modelde bu ağırlıkları dikkate alır.
- Blending: Stacking'in bir varyasyonu olan blending, çapraz doğrulama yerine bir hold-out validation seti kullanarak basit bir uygulama sağlar.
- Dynamic Stacking: Veri özelliklerine göre hangi temel modellerin kullanılacağını dinamik olarak seçer ve genelde daha karmaşık uygulamalarda tercih edilir (Van der Laan vd., 2007).

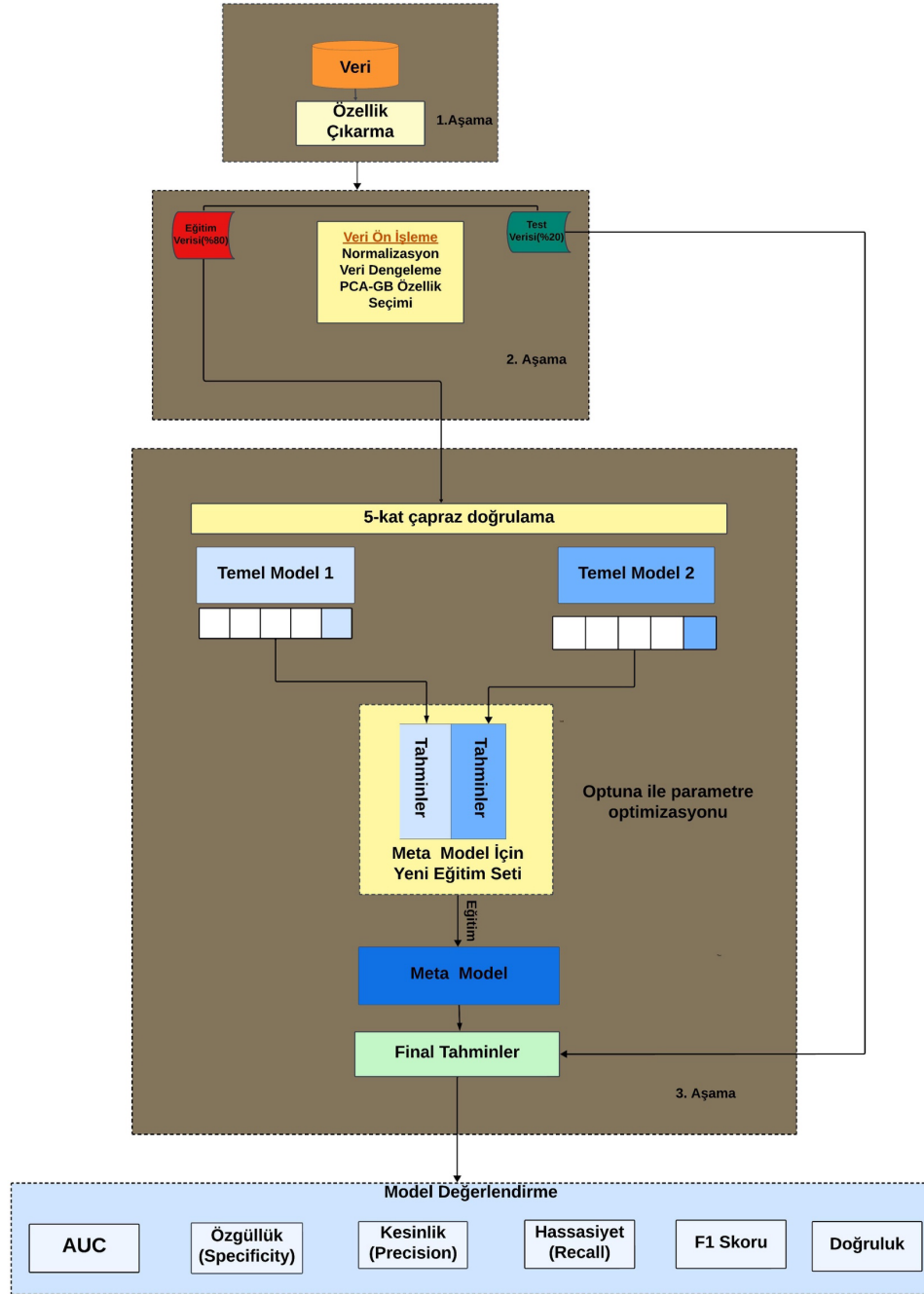
Bu çalışmada, 5 katlı çapraz doğrulama ile eğitilen, statik yapıda tek katmanlı hibrit stacking modeli kullanılmıştır. Kullanılan yaklaşımın detayları, bölüm 2 kısım 2.6.'da ele alınmıştır ve sürecin genel performansını artırmaya yönelik katkıları bu bölümde açıklanmıştır.

2. YAPILAN ÇALIŞMALAR

Bu çalışmada, veriye dayalı analiz ve modelleme süreci detaylı ve sistematik bir şekilde üç aşamada gerçekleştirilmiş olup, süreç Şekil 6'da görselleştirilmiştir. Sürecin ilk aşamasında, ham veri üzerinde kapsamlı bir özellik çıkarımı gerçekleştirilmiştir. Bu aşamada, ham veriden daha anlamlı ve açıklayıcı bilgiler türetmek için çeşitli sinyal işleme teknikleri uygulanmış; bu sayede, istatistiksel, zamansal, kinematik, frekans ve kepsral türlerde önemli özellikler elde edilmiştir. Bu süreç, veriye anlam kazandırarak analiz ve modelleme adımları için anlamlı ve özet bir veri temsili sağlamıştır.

İkinci aşamada, çıkarılan özellikler bir dizi veri ön işleme adımından geçirilmiştir. Bu adımlar arasında normalizasyon, veri dengeleme ve boyut indirgeme tekniklerinden biri olan PCA yöntemi yer almıştır. Ayrıca, GB algoritması ile özelliklerin önem dereceleri belirlenmiş ve bu bilgiler ışığında veri setindeki en anlamlı özellikler seçilmiştir. PCA ile GB yöntemlerinin birlikte kullanıldığı bu aşamanın amacı, veri kümesinde gereksiz veya düşük bilgi taşıyan özellikleri filtrelemek ve daha kompakt bir veri seti oluşturmaktır. Bu iki yöntemin birleştirilmesi hem boyut indirgeme hem de özellik seçimi süreçlerinin daha etkili bir şekilde uygulanmasını sağlamak için yapılmıştır. Bu aşamanın sonunda, yüksek boyutlu veri seti hem daha anlamlı hale getirilmiş hem de modelin işlem yükü optimize edilmiştir. Bu sürecin grafiksel bir temsili, Şekil 2.5'te detaylı olarak açıklanmıştır.

Üçüncü ve son aşamada, elde edilen optimize edilmiş verilerle 5-katlı çapraz doğrulama yöntemi uygulanmış ve temel modellerin çıktılarından meta model için yeni bir eğitim seti oluşturulmuştur. Modellerin parametre optimizasyonu, modern ve güçlü bir hiper parametre optimizasyon kütüphanesi olan Optuna kullanılarak gerçekleştirilmiştir.



Şekil 5. Tez kapsamında yapılan çalışmanın akış diyagramı

Çalışmanın nihai değerlendirme aşamasında, oluşturulan modellerin performansları ölçülmüş ve analiz edilmiştir. Değerlendirme metrikleri arasında AUC (Eğri Altındaki Alan), Özgüllük, Kesinlik, Hassasiyet, F1 Skoru ve Doğruluk gibi geniş bir yelpazede metrikler kullanılmıştır. Bu metrikler hem modelin genelleme yeteneğini hem de sınıflandırma doğruluğunu kapsamlı bir şekilde analiz etmek için kullanılmıştır.

2.1. EMOTHAW Veri Tabanı

EMOTHAW (EMOTion recognition from HAndWriting and draWing), Likforman-Sulem vd. (2017) tarafından oluşturulan, el yazısı ve çizimden duygusal durum tanıma için yeni bir veri tabanıdır. Bu veri tabanı, iki ana bileşenden oluşmaktadır: Biri DASS ölçeği, diğeri ise yedi farklı görev içeren bir değerlendirme setidir. DASS, üç duygusal durumu değerlendiren 42 maddelik bir anketten oluşur: Stres, Anksiyete ve Depresyon. Her biri 14 maddeden oluşan bu ölçek, bu duygusal durumlar arasındaki ayrımı maksimize eder ve 60 yılı aşkın süredir test edilen psikolojik ölçeklerden biridir (Lovibond vd., 1995). Değerlendirme setinin diğer kısmı, beş çizim ve iki yazılı görevden oluşur ve her bir duygusal durumu tanımlamak için özel olarak tasarlanmıştır.

Bu veri tabanı, insan özelliklerini içeren biyometrik veri türleri arasında sınıflandırılmıştır. Morfolojik biyometrik veriler arasında yüz tanıma, iris algılama ve parmak izi yer alırken; el yazısı ve çizim gibi veriler, davranışsal biyometri olarak adlandırılır. Bu alan hâlâ gelişim aşamasında olup, etiketli veri tabanları oluşturmanın zorluğu nedeniyle çok fazla açık kaynak bulunmamaktadır (De Stefano vd., 2019). Bu durum, EMOTHAW'ı çevrimiçi olarak mevcut birkaç kaynaktan biri yapmaktadır. EMOTHAW'ın temel uygulama alanları arasında e-güvenlik ve e-sağlık bulunmakta, özellikle hastalıkların ve zihinsel bozuklukların tespitinde kullanımı dikkat çekmektedir (Faundez-Zanuy vd., 2020).

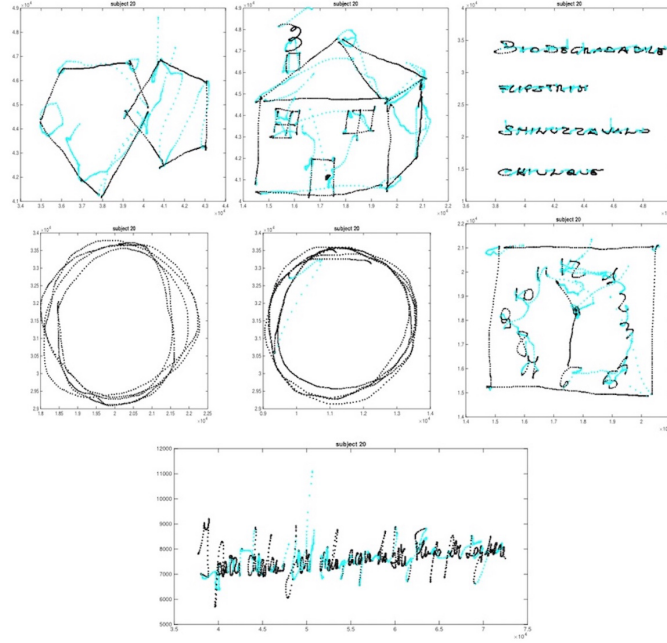
Yazarlar, Sensör veri alanındaki araştırmacılar için açık erişimli ve güvenilir bir alternatif sağlamak amacıyla EMOTHAW'ı geliştirmiştir. El yazısı yoluyla duygu algılama hâlâ erken aşamalarda ve mevcut kaynakların sınırlılığı EMOTHAW'ı daha da önemli hale getirmektedir (Nolazco-Flores vd., 2021).

2.1.1. Görevler Seti

El yazısı aktivitelerinden duygu tanıma, makine öğrenimi alanında nispeten yeni bir konudur (Nolazco-Flores vd., 2021). Ancak, özellikle depresyon ve anksiyete gibi belirli davranışların altında yatan nedenleri belirlemek için kullanılan doğrulanmış ve kanıtlanmış psikolojik yazılı ve çizimli testler mevcuttur (Likforman-Sulem vd., 2017). Bu testler ayrıca tıbbi referans olarak da değerlendirilmektedir. Beyin bozuklukları (ör. Parkinson, Alzheimer, Şizofreni) durumlarında, hastaların motor işlevlerini kaybetme eğiliminde

olduklarına dair kanıtlar bulunmuş ve bu durum el yazısı analizleriyle izlenebilmiştir (Faundez-Zanuy vd., 2013).

Bu çalışmada kullanılan yedi görev, bilişsel sorunları belirlemek için kullanılan üç farklı testten seçilmiştir: Mini Mental Durum Değerlendirme (MMSE) (Folstein vd., 1975), Ev-Ağaç-Kişi (HTP) (Kline, 2013) ve Saat Çizme Testi (CDT). Saat Çizme Testi daha sonra Saat Çizme Sistemi adıyla yeniden düzenlenmiş ve testin bir bilgisayarlı versiyonu oluşturulmuştur (Kim, 2013). Bu testler aracılığıyla EMOTHAW veri tabanı oluşturulmuştur. Saat, beşgen, daire ve ev çizimleri ile kelime ve cümle yazımı bu veri tabanının bir parçasıdır.



Şekil 6. Her katılımcı için gerçekleştirilen görevler ve tüm görevlerden toplanan yazı ve çizim örnekleri (Likforman-Sulem vd., 2017)

Şekil 6’da tüm görevlerden toplanan yazı ve çizim örnekleri gösterilmiştir. Bunlar, beşgenler, ev çizimleri, el yazısı yazılar, döngüler (sol el ve sağ el), saat çizimi ve cümle yazımı. Kalem yere temas ettiği ve havada olduğu veri noktaları sırasıyla siyah ve mavi renklerle gösterilmiştir.

Veri tabanı, tablet teknolojilerindeki ve yazılım sistemlerindeki gelişmeler sayesinde oluşturulabilmiştir. Özellikle INTOUS WACOM 4 serisi tabletler ve Intuos Inkpen gibi cihazlar, yüksek hassasiyet ve doğruluk sunmaktadır (Likforman-Sulem vd., 2017). Bu

gelişmeler, psikolojik testlerin dijital platformlarda uygulanmasını mümkün kılmakta ve bu sayede daha hızlı ve doğru teşhisler yapılmasına katkıda bulunmaktadır (Likforman-Sulem vd., 2017).

Bu görevlerin toplanması sırasında tablet, yedi farklı başlangıç niteliği kaydetmiştir: Yatay pozisyon $x(n)$, dikey pozisyon $y(n)$, zaman damgası (ts), kalem durumu $sq(n)$, azimut açısı $az(n)$, yükseklik $al(n)$ ve basınç $p(n)$. Her görev için bu yedi parametrenin değerleri kaydedilmiştir ve kayıtlar, Wacom'un svc formatında ASCII dosyaları olarak saklanmıştır. Şekil 7, bir kullanıcının beşgen çizim görevini tamamlaması sırasında oluşturulan bir SVC dosyasını temsil etmektedir. Bu detaylı veri toplama yöntemi, yazı dinamiklerini ve hareket özelliklerini derinlemesine analiz etmeye olanak tanımaktadır. Elde edilen yedi farklı başlangıç niteliğinden nasıl yeni özellikler türetildiği, bölüm 1 kısım 1.4.'te detaylı bir şekilde açıklanmıştır. Bu süreç, hareketlerin analizi için daha kapsamlı ve anlamlı veri setlerinin oluşturulmasını sağlamaktadır.

x pozisyonu	y pozisyonu	yükseklik	azimut	kalem durumu	zaman damgası	basınç
1796						
49076	34584	17606448	1	1870	560	45
49025	34608	17606456	1	1870	560	81
49009	34613	17606463	1	1870	560	157
48995	34614	17606478	1	1870	560	193
48993	34614	17606486	1	1870	560	219
48993	34614	17606493	1	1860	560	246
48993	34614	17606501	1	1860	550	284
...	...	...	...	...	...	...
50786	33795	17606756	1	1900	550	305
50727	33808	17606764	1	1900	540	130
50727	33808	17606771	0	1900	540	0
50640	33840	17606779	0	1900	540	0
50621	33860	17606786	0	1900	540	0
50619	33878	17606794	0	1900	540	0
...	...	...	...	...	...	...
51032	33781	17607320	0	1940	510	0
51032	33781	17607328	1	1940	510	84
51056	33773	17607336	1	1940	510	118
...	...	...	...	...	...	...

Şekil 7. Beşgen çizim göreviyle ilgili bir svc dosyasının bir bölümü

Dosya içerisinde toplamda 1796 nokta yer almakta olup, her bir nokta için yedi farklı nitelik kaydedilmiştir.

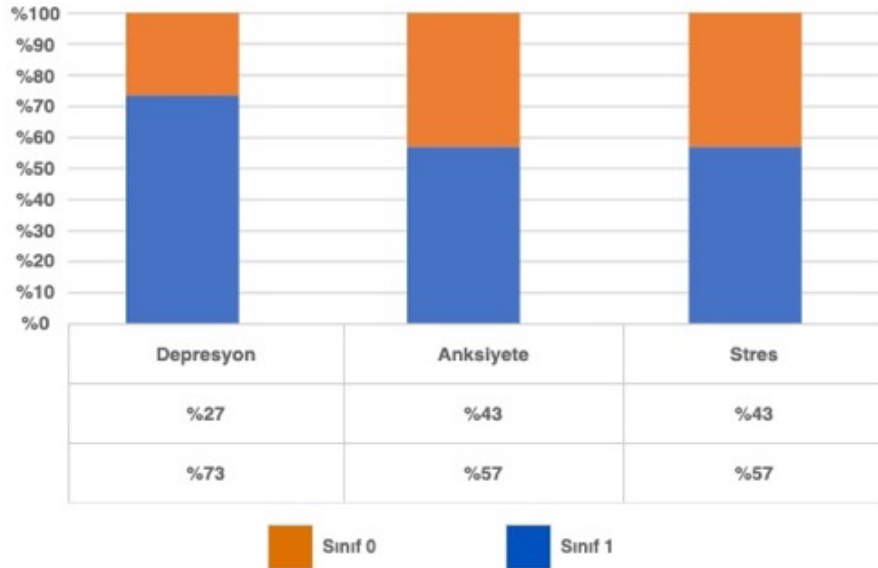
2.1.2. Katılımcılar

Çalışmaya, yaşları 21 ile 32 arasında değişen (ortalama yaş 24,8, standart sapması= 2,4) toplam 129 katılımcı dahil edilmiştir. Katılımcıların tamamı, Seconda Università di Napoli, Psikoloji Bölümü'nde öğrenim gören Lisans ve Yüksek Lisans öğrencilerinden oluşmaktadır. Katılımcıların demografik dağılımı 71 kadın ve 58 erkek şeklindedir. Yaş

aralığı, çalışmanın denekler arası değişkenliğini azaltmak amacıyla sınırlandırılmış ve yaş parametresi, bireylerin belirli bir duygusal duruma sahip olup olmadıklarının kontrol edilmesinde önemli bir faktör olarak değerlendirilmiştir. Diğer bir deyişle, yaşa bağlı değişkenlik, veri tabanına daha fazla karmaşıklık ve önyargı eklemiş olabilir.

Veri toplama süreci, katılımcıların öncelikle demografik bilgilerini içeren onam formunu doldurmaları ve ardından İtalyan DASS testini tamamlamaları ile başlamıştır. DASS testinden sonra, EMOTHAW veri tabanını oluşturmak amacıyla tasarlanan yedi el yazısı ve çizim görevini yerine getirmişlerdir. Bu görevlerden toplanan veri, bireylerin duygusal durumlarının yanı sıra yazı ve çizim davranışlarını analiz etmeye yönelik kapsamlı bir temel sağlamıştır.

Elde edilen veriler, üç duygusal durum (Depresyon, Anksiyete ve Stres) açısından incelendiğinde dengesiz bir dağılım ortaya koymuştur. Şekil 8’de görüldüğü üzere, anksiyete ve stres dağılımları denge problemine daha yakın bir yapıya sahipken, depresyon açık bir dengesizlik göstermektedir; burada sınıf 1 (sağlıklı) bireyleri, sınıf 0 (hasta) bireyleri temsil etmektedir. Sınıf 1’de sınıf 0’a göre çok daha fazla örnek bulunmaktadır. Bu durum, özellikle depresyon sınıfında veri dengesizliği sorununu göz önünde bulundurmaya gerektirmektedir. Ele alınan makine öğrenimi problemlerinde genelleştirmeyi daha da zorlaştıracak bir durum yaratmıştır.



Şekil 8. Veri tabanındaki örneklerin ikili sınıflandırma üzerindeki dağılımı (Tablo 2’ye bkz)

2.2. Özellik Analizi

Çalışmada, el yazısı ve çizim verilerinin analizi için toplamda 595 özellik çıkarılmış olup, bu özellikler zaman tabanlı, kinematik, istatistiksel ve spektral-kepstral özellikleri kapsayarak detaylı ve kapsamlı bir analiz çerçevesi sunmaktadır. Çalışmada elde edilen 595 özelliğin 170 tanesi yazma görevlerinden, geri kalan 425 tanesi ise çizim görevlerinden türetilmiştir. Zaman özellikleri, hareket sıklığını ve yazı sırasında oluşan duraksamaların etkilerini değerlendirmek için önemli ipuçları sunar (Likforman-Sulem vd.,2017).

Kinematik özelliklerin, verilen koşullar sağlandığında makine öğrenimi uygulamalarında oldukça faydalı olduğu kanıtlanmıştır (Impedovo, 2019). Bu çalışmada, kullanıcılar için kullanılan kinematik özellikler, parkinson hastalığıyla ilgili veri setlerinde kullanılan özelliklere benzer şekilde EMOTHAW veri tabanına da uygulanabilir. Bunun nedeni, her iki veri seti arasındaki yapısal benzerliklerdir (Nolazco-Flores et al., 2021).

İstatistiksel özellikler , verilerin genel dağılımını ve dinamiklerini anlamak amacıyla çıkarılmıştır. Literatürde yaygın olarak kullanılan tekniklere ek olarak, bu çalışmada 8 farklı istatistiksel yöntem uygulanmış ve böylece oldukça geniş kapsamlı bir özellik seti oluşturulmuştur.

Sinyal analizi için ise Fourier dönüşümü ve kepstral analiz yöntemleri tercih edilmiştir. Bu yöntemler sinyalin genel frekans düzenliliklerini analiz etmek için kullanılmıştır. Literatürde alternatif olarak değerlendirilen Mel-Frequency Cepstral Coefficients (MFCC) yöntemi belirli avantajlar sunmasına rağmen, çalışmamızın hedefleri doğrultusunda tercih edilmemiştir. MFCC, insan kulağının algısına uygun bir frekans ölçeği sunarak konuşma işleme gibi uygulamalarda ideal bir yöntemdir (Davis & Mermelstein,1980). Ancak bu çalışmada, yazım ve çizim hareketleri gibi fizyolojik sinyalleri analiz etmek hedeflendiği için, MFCC'nin mel ölçeği yerine genel frekans düzenliliklerini anlamaya odaklanan Fourier ve kepstral analiz yöntemleri tercih edilmiştir. Tüm özelliklerin tam tanımı ve kapsamı bölüm 1, kısım 1.4.'te ayrıntılı olarak açıklanmıştır.

2.3. Öznitelik Seçimi

Bölüm 1, kısım 1.5.1'de, avantajları ve yetenekleriyle birlikte açıklanan PCA, bu çalışmada önerilen özellik seçimi veri hattının ön işleme adımı olarak kullanılmıştır. Başka bir deyişle, özellik seti oluşturulduktan sonra, ilgili özelliklerin (veya bileşenlerin) seçilmesinden önce bir boyut azaltma tekniğiyle birleştirilmesi gerekmektedir. PCA, orijinal

veri setindeki değişkenler arasındaki kovaryans yapısını koruyarak, veriyi daha düşük boyutlu bir alt uzaya yansıtan bir boyut indirgeme tekniğidir.

PCA yöntemindeki giriş vektörüne göre sonuç vektörü, $\min(\text{len}(\text{gözlem sayısı}(n)), \text{len}(\text{özellik sayısı}(d)))$ ile sınırlandırılmıştır. Başka bir deyişle, problemde gözlem sayısına bağlıdır. Bu sınır, PCA'nın temelini oluşturan kovaryans matrisinin sıralamasına (rank) bağlıdır ve şu şekilde ifade edilir:

$$\text{rank}(C) = \min(n, d) \quad (31)$$

Burada C, veri setinin kovaryans matrisidir. PCA'nın çıktısı olan ana bileşenler (K), bu sınıra bağlı olarak en fazla $\min(n, d)$ boyutunda olabilir.

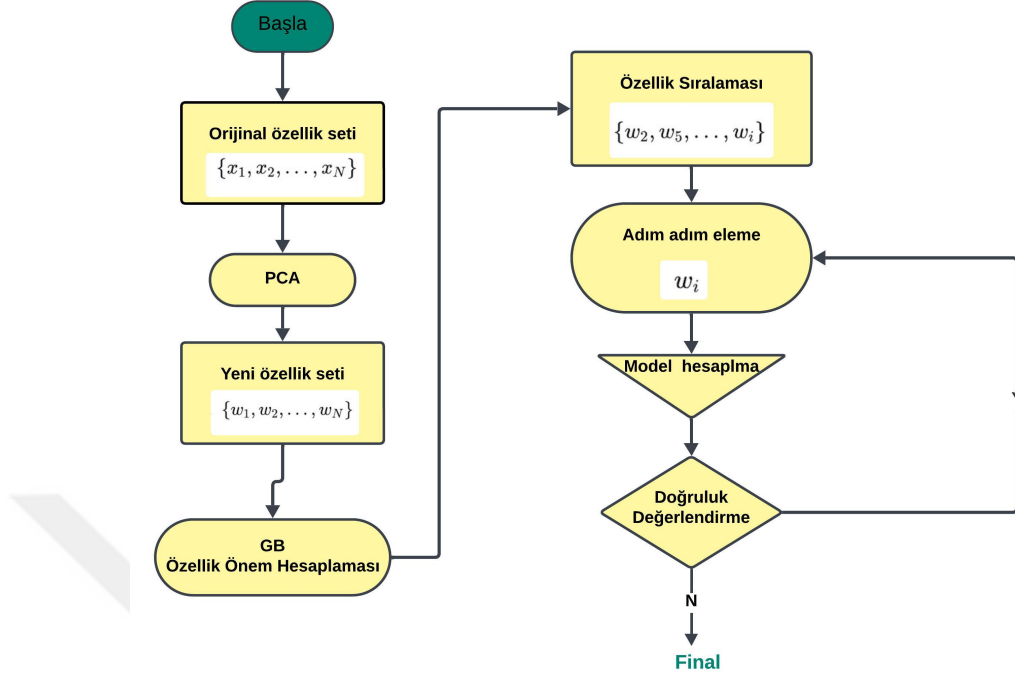
$$BV = PCA(\text{ÖS}, \%P) \quad \text{ve} \quad K \leq \min(n, d) \quad (32)$$

Burada:

- BV: PCA Uygulanarak elde edilen bileşenler vektörüdür.
- PCA: Temel bileşen analizi işlemini ifade eder.
- ÖS: Orijinal veri setinden türetilen özellik setidir.
- %P: PCA'nın hedeflediği açıklanan varyans yüzdesidir. Bu oran, PCA'nın kaç bileşeni koruyacağını belirler.
- K: Elde edilen yeni boyut sayısını ifade eder.

2.3.1. Gradyan Artırma ile Özellik Önem Derecesi Değerlendirme

Gradient Boosting (GB), daha sonra sınıflandırma başlığı altında (bölüm 2 kısım 2.5.3) detaylı açıklanacağı üzere, güçlü bir topluluk yöntemidir. Bu başlık altında ise GB'nin öznitelik seçimi için bir gömülü yöntem olarak kullanımına odaklanılmıştır. GB, model eğitimi sırasında özniteliklerin önem derecelerini belirleyerek düşük katkı sağlayan öznitelikleri elemektedir. Bu yaklaşım, çalışmamızda öznitelik seçimi sürecini daha sistematik hale getirmiş ve modelin performansını artırmaya katkı sağlamıştır.



Şekil 9. PCA-GB tabanlı özellik seçiminin akış şeması

Şekil 9’da da gösterildiği gibi bu çalışmada, PCA ve GB ‘un avantajlarını birleştirerek, probleme uyarlanmış bir PCA-GB hibrit tabanlı özellik seçimi yöntemi önerilmiştir. Bu yöntem, takip eden yığılma topluluk modeline bir veri ön işleme olarak kullanılmıştır. Literatürde genellikle önce öznitelik seçimi, ardından öznitelik çıkarımı uygulanırken, bu çalışmada tam tersine önce PCA ile öznitelik çıkarımı, ardından GB ile öznitelik seçimi gerçekleştirilmiştir. PCA’nın öncelikli olarak uygulanmasının temel nedeni, yüksek boyutlu verilerde bulunan gürültüyü azaltarak özellikler arasındaki gereksiz doğrusal ilişkileri ortadan kaldırmasıdır. Bu sayede GB yönteminin daha temiz, anlamlı ve kompakt bir öznitelik alt kümesini seçebilmesini sağlamaktır. Bu yöntem sıralamasıyla, PCA sonrası oluşan daha temiz ve anlamlı özellik uzayı üzerinde GB kullanılarak, özellik seçim performansının güçlendirilmesi ve modelin genel başarımına katkı sağlanması amaçlanmıştır.

Çalışmada, PCA’nın uygulanmasında toplam varyansın %95’ini açıklayan bileşenlerin seçilmesi tercih edilmiştir. Literatürde, toplam varyansın %95’ini açıklayan bileşenlerin seçimi, bilgi kaybını minimize etmek ve veri setinin temel yapısını korumak amacıyla yaygın bir yöntem olarak önerilmektedir (Jolliffe & Cadima, 2016; Zhang & Castelló, 2017;

Harrington vd., 2005). Çalışmada, PCA ile boyut indirgeme işlemi sonrasında, GB'un özellik önemini değerlendirme avantajı kullanılarak, yeni özellik setindeki daha az bilgi taşıyan özelliklerin filtrelenmesine sağlanmıştır.

2.4. Veri Dengeleme Yöntemi: ADASYN + Tomek Links Hibrit Yaklaşımı

ADASYN (Uyarlanabilir Sentetik Örnekleme), dengesiz veri setlerinde azınlık sınıfının öğrenilebilirliğini artırmak amacıyla, azınlık sınıfına yeni sentetik örnekler ekleyen etkili bir örnek çoğaltma yöntemidir. Bu yaklaşım, sentetik örneklerin yoğunluğunu azınlık sınıfının daha zor öğrenilen bölgelerine odaklayarak üretir ve bu sayede modelin azınlık sınıfını daha iyi tanımasını sağlar (Ahmed vd., 2022). Bu yöntem, SMOTE'un daha gelişmiş ve adaptif bir versiyonu olarak tasarlanmıştır.

ADASYN'nin temel fikri, öğrenilmesi daha zor olan azınlık sınıfı örneklerine daha fazla sentetik veri ekleyerek öğrenme zorluk seviyelerine göre farklı ağırlıklarda bir dağılım kullanmaktır. Bu, sınıf dengesizliğinden kaynaklanan önyargıyı azaltır ve sınıflandırma karar sınırını zor örneklerle uyarlanabilir şekilde kaydırarak öğrenme sürecini optimize eder (Tahsin vd., 2023). ADASYN'nin farkı, azınlık sınıfındaki her bir örneğin çevresindeki veri yoğunluğunu dikkate alarak, sentetik veri örneklerini dinamik bir şekilde üretmesidir. Yoğunluğu düşük olan (çoğunluk sınıfına daha yakın ve daha az temsil edilen) azınlık örnekleri için daha fazla sentetik veri oluşturulurken, yoğunluğu yüksek olan (daha iyi temsil edilen) azınlık örnekleri için daha az sentetik veri üretilir. Bu sayede, ADASYN, veri setinin yapısına daha iyi uyum sağlar, sınıf sınırlarını netleştirir ve dengesizlik sorununu etkin bir şekilde ele alır (He vd., 2008). ADASYN algoritmasının sınıf dengesizliğini gidermek için izlediği matematiksel adımlar aşağıdaki gibidir.

Girdi:

1. Eğitim veri seti D_T, m örnek içerir: (x_i, y_i) , $i = 1, \dots, m$. Burada x_i, n boyutlu özellik uzayındaki bir örnek; $y_i \in \{1, 0\}$ ise sınıf etiketidir. Azınlık sınıf örneklerinin sayısını m_s ve çoğunluk sınıf örneklerinin sayısını m_l olarak tanımlansın.

Bu durumda, $m_s < m_l$ ve $m_s + m_l = m$ olacaktır.

Prosedür:

1. Sınıf dengesizliği derecesini hesaplayın:

$$d = \frac{m_s}{m_l} \quad (33)$$

2. Eğer $d \leq d_h$ ise (d_h : sınıf dengesizliği oranı için maksimum tolerans eşiği) şu adımları yapın.

(a) Azınlık sınıfı için üretilmesi gereken toplam sentetik veri sayısını hesaplayın:

$$G = (m_l - m_s) \times \beta \quad (34)$$

Burada, $\beta \in [0,1]$ genel denge seviyesini temsil eder. ($\beta = 1$) Seçimi tamamen dengeli bir veri seti oluşturur.

(b) Her bir azınlık sınıf örneği için K en yakın komşuya dayalı oranı hesaplayın:

$$r_i = \frac{\Delta_i}{K}, \quad i = 1, \dots, m_s \quad (35)$$

Burada Δ_i, x_i örneğinin çoğunluk sınıfına ait en yakın komşularının sayısıdır.

(c) Normalize edilmiş yoğunluk oranını hesaplayın:

$$\hat{r}_i = \frac{r_i}{(\sum_{t=1}^{m_s} r_t)} \quad (36)$$

(d) Her azınlık sınıf örneği için üretilmesi gereken sentetik veri sayısını hesaplayın:

$$g_i = \hat{r}_i \times G \quad (37)$$

3. Her azınlık sınıf örneği için şu adımları tekrarlayın:

(a) Rastgele bir azınlık sınıf komşusu seçin, (x_{zi}).

(b) Sentetik veri örneğini şu şekilde oluşturun:

$$s_i = x_i + \lambda \times (x_{zi} - x_i) \quad (38)$$

Burada $\lambda \in [0,1]$ aralığında rastgele seçilen bir sayı.

4. 3 numaralı adımda tanımlanan döngüyü, $i = 1$ ' den g_i 'ye kadar çalıştırın.

Bu tez çalışmasında, dengesiz veri setlerinin neden olduğu sınıflandırma zorluklarını aşmak için ADASYN ve Tomek Links teknikleri birlikte kullanılmıştır. Öncelikle ADASYN, azınlık sınıf örneklerini çoğaltarak veri dengesini iyileştirirken; ardından Tomek Links yöntemiyle, sınıf sınırlarındaki gürültülü örnekler temizlenerek veri seti optimize edilmiştir. Bu sıralama, hem veri dengesizliğini azaltmak hem de sınıflar arasındaki ayrımı netleştirmek amacıyla tercih edilmiştir. Tomek Links yöntemi, bölüm 1, kısım 1.8.2’de detaylı olarak açıklanmıştır. Bu iki yöntemin birlikte kullanılmasıyla, veri setindeki dengesizlikle başa çıkılırken modelin doğruluğu ve genellenebilirliği artırılarak sınıflandırma performansının optimize edilmesi hedeflenmiştir.

2.5. Sınıflandırma

Sınıflandırma, modelin belirli bir girdi verisine ait doğru etiket veya kategori tahmin ettiği, denetimli bir makine öğrenimi yöntemidir. Bu süreçte, model etiketlenmiş eğitim verileriyle eğitilir ve ardından yeni, görülmemiş veriler üzerinde tahmin yapmak üzere test edilir. Sınıflandırma, hedef değişkenin ayrık olduğu durumlarda uygulanır ve spam filtreleme, görüntü sınıflandırma, dolandırıcılık tespiti, duygu analizi gibi birçok alanda yaygın olarak kullanılır. Bu yöntem, test verilerini sınıflandırmak için eğitim aşamasında öğrenilen ilişkiler ve özelliklerden yararlanarak verileri kategorilere ayırır (Singh vd., 2016).

2.5.1. Lojistik Regresyon

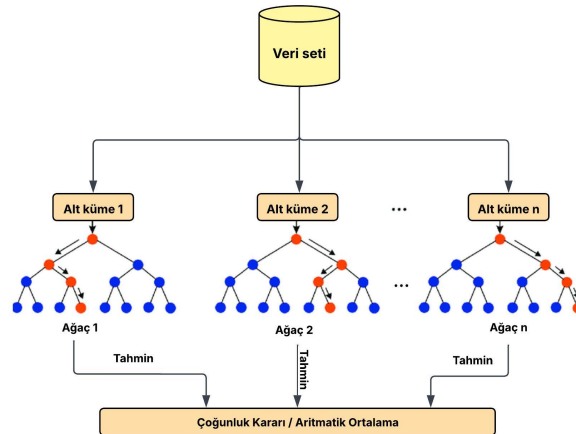
Lojistik regresyon (LR), istatistiksel modelleme ve makine öğrenimi alanlarında, özellikle iki sınıflı (binary) sınıflandırma problemlerinde bağımlı değişkenin, bir veya daha fazla bağımsız değişkenle olan ilişkisini modellemek için kullanılan temel bir yöntemdir. Modelin matematiksel temeli, doğrusal regresyonun sınırlamalarını aşmak için tasarlanan lojistik (sigmoid) fonksiyona dayanır. Bu fonksiyon, bir olayın gerçekleşme olasılığını tahmin etmek için kullanılabilir ve şu şekilde ifade edilir:

$$P(y = 1|X) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 X_1 + \dots + \beta_k X_k)}} \quad (39)$$

Burada β_0 , sabit terimi; $\beta_1, \beta_2, \dots, \beta_k$ ise bağımsız değişkenlerin katsayılarını temsil eder. Lojistik regresyonun teorik altyapısı, olasılık ve bilgi teorisi kavramlarına dayanmaktadır ve özellikle maksimum olabilirlik tahmini yöntemiyle parametrelerin hesaplanması üzerine kuruludur (Hosmer vd., 2013; Peng vd., 2002). Fonksiyon, giriş verilerine (bağımsız değişkenlere) karşılık gelen katsayıları ($\beta_1, \beta_2, \dots, \beta_k$) öğrenerek her bir örnek için 0 ile 1 arasında bir olasılık değeri üretir. Bu olasılık değeri belirli bir eşik değeri (örneğin 0.5) ile karşılaştırılarak sınıf ataması yapılır. Bu sayede, model doğrulukla sınıf ayrımı yapabilir. Lojistik regresyon modeli, bu yapıyla özellikle tıbbi tanı, dolandırıcılık tespiti ve duygu analizi gibi birçok alanda başarılı sınıflandırma sonuçları vermektedir. Bu yöntem, sınıf ayrımını güçlendirmenin yanı sıra, modelin veriye en iyi uyumu sağlayacak parametreleri seçmesini de mümkün kılar.

2.5.2. Rastgele Orman Algoritması

Rastgele Orman (Random Forest - RF), geleneksel sınıflandırma algoritmalarından biri olarak, ağaç tabanlı bir topluluk sınıflandırma yöntemidir ve ağaçların tam kapasiteyle büyütüldüğü örnekleme tabanlı toplulaştırma (bagging, bootstrap aggregation) tekniğini kullanır (Breiman, 2001). Bu algoritmada her bir karar ağacı, verilerin bootstrap yöntemiyle alınan rastgele bir örneğiyle bağımsız olarak eğitilir ve her düğümde rastgele seçilen özellikler üzerinden yapılandırılır. Her bir ağaç, zayıf öğrenici olarak değerlendirilse de topluluk halinde çalışarak güçlü ve kararlı bir model oluşturur. RF'nin nihai tahmini, tüm zayıf öğrenicilerden alınan çoğunluk oyuna veya regresyon problemlerinde ortalama tahmine dayanır.



Şekil 10. Rastgele orman sınıflandırıcısının çalışma mimarisi

RF algoritmasının matematiksel olarak işleyişi aşağıdaki gibi tanımlanır:

Girdi:

- Eğitim veri kümesi: $D = \{(x_i, y_i)\}_{i=1}^n$
- Toplam ağaç sayısı: N

Algoritma adımları:

1. Yeniden örnekleme (Bootstrap) : Her $k = 1, 2, \dots, N$ için, D kümesinden rastgele örnekleme ile bir alt küme D_k oluşturulur.
2. Karar ağacı eğitimi: Her D_k üzerinde, rastgele seçilmiş özellik alt kümesiyle bir karar ağacı $h(x, \theta_k)$ eğitilir. Burada θ_k , k 'inci ağacın parametrelerini (bölünme kriterleri, özellik seçimleri vb.) ifade eder.
3. Son olarak tüm N adet karar ağacından oluşan topluluk tahmini aşağıdaki şekilde hesaplanır:

- Sınıflandırma problemi için çoğunluk oylamasıyla:

$$f(x) = \arg \max_y \sum_{k=1}^N I(h_k(x, \theta_k) = y) \quad (40)$$

- Regresyon problemi için tahminlerin ortalamasıyla:

$$f(x) = \frac{1}{M} \sum_{k=1}^M h_k(x, \theta_k) \quad (41)$$

Burada, $h_k(x, \theta_k)$ k -inci ağacın tahminini, θ_k ise ağacın parametrelerini ifade eder. RF'nin temel avantajları arasında gürültü ve aşırı uyuma (overfitting) karşı dayanıklılığı, büyük veri setlerini verimli bir şekilde işleyebilmesi ve çok sınıflı problemler için doğal olarak uyum sağlaması yer alır (Robnik-Šikonja, 2004).

2.5.3. Gradyan Artırma Algoritması

GB, bir dizi zayıf öğreniciyi birleştirerek güçlü bir model oluşturmayı amaçlayan bir topluluk öğrenme yöntemidir. Bu yöntemin temel fikri, ardışık olarak yeni modeller ekleyerek, önceki modellerin hatalarını azaltmaya odaklanan bir süreçtir (Friedman, 2001; Natekin&Knoll, 2013). Gradyan artırma yöntemleri, istatistiksel çerçevede kayıp fonksiyonunun gradyanına dayanan bir formülasyon ile tanımlanır. Bu formülasyon, gradyan iniş algoritmasını kullanarak model parametrelerini günceller ve böylece modelin hedef fonksiyonunu optimize eder. Her bir adımda yeni bir model oluşturulur ve bu model, topluluğun genel hatalarını azaltmaya odaklanır. Temel öğreniciler, kayıp fonksiyonunun negatif gradyanına en iyi uyacak şekilde eğitilerek, topluluk modeli güçlendirilir. Bu süreç, topluluğun genel performansını artırarak güçlü bir model elde edilmesini sağlar. Friedman'ın Gradyan Artırma algoritmasının uygulama adımları aşağıda gösterilmiştir.

Girdi:

- Eğitim verisi $\{x_i, y_i\}_{i=1}^N$
- İterasyon sayısı M
- Kayıp fonksiyonunun seçimi $L(y, f(x))$
- Temel öğrenci modelinin seçimi $h(x, \theta)$

Algoritma:

1. $f_0(x)$ Değerini bir sabit ile başlat.
2. $m = 1$ 'den M 'ye kadar şu adımları yap.
 - (a) Negatif gradyanı hesapla:

$$g_i = - \left[\frac{\partial L(y_i, f(x_i))}{\partial f(x_i)} \right]_{f(x)=f_{m-1}(x)} \quad (42)$$

(b) g_i üzerinde temel öğrenciyi $h(x, \theta_m)$ ile eğit.

(c) En iyi gradyan azalma adım boyutunu bul:

$$\rho_m = \arg \min_{\rho} \sum_{i=1}^N L(y_i, f_{m-1}(x_i) + \rho h(x_i, \theta_m)) \quad (43)$$

(d) Fonksiyon tahminini güncelle:

$$f_m(x) = f_{m-1}(x) + \rho_m h(x, \theta_m) \quad (44)$$

3. Döngüyü sonlandır.

2.5.4. Aşırı Gradyan Artırma Algoritması

Extreme Gradient Boosting (XGB, XGBoost), Chen & Guestrin (2016) tarafından sınıflandırma ve regresyon problemleri için geliştirilmiş, ağaç tabanlı bir gradyan artırma algoritmasıdır. XGB, birden fazla yenilikçi algoritmayı bir araya getirerek hem doğruluk hem de işlem hızı açısından önemli avantajlar sağlar. İlk versiyonlarında Exact Greedy (Tam Açgözlü) ve Approximate (Yaklaşık) algoritmalarını sunan XGB, daha sonra histogram tabanlı bir bölme algoritması geliştirmiştir. Bu algoritma, özellik değerlerini gruplandırarak hesaplama karmaşıklığını azaltır ve büyük veri kümelerinde eğitim sürecini önemli ölçüde hızlandırır.

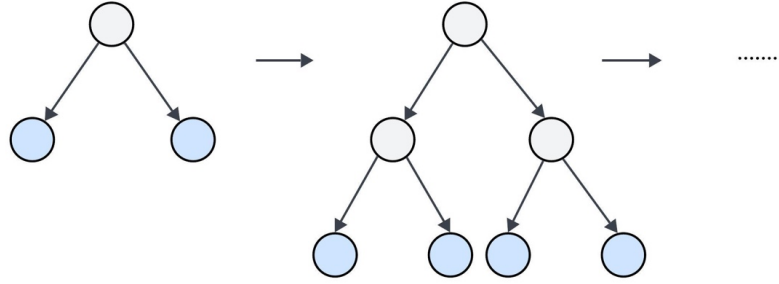
XGB'nin başarısında önemli bir diğer yenilik ise Gradyan Tabanlı Optimizasyon yöntemidir. Bu yöntem, kayıp fonksiyonunun optimizasyonunda ikinci dereceden Taylor genişlemesini kullanarak modelin doğruluğunu artırır ve eğitimi hızlandırır. Ayrıca, seyrek veri ile uyumlu bölme algoritması, sıfır değerli özelliklerin tamamını kaybı en çok azaltan tarafa yerleştirerek seyrek veri kümeleri üzerinde etkili bir şekilde çalışmasını sağlar. Bunun yanı sıra, yaprak bazlı bölme stratejisi, ağaç yapısında daha az bölme gerçekleştirerek yüksek kaliteli modeller oluşturur ve aşırı öğrenme (overfitting) riskini azaltır.

Paralel işlem yapısını optimize eden XGB, çok iş parçacıklı çalışmayı kolaylaştırarak büyük veri kümelerinde eğitim sürelerini kısaltır ve kullanıcıya esnek, ölçeklenebilir bir çözüm sunar (Chen & Guestrin, 2016).

XGBoost'un Gradient Boosting (GB) algoritmasından farkları şu şekilde özetlenebilir:

- Taylor Serisi Genişlemesi: XGBoost hem gradyanı hem de ikinci türevi kullanarak daha doğru ve hızlı optimizasyon yapar. Bu yaklaşım daha hassas kararlar almayı sağlar. GB algoritması yalnızca birinci derece gradyan bilgisine dayanır.
- Histogram Tabanlı Bölme: GB'de özellikler doğrudan kullanılırken, XGBoost, isteğe bağlı olarak histogram tabanlı bölme yöntemi kullanabilir. Bu yöntemle öznitelik değerleri aralıklara ayrılarak grup halinde değerlendirilir, bu da hesaplama süresini azaltır ve büyük veri kümeleriyle çalışmayı kolaylaştırır.

- **Exact Greedy ve Approximate Algoritmaları:** XGBoost, veri setini dallara ayırırken her bölme için tüm örnekleri değerlendirerek Exact Greedy (kesin bölme) adı verilen bir yöntemle çalışır. Bu, küçük ve orta boyutlu veri kümelerinde doğruluk avantajı sunar. Büyük veri kümelerinde ise Approximate (yaklaşık bölme) yöntemle daha hızlı çalışabilir.
- **Seyrek Veri Desteği:** GB algoritmasında eksik veya sıfır veriler için ön işlem şarttır. XGBoost bu tür verileri otomatik olarak en uygun dala yönlendirebilir. Bu özellikle sıfır değerli ve eksik özelliklerle karşılaşıl原因 veri setlerinde büyük avantaj sağlar.
- **Aşırı Öğrenmeye Karşı Yerleşik Önlemler:** XGBoost, GB'ye kıyasla düzenlileştirici terimler (L1 ve L2) ile birlikte gelir. Bu sayede modelin aşırı öğrenme ihtimali azaltılır. Ayrıca erken durdurma gibi mekanizmalarla eğitim süreci daha kontrollü yürütülür.



Şekil 11. Aşırı gradyan artırma algoritmasının çalışma mantığı

2.5.5. Hafif Gradyan Artırma Makinesi

Light Gradient Boosting Machine (LightGBM-LGB), Ke vd. (2017) tarafından geliştirilen, gradyan artırma algoritmalarını optimize eden bir yöntemdir. LightGBM, birkaç temel algoritma ve yenilikçi stratejiye dayanmaktadır:

Yaprak Tabanlı Karar Ağacı (Leaf-wise Tree Growth): LightGBM, yaprak bazlı karar ağacı stratejisi ile çalışır ve ağacı yatay yerine dikey büyütürken en iyi bölünme noktalarını kaybı en çok azaltan yaprağı seçer. Bu yenilikçi strateji, modelin doğruluğunu artırırken işlem maliyetlerini düşürür.

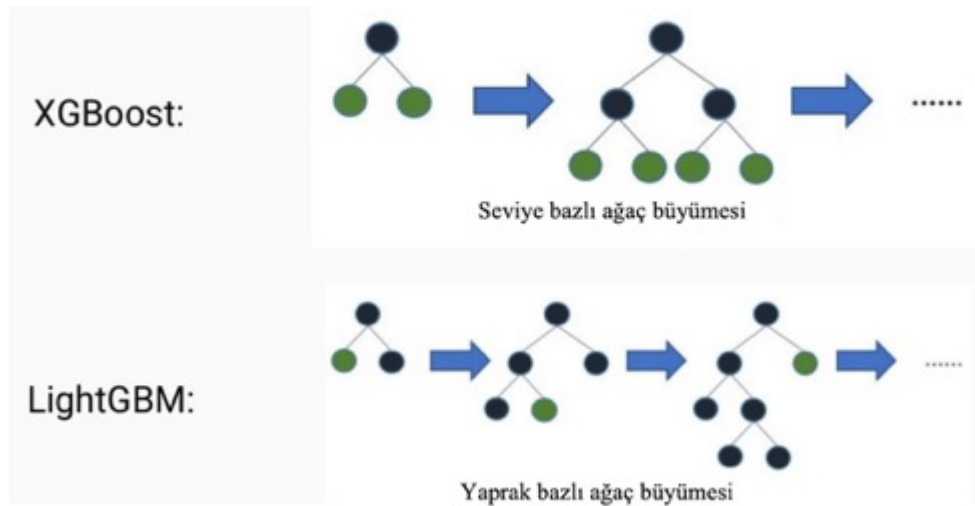
Gradyan Tabanlı Tek Taraflı Örnekleme (Gradient-based One-Side Sampling - GOSS):GOSS algoritması, yüksek gradyanlı örneklere öncelik vererek düşük gradyanlı örneklerin ağırlıklarını artırır ve böylece dengesiz veri kümelerine karşı hassasiyet azaltır. Böylece model hem doğru hem de hızlı bir şekilde büyük veri kümelerini işleyebilir.

Özel Özellik Paketleme (Exclusive Feature Bundling - EFB): LightGBM, veri boyutsallığını azaltmak için EFB algoritmasını kullanır. Bu algoritma, Özellikler arasında dışlayıcı olanları birleştirerek (örneğin, aynı anda aktif olmayan özellikler) bilgi kaybı olmadan özellik sayısını optimize eder.

Histograma Dayalı Karar Ağaçları: LightGBM, histograma dayalı karar ağacı algoritmasını kullanarak sürekli özellikleri histogramlara dönüştürerek özellik bölünmelerini daha hızlı bir şekilde gerçekleştirir. Ayrıca, histogram farklarını hızlandırma gibi yenilikler, eğitim süresini daha da kısaltır.

Kategorik Özelliklere Doğrudan Destek: LightGBM, kategorik veriler üzerinde otomatik işlem yaparak manuel dönüştürme gereksinimini ortadan kaldırır. Bu özellik, özellikle büyük ölçekli veri kümelerinde işlemleri basitleştirir.

Paralel Eğitim ve Büyük Veri Desteği: LightGBM, paralel hesaplama desteği ve düşük bellek tüketimiyle büyük veri kümelerini kolayca işleyebilir. Bu yenilikler sayesinde LightGBM, hızlı, esnek ve güçlü bir modelleme aracı olarak öne çıkar.



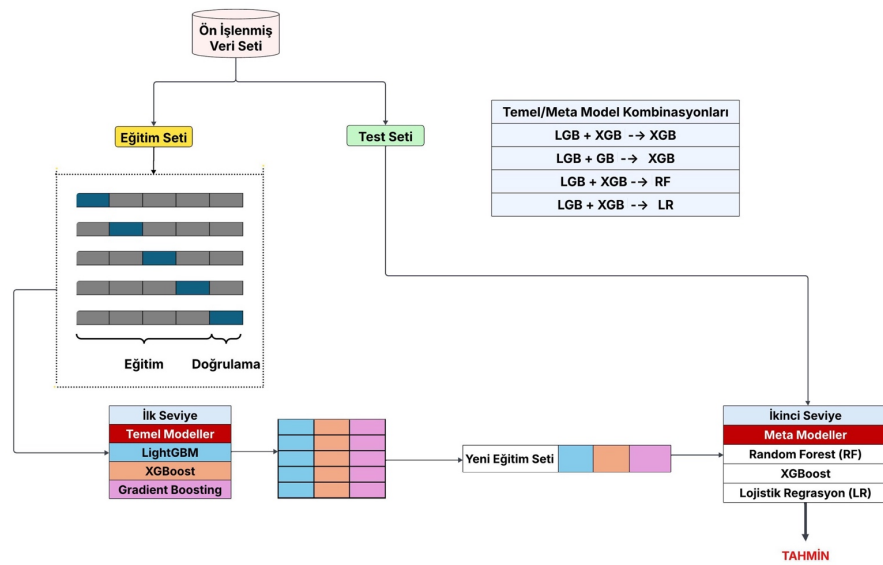
Şekil 12. XGBoost ve LightGBM algoritmalarında ağaç büyüme stratejileri

2.6. İstifleme Modeller

Bu çalışmada 5 katlı Statik Tek Katmanlı Hibrit İstifleme model yaklaşımı uygulanmıştır. Statik yapı, tüm model kombinasyonlarının ve işlem adımlarının sabit ve tutarlı bir şekilde belirlenmiş olduğu anlamına gelirken, tek katmanlı tasarım, temel modellerden elde edilen çıktılarla yalnızca bir meta modelin eğitildiğini ifade eder. Tek katmanlı tasarım, modelin basit ve verimli olmasını sağlarken, aşırı uyum riskini azaltır ve genelleme performansını korur.

Modelleme sürecinde, temel ve meta modellerde yer alan farklı algoritma türlerini bir araya getirerek hibrit bir yapı oluşturulmuştur. Temel modellerde, boosting tabanlı algoritmalar olan LightGBM, XGBoost ve GB kullanılmış, bu şekilde modelin, veri setindeki doğrusal olmayan ve karmaşık ilişkileri öğrenme kapasitesine sahip olması hedeflenirken, meta modellerde ise hem bagging tabanlı RF hem doğrusal sınıflandırmaya dayalı LR hem de boosting tabanlı algoritma olan XGBoost kullanılarak heterojen ve homojen bir perspektif sağlanmıştır.

Homojen kombinasyonlar ile benzer tür algoritmaların performansını optimize etmek amaçlanırken; heterojen kombinasyonlar ile farklı algoritmaların güçlü yönlerinin bir araya getirilmesi hedeflenmiştir. Temel model olarak iki model kullanımı ise, algoritmaların güçlü yönlerini birleştirirken, sistemin karmaşıklığını kontrol altında tutarak aşırı uyum riskini optimize etmek amacıyla tercih edilmiştir.



Şekil 13. 5- katlı statik tek katmanlı hibrit istifleme mimarisi

Kullanılan bu hibrit yöntem ile boosting'in hataları iteratif düzeltme yeteneği ile bagging'in varyansı azaltma kapasitesi birleştirilmiştir. Ayrıca stacking yaklaşımlarında meta modelin (üst model) seçimi, temel modellerin korelasyon düzeyine ve verinin yapısına bağlı olarak değişiklik gösterebilir. Eğer temel modeller benzer örüntüler gösteriyor veya yüksek korelasyon içeriyorsa, doğrusal sınıflandırıcılar (örneğin Logistic Regression gibi) aşırı öğrenmeyi en aza indirerek etkili bir biçimde bu modelleri birleştirebilir. Bu doğrultuda LR meta model olarak seçilerek tahmin doğruluğu ve genelleme performansını artırmak hedeflenmiştir.

2.7. Optuna: Dinamik Hiperparametre Optimizasyonu

Optimizasyon, bir problemin çözümünde belirli bir hedef fonksiyonun en iyi değerine ulaşmayı amaçlayan bir süreçtir. Makine öğrenimi bağlamında bu, bir modelin performansını artırmak için en uygun hiper parametre kombinasyonlarını bulmayı ifade eder. Matematiksel olarak bu süreç, bir hedef fonksiyonun optimize edilmesiyle açıklanabilir:

$$f(\theta) \rightarrow \min_{\theta \in \Theta} f(\theta) \quad (45)$$

Hiper parametreler, modelin öğrenme sürecini etkileyen ve önceden tanımlanması gereken değerlerdir; öğrenme oranı, düzenleme parametreleri veya ağ katmanlarının sayısı gibi. Bu parametrelerin doğru bir şekilde ayarlanması, modelin genelleştirme yeteneğini ve doğruluğunu önemli ölçüde artırabilir. Geleneksel hiper parametre optimizasyonu yöntemleri, grid search ve random search gibi, belirli bir arama alanında tüm olasılıkları sistematik veya rastgele bir şekilde denemeyi içerir. Ancak bu yöntemler, arama alanının boyutunun hızla büyümesiyle yüksek maliyetli ve zaman alıcı hale gelir (Bergstra ve Bengio, 2012).

Optuna, bu sorunu çözmek için geliştirilmiş modern ve dinamik bir hiperparametre optimizasyon çerçevesidir. Optuna'nın öne çıkan özelliği, dinamik bir arama stratejisi olan Tree-structured Parzen Estimator (TPE) algoritmasını kullanarak, hiperparametre aramasını daha verimli hale getirmesidir. Arama uzayının boyutu aşağıdaki şekilde gösterilir:

$$\text{Arama alanı boyutu} = \prod_{i=1}^n |\Theta_i| \quad (46)$$

Burada n , hiperparametre sayısını; Θ_i , her bir hiperparametrenin alabileceği değer kümesini ifade eder. Bu yapı, özellikle yüksek boyutlu hiperparametre uzaylarında Optuna'nın daha esnek ve etkin bir arama yapmasını sağlar (Akiba vd., 2019).

Optuna'nın seçim stratejilerinden biri, fayda/maliyet oranına dayalı bir edinim fonksiyonudur. Bu strateji, denenecek hiperparametre kombinasyonunun seçiminde aşağıdaki ifadeye göre karar verir:

$$x_{\text{sonraki}} = \arg \max_x \frac{l(x)}{g(x)} \quad (47)$$

Bu formülde:

- x_{sonraki} , bir sonraki denenecek hiperparametre kombinasyonudur.
- $l(x)$, o hiperparametre kombinasyonunun modelin başarı fonksiyonuna katkısını (örneğin doğruluk artışı gibi) temsil eder.
- $g(x)$, aynı kombinasyonun değerlendirilme maliyetini (örneğin zaman veya hesaplama yükü) temsil eder.

Bu oran, en yüksek faydayı en düşük maliyetle sağlayacak hiperparametrelerin seçilmesine olanak tanır. Böylece arama süreci, rastgeleliğe dayalı yöntemlere kıyasla çok daha verimli şekilde yürütülür.

Ayrıca Optuna, düşük başarıma sahip denemeleri önceden tahmin ederek devre dışı bırakabilen bir erken durdurma mekanizması da sunar. Bu mekanizma, aşağıdaki formülle ifade edilir:

$$E[f(\theta_{t+k}) \mid \theta_t] < T \quad (48)$$

Burada:

- θ_t , şu anki hiperparametre durumu,
- $\theta_{(t+k)}$, gelecekte denenmesi planlanan hiperparametre kombinasyonu,
- $E[f(\theta_{t+k}) \mid \theta_t]$, şu anki bilgiye göre gelecekteki kombinasyonun beklenen başarıma değeri,

- T , kabul edilebilir başarımlar eşiğini temsil eder.

Bu eşitsizlik sağlanmıyorsa, yani beklenen başarımlar düşükse, ilgili deneme gerçekleştirilmez. Bu sayede zaman ve kaynak israfı önlenmiş olur. Bu çalışmada, hiper parametre optimizasyonu sürecinde Optuna'nın sağladığı bu modern ve dinamik mekanizmalardan yararlanılmıştır. Optuna ile gerçekleştirilen bu süreçte, LightGBM, XGBoost, Random Forest (RF) ve Gradient Boosting (GB) algoritmalarına ait temel hiperparametreler üzerinde yoğunlaşmıştır. Bu kapsamda; tahminci sayısı ($n_estimators$), öğrenme oranı ($learning_rate$), maksimum derinlik (max_depth), yaprak sayısı (num_leaves), örnekleme oranları ($subsample$, $colsample_bytree$), minimum örnek sayıları ($min_child_samples$, $min_samples_split$, $min_samples_leaf$) ve düzenleme parametreleri ($alpha$, $lambda$, $gamma$) gibi performansı doğrudan etkileyen parametreler, uygun aralıklarla optimize edilmiştir.

2.8. Performans Değerlendirme

Bu çalışmada, modellerin genelleme yeteneği doğruluk, AUC, hassasiyet, kesinlik, özgüllük ve F1 skoru gibi metriklerle analiz edilmiştir. Bu metrikler, sınıflandırma doğruluğunu ve farklı sınıflardaki performansı değerlendirmek için seçilmiş ve modellerin güçlü ve zayıf yönlerini belirlemede kullanılmıştır.

Bu metrikler, sınıflandırma sonuçlarını özetlemek ve analiz etmek için karışıklık matrisinden hesaplanmıştır. Karışıklık matrisi, tahminleri TP, FP, TN ve FN kategorilerine ayırır.

- Gerçek Pozitif (TP): Pozitif olarak tahmin edilen ve gerçekte de pozitif olan örnekler.
- Yanlış Pozitif (FP): Pozitif olarak tahmin edilen ancak gerçekte negatif olan örnekler.
- Gerçek Negatif (TN): Negatif olarak tahmin edilen ve gerçekte de negatif olan örnekler.
- Yanlış Negatif (FN): Negatif olarak tahmin edilen ancak gerçekte pozitif olan örnekler (kaçırılan örnekler).

		GERÇEK	
		Pozitif	Negatif
TAHMİN	Pozitif	TP	FP
	Negatif	FN	TN

Şekil 14. Karışıklık matrisi

2.8.1. Doğruluk

Tüm doğru tahminlerin (TP ve TN) toplam örnek sayısına oranıdır. Modelin tüm sınıflar üzerinde doğru tahmin oranını ölçer. Ancak, dengesiz veri setlerinde yanıltıcı olabilir.

$$\text{Doğruluk} = \frac{TP + TN}{TP + TN + FP + FN}$$

2.8.2. Özgüllük

Negatif sınıfların doğru tahmin edilme oranını ölçer. Yanlış pozitiflerin önlenmesi gereken durumlarda önemlidir.

$$\text{Özgüllük} = \frac{TN}{TN + FP}$$

2.8.3. Hassasiyet

Gerçek pozitiflerin ne kadarının doğru tahmin edildiğini ölçer. Pozitif sınıfı kaçırmamak önemliyse bu metrik önemlidir.

$$\text{Hassasiyet} = \frac{TP}{TP + FN}$$

2.8.4. Kesinlik

Pozitif tahminlerin ne kadarının doğru olduğunu ölçer.

$$\text{Kesinlik} = \frac{TP}{TP + FP}$$

2.8.5. F1 Skoru

Hassasiyet ve kesinlik arasında bir denge sağlar, dengesiz veri setlerinde ideal bir metrik olarak öne çıkar.

$$F1 = 2 \cdot \frac{\text{Kesinlik} \cdot \text{Hassasiyet}}{\text{Kesinlik} + \text{Hassasiyet}}$$

2.8.6. Eğri Altındaki Alan (AUC)

AUC, ROC eğrisinin altında kalan alanı ölçer ve modelin eşik bağımsız performansını temsil eder. Yüksek AUC, modelin hem pozitif hem de negatif sınıfları iyi ayırt edebildiğini gösterir.

$$AUC = \int_0^1 \text{TPR}(\text{FPR}^{-1}(x)) dx$$

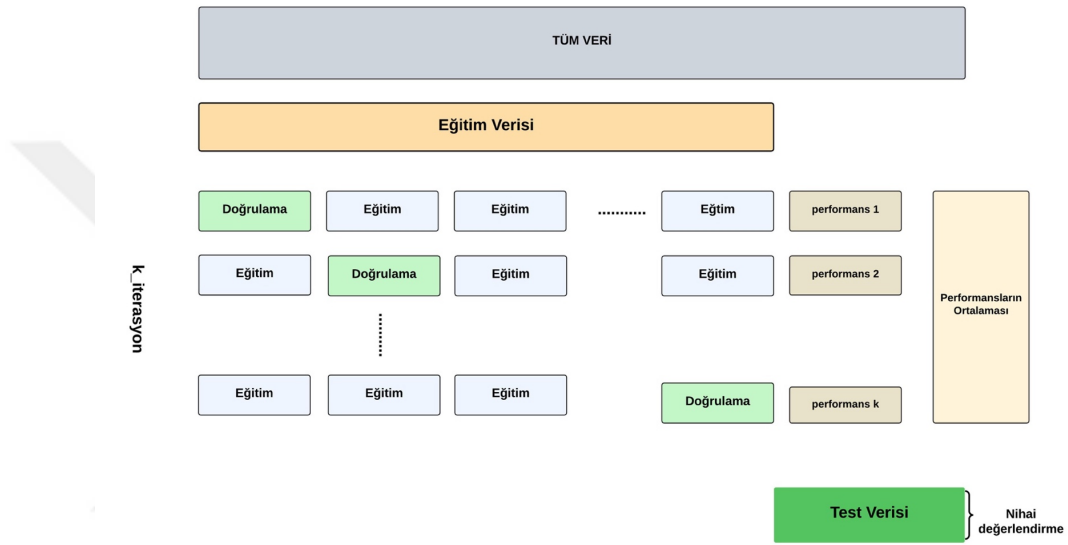
Burada, TPR(Doğru pozitif oranı) ve FPR (Yanlış pozitif oranı) ROC eğrisinde hesaplanan oranlardır.

2.9. k-Katlı Çapraz Doğrulama ile Eğitim-Test-Doğrulama Setlerinin Kullanımı

Cross-validation (çapraz-doğrulama, CV), makine öğrenimi modellerinin performansını değerlendirmek ve bu performansı iyileştirmek için kullanılan bir yöntemdir. CV, özellikle veri setlerinin sınırlı olduğu durumlarda, modelin genelleştirme yeteneğini artırmak amacıyla uygulanır.

Veri setini farklı eğitim ve test alt kümelerine ayırarak modelin hem aşırı öğrenme hem de yetersiz öğrenme eğilimlerini tespit etmeye olanak tanır. En yaygın kullanılan

tekniklerden biri olan k-katlamalı çapraz doğrulama (k-CV), modelin farklı veri segmentlerinde performansını değerlendirmek için veri setini k eşit parçaya böler ve her yinelemede, bir parça doğrulama seti olarak kullanılırken, kalan $k-1$ katlama eğitim seti olarak kullanılır. Bu işlem k kez tekrarlanır ve her katlama bir kez doğrulama seti olarak kullanılmış olur. Son olarak bu k başarımlar ortalaması alınarak modelin genel başarımı hesaplanmıştır. Bu süreç tüm veri kümesi üzerinde tekrarlanarak model doğruluğunun güvenilir bir tahminini sağlar (Anguita vd., 2012).



Şekil 15. k-Kat çapraz doğrulama ve eğitim-test ayrımının gösterimi

3. BULGULAR VE TARTIŞMA

Bu çalışmada, yazı ve çizim verilerinden depresyon, anksiyete ve stres sınıflandırmalarını iyileştirmek amacıyla yedi farklı veri işleme ve dört farklı modelleme yaklaşımı değerlendirilmiştir.

İlk olarak, ham veriden çıkarılan özellikler üzerinde herhangi bir ön işleme uygulanmaksızın seçilen sınıflandırma modellerinin bireysel performans analizleri yapılmıştır. Ardından, meta model olarak Lojistik Regresyon (LR) ve XGBoost kullanılarak, meta model seçimi üzerine bir karşılaştırma gerçekleştirilmiştir. Lojistik Regresyon, istatistiksel modellemede yaygın olarak kullanılan basit ve yorumlanabilir bir yöntemdir. Öte yandan, XGBoost (eXtreme Gradient Boosting), daha karmaşık veri yapılarıyla başa çıkabilen, yüksek performanslı bir makine öğrenmesi algoritmasıdır.

Daha sonraki aşamalarda, yöntemsel çeşitlilik sağlamak amacıyla farklı istifleme (stacking) kombinasyonları denenmiştir. Bu kapsamda, ham veriden çıkarılan özellikler üzerine PCA ve GBC teknikleri ayrı ayrı uygulanmış, bazı senaryolarda ise veri artırma yöntemiyle birlikte entegre edilmiştir. Her bir durum için performans karşılaştırmaları gerçekleştirilmiştir. Ardından PCA ve GBC hibrit olarak uygulandı ve bireysel olarak uygulanan PCA ve GB ile performans karşılaştırmaları yapıldı. Son olarak, özellik seçiminin hangi aşamada yapılmasının daha etkili olduğu (özellik seçiminden önce mi sonra mı) analiz edilmiştir.

Çalışmada yapılan analizlerin özeti:

- Model 1: Hiçbir Ön İşlem Yok + Sınıflandırma Modellerinin Bireysel Performans Analizi
- Model 2: Hiçbir Ön İşlem Yok + Meta Model = LR
- Model 3: Hiçbir Ön İşlem Yok + Meta Model = XGBoost
- Model 4: Veri Artırma + GBC
- Model 5: Veri Artırma + PCA
- Model 6: PCA + GBC + Özellik Seçiminden Sonra Veri Artırma
- Model 7: Özellik Seçimi Yok + Veri Artırma
- Model 8: Veri Artırma Olmadan + PCA + GBC
- Model 9: Veri Artırma + PCA + GBC

3.1. Sınıflandırma Modelleri

Tablo 8. Modellerin bireysel performans analizleri

Duygular	Görevler	Modeller	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	XGB	82.11	55.19	30.95	28.57	41.26	67.69
		LGB	76.84	46.09	18.67	18.57	21.33	61.15
		GB	77.37	49.25	19.07	18.57	21.05	61.54
		RF	81.58	48.72	24.91	21.43	32.52	65.38
		LR	64.21	35.16	16.33	18.57	15.30	51.92
	ÇİZME	XGB	77.89	51.13	21.36	20.00	23.79	62.31
		LGB	75.26	53.38	24.44	22.86	27.87	61.15
		GB	73.68	44.66	17.53	17.14	18.92	58.46
		RF	88.42	55.41	17.23	12.86	28.17	68.08
		LR	67.37	45.79	28.48	31.43	28.54	57.69
	YAZMA	XGB	77.89	53.08	28.98	27.14	36.46	64.23
		LGB	77.89	51.35	23.28	22.86	24.45	63.08
		GB	74.74	51.95	21.45	21.43	23.17	60.38
		RF	84.21	43.80	12.55	10.00	22.33	64.23
		LR	68.95	52.78	28.92	31.43	28.29	58.85
ANKŞİYETE	TÜMÜ	XGB	63.33	48.67	34.78	33.64	37.56	50.77
		LGB	62.67	50.67	45.22	44.55	47.06	55.00
		GB	63.33	56.12	39.63	38.18	42.24	52.69
		RF	71.33	50.42	37.98	34.55	49.82	55.77
		LR	60.67	54.00	41.68	42.73	42.74	53.08
	ÇİZME	XGB	60.67	51.33	43.44	43.64	45.01	53.46
		LGB	63.33	53.70	42.93	41.82	46.17	54.23
		GB	72.67	54.91	44.97	40.91	51.46	59.23
		RF	66.67	43.45	27.74	24.55	36.67	48.85
		LR	61.33	54.73	47.48	49.09	48.22	56.15
	YAZMA	XGB	66.00	58.30	44.37	42.73	47.93	56.15
		LGB	58.67	53.70	42.63	42.73	43.46	51.92
		GB	57.33	49.39	42.22	43.64	42.08	51.54
		RF	73.33	52.73	31.35	26.36	43.50	53.46
		LR	62.00	53.39	44.91	44.55	46.70	54.62
STRES	TÜMÜ	XGB	57.33	46.73	36.06	35.45	37.82	48.08
		LGB	56.67	52.85	44.94	48.18	44.03	53.08
		GB	66.67	48.55	34.10	30.00	40.48	51.15
		RF	76.00	52.12	35.61	30.00	45.77	56.54
		LR	56.00	50.12	40.55	41.82	41.04	50.00
	ÇİZME	XGB	54.67	45.76	34.33	33.64	36.06	45.77
		LGB	52.67	46.79	38.95	40.91	38.44	47.69
		GB	58.67	38.67	27.16	25.45	30.27	44.62
		RF	79.33	45.82	23.32	18.18	36.92	53.46
		LR	62.67	43.82	32.61	30.00	37.03	48.85
	YAZMA	XGB	68.67	56.52	45.02	42.73	49.05	57.69
		LGB	67.33	55.79	41.54	40.91	45.31	56.15
		GB	62.67	53.21	40.60	40.00	43.50	53.08
		RF	72.00	49.88	33.57	28.18	43.95	53.46
		LR	60.67	55.09	44.62	46.36	45.81	54.62

Tablo 8, farklı sınıflandırma modellerinin depresyon, anksiyete ve stres kategorilerindeki performansını detaylı şekilde farklı görevler (tümü, çizme, yazma) üzerinden incelenmiş ve karşılaştırmalı olarak sunmaktadır. Tablo, özgüllük, AUC, F1 skoru, hassasiyet, kesinlik ve doğruluk metrikleri üzerinden her modelin güçlü ve zayıf yönlerinin değerlendirilmesini göstermektedir.

Depresyon sınıflandırmasında tüm veri üzerinden en yüksek doğruluğu %67,69 XGB modeli sağlarken, en düşük doğruluğun %51,92 LR modeli tarafından elde edildiği görülmüştür. Görev türleri bazında incelendiğinde, yazma verileri ile en yüksek doğruluk %64,23 ile yine XGB tarafından sağlanmış, ancak bu görevde LR modeli %58,08 ile en düşük performansı göstermiştir. Çizme görevinde ise RF %68,08 en iyi sonucu sağlarken, LR modeli %57,69 ile sınırlı kalmıştır. Bu durum, depresyon tanımda XGB modelinin her üç görevde de üstünlüğünü açıkça ortaya koymaktadır.

Anksiyete sınıflandırmasında modellerin performansı genel olarak depresyon sınıflandırmasına kıyasla daha düşük ve daha tutarlı değerler göstermiştir. Tüm veri ile yapılan sınıflandırmada RF modeli %55,77 doğrulukla en iyi sonucu verirken, en düşük performans XGB modeline ait olup %50,77'de kalmıştır. Görev bazında ayrı ayrı bakıldığında, çizme görevinde GB %59,23 doğruluk ile en iyi performansı sergilerken, RF %48,85 ile en düşük doğruluğa sahiptir. Yazma görevinde ise XGB modeli yine en iyi sonucu %56,15 verirken, GB modeli %51,54 doğruluk ile sınırlı kalmıştır. Bu sonuçlar, anksiyete verilerinin görevlere göre performansının daha dengeli ancak sınırlı olduğunu göstermektedir. Ayrıca XGB'nin farklı görevlerde tutarlı şekilde en iyi performansı göstermesi dikkate değer bir bulgudur.

Stres sınıflandırmasında ise modellerin performans değerlerinin genel olarak daha düşük seviyelerde olduğu ve görevlere göre farklılıklar içerdiği görülmektedir. Tüm veriler (yazma + çizme) kullanıldığında en yüksek doğruluk %56,54 ile RF modeline ait iken, en düşük doğruluk %50,00 ile LR tarafından elde edilmiştir. Çizme görevinde ise RF modeli %53,46 ile en iyi sonucu sağlamış fakat doğruluk %44,62 ile GB modelinde en düşük seviyeye düşmüştür. Yazma görevinde XGB %57,69 ile en yüksek performansı gösterirken, GB modeli %53,08 ile sınırlı kalmıştır. Bu durum, stres sınıflandırmasında özellikle yazma verilerinde XGB modelinin üstünlüğüne işaret ederken, çizme görevinde RF'nin daha uygun olduğunu göstermektedir. Bu da stres verilerinde görev türlerinin ve model seçimlerinin önemini açıkça ortaya koymaktadır.

Genel olarak modellerin F1 skoru ve hassasiyet değerleri düşük olup, pozitif sınıfların doğru tahmin edilmesinde sınırlamalar olduğunu ortaya koymaktadır. Buna karşın RF ve XGB modelleri özellikle bazı görev ve duygu kategorilerinde daha dengeli ve yüksek performans sergilemiştir. Bu sonuçlar, depresyon, anksiyete ve stres duygularının sınıflandırılmasında seçilen görev türünün (çizme, yazma veya tümü) ve kullanılan modellerin performans üzerinde önemli ve dikkate değer bir etkiye sahip olduğunu göstermektedir. Elde edilen bulgular, sınıflandırma performansının artırılması için özellik seçimi ve veri artırma gibi gelişmiş veri ön işleme tekniklerinin önemine işaret etmektedir.

3.2. Meta Model Seçimi Karşılaştırılması

Tablo 9. Ön işleme uygulanmadan, LGB + XGB temel modelleri ve LR meta model ile elde edilen sonuçlar

Duygular	Görevler	Özgüllük	AUC	F1 Skoru	Hassasiyet	Keskinlik	Doğruluk
DEPRESYON	TÜMÜ	45.79	54.21	37.18	57.14	27.84	48.85
	ÇİZME	40.53	38.05	25.12	37.14	20.75	39.62
	YAZMA	48.42	39.47	22.81	34.29	17.35	44.62
ANKSİYETE	TÜMÜ	57.33	49.21	44.09	46.36	43.53	52.69
	ÇİZME	45.33	43.94	40.66	44.55	37.93	45.00
	YAZMA	45.33	48.30	43.95	50.91	39.89	47.69
STRES	TÜMÜ	52.00	49.70	42.96	48.18	41.78	50.38
	ÇİZME	42.67	44.36	36.08	40.00	33.10	41.54
	YAZMA	52.67	48.85	39.43	41.82	37.97	48.08

Tablo 9'da görüldüğü üzere, hiçbir ön işlem uygulanmadan yapılan bu ilk analiz, depresyon, anksiyete ve stres sınıflandırmaları için yetersiz bir performans ortaya koymuştur. Tablo 9 'da görülen düşük F1 skorları ve özgüllük değerleri, modelin pozitif ve negatif sınıfları tahmin etme konusunda ciddi eksiklikler yaşadığını açıkça göstermektedir. Bu analiz, LGB + XGB → LR kombinasyonunun sınıflandırma görevlerinde düşük bir performans sergilediğini ortaya koymaktadır. Özellikle çizme ve yazma kategorilerinde depresyon ve stres sınıflandırmalarında sonuçlar oldukça sınırlıdır. Bununla birlikte, yazma kategorisinde anksiyete sınıflandırmasında daha dengeli bir performans elde edilmiştir. Meta model olarak kullanılan LR 'nin performansı, Tablo 10 ile karşılaştırıldığında oldukça düşük seviyede kalmıştır. LR, basitliği ve yorumlanabilirliği nedeniyle yaygın olarak kullanılır. Ancak, doğrusal olmayan ilişkileri yakalamada sınırlı kalabilir.

Tablo 10. Ön işleme uygulanmadan, LGB + XGB temel modelleri ve XGBoost meta model ile elde edilen sonuçlar

Duygular	Görevler	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	91.05	55.98	24.21	17.14	50.33	71.15
	ÇİZME	89.47	51.69	20.92	15.71	35.00	69.62
	YAZMA	91.05	56.17	23.34	18.57	38.55	71.54
ANKSİYETE	TÜMÜ	70.67	49.67	36.11	31.82	42.98	54.23
	ÇİZME	60.67	41.15	33.19	31.82	36.00	48.46
	YAZMA	66.00	51.58	37.85	35.45	42.53	53.08
STRES	TÜMÜ	64.67	47.97	38.65	35.45	43.02	52.31
	ÇİZME	65.33	49.36	42.25	40.91	46.02	55.00
	YAZMA	61.33	54.55	46.45	46.36	47.61	55.00

Tablo 10, hiçbir ön işlem uygulanmadan meta model olarak XGBoost kullanılarak gerçekleştirilen sınıflandırma analiz sonuçlarını sunmaktadır. Bu analizlerde depresyon, anksiyete ve stres sınıflandırmaları; tümü, çizme ve yazma görevleri üzerinden ayrı ayrı değerlendirilmiştir.

Depresyon için en iyi doğruluk, tümü %71,15 ve yazma %71,54 kategorilerinde elde edilmiştir. Modellerin bireysel analiz performansları ile kıyaslandığında, depresyon için tüm görevlerde (yazma, çizme, tümü) performans artışı sağlanmıştır.

Stres sınıflandırmasında, yazma kategorisinde %55,00 doğruluk ve %46,45 F1 skoru ile model daha iyi bir sonuç sunmuştur. Stres sınıflandırmasında yazma ve çizme kategorileri, pozitif ve negatif sınıflar arasındaki dengeli tahmin yeteneğini yansıtmaktadır. Ancak, depresyon ve anksiyete sınıflandırmalarında pozitif sınıfları tanımadaki başarısızlık, modelin performansını sınırlamaktadır. Tablo 10 'da yer alan XGBoost modeli ile bir önceki Tablo 9'da yer alan lojistik regresyon modeli ile karşılaştırıldığında, XGB meta modeli (LGB + XGB → XGB), genel doğruluk ve pozitif sınıfları tanımadaki dengeli performansı ile LR modeline (LGB + XGB → LR) göre daha iyi sonuçlar sunmuştur. Bununla birlikte, her iki modelde de F1 skorlarının genelde düşük olması, pozitif sınıfların yeterince iyi ayırt edilemediğini, veri artırma tekniklerinin uygulanması ve öz nitelik seçimi gibi iyileştirmeler yapılması gerektiğini göstermektedir.

3.3. Temel Bileşen Analizi ile Boyut İndirgeme Yaklaşımı

Bu aşamada veri artırma, boyut indirgeme, özellik seçimi ve farklı kombinasyonların sınıflandırma performansı üzerindeki avantaj ve dezavantajlarını görmek için uygulanmıştır.

Tablo 11. PCA ile model kombinasyonlarının sınıflandırma performansları

Duygular	Görevler	Modeller	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	LGB+XGB→XGB	86.84	75.79	68.98	71.43	67.56	82.69
		LGB+GB→XGB	84.21	69.96	65.70	67.14	67.86	79.62
		LGB+GB→RF	85.26	64.44	52.63	51.43	57.73	76.15
	ÇİZME	LGB+XGB→XGB	89.47	70.00	64.39	61.43	70.25	81.92
		LGB+GB→XGB	85.26	70.71	64.59	65.71	65.92	80.00
		LGB+GB→RF	90.53	68.68	52.76	45.71	64.88	78.46
	YAZMA	LGB+XGB→XGB	90.00	71.50	67.47	65.71	74.58	83.46
		LGB+GB→XGB	87.89	74.62	65.91	67.14	67.87	82.31
		LGB+GB→RF	84.21	70.49	57.20	55.71	61.57	76.54
ANKSİYETE	TÜMÜ	LGB+XGB→XGB	84.00	82.48	79.70	80.91	80.06	82.69
		LGB+GB→XGB	80.67	80.30	78.03	80.91	76.31	80.77
		LGB+GB→RF	76.67	78.00	74.38	78.18	71.92	77.31
	ÇİZME	LGB+XGB→XGB	74.00	81.39	78.70	87.27	73.69	79.62
		LGB+GB→XGB	80.67	80.76	79.24	82.73	76.59	81.54
		LGB+GB→RF	80.00	84.39	76.47	79.09	75.14	79.62
	YAZMA	LGB+XGB→XGB	77.33	79.52	76.40	80.91	72.91	78.85
		LGB+GB→XGB	81.33	77.00	74.77	75.45	75.70	78.85
		LGB+GB→RF	80.67	81.45	74.50	75.45	77.18	78.46
STRES	TÜMÜ	LGB+XGB→XGB	77.33	83.18	77.87	83.64	74.52	80.00
		LGB+GB→XGB	80.00	85.52	78.86	82.73	76.56	81.15
		LGB+GB→RF	77.33	80.58	76.76	81.82	74.86	79.23
	ÇİZME	LGB+XGB→XGB	81.33	81.21	78.11	80.00	78.06	80.77
		LGB+GB→XGB	78.00	76.61	75.17	78.18	73.33	78.08
		LGB+GB→RF	75.33	75.24	71.98	75.45	69.52	75.38
	YAZMA	LGB+XGB→XGB	80.00	77.21	75.93	78.18	75.52	79.23
		LGB+GB→XGB	84.67	78.48	79.06	79.09	80.21	82.31
		LGB+GB→RF	79.33	78.39	72.42	73.64	72.50	76.92

Tablo 11, veri artırma , boyut indirgeme (PCA) ve farklı model kombinasyonlarının sınıflandırma performansı üzerindeki etkileri detaylı bir şekilde incelenmiştir. Depresyon, anksiyete ve stres sınıflandırmaları; tümü, çizme ve yazma görevlerine dayalı olarak analiz edilmiştir. Elde edilen bulgulara göre, PCA özellikle XGB meta modeli ile birlikte kullanıldığında depresyon ve stres sınıflandırmalarında anlamlı performans artışı sağlamıştır. Örneğin, depresyon sınıflandırmasında LGB+XGB→XGB kombinasyonuna PCA uygulandığında doğruluk oranı %82,69'a, F1 skoru ise %68,98'e yükselmiştir. Benzer şekilde, stres sınıflandırmasında da PCA'nın katkısı ile doğruluk %82,31'e kadar ulaşmıştır. Anksiyete sınıflandırmasında ise, tüm kategorilerde ve model kombinasyonlarında daha dengeli sonuçlar elde edilmiş, F1 skorları %74,38 ile %79,70 arasında değişmiştir. Bununla birlikte, PCA'nın etkisinin sınırlı kaldığı bazı senaryolar da mevcuttur; özellikle RF içeren kombinasyonlarda performans artışı daha sınırlı kalmıştır.

PCA ile yapılan boyut indirgeme işlemi sonucunda, 595 boyutlu orijinal öznetelik kümesi yaklaşık 61 boyuta düşüş sağlanmış; bu sayede hem hesaplama maliyeti azaltılmış

hem de daha anlamlı özniteliklerin modele dâhil edilmesi mümkün olmuştur. Performans artışı yalnızca boyutun küçülmesinden değil, aynı zamanda PCA'nın veri setindeki gürültülü ve ayırt edici olmayan özniteliklerin elimine edilmesiyle ilişkilidir. Bu bulgular, yüksek boyutlu veri ile çalışırken etkili boyut indirgeme yöntemlerinin önemini ortaya koymaktadır. Ayrıca model karşılaştırmalarında, XGB meta modelinin RF'ye kıyasla depresyon, anksiyete ve stres sınıflandırmalarının tümünde daha yüksek doğruluk ve F1 skorlarına ulaştığı gözlemlenmiştir. Genel olarak PCA, özellikle yüksek boyutlu ve kompleks modellerle birleştirildiğinde anlamlı bir katkı sağlamış, performans üzerindeki etkisi modele ve duygu türüne göre değişkenlik göstermiştir.

3.4. Gradyan Artırma ile Özellik Seçimi Yaklaşımı

Tablo 12. GBC ile model kombinasyonlarının sınıflandırma performansları

Duygular	Görevler	Modeller	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	LGB+XGB→XGB	88.42	72.33	64.46	62.86	69.38	81.54
		LGB+GB→XGB	85.26	65.64	57.25	57.14	59.21	77.69
		LGB+GB→RF	83.68	60.94	54.94	54.29	57.29	75.77
	ÇİZME	LGB+XGB→XGB	93.68	66.77	65.33	57.14	79.33	83.85
		LGB+GB→XGB	93.68	71.35	73.76	68.57	88.33	86.92
		LGB+GB→RF	90.00	69.17	58.85	54.29	75.46	80.38
	YAZMA	LGB+XGB→XGB	85.26	72.86	64.92	67.14	65.51	80.38
		LGB+GB→XGB	81.58	73.01	67.22	75.71	62.47	80.00
		LGB+GB→RF	84.21	62.41	55.42	54.29	61.26	76.15
ANKSİYETE	TÜMÜ	LGB+XGB→XGB	70.00	73.76	74.77	83.64	69.22	75.77
		LGB+GB→XGB	78.67	80.55	76.65	80.00	75.07	79.23
		LGB+GB→RF	78.67	76.82	71.44	71.82	71.63	75.77
	ÇİZME	LGB+XGB→XGB	77.33	73.82	77.24	81.82	76.64	79.23
		LGB+GB→XGB	78.67	73.61	74.95	77.27	74.71	78.08
		LGB+GB→RF	80.67	74.21	72.00	70.91	74.19	76.54
	YAZMA	LGB+XGB→XGB	82.67	78.58	76.50	76.36	78.07	80.00
		LGB+GB→XGB	70.00	73.55	74.10	82.73	68.65	75.38
		LGB+GB→RF	72.00	77.61	74.58	81.82	69.90	76.15
STRES	TÜMÜ	LGB+XGB→XGB	81.33	74.36	77.09	78.18	77.03	80.00
		LGB+GB→XGB	68.00	72.24	70.69	78.18	64.97	72.31
		LGB+GB→RF	74.67	75.06	69.79	71.82	69.17	73.46
	ÇİZME	LGB+XGB→XGB	76.00	76.79	73.33	76.36	71.22	76.15
		LGB+GB→XGB	79.33	74.12	72.22	72.73	73.64	76.54
		LGB+GB→RF	69.33	71.76	71.29	79.09	66.16	73.46
	YAZMA	LGB+XGB→XGB	78.67	76.39	73.23	74.55	73.45	76.92
		LGB+GB→XGB	79.33	80.97	78.36	82.73	76.44	80.77
		LGB+GB→RF	78.00	78.45	73.31	75.45	71.98	76.92

Tablo 12'deki analizler, veri artırma , GB sınıflandırıcısı ile özellik seçimi ve farklı model kombinasyonlarının depresyon, anksiyete ve stres sınıflandırmalarında kategori bazında öne çıkan performansları incelenmiştir.

Depresyon sınıflandırmasında, LGB + XGB→XGB modeli tümü (hem yazma hem çizme) kategorisinde %81,54 doğruluk, %64,46 F1 skoru ve %88,42 özgüllük ile en yüksek performansı sergilemiştir.

Çizme kategorisinde ise LGB + GB → XGB kombinasyonu %86,92 doğruluk, %93,68 özgüllük ve %88,33 kesinlik ile depresyonu en başarılı şekilde sınıflandırmıştır.

Anksiyete sınıflandırmasında, modeli tümü (hem yazma hem çizme) kategorisinde yine LGB+GB→XGB kombinasyonu %79,23 doğruluk, %76,65 F1 skoru ve %80,00 hassasiyet ile pozitif sınıfları dengeli bir şekilde tahmin etmiştir.

Stres sınıflandırmasında ise yazma kategorisinde LGB + GB→XGB modeli %80,77 doğruluk, %78,36 F1 skoru ve %82,73 hassasiyet ile en iyi sonuçları elde etmiştir.

3.5. Özellik Seçimi ve Veri Artırma Sırasının Performansa Etkisi

Tablo 13. PCA+GBC özellik seçiminden sonra veri artırma, LGB + XGB + XGB modeli ile elde edilen sonuçlar

Duygular	Görevler	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	78.95	53.38	20.89	18.57	24.45	62.69
	ÇİZME	76.32	59.21	31.96	34.29	33.86	65.00
	YAZMA	70.00	48.08	29.20	31.43	29.14	59.62
ANKSİYETE	TÜMÜ	56.67	56.03	46.11	48.18	45.34	53.08
	ÇİZME	59.33	54.61	48.62	50.00	47.92	55.38
	YAZMA	60.00	53.55	45.49	46.36	45.94	54.23
STRES	TÜMÜ	56.67	52.30	42.46	43.64	43.35	51.15
	ÇİZME	58.00	54.09	46.35	50.00	45.04	54.62
	YAZMA	61.33	59.67	46.38	48.18	46.15	55.77

Tablo 13'te yer alan analizler, özellik seçiminden (PCA+GBC kaskat yaklaşımı) sonra tüm duygular üzerinde veri artırma yapılan sınıflandırma analiz sonuçlarını göstermektedir.

Tümü kategorisinde, anksiyete sınıflandırması, %46,11 F1 skoru ve %53,08 doğruluk ile depresyon ve stres sınıflarına kıyasla pozitif sınıf için daha iyi bir performans göstermiştir. Ancak tüm sınıflarda pozitif tahmin başarısı kötüdür.

Çizme kategorisinde, depresyon sınıflandırması, %65,00 doğruluk ve %76,32 özgüllük ile negatif sınıfları ayırt etmede başarılıdır. Anksiyete sınıflandırmasında %48,62 F1 skoru ile pozitif sınıflarda diğerlerine göre kısmen başarı sağlanmıştır.

Yazma kategorisinde ise stres sınıflandırması, %46,38 F1 skoru ve %48,18 hassasiyet ile yazma kategorisinde en iyi performansı göstermiştir. Genel olarak, pozitif ve negatif sınıfların tahmini tüm kategorilerde düşük performans göstermiştir.

Tablo 14. PCA+GBC özellik seçiminden önce veri artırma, LGB + XGB + XGB modeli ile elde edilen sonuçlar

Duygular	Görevler	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	91.05	68.01	60.19	54.29	72.89	81.15
	ÇİZME	82.63	67.26	61.88	65.71	69.85	78.08
	YAZMA	88.42	79.44	72.70	75.71	72.44	85.00
ANKSİYETE	TÜMÜ	72.67	77.03	75.26	82.73	69.56	76.92
	ÇİZME	76.67	77.45	76.64	81.82	73.43	78.85
	YAZMA	76.00	76.52	76.31	81.82	72.11	78.46
STRES	TÜMÜ	76.00	78.70	76.31	81.82	72.01	78.46
	ÇİZME	71.33	75.94	76.38	85.45	69.71	77.31
	YAZMA	84.67	79.48	77.54	76.36	79.83	81.15

Tablo 13 ve Tablo 14'te görüldüğü üzere, PCA + GBC özellik seçimi öncesinde ve sonrasında yapılan veri artırma işlemleri, sınıflandırma performansında belirgin farklılıklar ortaya koymuştur.

PCA + GBC özellik seçimi öncesinde elde edilen sonuçlar, genel olarak tüm kategorilerde daha yüksek doğruluk ve F1 skoru sağlamış, özellikle yazma kategorisinde stres sınıflandırmasında %81,15 ve depresyon için %85,00 doğruluk ile en iyi performansı göstermiştir. Bu durum, PCA + GBC yönteminin pozitif ve negatif sınıfları ayırt etmede daha güçlü olduğunu göstermektedir.

Ancak PCA + GBC sonrası yapılan veri artırma analizinde, anksiyete sınıflandırması tüm kategorilerde %46,11 F1 skoru ve %53,08 doğrulukla pozitif sınıf performansı açısından düşük kalmıştır. Ayrıca, yazma kategorisinde depresyon sınıflandırması %59,62 doğruluk ve %29,20 F1 skoru ile pozitif sınıfları tanımlamada yetersiz kalmıştır.

3.6. Veri Artırma Olmadan Model Performansı

Tablo 15. PCA+GBC veri artırma olmadan, LGB + XGB + XGB modeli ile elde edilen sonuçlar

Duygular	Görevler	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	83.16	46.69	12.64	10.00	17.67	63.46
	ÇİZME	81.58	56.47	20.19	18.57	30.94	64.62
	YAZMA	82.11	48.46	15.40	12.86	20.92	63.46
ANKSİYETE	TÜMÜ	60.00	50.64	40.56	40.91	42.57	51.92
	ÇİZME	66.67	45.97	32.21	28.18	40.18	50.38
	YAZMA	62.67	53.82	44.69	44.55	47.42	55.00
STRES	TÜMÜ	66.67	50.55	37.80	34.55	44.17	53.08
	ÇİZME	61.33	45.48	29.57	26.36	34.97	46.54
	YAZMA	62.00	52.55	37.44	35.45	40.37	50.77

Tablo 14 ve Tablo 15 karşılaştırmalı olarak incelendiğinde, PCA ve GBC ile yapılan özellik seçimi öncesinde veri artırma uygulanmış senaryoların genel olarak daha yüksek doğruluk ve F1 skoru sağladığı gözlemlenmiştir. Tablo 14'te yer alan sonuçlara göre, veri artırma sonrası yazma görevinde depresyon sınıflandırmasında doğruluk oranı %85,00, F1 skoru ise %72,70 olarak gerçekleşmiştir. Oysa veri artırma olmadan elde edilen Tablo 15'te bu değerler sırasıyla %63,46 ve %15,40'a düşmektedir. Benzer şekilde, anksiyete sınıflandırmasında da çizme göreviyle veri artırma uygulandığında %78,85 doğruluk ve %76,64 F1 skoru elde edilirken, veri artırmamasız durumda doğruluk %50,38, F1 skoru %32,21 olarak kaydedilmiştir.

Stres sınıflandırması açısından da yazma görevinde veri artırmanın etkisi açıkça görülmektedir: F1 skoru %77,54'ten %37,44'e, doğruluk ise %81,15'ten %50,77'ye düşmüştür. Bu farklılıklar, veri artırma uygulamasının sınıflandırma başarımı üzerinde belirgin bir etkisi olduğunu, özellikle pozitif sınıfların tanınmasında modelin genel performansını iyileştirdiğini ortaya koymaktadır.

Tüm görevlerin birlikte değerlendirildiği genel durumda, veri artırma yapılan tabloda tüm duygular için doğruluk ve F1 skorları belirgin şekilde daha yüksek olup, örneğin depresyonda %81,15 doğruluk ve %60,19 F1 skoru elde edilmiştir. Buna karşın, veri artırma uygulanmayan durumda aynı kategoride doğruluk %63,46 ve F1 skoru %12,64 seviyesinde kalmıştır. Bu bulgular, veri artırma stratejisinin sınıflandırma modellerinin dengesini

iyileştirme ve özellikle azınlık sınıfların daha doğru tahmin edilmesini sağlama açısından önemli katkılar sunduğunu göstermektedir.

3.7. Temel Bileşen Analizi ve Gradyan Artıma Yaklaşımının Kullanımı

Tablo 16. PCA + GBC ile istifleme modellerinin sınıflandırma performansları

Duygular	Görevler	Modeller	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	LGB+XGB→XGB	91.05	68.01	60.19	54.29	72.89	81.15
		LGB+GB→XGB	86.32	68.53	63.44	62.86	65.64	80.00
		LGB+GB→RF	84.74	62.07	46.96	45.71	51.83	74.23
	ÇİZME	LGB+XGB→XGB	82.63	67.26	61.88	65.71	69.85	78.08
		LGB+GB→XGB	83.16	67.67	62.78	67.14	61.64	78.85
		LGB+GB→RF	86.32	64.70	52.87	52.86	61.73	77.31
	YAZMA	LGB+XGB→XGB	88.42	79.44	72.70	75.71	72.44	85.00
		LGB+GB→XGB	87.37	74.21	66.75	67.14	67.42	81.92
		LGB+GB→RF	87.37	71.80	64.20	64.29	66.10	81.15
ANKSİYETE	TÜMÜ	LGB+XGB→XGB	72.67	77.03	75.26	82.73	69.56	76.92
		LGB+GB→XGB	74.67	74.03	76.37	82.73	72.53	78.08
		LGB+GB→RF	72.67	73.67	69.54	74.55	66.36	73.46
	ÇİZME	LGB+XGB→XGB	76.67	77.45	76.64	81.82	73.43	78.85
		LGB+GB→XGB	76.67	75.09	76.09	80.91	73.09	78.46
		LGB+GB→RF	77.33	78.15	76.65	81.82	73.60	79.23
	YAZMA	LGB+XGB→XGB	76.00	76.52	76.31	81.82	72.11	78.46
		LGB+GB→XGB	77.33	77.88	78.07	83.64	73.80	80.00
		LGB+GB→RF	76.00	73.09	73.17	76.36	71.01	76.15
STRES	TÜMÜ	LGB+XGB→XGB	76.00	78.70	76.31	81.82	72.01	78.46
		LGB+GB→XGB	78.00	74.15	72.65	74.55	72.22	76.54
		LGB+GB→RF	77.33	76.97	73.76	76.36	71.94	76.92
	ÇİZME	LGB+XGB→XGB	71.33	75.94	76.38	85.45	69.71	77.31
		LGB+GB→XGB	79.33	78.55	77.52	80.91	74.86	80.00
		LGB+GB→RF	74.67	78.45	77.82	85.45	72.10	79.23
	YAZMA	LGB+XGB→XGB	84.67	79.48	77.54	76.36	79.83	81.15
		LGB+GB→XGB	80.00	78.09	77.35	80.00	75.33	80.00
		LGB+GB→RF	82.00	78.06	73.45	72.73	75.45	78.08

Tümü (hem yazma hem çizme) kategorisinde, depresyon sınıflandırmasında LGB + XGB → XGB kombinasyonu doğruluk (%81,15), özgüllük %91,05 ile negatif sınıfları başarılı bir şekilde ayırt ederken pozitif sınıfları tahmin etmede %54,29 hassasiyet, %60,19 F1 skoru) yetersiz kalmıştır. Anksiyete sınıflandırmasında, aynı model doğruluk %76,92, F1 skoru %75,26 ve hassasiyet %82,73 ile pozitif sınıfları yüksek doğrulukla tahmin ederken, aynı model stres sınıflandırmasında pozitif sınıfları oldukça başarılı bir şekilde tahmin ederken %81,82 hassasiyet, doğruluk %77,31, F1 skoru %76,31 ile model pozitif ve negatif sınıflar arasında dengeli bir performans sergilemiştir.

Çizme kategorisinde, depresyon sınıflandırmasında LGB + GB $\rightarrow$ XGB modeli %78,85 doğruluk, özgüllük %83,16 ve F1 skoru %62,78 ile negatif sınıfları ayırt etmede başarılı olmasına rağmen pozitif sınıfları tahmin etme başarısının sınırlı olduğunu ortaya koymaktadır. Anksiyete sınıflandırmasında ise aynı model doğruluk %78,46, F1 skoru %76,09 ve hassasiyet %80,91 ile hem pozitif hem de negatif sınıflar arasında dengeli bir performans göstermiştir. Stres sınıflandırmasında doğruluk (%77,31), F1 skoru (%75,94) ve hassasiyet %85,45 pozitif ve negatif sınıfları iyi bir şekilde tahmin edildiğini göstermektedir.

Yazma kategorisinde, depresyon sınıflandırmasında LGB + GB $\rightarrow$ XGB modeli hem pozitif %76,31 hassasiyet hem de negatif sınıfları %88,42 özgüllük dengeli bir şekilde ayırt edebilmiştir. F1 skoru %76,52 ve doğruluk oranı %88,42 modelin oldukça başarılı bir sınıflandırma sunduğunu göstermektedir. Anksiyete sınıflandırmasında ise aynı model hem pozitif %83,64 hassasiyet hem de negatif sınıfları %77,33 özgüllük iyi bir şekilde sınıflandırmıştır. F1 skoru %78,88 ve doğruluk oranı %80,00 modelin güçlü bir performans sergilediğini göstermektedir. Stres sınıflandırmasında LGB + XGB $\rightarrow$ RF modeli hem pozitif %80,00 hassasiyet hem de negatif sınıflar %80,00 özgüllük arasında dengeli bir şekilde sınıflandırmıştır. F1 skoru %77,35 ve doğruluk oranı %80,00 modelin pozitif ve negatif sınıflar arasında iyi bir performans sergilediğini göstermektedir.

3.8. Sadece Veri Artırma ile Model Performansının Değerlendirilmesi

Tablo 17: Özellik seçimi uygulanmadan, yalnızca veri artırma ile elde edilen model sonuçları

Duygular	Görevler	Özgüllük	AUC	F1 Skoru	Hassasiyet	Kesinlik	Doğruluk
DEPRESYON	TÜMÜ	91.58	54.66	16.36	11.43	33.33	70.00
	ÇİZME	90.53	52.97	24.22	18.57	42.00	71.15
	YAZMA	85.79	52.56	19.17	15.71	26.10	66.92
ANKSİYETE	TÜMÜ	66.00	55.15	39.28	36.36	43.65	53.46
	ÇİZME	66.67	47.79	40.07	38.18	44.09	54.62
	YAZMA	65.33	49.58	38.23	36.36	43.43	53.08
STRES	TÜMÜ	63.33	49.70	38.20	37.27	41.32	52.31
	ÇİZME	64.00	54.18	41.67	40.00	45.66	53.85
	YAZMA	57.33	47.61	38.01	37.27	39.89	48.85

Tablo 17’da yalnızca veri artırma (data augmentation) yöntemi uygulanarak yazma, çizme ve her iki işlem üzerinden depresyon, anksiyete ve stres sınıflarına etkisi değerlendirilmiştir. Depresyon sınıfında en yüksek doğruluk %71,15 ile çizme işleminde, en

düşük doğruluk ise %66,92 ile yazma işleminde elde edilmiştir. Bununla birlikte, F1 skoru ve hassasiyet değerlerinin düşük olması, özellikle veri artırma yöntemiyle dengesiz sınıfların hâlâ tam olarak optimize edilemediğini göstermektedir. Anksiyete sınıfında doğruluk değerleri %54,62 (çizme) ile %53,08 (yazma) arasında sınırlı bir performans göstermektedir. Stres sınıfında ise en düşük doğruluk %48,85 ile yazma işleminde kaydedilmiştir. Özellikle strese yönelik sonuçlar, diğer duygulara kıyasla daha düşük performans göstermiştir. Bu sonuçlar, yalnızca veri artırma yönteminin performans açısından sınırlı kaldığını ortaya koymaktadır.

Ayrıca, bu tezde her model birden fazla kez çalıştırılmış ve her koşumdan elde edilen precision, recall ve F1-score değerlerinin ortalaması alınmıştır. Bu yöntem, genel eğilimi yansıtmada pratik ve yaygın bir yaklaşımdır. Ancak bu hesaplama biçimi, bazı tablolarda F1 skorunun teorik beklentilerin (precision ve recall değerleriyle ilişkili) altında görünmesine neden olabilmektedir. Bu durum, metriklerin doğrudan ortalamasının alınmasından kaynaklanmakta olup, değerlendirme sürecine açıklık kazandırmak amacıyla burada belirtilmiştir.

4. SONUÇLAR

Yapılan analizlerde modellerin doğruluk oranları incelendiğinde, her duygu durumunun belirli kategorilerde diğerlerine kıyasla öne çıktığı görülmektedir. Modeller arasında LGB + GB → XGB ve LGB + XGB → XGB kombinasyonlarının en yüksek doğruluk oranlarına ulaşması, bu kombinasyonların sınıflandırma performansı üzerindeki etkisini göstermektedir.

Tablo 18. EMOTHAW veri seti kullanılarak yapılan çalışma sonuçlarının karşılaştırılması

Duygu	Görev	Doğruluk (%)						
		Likforman-Sulem vd. RF	Nolazco-Flores vd. SVM	Rahman & Halim BİLSTM	Khan-Xia vd. Transformer	Nolazco-Flores vd. AutoML	Nolazco-Flores vd. PCA-mFCBF	Önerilen Metot
Depresyon	Çizme	72.80	75.59	83.28	86.15	-	-	78.08
	Yazma	67.80	80.31	89.21	91.39	-	-	85.00
	Tümü	71.20	74.01	87.11	92.64	88.60	92.10	81.15
Anksiyete	Çizme	60.50	67.71	76.12	79.51	-	-	78.85
	Yazma	56.30	68.50	74.54	77.38	-	-	80.00
	Tümü	60.00	72.44	80.03	83.22	81.58	85.96	78.08
Stres	Çizme	60.10	67.71	75.39	78.76	-	-	80.00
	Yazma	51.20	67.71	75.17	79.41	-	-	81.15
	Tümü	60.20	70.07	74.38	78.05	81.58	88.59	78.46

Tablo 18, EMOTHAW veri seti kullanılarak gerçekleştirilen farklı çalışmaların doğruluk oranlarının önerilen yöntemle karşılaştırmasını ve yazma, çizme, tümü verilerinin duyguları sınıflandırma karşılaştırmalarını göstermektedir. Çizme, yazma ve tüm veri kategorilerinde depresyon, anksiyete ve stres sınıflandırmaları ayrı ayrı değerlendirilmiştir. Önerilen yöntem, birçok kategoride ve duygu sınıflandırmasında önceki çalışmalara kıyasla güçlü bir performans sergilemiştir. Yazma verileri, depresyon ve anksiyete sınıflarında yüksek doğruluk oranları (%85,0 ve %80,0) ile dikkat çekerken, stres sınıfında da güçlü bir performans %81,15 sergilemiştir. Çizme verileri, stres sınıflandırmasında %80,0 doğruluk oranıyla etkili olurken, depresyon sınıfında %78,85 ile yazma verilerinden bir miktar düşük kalmıştır. Tümü verileri ise her üç duygusal durum için daha dengeli bir performans sağlamış ve genel olarak (%78,08-%81,15) arasında doğruluk oranına ulaşmıştır. Çizme ve yazma verileri birleştirilerek oluşturulan veri kümesinde elde edilen performansın, her bir veri türünün tek başına elde ettiği en yüksek performanstan düşük çıkması beklenmedik gibi görünse de bu durum, farklı veri türlerinin birleşiminden kaynaklanan heterojen özellik etkileşimlerinden ve artan veri karmaşıklığından kaynaklanmış olabilir. Önerilen yöntem,

depresyon sınıflandırmasında çizme verilerinde %78,85, yazma verilerinde %85,0 ve tüm verilerde %81,15 doğruluk oranıyla, özellikle Rahman ve Halim'in çalışmasına yakın bir performans sergilemiştir. Anksiyete sınıfında yazma verilerinde %80,0 doğruluk oranı ile diğer yöntemlere kıyasla daha iyi sonuç bir sonuç elde eden önerilen yöntem, tüm verilerde %78,46 doğruluk oranıyla genel anlamda güçlü bir performans ortaya koymuştur. Stres sınıflandırmasında ise çizme ve yazma verilerinde sırasıyla %80,0 ve %81,15 doğruluk oranları diğer yöntemlere kıyasla öne çıkmıştır. Ayrıca, önerilen yöntemin diğer çalışmalarla rekabet edebilir bir doğruluk sağladığı ve bazı durumlarda bu yöntemleri geride bıraktığı ortaya konmuştur.

Bu analizler, artırma ve torbalama tabanlı model kombinasyonlarının ve kategori bazlı veri işleme yöntemlerinin, farklı duygu sınıflarını sınıflandırmada başarılı sonuçlar sergilediğini göstermektedir. Genel olarak, önerilen yöntem, rekabetçi bir sınıflandırma modeli sunmuş ve önceki çalışmalarla uyumlu sonuçlar elde ettiğini göstermiştir.

Çalışmanın literatürdeki benzer yaklaşımlar karşısındaki konumunu daha net ortaya koymak adına, benzer çalışma olan Nolzco-Flores vd. (2022) ile önerilen yöntemin farkları aşağıda özetlenmiştir:

Nolzco-Flores vd. çalışması veri dengesizliği problemini sınırlı düzeyde ele almıştır. Bu tezde, aşırı örnekleme (oversampling) ve az örnekleme (undersampling) teknikleri birlikte uygulanarak sınıflar dengelenmiştir.

İlgili çalışmada sınıflandırma süreci AutoML tabanlı bir yapı ile gerçekleştirilmiştir. Bu tezde ise istifleme modeli uygulanarak farklı taban öğrenicilerin (LGB, RF, XGB) çıktıları birleştirilmiş ve genel performans artırılmaya çalışılmıştır.

AutoML, model seçimi ve hiperparametre ayarlarını otomatikleştirmesiyle birçok kolaylık sunmaktadır. Ancak bu süreçte yapılan işlemler genellikle sınırlı açıklanabilirliğe sahiptir. Bu çalışmada ise model seçimi, öznitelik mühendisliği ve parametre aralıkları kullanıcı tarafından belirlenmiş, yorumlanabilir ve tekrarlanabilir bir yapı oluşturulmuştur.

Öznitelik seçimi sürecinde, ilgili çalışmada PCA ve mFCBF birlikte kullanılmışken; bu tezde PCA ile boyut indirgeme ardından gradyan artırma (Gradient Boosting Classifier, GBC) ile öznitelik sıralaması yapılmış,

AutoML ile parametre seçimi otomatik gerçekleştirilmişken; bu çalışmada ise her modelin hiperparametreleri Optuna kütüphanesiyle manuel olarak kontrol edilen bir hiperparametre optimizasyon süreci kullanılmıştır.

Bu alıřmada, performans karřılařtırmaları duygu t, grev t ve iřleme adımları (veri artırma/znitelil seimi ncesi-sonrası/znitelik seimi var-yok) zerinden detaylı analiz edilmiřtir. Literatrde bu derece ok ynl ve sistematik bir analiz sunan sınırlı sayıda alıřma bulunmaktadır.

Ayrıca bu tezde, EMOTHAW veri setinden tretilen temel zniteliklere ek olarak kalemin konum bilgisi, azimut, eėim, basın gibi sensr verilerinden ek znitelikler tretilmiř ve sınıflandırmaya dahil edilmiřtir.



5. ÖNERİLER

Bu çalışmada elde edilen bulgular doğrultusunda, önerilen yöntemlerin geliştirilmesi ve gelecekteki çalışmalar için aşağıdaki öneriler sunulabilir:

Öncelikle, veri artırma tekniklerinin farklı varyasyonlarının (örneğin , gürültü ekleme veya sentetik veri oluşturma yöntemleri) uygulanması, sınıflandırma performansının artırılmasına katkı sağlayabilir.

Bunun yanı sıra, farklı model kombinasyonlarının ve yöntemlerinin (örneğin, torbalama, artırma, derin öğrenme yöntemleri ve basit doğrusal modeller gibi) daha geniş bir şekilde ve alternatif istifleme mimarileri ile (örneğin, çok seviyeli istifleme veya model ağırlıklandırma teknikleri) test edilmesi, sonuçların genelleştirilebilirliğini güçlendirebilir.

Model seviyesinde, LSTM ve CNN gibi derin öğrenme yaklaşımlarının uygulanabilirliği incelenebilir. Bu bağlamda, hibrit model tasarımlarının veya çok seviyeli istifleme mimarilerinin kullanımıyla daha yüksek doğruluk oranları elde edilmesi mümkün olabilir. Bu iyileştirmeler, duygusal durumların sınıflandırılmasında daha yüksek doğruluk oranları elde edilmesine yardımcı olabilir.

Bu çalışmada frekans tabanlı özellikler X ve Y eksenleri için ayrı ayrı hesaplanmıştır. Ancak bu yöntem, sinyalin uzaydaki yönüne (örneğin kâğıdın döndürülmesi gibi) duyarlı olup analiz sonuçlarının değişmesine neden olabilmektedir. Gelecek çalışmalarda, X ve Y eksenlerine ait konum bilgilerinin birlikte değerlendirilerek tek boyutlu bir yapıya dönüştürülmesi önerilmektedir. Bu sayede, farklı yönelimlere karşı daha kararlı ve genellenebilir analiz sonuçları elde edilebilir.

Çizme ve yazma verilerinin birlikte kullanıldığı durumda, her iki veri türünün tek başına ulaştığı en yüksek performans değerlerinden daha düşük sonuçlar gözlemlenmiştir. Bu durum, heterojen veri kaynaklarının birleştirilmesinde mevcut yöntemlerin yetersiz kalabileceğini göstermekte olup, gelecekteki çalışmalarda heterojen veri birleşimlerinde daha gelişmiş özellik füzyonu ve entegrasyon yöntemlerinin kullanılması önerilmektedir.

6. KAYNAKLAR

- Ayzeren, Y. B., Erbilek, M., & Çelebi, E. (2019). Emotional state prediction from online handwriting and signature biometrics. *IEEE Access*, 7, 164759-164774.
- Abdi, H., & Williams, L. J. (2010). Principal component analysis. *Wiley interdisciplinary reviews: computational statistics*, 2(4), 433-459.
- Ahmed, G., Er, M. J., Fareed, M. M. S., Zikria, S., Mahmood, S., He, J., ... & Aslam, M. (2022). Dad-net: Classification of alzheimer's disease using adasyn oversampling technique and optimized neural network. *Molecules*, 27(20), 7085.
- Akiba, T., Sano, S., Yanase, T., Ohta, T., & Koyama, M. (2019). Optuna: A next-generation hyperparameter optimization framework. *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 2623–2631.
- Anguita, D., Ghelardoni, L., Ghio, A., Oneto, L., & Ridella, S. (2012). The K'in K-fold Cross Validation. In *ESANN* (Vol. 102, pp. 441-446).
- Bhattacharya, S., Islam, A., & Shahnawaz, S. (2022, March). TEemoDec: emotion detection from handwritten text using agglomerative clustering. In *2022 First International Conference on Artificial Intelligence Trends and Pattern Recognition (ICAITPR)* (pp. 1-6). IEEE.
- Beck, A. T., & Beamesderfer, A. (1974). Assessment of depression: The depression inventory. S. Karger.
- Baum, A. (1990). Stress, intrusive imagery, and chronic distress. *Health Psychology*, 9(6), 653.
- Barandela, R., Sánchez, J. S., García, V., & Rangel, E. (2003). Strategies for learning in class imbalance problems. *Pattern Recognition*, 36(3), 849-851.
- Batista, G. E., Prati, R. C., & 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.
- Breiman, L. (1996). Bagging predictors. *Machine Learning*, 24(2), 123–140.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
- Bergstra, J., & Bengio, Y. (2012). Random search for hyper-parameter optimization. *Journal of Machine Learning Research*, 13(1), 281–305.
- Christensen, A. V., Dixon, J. K., Juel, K., Ekholm, O., Rasmussen, T. B., Borregaard, B., ... & Berg, S. K. (2020). Psychometric properties of the Danish Hospital Anxiety and

- Depression Scale in patients with cardiac disease: results from the DenHeart survey. *Health and quality of life outcomes*, 18, 1-13.
- Chitlangia, A., & Malathi, G. (2019). Handwriting analysis based on histogram of oriented gradient for predicting personality traits using SVM. *Procedia computer science*, 165, 384-390.
- Cordasco, G., Scibelli, F., Faundez-Zanuy, M., Likforman-Sulem, L., & Esposito, A. (2019). Handwriting and drawing features for detecting negative moods. *Quantifying and Processing Biomedical and Behavioral Signals* 27, 73-86.
- Costello, C., & Comrey, A. L. (1967). Scales for measuring depression and anxiety. *The Journal of Psychology*, 66(2), 303–313.
- Christensen, A. V., Dixon, J. K., Juel, K., Ekholm, O., Rasmussen, T. B., Borregaard, B., Mols, R. E., Thrysøe, L., Thorup, C. B., & Berg, S. K. (2020). Psychometric properties of the Danish hospital anxiety and depression scale in patients with cardiac disease: Results from the DenHeart survey. *Health and Quality of Life Outcomes*, 18(1), 1–13.
- Crawford, J. R., & Henry, J. D. (2003). The depression anxiety stress scales (DASS): Normative data and latent structure in a large non-clinical sample. *British Journal of Clinical Psychology*, 42(2), 111–131.
- Cao, L. J., Chua, K. S., Chong, W. K., Lee, H. P., & Gu, Q. M. (2003). A comparison of PCA, KPCA and ICA for dimensionality reduction in support vector machine. *Neurocomputing*, 55(1-2), 321-336.
- Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: Synthetic minority over-sampling technique. *Journal of Artificial Intelligence Research*, 16, 321–357.
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785–794.
- Drotár, P., Mekyska, J., Rektorová, I., Masarová, L., Smékal, Z., & Faundez-Zanuy, M. (2014). Analysis of in-air movement in handwriting: A novel marker for Parkinson's disease. *Computer Methods and Programs in Biomedicine*, 117(3), 405–411.
- Dash, M., & Liu, H. (1997). Feature selection for classification. *Intelligent Data Analysis*, 1(1-4), 131–156.
- Devi, D., & Purkayastha, B. (2017). Redundancy-driven modified Tomek-link based undersampling: A solution to class imbalance. *Pattern Recognition Letters*, 93, 3-12.
- Dietterich, T. G. (2000). Ensemble methods in machine learning. In *International workshop on multiple classifier systems* (pp. 1–15). Springer.

- De Stefano, C., Fontanella, F., Impedovo, D., Pirlo, G., & Di Freca, A. S. (2019). Handwriting analysis to support neurodegenerative diseases diagnosis: A review. *Pattern Recognition Letters*, 121, 37–45.
- Davis, S., & Mermelstein, P. (1980). Comparison of parametric representations for monosyllabic word recognition in continuously spoken sentences. *IEEE transactions on acoustics, speech, and signal processing*, 28(4), 357-366.
- Elhassan, T., & Aljurf, M. (2016). Classification of imbalance data using tomes link (t-link) combined with random under-sampling (rus) as a data reduction method. *Global J Technol Optim S*, 1, 2016.
- Esfahani, M. S., & Dougherty, E. R. (2014). Effect of separate sampling on classification accuracy. *Bioinformatics*, 30(2), 242–250.
- Efron, B., & Tibshirani, R. J. (1994). *An introduction to the bootstrap*. Chapman and Hall/CRC.
- Faundez-Zanuy, M., Fierrez, J., Ferrer, M. A., Diaz, M., Tolosana, R., & Plamondon, R. (2020). Handwriting biometrics: Applications and future trends in e-security and e-health. *Cognitive Computation*, 12(5), 940–953.
- Faundez-Zanuy, M., Hussain, A., Mekyska, J., Sesa-Nogueras, E., Monte-Moreno, E., Esposito, A., ... & Lopez-de-Ipiña, K. (2013). Biometric applications related to human beings: there is life beyond security. *Cognitive Computation*, 5, 136-151.
- Folstein, M. F., Folstein, S. E., & McHugh, P. R. (1975). “Mini-mental state”: a practical method for grading the cognitive state of patients for the clinician. *Journal of psychiatric research*, 12(3), 189-198.
- Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *Annals of Statistics*, 29(5), 1189–1232.
- Greco, C., Raimo, G., Amorese, T., Cuciniello, M., Mcconvey, G., Cordasco, G., ... & Esposito, A. (2023). Discriminative power of handwriting and drawing features in depression. *International Journal of Neural Systems*. 2024; 34 (2): 2350069.
- Guyon, I., Gunn, S., Nikravesh, M., & Zadeh, L. A. (2008). *Feature extraction: Foundations and applications* (Vol. 207). Springer.
- Guyon, I., & Elisseeff, A. (2003). An introduction to variable and feature selection. *Journal of Machine Learning Research*, 3(Mar), 1157–1182.
- Güzel, K. (2020, December 28). *Boosting nedir? Adım adım AdaBoost algoritması*. Medium. <https://kadirguzel.medium.com/boosting-nedir-adim-adim-adaboost-algoritması-439cce20ab9a>

- He, H., & Ma, Y. (2013). *Imbalanced learning: Foundations, algorithms, and applications*.
- Harrington, P. D. B., Vieira, N. E., Espinoza, J., Nien, J. K., Romero, R., & Yergey, A. L. (2005). Analysis of variance–principal component analysis: A soft tool for proteomic discovery. *Analytica chimica acta*, 544(1-2), 118-127.
- He, H., Bai, Y., Garcia, E. A., & Li, S. (2008, June). ADASYN: Adaptive synthetic sampling approach for imbalanced learning. In *2008 IEEE international joint conference on neural networks (IEEE world congress on computational intelligence)* (pp. 1322-1328). IEEE.
- Hosmer Jr, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). *Applied logistic regression*. John Wiley & Sons.
- Impedovo, D. (2019). Velocity-based signal features for the assessment of parkinsonian handwriting. *IEEE Signal Processing Letters*, 26(4), 632–636.
- Jolliffe, I. T., & Cadima, J. (2016). Principal component analysis: a review and recent developments. *Philosophical transactions of the royal society A: Mathematical, Physical and Engineering Sciences*, 374(2065), 20150202.
- Khan, Z. A., Xia, Y., Aurangzeb, K., Khaliq, F., Alam, M., Khan, J. A., & Anwar, M. S. (2024). Emotion detection from handwriting and drawing samples using an attention-based transformer model. *PeerJ Computer Science*, 10, e1887.
- Khalid, S., Khalil, T., & Nasreen, S. (2014). A survey of feature selection and feature extraction techniques in machine learning. In *2014 Science and Information Conference* (pp. 372–378). IEEE.
- Kline, P. (2013). *Handbook of psychological testing*. Routledge.
- Kim, H. (2013). *The ClockMe system: computer-assisted screening tool for dementia* (Doctoral dissertation, Georgia Institute of Technology).
- Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., ... & Liu, T.-Y. (2017). LightGBM: A highly efficient gradient boosting decision tree. *Advances in Neural Information Processing Systems*, 30, 3146–3154.
- Likforman-Sulem, L., Esposito, A., Faundez-Zanuy, M., Cléménçon, S., & Cordasco, G. (2017). EMOTHAW: A novel database for emotional state recognition from handwriting and drawing. *IEEE Transactions on Human-Machine Systems*, 47(2), 273-284.
- Lovibond, P. F., & Lovibond, S. H. (1995). The structure of negative emotional states: Comparison of the depression anxiety stress scales (DASS) with the Beck depression and anxiety inventories. *Behaviour Research and Therapy*, 33(3), 335–343.

- Mohammed, R., Rawashdeh, J., & Abdullah, M. (2020). Machine learning with oversampling and undersampling techniques: Overview study and experimental results. In *2020 11th International Conference on Information and Communication Systems (ICICS)* (pp. 243–248). IEEE.
- Nolazco-Flores, J. A., Faundez-Zanuy, M., Velázquez-Flores, O. A., Cordasco, G., & Esposito, A. (2021). Emotional state recognition performance improvement on a handwriting and drawing task. *IEEE Access*, 9, 28496-28504.
- Nolazco-Flores, J. A., Faundez-Zanuy, M., Velázquez-Flores, O. A., Del-Valle-Soto, C., Cordasco, G., & Esposito, A. (2022). Mood state detection in handwritten tasks using PCA–mFCBF and automated machine learning. *Sensors*, 22(4), 1686.
- Nolazco-Flores, J. A., Faundez-Zanuy, M., De La Cueva, V. M., & Mekyska, J. (2021). Exploiting spectral and cepstral handwriting features on diagnosing Parkinson's disease. *IEEE Access*, 9, 1–1.
- Natekin, A., & Knoll, A. (2013). Gradient boosting machines, a tutorial. *Frontiers in Neurorobotics*, 7, 21.
- Partridge, M., & Calvo, R. (1997). Fast dimensionality reduction and simple PCA. *Intelligent data analysis*, 2(3), 292-298.
- Pereira, R. M., Costa, Y. M., & Silla Jr, C. N. (2020). MLTL: A multi-label approach for the Tomek Link undersampling algorithm. *Neurocomputing*, 383, 95-105.
- Polikar, R. (2006). Ensemble based systems in decision making. *IEEE Circuits and Systems Magazine*, 6(3), 21–45.
- Peng, C. Y. J., Lee, K. L., & Ingersoll, G. M. (2002). An introduction to logistic regression analysis and reporting. *The journal of educational research*, 96(1), 3-14.
- Rahman, A. U., & Halim, Z. (2022). Predicting the big five personality traits from handwritten text features through semi-supervised learning. *Multimedia Tools and Applications*, 81(23), 33671-33687.
- Rahman, A. U., & Halim, Z. (2023). Identifying dominant emotional state using handwriting and drawing samples by fusing features. *Applied Intelligence*, 53(3), 2798-2814.
- Robnik-Šikonja, M. (2004). Improving random forests. In *Proceedings of the European Conference on Machine Learning* (pp. 359-370).
- Schlenker, B. R., & Leary, M. R. (1982). Social anxiety and self-presentation: A conceptualization model. *Psychological bulletin*, 92(3), 641.
- Spielberg, C. D. (2010). State-trait anxiety inventory. *The Corsini Encyclopedia of Psychology*, 1–1.
- Selye, H. (1956). *The stress of life*. McGraw-Hill.

- Shaw, B., & Jebara, T. (2009). Structure preserving embedding. In *Proceedings of the 26th Annual International Conference on Machine Learning* (pp. 937–944).
- Sammut, C., & Webb, G. I. (2011). *Encyclopedia of machine learning*. Springer Science & Business Media.
- Seera, M., & Lim, C. P. (2014). A hybrid intelligent system for medical data classification. *Expert Systems with Applications*, 41(5), 2239–2249.
- Schapire, R. E. (1990). The strength of weak learnability. *Machine Learning*, 5(2), 197–227.
- Sagi, O., & Rokach, L. (2018). Ensemble learning: A survey. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 8(4), e1249.
- Singh, A., Thakur, N., & Sharma, A. (2016). A review of supervised machine learning algorithms. *2016 3rd International Conference on Computing for Sustainable Global Development (INDIACom)*, New Delhi, India, 1310-1315.
- Turner, A. I., Smyth, N., Hall, S. J., Torres, S. J., Hussein, M., Jayasinghe, S. U., ... & Clow, A. J. (2020). Psychological stress reactivity and future health and disease outcomes: A systematic review of prospective evidence. *Psychoneuroendocrinology*, 114, 104599.
- Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society: Series B (Methodological)*, 58(1), 267–288.
- Ting, K. M., & Witten, I. H. (1997). Stacking bagged and dagged models. *Proceedings of the Fourteenth International Conference on Machine Learning*, 367–375.
- Tasin, I., Nabil, T. U., Islam, S., & Khan, R. (2023). Diabetes prediction using machine learning and explainable AI techniques. *Healthcare Technology Letters*, 10(1-2), 1-10.
- Van Der Maaten, L., Postma, E. O., & Van Den Herik, H. J. (2009). Dimensionality reduction: A comparative review. *Journal of machine learning research*, 10(66-71), 13.
- Van der Laan, M. J., Polley, E. C., & Hubbard, A. E. (2007). Super learner. *Statistical applications in genetics and molecular biology*, 6(1).
- Wolpert, D. H. (1992). Stacked generalization. *Neural Networks*, 5(2), 241–259.
- Yu, L., & Liu, H. (2003). Feature selection for high-dimensional data: A fast correlation-based filter solution. In *Proceedings of the 20th International Conference on Machine Learning (ICML-03)* (pp. 58–63).
- Zhou, Z. H. (2012). *Ensemble methods: foundations and algorithms*. CRC press.
- Zhang, Z., & Castelló, A. (2017). Principal components analysis in clinical studies. *Annals of translational medicine*, 5(17).

ÖZGEÇMİŞ

Semra BAYRAK, İlköğrenimini bitirdikten sonra lise eğitimini 2011 yılında Gümüşhane Lisesi'nde tamamladı. 2011-2012 öğretim yılında Karadeniz Teknik Üniversitesi Mühendislik Fakültesi Bilgisayar Mühendisliği bölümünde lisans eğitimine başladı. 2019 yılında bu bölümden mezun oldu. 2021 yılında Karadeniz Teknik Üniversitesi, Fen Bilimleri Enstitüsü, Bilgisayar Bilimleri Anabilim dalında tezli yüksek lisans programına başladı. İyi derece İngilizce bilmektir.

