



Article Info/Makale Bilgisi

✓Received/Geliş:01.09.2025 ✓Accepted/Kabul:16.10.2025

DOI:<https://doi.org/10.30794/pausbed.1775480>

Research Article/Araştırma Makalesi

Kocarik Gacar, B. (2025). "EEG Tabanlı Nöbet Sınıflandırması: PCA'nın Etkisi ve Lojistik Regresyon, SVM, XGBoost Karşılaştırması", *Pamukkale Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, Sayı:71 EYS'25 Özel Sayısı, 641-654.

EEG TABANLI NÖBET SINIFLANDIRMASI: PCA'NIN ETKİSİ VE LOJİSTİK REGRESYON, SVM, XGBOOST KARŞILAŞTIRMASI*

Burcu KOCARIK GACAR*

Öz

Epilepsi, bireylerin sosyal ilişkilerini, toplumsal uyumunu ve yaşam kalitesini olumsuz yönde etkileyen bir hastalıktır. Bu çalışmada, çok öznitelikli (16 kanal), çok sınıflı (sağlıklı, jeneralize nöbet, fokal nöbet, nöbet aktivitesi) ve dengeli Bangalore Epilepsi veri kümesi üzerinde Lojistik Regresyon, Destek Vektör Makineleri (SVM) ve Aşırı Gradyan Artırma (XGBoost) modelleri; Temel Bileşenler Analizi (PCA) ile ve PCA'sız iki işlem hattında değerlendirilmiştir. Doğruluk-çalışma süresi dengesi bağlamında PCA'nın marjinal katkısı nicel olarak gösterilmiş; sınıflar arası ayrışmanın zorlaştığı durumlarda hangi modelin daha tutarlı ve etkin performans gösterdiği incelenmiştir. Bulgular, öznitelik korelasyonlarının yüksek olduğu veri kümelerinde XGBoost'un doğruluk, F1, ROC-AUC açısından SVM ve lojistik regresyona kıyasla daha iyi performans sergilediğine işaret etmektedir. Bu durum epileptik nöbet tespitinde PCA'nın her durumda model başarımını artırmadığını; dolayısıyla veri ön-işleme stratejilerinin model tipine göre dikkatle seçilmesi gerektiğini göstermektedir. Elde edilen sonuçlar, hem bilişsel sinyal işleme literatürüne metodolojik bir karşılaştırma katkısı sunmakta hem de model seçimi konusunda uygulayıcılara yol göstermektedir. Standart metrikler ile on katlı çapraz doğrulama ve farklı oranlarda ayrılmış test kümelerine dayalı bir kıyaslama sunması çalışmanın başlıca katkılarındandır.

Anahtar kelimeler: *Epileptik Nöbet Sınıflandırma, Bangalore Epilepsi verisi, Temel Bileşenler Analizi, Makine Öğrenmesi Teknikleri.*

EEG-BASED SEIZURE CLASSIFICATION: THE EFFECT OF PCA AND A COMPARISON OF LOGISTIC REGRESSION, SVM, AND XGBOOST

Abstract

Epilepsy is a disorder that adversely affects individuals' social relationships, societal integration, and overall quality of life. In this study, Logistic Regression, Support Vector Machines (SVM), and XGBoost models were evaluated on the multivariate (16-channel), multiclass (healthy, generalized seizure, focal seizure, seizure activity), and balanced Bangalore EEG Epilepsy Dataset using two processing pipelines—with and without Principal Component Analysis (PCA). The marginal contribution of PCA was quantitatively examined in terms of the accuracy–runtime trade-off, and the consistency and efficiency of each model were analyzed in cases where class separation became more challenging. The findings indicate that, in datasets with high feature correlations, XGBoost outperformed SVM and Logistic Regression in terms of accuracy, F1, and ROC-AUC metrics. This suggests that PCA does not necessarily improve model performance in epileptic seizure detection, emphasizing that data pre-processing strategies should be carefully chosen according to the model type. The results provide a methodological comparison that contributes to the cognitive signal processing literature and offer guidance for practitioners in model selection. Presenting a benchmark based on standard metrics, ten-fold cross-validation, and test sets with different split ratios constitutes one of the main contributions of this study.

Keywords: *Epileptic Seizure Classification, BEED, Principal Component Analysis, Machine Learning.*

*Bu çalışma, EYS'25 Yapay Zeka modülünde yer alan Python Destekli Denetimli Makine Öğrenmesi dersi kapsamında öğretilen teknikler kullanılarak hazırlanmıştır.

**Dr. Öğr. Üyesi, Manisa Celal Bayar Üniversitesi, İktisadi ve İdari Bilimler Fakültesi, Ekonometri Bölümü, İstatistik Ana Bilim Dalı, MANİSA. e-posta: burcu.kocarik@cbu.edu.tr, (<https://orcid.org/0000-0001-5944-4456>)

1. GİRİŞ

Epilepsi tanısında EEG (Elektroensefalografi) temelli otomatik nöbet sınıflandırması, tanı doğruluğunu artırması ve gerçek zamanlı klinik karar desteğini mümkün kılması bakımından kritik önemdedir. Derin ve klasik makine öğrenmesindeki son gelişmeler, EEG sinyallerinin daha isabetli ve verimli biçimde işlenmesini sağlayarak nöbet başlangıcı örüntülerine ilişkin içgörülerini derinleştirmiş ve sınıflandırma başarımını artırmıştır. Bununla birlikte, farklı veri kümeleri ve hasta demografileri arasında genellenebilirlik hala temel bir güçlük olarak sürmektedir. Bu nedenle literatürde, modellerin sağlamlığını artırmak amacıyla çeşitli ön işleme, öznelilik çıkarımı ve sınıflandırma stratejileri sistematik biçimde incelenmektedir.

Son yıllarda EEG tabanlı nöbet tespiti çalışmalarında öznelilik uzayının boyutu, sinyal gürültüsü ve kanallar arası yüksek korelasyon gibi sorunlar, sınıflandırma modellerinin performansını doğrudan etkileyen unsurlar olarak öne çıkmaktadır. Bu kapsamda, boyut indirgeme yöntemleri- özellikle PCA - hem hesaplama yükünü azaltmak hem de çoklu doğrusal bağlantı problemini gidermek için sıklıkla kullanılmaktadır. PCA'nın her durumda sınıflandırma başarımını artırıp artırmadığına dair bulgular tutarsızdır; bazı çalışmalar PCA'nın model performansını iyileştirdiğini bildirirken, bazıları önemli bilgi kayıplarına yol açabileceğini göstermektedir. Bu sebeple, EEG verilerinde PCA'nın farklı makine öğrenmesi algoritmaları üzerindeki etkisinin sistematik ve karşılaştırmalı biçimde değerlendirilmesinin literatürde devam eden bir ihtiyaç olduğu görülmektedir. Özellikle Lojistik Regresyon, SVM ve XGBoost gibi klasik modellerin PCA ön işleme adımıyla olan etkileşimini doğrudan karşılaştıran çalışmalar sınırlıdır. Mevcut literatürdeki bu boşluk, EEG tabanlı nöbet sınıflandırmasında yöntemsel optimizasyonu güçleştirmekte ve model seçimi konusunda standart bir çerçevenin oluşmasını engellemektedir.

Bu çalışma, söz konusu araştırma boşluğunu doldurmayı amaçlamakta; PCA uygulanan ve uygulanmayan iki farklı işlem hattı altında, çeşitli makine öğrenmesi modellerinin performansını karşılaştırmalı olarak değerlendirmektedir. Böylece, EEG sinyallerinin sınıflandırılmasında hangi modelin hangi koşullarda daha yüksek doğruluk ve tutarlılık sergilediği nicel olarak ortaya konmaktadır. Çalışmanın özgünlüğü, PCA'nın etkisini yalnızca bir ön işleme adımı olarak değil, modelin genel başarımına olan marjinal katkısı üzerinden değerlendirmesinde yatmaktadır. Sonuç olarak, bu araştırma EEG tabanlı nöbet sınıflandırmasında model seçimi, öznelilik boyutu yönetimi ve veri ön işleme stratejilerinin birlikte ele alınmasına yönelik bütüncül bir yaklaşım önermekte hem metodolojik hem de uygulamalı literatüre katkı sağlamaktadır.

2. LİTERATÜR ARAŞTIRMASI

Son yıllarda EEG tabanlı epileptik nöbet sınıflandırması alanında hem klasik makine öğrenmesi yöntemleri (lojistik regresyon, k-NN, SVM, karar ağaçları, rastgele orman, yapay sinir ağları, XGBoost) hem de derin öğrenme yaklaşımları (CNN – Convolutional Neural Network (Evrimsel Sinir Ağı) tabanlı mimariler; RNN – Recurrent Neural Network (Yinelemeli Sinir Ağı) / GRU – Gated Recurrent Unit (Kapalı Yinelemeli Birim) / LSTM – Long Short-Term Memory (Uzun Kısa Dönem Bellek Ağı) ve Transformer (Dönüştürücü Mimarisi) tabanlı modeller) yoğun biçimde araştırılmaktadır. Bu bölümde, özellikle çoklu doğrusal bağlantı, yüksek boyut ve ön işleme/boyut indirgeme gereksinimleri bağlamında, yapay zekâ destekli sınıflandırmaya yönelik çalışmalardan öne çıkanlar sistematik olarak özetlenmiştir.

Wei ve Mooney (2020), TUH-EEG Corpus'tan EEG'lerde nöbetleri tespit etmek için XGBoost tabanlı bir yöntem kullanmıştır. Yöntemi eğitmek için 4597 EEG dosyasından 64 öznelilik seçilmiştir. SMOTE (Sentetik Azınlık Aşırı Örneklemme Tekniği) kullanılmıştır. Önerilen XGBoost tabanlı yöntem, test setinde %20 duyarlılık ve 24 saatte 15.59 yanlış alarm elde etmiştir. Sonuç olarak yöntem klinik EEG kayıtlarındaki nöbetleri otomatik olarak analiz etmeye yardımcı olma potansiyeline sahip bulunmuştur. Wu vd. (2020), epileptik nöbetlerin doğru bir şekilde tespiti için CEEMD (Complete Ensemble Empirical Mode Decomposition) ve XGBoost'u entegre eden CEEMD-XGBoost adlı otomatik tespit yaklaşımını önermiştir. Bonn ve CHB-MIT (Boston Çocuk Hastanesi-Massachusetts Teknoloji Enstitüsü) olarak bilinen iki referans epilepsi EEG veri kümesi üzerinde deneyler yürütülmüş; sonuçlar, CEEMD-XGBoost'un duyarlılık, özgüllük ve doğruluk açısından nöbet tespitini iyileştirdiğini göstermiştir. Rahman vd. (2021), meme kanserinin erken tanısında düzenli kan analizleri ve antropometrik ölçümlerden elde edilebilecek biyobelirteçleri belirlemeyi ve bilgisayar destekli tanı (CAD) sistemlerinin performansını iyileştirmeyi amaçlamıştır. Dokuz antropometrik ve klinik özellik içeren güncel bir veri kümesiyle standart makine öğrenmesi modelleri eğitilmiştir. RBF çekirdeğine sahip SVM en başarılı sınıflandırıcı olmuştur. Sınıflandırma doğruluğu, duyarlılık ve özgüllük sırasıyla %93,9, %95,1 ve %94,0 olarak bulunmuştur. Chan vd. (2022), doğrusal bağlantıyı indirmek için öznelilik seçim yöntemleri ve düzenlenmiş tahmin yöntemlerini kullanmıştır. Çalışmada, değişken seçim yöntemlerinin bazı değişkenleri modelden çıkarması nedeniyle veri toplama çabalarını boşa çıkarabileceği; bu nedenle çoklu doğrusal bağlantı içeren verilerin, istatistiksel tahmin edicilere kıyasla makine öğrenmesine dayalı optimizasyon yaklaşımlarıyla daha iyi ele alınabildiği belirtilmiştir. Islam vd. (2022), epileptik nöbetleri tespit

etmek için derin öğrenme modeli (Epileptic-Net) kullanan dinamik bir yöntem önermiştir. Bonn Üniversitesi EEG veri kümesi üzerinde önerilen Epileptic-Net, iki sınıflı sınıflandırmada %99,95, üç sınıflı %99,98, dört sınıflı %99,96 ve beş sınıflı karmaşık problemi sınıflandırmada %99,96 doğruluk sağlamıştır. Sonuç olarak, önerilen yaklaşımın epileptik nöbetleri tespit etmede mevcut modellere göre rekabetçi sonuçlar verdiği görülmüştür.

Diler ve Demir (2024), Lojistik Regresyon, Naive Bayes, k-NN, SVM ve XGBoost algoritmalarını kullanmıştır. Çalışmada koşul indeksi ile çoklu doğrusal bağlantının derecesi belirlenmiştir. Benzetim çalışmaları ile performanslar gözlenmiş sonuçlar orjinal veri kümeleri ile karşılaştırılmıştır. Sonuç olarak, çoklu doğrusal bağlantı varlığında XGBoost algoritmasının en yüksek, Naive Bayes algoritmasının ise en düşük performansa sahip olduğu gözlenmiştir. Kong vd. (2024), Ayrık Dalgacık Dönüşümü (DWT) ve Welch tabanlı, yüksek korelasyon gösteren 35 öznitelik üzerinde Parçacık Sürü Optimizasyonu (PSO) ve Pearson korelasyon analizini bir araya getirmiştir. Epileptik nöbet tespit modelleri oluşturmak için SVM, Yapay Sinir Ağları (YSA), Rastgele Orman (RF) ve XGBoost sınıflandırıcıları kullanılmıştır. Sonuç olarak önerilen yöntem %99,32 doğruluk, %99,64 özgüllük ve %99,29 duyarlılık elde etmiştir. Berrich ve Guennoun (2025), CNN ve Derin Sinir Ağları (DNN) ile SVM'nin entegrasyonunu, PCA'nın öznitelik boyutunu sadeleştirmedeki rolüne özellikle vurgu yaparak incelemiştir. Sağlamlığın değerlendirilmesi amacıyla Epileptic Seizure Recognition ve BONN veri kümeleri karşılaştırılmıştır. Çalışmada CNN-SVM-PCA modelinin doğruluğunun Epileptic Seizure Recognition veri kümesinde %99,42'ye, BONN veri kümesinde ise %99,96'ya ulaştığı raporlanmıştır. Ayrıca PCA'nın DNN-SVM modeline entegrasyonu, doğrulukta Epileptic Seizure Recognition için %0.36, BONN için %3,07 artış sağlamıştır. Li vd. (2025), Evrişimli Sinir Ağı (CNN) ile Vision Transformer (ViT)'i birleştiren Çok Akışlı Öznitelik Füzyonu (MSFF) stratejisine dayalı CMFViT adlı epilepsi tespit modelini önermiştir. Modelin etkinliği, CHB-MIT ve Kaggle'da yer alan 121 katılımcılı epilepsi veri kümesi üzerinde yapılan deneysel değerlendirmelerle doğrulanmıştır. CMFViT, CHB-MIT üzerinde özne-içi deneylerde %98,85 doğruluk ve diğer metriklerde de yüksek değerler elde etmiş; Kaggle veri kümesinde öznelere arası deneylerde de güçlü performans sergilemiştir. Najmusseher ve Banu (2025), Hızlı Fourier Dönüşümü (FFT) ile Uniform Manifold Yaklaştırması ve Projeksiyonu (UMAP) kullanarak spektral ve zamansal özniteliklerin birleştirilmesini ve etkili nöbet sınıflandırmasını sağlamak üzere, makine öğrenmesi ile derin öğrenmeyi entegre eden Sıralı Güçlendirme Ağı (SeqBoostNet)'i tanıtmıştır. UCI Machine Learning Repository üzerinden temin edilen BONN ve BEED veri kümeleri üzerinde karşılaştırmalı olarak doğrulanan yaklaşım, fokal ve jeneralize nöbet başlangıçlarını %95,91 doğrulukla ayırt etmiş; ayrıca BEED'de %96,71 ve BONN'da %97,11 ortalama doğruluk sağlamıştır. Carvajal-Dossman vd. (2025), derin öğrenme ağları dahil olmak üzere çok sayıda modelin EEG'lerden nöbet tespitindeki doğruluğunu araştırmıştır. Karşılaştırmalı değerlendirme, eğitim ve ilk test için üç farklı veri kümesini ve dış doğrulama için yerel bir hastadan manuel olarak kaydedilmiş EEG'yi içermiştir. RF ve CNN en iyi sonuçları elde etmiştir.

Literatürdeki önceki çalışmaların önemli bir bölümü, EEG tabanlı nöbet sınıflandırmasında yalnızca tek bir model veya derin öğrenme tabanlı mimariler üzerinde yoğunlaşmış, klasik makine öğrenmesi yöntemlerini sistematik biçimde karşılaştırmamıştır. Ayrıca PCA'nın bu modeller üzerindeki etkisi genellikle tek bir algoritma özelinde incelenmiş; farklı sınıflandırıcılar arasında performans farkları bütüncül biçimde değerlendirilmemiştir. Bu çalışmada, çoklu doğrusal bağlantı içeren EEG öznitelikleri altında klasik makine öğrenmesi sınıflandırıcılarının performanslarını karşılaştırmaktadır. Lojistik Regresyon, SVM ve XGBoost'un PCA'lı ve PCA'sız denemelerinin (PCA-lojistik regresyon, PCA-SVM, PCA-XGBoost ve karşılıkları) epileptik nöbet örüntülerini doğru belirlemedeki başarımları değerlendirilmiş; böylece literatüre nicel ve karşılaştırılabilir bulgular sunulmuştur. Bu çalışma, PCA'nın klasik sınıflandırma modelleri üzerindeki marjinal katkısını doğrudan karşılaştırarak, model başarısındaki varyasyonları istatistiksel olarak ortaya koyması bakımından literatürdeki boşluğu doldurmaktadır. Özellikle farklı test oranları ve çapraz doğrulama yaklaşımlarının birlikte değerlendirilmesi, PCA'nın etkisine ilişkin daha güvenilir ve genellenebilir sonuçlar sağlamıştır. Böylece, EEG tabanlı nöbet sınıflandırmalarında model seçimi ve ön işleme stratejilerinin etkileşimini açıklayan, nicel olarak karşılaştırılabilir bir çerçeve önerilmiştir.

3. YÖNTEM

Bu bölümde, EEG temelli epileptik nöbet sınıflandırmasında kullanılan veri seti, ön işleme adımları ve uygulanan makine öğrenmesi yöntemleri ayrıntılı biçimde açıklanmaktadır. Çalışmanın temel amacı, çoklu doğrusal bağlantı koşulunda klasik sınıflandırma algoritmalarının tutarlılık ve etkinliğini değerlendirmek; farklı model türlerinin PCA uygulaması altında nasıl bir performans farkı gösterdiğini sistematik biçimde incelemektir. Bu amaç doğrultusunda, literatürde sıkça kullanılan ancak farklı veri yapılarında tutarsız sonuçlar üretebilen PCA'nın model başarımına katkısı, çeşitli algoritmalar üzerinden karşılaştırmalı olarak analiz edilmiştir. Seçilen yöntemlerin hem doğrusal hem doğrusal olmayan ilişkileri temsil edebilme özellikleriyle, çalışmanın araştırma sorusuna doğrudan yanıt verecek nitelikte olduğu düşünülmektedir.

Bu çalışmada, veri kalitesini olumsuz etkileyen çoklu doğrusal bağlantı (multicollinearity) koşulunda sınıflandırma algoritmalarının performansını incelemeyi temel motivasyon olarak benimsenmekte; söz konusu bağımlılığın varlığında yöntemlerin tutarlılık ve etkinliğini nicel olarak karşılaştırmak amaçlanmaktadır. Bu amaçla, 16 kanallı, dört sınıflı (sağlıklı, jeneralize nöbet, fokal nöbet, nöbet aktivitesi) ve dengeli BEED veri kümesi üzerinde makine öğrenmesi yöntemlerinden Lojistik Regresyon, SVM (RBF) ve XGBoost modelleri PCA'lı ve PCA'sız iki ayrı işlem hattında (pipeline) sınanmıştır. Bu algoritmaların seçilmesindeki temel sebep, farklı veri yapılarına uygunlukları ve regresyon problemlerindeki tahmin performanslarının literatürde sıklıkla kanıtlanmış olmasıdır. Kullanılan BEED veri kümesine UCI Machine Learning Repository üzerinden erişilmiştir. Çalışmada gerçekleştirilen analizler Python programlama dili ve scikit-learn kütüphanesi kullanılarak yapılmıştır.

Doğrulama düzeninde eğitim verisi üzerinde on katlı çapraz doğrulama ve %30, %20, %10 test paylarına sahip üç ayrı test senaryosu kullanılmıştır. Tüm ön işleme adımları (ölçekleme, PCA) yalnızca eğitim katlarında öğrenilip test dilimlerine uygulanmıştır. Bu çerçevede, PCA'nın eklenmesi/çıkarılmasının model performansı ve kaynak kullanımı üzerindeki etkisi nicel olarak değerlendirilmiş; doğruluk, F1, ROC-AUC ve çalışma süresi ölçütleri raporlanmıştır.

3.1. Temel Bileşenler Analizi

Makine öğrenmesi yöntemlerinde çok sayıda öznelik hesaplama maliyetini artırmakta, aşırı öğrenme (overfitting) riskini yükseltmekte ve çoklu doğrusal bağlantı problemine yol açmaktadır. PCA ile öznelikler birkaç temel bileşene indirgenerek, veri boyutu küçültülürken varyansın büyük bir kısmı korunmaktadır. Böylece önemli bilgiler korunarak gürültü azaltılıp boyut indirgenebilir. Bağımsız yeni bileşenler algoritmaların performansını artırabilmektedir. PCA, özellikle kümeleme (clustering) ve sınıflandırma (classification) algoritmaları ile veri keşfi (exploratory analysis) aşamalarında sıklıkla kullanılmaktadır. Bazı algoritmalarda (örneğin, lineer regresyon, lojistik regresyon, SVM vb.) PCA sonrası model daha iyi genelleşmektedir. Özetle, PCA, makine öğrenmesinde boyut indirgeme, gürültü azaltma ve ön işleme aracı olarak kullanılmaktadır (Bishop, 2006; Hastie vd., 2009; Murphy, 2012).

Çoklu doğrusal bağlantı, açıklayıcı değişkenler arasında yüksek korelasyon yapısını ifade etmektedir. Bu ilişki katsayıların hatalı hesaplanmasına yol açtığı ve standart hatalarını büyüttüğü için model tahminlerini olumsuz etkilemektedir. Ayrıca modelde hangi değişkenin daha etkili olduğunu ayırt etmek zorlaşır. Bu sebeple regresyon modellemesi açıklayıcı değişkenlerin birbiri ile ilişkisiz olduğu varsayımına dayalıdır. Genellikle uygulamalarda bu durum ihmal edilmekte ve açıklayıcı değişkenler birbirleri ile ilişkili çıkabilmektedir (Alpar, 2013; Chan vd., 2022). Benzer şekilde makine öğrenmesi araştırmalarında da öznelikler arasında çoklu doğrusal bağlantıyla karşılaşılabilir. Problemin tespiti amacıyla korelasyon matrisinde değişkenler arasındaki ikili korelasyon değerleri incelenir. Katsayının 0.80'den yüksek çıkması çoklu doğrusal bağlantı göstergesidir (Mason ve Perreault, 1991).

PCA, çoklu doğrusal bağlantıyı çözmek için sıkça kullanılan güçlü bir yöntemdir. PCA, Denklem (1)'de olduğu gibi birbirine yüksek korelasyonlu öznelikleri ($X_1, X_2, X_3, \dots$) ortogonal (bağımsız) yeni bileşenlere ($PC_1, PC_2, PC_3, \dots$) dönüştürmektedir. Bu yeni bileşenler temel bileşenler (principal components, PC) olarak adlandırılmaktadır. Her biri orijinal değişkenlerin doğrusal kombinasyonudur. Böylece korelasyon problemi ortadan kalkar; çünkü PC'ler birbirinden bağımsızdır. Özellikle yüksek boyutlu verilerde (çok sayıda değişken varsa) sadeleşmiş bir model sunmaktadır. Modelin tahmin performansını artırmaktadır. Bu dönüşüm birleştirme işleminden dolayı yorum kolaylığını azaltmaktadır, dolayısıyla amaç sadece tahmin yapmak ise PCA uygun bir yöntem olmaktadır (Abdi ve Williams, 2010; Jolliffe, 2002; Shlens, 2014).

$$\begin{aligned} PC_1 &= v_{11}X_1 + v_{12}X_2 + \dots + v_{1p}X_p \\ PC_2 &= v_{21}X_1 + v_{22}X_2 + \dots + v_{2p}X_p \end{aligned} \tag{1}$$

...

Denklem (1)'de v katsayıları özvektör bileşenleridir. Özvektör, temel bileşenin yönünü göstermektedir. Özdeğer ise ilgili bileşenin açıkladığı varyansı göstermektedir. Özdeğeri 1'den büyük olan bileşenler seçilerek sezgisel olarak bileşen sayısına karar verilebilir. Birinci bileşen en fazla varyansı açıklamaktadır, ikinci ondan daha az vb. şeklinde değerlendirilmektedir. Bu işlem toplam varyansın %80-95'ini açıklayan ilk k bileşeni seçmek şeklinde değerlendirilmektedir (Jolliffe, 2002; Hair vd., 2019).

3.2. Makine Öğrenmesi Sınıflandırma Algoritmaları

Bu bölümde, EEG temelli nöbet sınıflandırmasında kullanılan üç temel makine öğrenmesi algoritması açıklanmaktadır: Lojistik Regresyon, SVM ve XGBoost. Bu modeller hem teorik temelleri hem de farklı veri yapılarıyla uyumlu olmaları bakımından literatürde sıkça tercih edilmektedir. Lojistik Regresyon, doğrusal ilişkilerin modellenmesi açısından temel bir referans noktası sunarken; SVM, çekirdek (kernel) fonksiyonu sayesinde doğrusal olmayan karar sınırlarını da yakalayabilmektedir. XGBoost ise çok sayıda zayıf sınıflandırıcının birleşiminden oluşan topluluk (ensemble) yapısı ile yüksek doğruluk ve genellebilirlik sağlamaktadır. Bu algoritmaların seçilmesindeki amaç EEG verisinin karmaşık yapısını temsil etme güçlerini ve PCA uygulamasının bu modeller üzerindeki etkilerini karşılaştırmalı olarak değerlendirmektir.

3.2.1. Lojistik Regresyon (Logistic Regression)

Lojistik Regresyon genellikle belirli bir sınıfa ait olma olasılığını tahminlemek için kullanılır. Örneğin, ikili sınıflandırmada model olasılık tahmini 0.50'den büyüğe gözlemin o sınıfa ait olduğunu (pozitif sınıf, "1" olarak etiketlenir) ve aksi takdirde olmadığını (negatif sınıf, "0" olarak etiketlenir) tahmin eder. Lojit (logit), 0 ile 1 aralığındaki bir olasılığı gerçek sayılar aralığına dönüştüren bir fonksiyondur. Ters fonksiyonu olan lojistik (sigmoid) fonksiyon ise tam tersini yapar ve herhangi bir gerçek sayıyı (0,1) açık aralığına dönüştürür. Bu sürekli olasılığı ikili etikete dönüştürmek için bir eşik değeri (en yaygın olanı 0.5'tir) belirlenir. Model bu olasılığı kullanarak ikili bir karar (0 ya da 1) vermektedir. Pozitif sınıf için tahmini olasılık ile negatif sınıf için tahmini olasılık arasındaki oranın lojiti olduğundan dolayı log-olasılık oranı olarak da adlandırılır. Bir olayın gerçekleşme olasılığının gerçekleşmeme olasılığına oranı ise odds oranı olarak ifade edilmektedir. Lojistik regresyon, hedef değişken kategorik verilerden oluştuğu için geliştirilmiş doğrusal modellerin (GLM) bir üyesidir. Dolayısıyla ayırma karar sınırı doğrusal olarak belirlenir. Lojistik regresyon modeli, birden fazla sınıfı destekleyecek biçimde geliştirilebilir. Bu modele çok sınıflı (multinomial/softmax) lojistik regresyon adı verilmektedir (Friedman vd., 2001; Hosmer vd., 2013; Lewis, 2017; Montgomery vd., 2021). Karar sınırları, etkileşim terimleri eklendiğinde veya çekirdek yöntemler kullanıldığında doğrusal olmayan sınırlara dönüşebilir. Her sınıf için lojit skor hesaplandıktan sonra, bu skorlar softmax fonksiyonundan geçirilerek k sınıfına ait olma olasılığı $\hat{p}_k$, Denklem (2)'de gösterildiği üzere tahmin edilmektedir (Geron, 2019).

$$\hat{p}_k = \sigma(s(x))_k = \frac{\exp(s_k(x))}{\sum_{j=1}^k \exp(s_j(x))} \quad (2)$$

Denklem (2)'de k sınıf sayısı; $s_k(x)$ her sınıfın lojitini içeren vektör; $\sigma(s(x))_k$ ise k sınıfına ait olma olasılığının tahminini göstermektedir. Bu fonksiyon ile her skorun üstel değeri tüm üstellerin toplamına bölünmektedir. Lojistik regresyon sınıflandırıcısı, Denklem (3)'te gösterildiği üzere, tahmin edilen olasılıklar arasından en yüksek değere sahip sınıfı (argmax) seçmektedir.

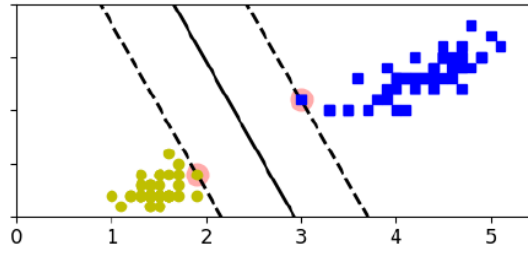
$$\hat{y} = \operatorname{argmax}_k \sigma(s(x))_k = \operatorname{argmax}_k s_k(x) = \operatorname{argmax}_k x^T \theta^k \quad (3)$$

Denklem (3)'te *argmax*, fonksiyonu maksimum yapan bir değişkenin değerini döndüren operatördür. Burada, $s_k(x) = x^T \theta^k$ olup, θ^k k. sınıfın parametre vektörünü göstermektedir.

3.2.2. Destek Vektör Makineleri (Support Vector Machine)

Doğrusal ve doğrusal olmayan ilişkileri modelleyebilme yeteneğiyle regresyon ve sınıflandırma problemlerinde kullanılan SVM oldukça etkili bir yöntemdir. Ayrıca aykırı değer tespiti yapabilen güçlü ve çok yönlü bir algoritma sunmaktadır (Moguerza ve Muñoz, 2006). SVM yöntemi ilk olarak Cortes ve Vapnik (1995) tarafından geliştirilmiş olup, maksimum marj prensibine dayalı doğrusal bir sınıflandırma yaklaşımı ileri sürmektedir (Cortes ve Vapnik, 1995). SVM, istatistiksel öğrenme teorisi temelinde bir yöntem olup örüntüleri tespit etme sürecinde en geniş marjı belirlemektedir (Kecman, 2005). SVM algoritması, tahmin doğruluğunu artırmak için bir hata toleransı belirlemektedir. Böylece belirli ölçüde yanlış sınıflandırmalara izin veren ceza parametresi ile hata-marj dengesi kurmaktadır. Model, bu denge altında en büyük marjınli ayırıcı hiperdüzlemi belirlemektedir. Şekil 1'de doğrusal SVM için karar hiperdüzlemi (siyah çizgi), marjin (paralel kesikli çizgiler) ve destek vektörleri (içi dolu işaretçiler) gösterilmektedir. Başka bir ifadeyle sınıflar (mavi kare ve sarı daire) arasındaki mümkün olan en geniş aralığa marjin denilmekte; marjini tanımlayan paralel kesikli çizgilere dokunan örnekler ise destek vektörleri olarak adlandırılmaktadır. SVM, karar hiperdüzlemi etrafındaki marjini (iki paralel destek hiperdüzlemi arasındaki uzaklık) maksimize etmektedir. Hard-margin'de bu destek hiperdüzlemlerine

temas eden örnekler destek vektörleridir; soft-margin'de ise marjin içinde kalan veya hatalı sınıflanan ve modele etkisi olan örnekler de destek vektörüdür (Geron, 2019).



Şekil 1: Destek Vektörleri

Kaynak: Geron, 2019.

Karar sınırı yalnızca sınıfları ayırmakla kalmamakta; aynı zamanda mümkün olduğunca en yakın eğitim örneklerinden de uzak kalmaktadır. Çekirdek fonksiyonları yardımı ile doğrusal olmayan sınıflandırma yapabilmektedir. Yaygın olarak kullanılan çekirdek fonksiyonlar arasında doğrusal, polinomial, sigmoid ve radyal tabanlı fonksiyonlar (RBF) bulunmaktadır. SVM, büyük veri setlerinde yüksek doğruluk sağlarken, az sayıda destek vektörü kullandığı için bellek verimliliği de sunmaktadır (Smola ve Schölkopf, 2004). Modelin genelleştirme yeteneği hiperparametrelerden etkilendiği için bu parametrelerin optimize edilmesi önem kazanmaktadır. En uygun hiperparametre değerlerini belirlemek için yaygın bir yaklaşım ızgara aramasıdır (grid search). SVM'ler öznitelik ölçeklerine duyarlıdır; bu sebeple modelleme öncesinde özniteliklerin ölçeklenmesi/standartlaştırılması gerekmektedir. Veri sızıntısını önlemek için ölçekleme işlem hattı (pipeline) içinde GridSearchCV ile gerçekleştirilmektedir.

3.2.3. Aşırı Gradyan Artırma Algoritması (EXtreme Gradient Boosting)

XGBoost, gradient boosting (gradyan artırma) algoritmasına dayanan bir makine öğrenmesi algoritmasıdır. Boosting yaklaşımı zayıf öğrencileri birleştirerek güçlü bir model oluşturmayı hedeflemektedir. Doğrusal olmayan ilişkileri modellemek amacıyla karar ağaçlarını art arda birleştirmektedir. Her yeni ağaç, önceki modelin yaptığı hataları öğrenerek düzeltmeye çalışmaktadır. Bu işlem birçok kez tekrarlanmakta ve modeller üst üste eklenmektedir. Böylece tahmin doğruluğunu artmaktadır. Tahmin güncellenmesi Denklem (4) ile gösterilmektedir (Friedman, 2001; Hastie vd., 2009).

$$\hat{y}_i^{(t)} = \hat{y}_i^{(t-1)} + \eta f_t(x_i) \quad (4)$$

Denklem (4)'te η öğrenme oranını (eta), $f_t(x_i)$ t. ağacı göstermektedir.

XGBoost, klasik gradient boosting'e birkaç güçlü iyileştirme getirmektedir. Modelin arka planında uygulanan optimizasyon teknikleri hesaplama hızını artırırken, aşırı öğrenmeyi kontrol altına almak için düzenleme (regularization) terimlerini de içerebilmektedir. Özellikle LASSO ve ridge regularization ile ceza terimlerinin kullanımı, XGBoost'u diğer gradient boosting modellerine kıyasla daha dayanıklı, güçlü ve esnek hale getirir. XGBoost, hızlı hesaplama yeteneği ve büyük veri setleri üzerindeki yüksek doğruluk performansı sebebiyle makine öğrenmesi problemlerinde genelleme imkanı sunmaktadır (Chen ve Guestrin, 2016). Ayrıca hangi değişkenlerin modele en çok katkı sağladığı ölçülmektedir. Veride eksik değer söz konusu ise otomatik olarak en uygun dalı öğrenmektedir. Ayrıca, modelin hiperparametrelerinin geniş açıda ayarlanabilir olması, farklı veri yapılarına ve problemlere uyarlanmasına olanak tanımaktadır (Chen ve He, 2020).

3.2.4. Model Performansını Değerlendirme

Tahmin performanslarını karşılaştırmak amacıyla kesinlik (precision), duyarlılık (recall-sensitivity), doğruluk (accuracy) ve F1 kriterleri kullanılmaktadır. Performans ölçütleri, karmaşıklık matrisi değerlerinden elde edilmektedir. İkili sınıflandırma için (2x2) boyutlu karmaşıklık matrisinin bir tanımı Tablo 1 ile sunulmuştur.

Tablo 1: Karmaşıklık Matrisi

		Tahmin	
Gerçek	(pozitif)	(pozitif)	(negatif)
	(negatif)	Doğru Pozitifler (TP)	Yanlış Negatifler (FN)
		Yanlış Pozitifler (FP)	Doğru Negatifler (TN)

Performans ölçütlerinin hesaplanma yöntemleri, Denklem (5)–(8)’de gösterilmektedir (Gupta ve Goel, 2023). Tek bir ölçüt ile değerlendirme yapmak yeterli olmamaktadır. Araştırmanın amacına göre yanlış tahminlerin maliyetinin önemi farklılık gösterebilmektedir. Yüksek duyarlılık değeri, düşük miktarda yanlış negatif durumu gösterirken; yüksek kesinlik değeri, düşük miktarda yanlış pozitif durumu göstermektedir. Eğer yanlış pozitifleri en aza indirmek önemliyse, kesinlik tercih edilen ölçüt olabilmektedir. Eğer yanlış pozitifler çok fazla sayıda ise kesinlik düşmektedir. Kesinlik değeri, doğru pozitiflerin tüm pozitif tahminlere oranını vermektedir. Eğer yanlış pozitifleri kontrol etmek önemliyse (örneğin, sahte alarm istenmiyorsa), o zaman kesinlik uygun bir ölçüt olmaktadır (Shalev-Shwartz ve Ben-David, 2014; Kendirkıran ve Doğan, 2024; Kıvrak, 2025).

$$\text{Accuracy (Doğruluk)} = \frac{TP + TN}{TP + TN + FP + FN} \quad (5)$$

$$\text{Recall – Sensitivity (Duyarlılık)} = \frac{TP}{TP + FN} \quad (6)$$

$$\text{Precision (Kesinlik)} = \frac{TP}{TP + FP} \quad (7)$$

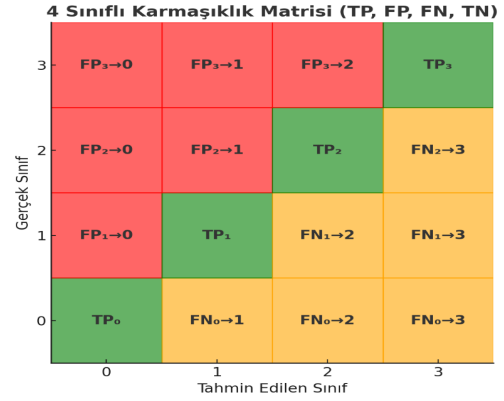
$$F1 = 2 * \frac{\text{precision} * \text{recall}}{\text{precision} + \text{recall}} \quad (8)$$

Öte yandan yüksek sayıda yanlış negatif durumda duyarlılık ölçütü tercih edilmektedir. Yanlış negatifler çok olduğunda duyarlılık düşmektedir. Yanlış negatifler azaldıkça, doğru pozitiflerin yakalanma oranı yani duyarlılık artmaktadır. Eğer bir problemde yanlış negatifleri azaltmak önemliyse (örneğin, kanserli hastayı hasta değil olarak sınıflandırmamak), duyarlılık kritik bir ölçüt olmaktadır. Yanlış negatif ve yanlış pozitif sayıları dengelendiğinde doğruluk metriği daha anlamlı hale gelebilmektedir. Ancak dengesiz (imbalanced) sınıf dağılımlarında yalnızca doğruluk tek başına güvenilir bir ölçüt olmayabilir. Örneğin, pozitif sınıfın oranının %95 olduğu dengesiz bir veri kümesinde, tüm örnekleri “pozitif” sınıflandıran bir modelin doğruluğu yaklaşık %95 görünse de model negatif sınıfı ayırt edememekte; dolayısıyla doğruluk tek başına yanıltıcı olmaktadır. F1 değeri kesinlik ile duyarlılığın harmonik ortalaması olarak hesaplanmaktadır. Yüksek bir F1 değeri hem yanlış pozitiflerin hem de yanlış negatiflerin düşük olduğu, yani modelin daha dengeli bir performans gösterdiği anlamına gelmektedir. Harmonik ortalama, iki değerden küçük olanına daha fazla ağırlık vermektedir; bu sebeple kesinlik ve duyarlılığın biri çok düşükse F1 de düşük olmaktadır (Fawcett, 2006; Powers, 2011; Han vd., 2012).

Çok sınıflı sınıflandırma, karmaşıklık matrisi k sınıf için k×k boyutunda genelleştirilir. Matriste i hücresi, gerçek sınıfın i olup modelin j sınıfını tahmin ettiği veri sayısını gösterir. Ana köşegen doğru sınıflamaları, köşegen dışı hücreler ise hataları temsil eder. Tablo 2’de araştırmada kullanılan k = 4 durumu için y = {0, 1, 2, 3} sınıflarına karşılık gelen karmaşıklık matrisi hücreleri gösterilmektedir. Buna göre köşegen değerleri (TP₀, TP₁, TP₂, TP₃) doğru tahmin edilenler, köşegen dışındakiler yanlış sınıflandırmalardır. FN, satırdaki hatalı tahminlerdir. Örneğin, FN₀ olan hücreler gerçek sınıf 0 iken tahmin sınıfı 1, 2 ya da 3 olanlar olup FN₀→1, FN₀→2 ya da FN₀→3 şeklinde gösterilmiştir. FP, sütundaki hatalı tahminlerdir. Örneğin, FP₀ olan hücreler gerçek sınıf 1, 2 ya da 3 iken tahmin sınıfı 0 olanlardır. Bu durumda TN, matriste ilgili satır ve sütun dışındaki tüm hücreler (geri kalanlar) şeklinde değerlendirilmektedir. Şekil 2’de gösterilen yapı ile performans hesapları her sınıf için ayrı ayrı hesaplanmaktadır. Ardından makro ortalama, mikro ortalama, ağırlıklı ortalama alınarak genel metrikler elde edilmektedir (Powers, 2011; Han vd., 2012; Provost ve Fawcett, 2013).

Tablo 2: k = 4 için Karmaşıklık Matrisi

		Tahmin			
		(0)	(1)	(2)	(3)
Gerçek	(0)	TP ₀	FN ₀	FN ₀	FN ₀
	(1)	FP ₀	TP ₁	FN ₁	FN ₁
	(2)	FP ₀	FP ₁	TP ₂	FN ₂
	(3)	FP ₀	FP ₁	FP ₂	TP ₃



Şekil 2: k = 4 Karmaşıklık Matrisi

Makine öğrenmesi ve istatistiksel sınıflandırma problemlerinde, modellerin ayırıştırma gücünü değerlendirmek için kullanılan en temel ölçütlerden biri de ROC (Receiver Operating Characteristic) eğrisi ve bu eğrinin altında kalan alanı ifade eden AUC (Area Under the Curve) değeridir. ROC eğrisi, modelin doğru pozitif oranı (TP Rate) ile yanlış pozitif oranı (FP Rate) arasındaki ilişkiyi gösteren bir grafikdir. ROC eğrisinin her noktası, modelin farklı eşik değerleri için elde ettiği TPR ve FPR ikililerini temsil etmektedir. Bu nedenle ROC eğrisi, modelin tahminlerinin olasılık dağılımını dikkate alarak performansını tüm olası eşik değerleri üzerinden değerlendirmektedir (Bradley, 1997; Fawcett, 2006; Saito ve Rehmsmeier, 2015). AUC değeri modelin rastgele sınıflandırma yapma olasılığına göre ne kadar iyi performans gösterdiğini ölçmektedir. AUC değeri 0 ile 1 arasında değişmekte; 0.5 değeri rastgele tahmin yapan bir model anlamına gelirken, 1.0 değeri mükemmel ayırıştırma başarısını göstermektedir. İstatistiksel olarak AUC, bir modelin rastgele seçilen bir pozitif örneğe negatif bir örnekten daha yüksek olasılık skoru verme ihtimali olarak da yorumlanabilmektedir. Bu özellik, ROC-AUC ölçütünü özellikle sınıf dengesizliği içeren problemlerde, tek başına doğruluk değerinden daha güvenilir hale getirmektedir. ROC-AUC, çok sınıflı durumlara genelleştirilebilmektedir. Böylece her sınıfın diğerlerine karşı ayırıştırma gücü belirlenebilmektedir (Hanley ve McNeil, 1982). Bu yönüyle ROC-AUC, model seçimi ve hiperparametre optimizasyonu süreçlerinde yaygın olarak kullanılan, dengeli bir performans ölçütü olmaktadır.

4. VERİ SETİ

Çalışma kapsamında, UCI makine öğrenmesi platformundan sağlanan epileptik nöbet tespiti ve sınıflandırması için kapsamlı bir EEG koleksiyonu olan Bangalore EEG Epilepsi Veri Seti (BEED) kullanılmıştır (Najmusseher ve Banu, 2024: <https://archive.ics.uci.edu/dataset>). BEED, Creative Commons Attribution 4.0 International (CC BY 4.0) lisansı kapsamında araştırmacıların kullanımına sunulmuştur. Bu veri seti Hindistan Bangalore'daki bir nörolojik araştırma merkezinde kaydedilen 8000 gözlem verisinden oluşan özgün ve güncel bir veri seti olup elektrot sistemi kullanılarak yakalanan yüksek kaliteli EEG sinyallerini içermektedir. Bu sinyaller beyin hücrelerinin elektriksel aktivitesini ölçmek amacıyla kafatası üzerine yerleştirilen elektrotlar (standart 10-20 elektrot sistemi) sayesinde nöronların sinyallerini kaydetmektedir.

BEED, 16 öznitelik ve bir hedef değişkenden oluşmaktadır. Öznitelikler farklı beyin bölgelerine karşılık gelen 256 Hz örnekleme hızına sahip 16 EEG kanalından gelen ham dalga formu sinyalleridir (X_1-X_{16}). Hedef değişken, epilepsi nöbeti varlığını veya yokluğunu ifade etmekte olup dört kategoriden oluşmaktadır. Tablo 3'te özetlendiği üzere sınıflar sırasıyla sağlıklı denekler (0), jeneralize nöbetler (1), fokal nöbetler (2) ve nöbet aktivitesinin göz kırpmaya, tırnak yeme veya dik dik bakma gibi olaylarla birlikte meydana geldiği nöbet olayları (3) şeklindedir (Najmusseher ve Banu, 2025). Her kategori eşit cinsiyet temsiline sahip yetişkin denekler (yaşları 21-55) 20 saniyelik EEG kaydı içermektedir. Dört kategoriye eşit olarak dağıtılmış -her kategori 2000 gözlem- dengeli bir veri seti söz konusudur.

Tablo 3: BEED Sınıf Etiketleri ve Tanımları

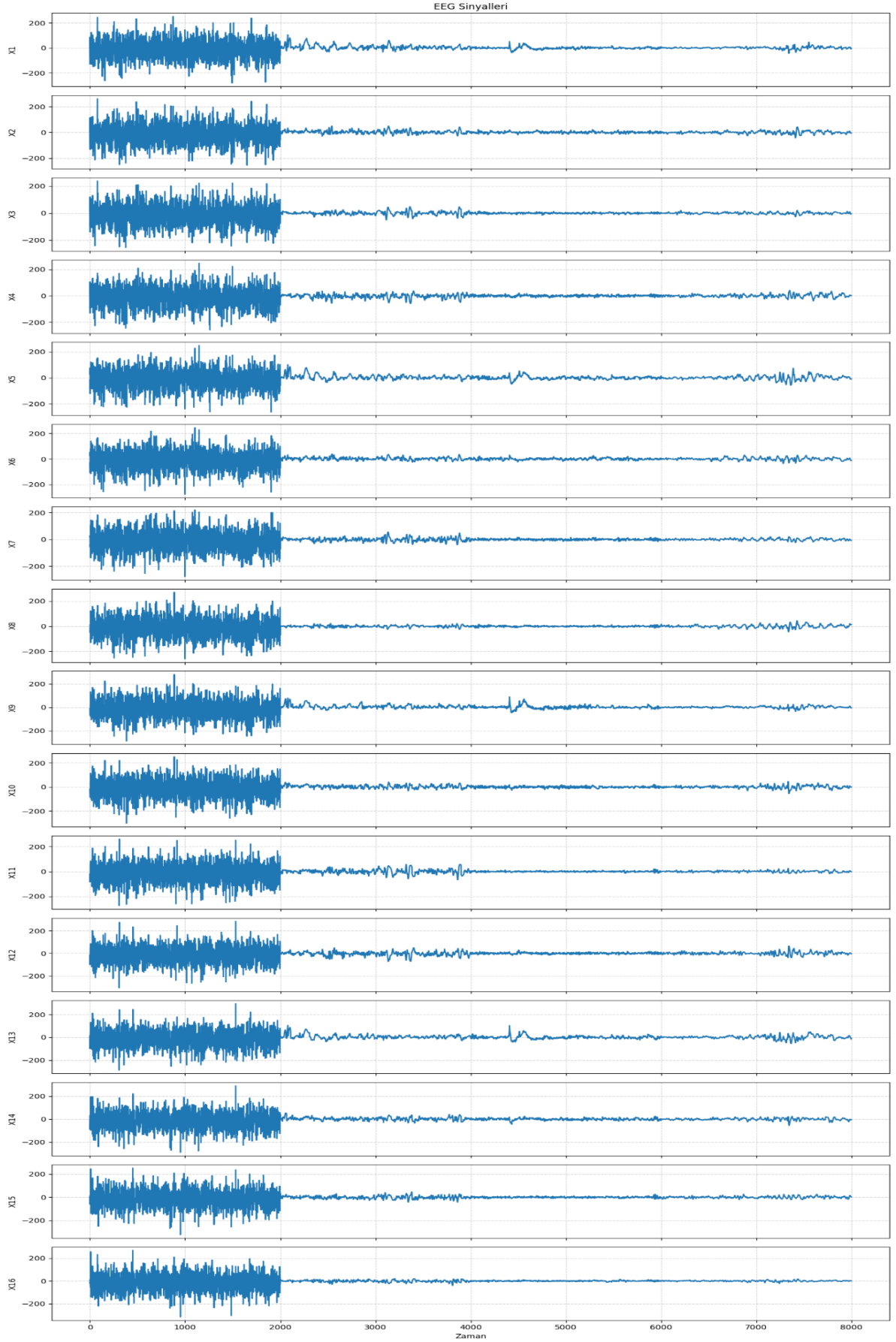
Nöbet Sınıfı	Tanım
Sağlıklı (0)	Epilepsi tanısı olmayan, nöbet geçirmeyen bireylerden alınan kayıtlardır.
Jeneralize nöbet (1)	Beynin her iki yarım küresinde görülen nöbet kayıtlarıdır. Nöbet aktivitesi beynin büyük bir bölümünü etkiler. EEG’de yaygın bozulmalar gözlenir.
Fokal nöbet (2)	Beynin belirli bir bölgesinde başlayan ve sınırlı kalan nöbet kayıtlarıdır. Zamanla genel nöbete dönüşebilir, başlangıç noktası önemlidir.
Nöbet aktiviteleri (3)	Fiziksel hareketler sırasında kaydedilen EEG verileridir. Bu grup, göz kırpması, boşluğa bakma gibi nöbete eşlik eden davranışsal olayları içerir.

5. BULGULAR

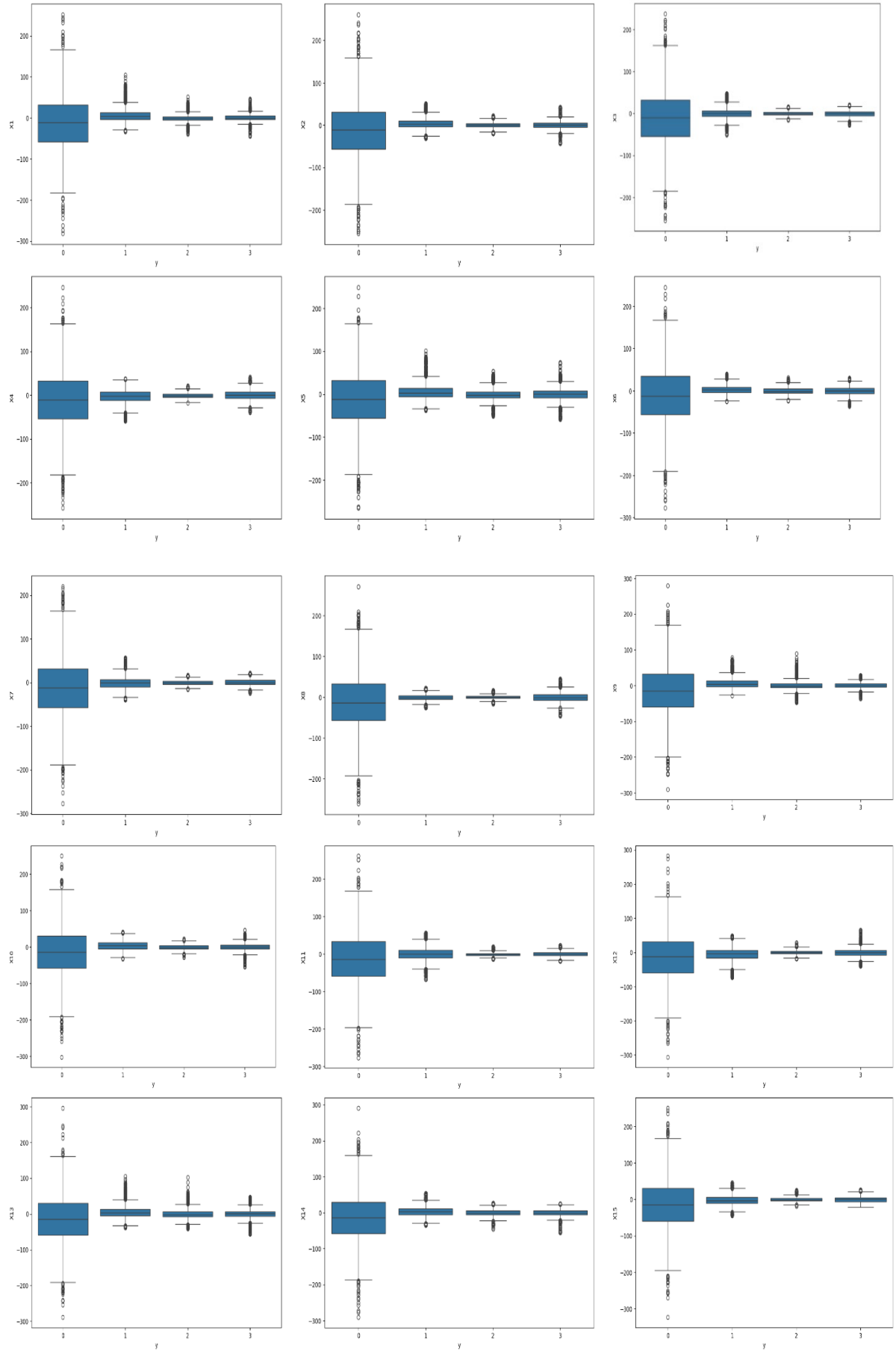
Bu çalışmada, epilepsi nöbeti varlığının veya yokluğunun tahmin edilmesi amacıyla kullanılan makine öğrenmesi yöntemleri Lojistik regresyon, SVM ve XGBoost’tur. Veriler eğitim ve test olmak üzere iki alt kümede düzenlenmiştir. Nesnel tahmin performansını değerlendirmek amacıyla üç ayrı deney yürütülmüş; test payı sırasıyla %30, %20 ve %10 olarak belirlenmiştir. Modeller eğitim veri setleri kullanılarak tahmin edilmiş ve sonrasında, modellerin öngörü performansları test veri setleri üzerinden ölçülmüştür.

BEED veri kümesindeki 16 kanallı sinyal örneği Şekil 3’te görselleştirilmiştir. Şekil 3’te, BEED veri kümesine ait 16 kanallı EEG sinyalleri zaman ekseninde gösterilmiştir. İlk segmentlerde gözlenen yüksek genlikli ve düzensiz salınımlar, epileptik nöbet aktivitesine karşılık gelirken, daha sonraki düzleşen kısımlar sağlıklı veya nöbet dışı dönemleri temsil etmektedir. Kanallar arası örüntüler genel olarak benzer biçimde ilerlese de genlik ve frekans bileşenleri açısından farklılıklar görülmektedir; bu da her kanalın nöbet tespitine özgü bilgi taşıdığını göstermektedir.

Her bir kanalın (X_1-X_{16}) dört sınıfa ($y = 0$ sağlıklı, 1 jeneralize nöbet, 2 fokal nöbet, 3 nöbet aktivitesi) göre kutu grafikleri Şekil 4’te sunulmuştur. Görüldüğü üzere özellikle nöbet sınıflarına (1, 2, 3) ait dağılımlar, sağlıklı sınıfa (0) kıyasla daha dar aralıkta ve negatif medyan (kutu orta çizgi) değerlerde yoğunlaşmaktadır. Bu durum, nöbet sırasında EEG sinyallerinin genlik varyansının azaldığını ve belirli kanallarda baskın bir yönelme gösterdiğini ortaya koymaktadır. Ayrıca bazı kanallarda aykırı değerlerin (outlier) sayısının artması, epileptik aktiviteye özgü ani potansiyel değişimlerinin varlığına işaret etmektedir. Genel olarak kutu grafikler, sınıflar arasında anlamlı genlik farklarının bulunduğunu ve bu farkların sınıflandırma algoritmaları açısından ayırt edici öznitelikler sunduğunu göstermektedir. Kutu grafiklerindeki bıyıkların (genellikle $Q_1-1.5(Q_3-Q_1)$ ile $Q_3+1.5(Q_3-Q_1)$ aralığı) dışında kalan gözlemler aykırı değer olarak değerlendirilmekte ve grafikte yuvarlak noktalar ile simgelenmektedir. İlk aşamada aykırı değerler modelden çıkarılmamıştır; sınıflandırma algoritmalarının bu gözlemlerden de örüntü öğrenebilme olasılığı göz önünde bulundurulmuştur.



Şekil 3: EEG Sinyalleri ($X_1 - X_{16}$)



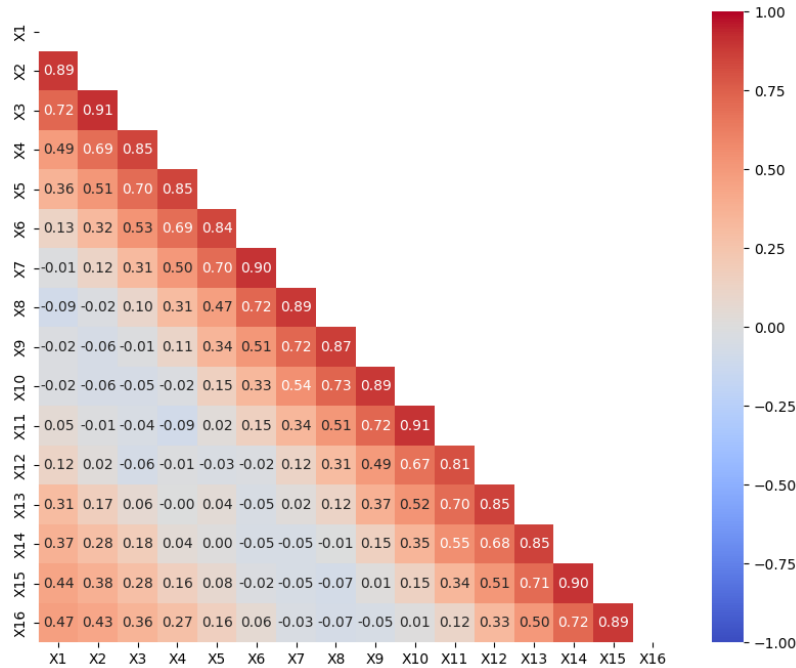
Şekil 4: Kutu Grafikleri

Veri kümesinde yer alan özniteliklere ilişkin temel istatistikler Tablo 4'te sunulmuştur. Veri kümesinde kayıp/eksik gözlem bulunmamaktadır. Tüm öznitelikler sayısal (nümerik) yapıdadır. Özet istatistikler (ortalama, standart sapma, çeyreklikler, minimum–maksimum) incelendiğinde, bazı öznitelikler için aykırı değerlerin bulunduğu görülmektedir. Analizin başlangıcında öznitelikler StandardScaler ile standartlaştırılmıştır (sızıntıyı önlemek için ölçekleme yalnızca eğitim verisi üzerinde öğrenilmiş, test verisine uygulanmıştır).

Tablo 4: Özet İstatistikler

Öznitelik	Tür	Ortalama	Std. sapma	Min	%25	%50	%75	Max
X ₁	Nümerik	-1.49	36.82	-281	-7.00	0.0	8.00	252
X ₂	Nümerik	-2.19	36.11	-255	-7.00	0.0	8.00	261
X ₃	Nümerik	-3.24	35.80	-255	-7.00	-1.0	5.00	238
X ₄	Nümerik	-4.12	36.29	-257	-10.00	-1.0	7.00	246
X ₅	Nümerik	-1.82	37.62	-264	-10.00	0.0	10.00	249
X ₆	Nümerik	-2.31	36.31	-277	-8.00	0.0	8.00	245
X ₇	Nümerik	-3.40	36.36	-277	-8.00	-1.0	6.00	220
X ₈	Nümerik	-3.45	36.52	-260	-7.00	-1.0	5.00	271
X ₉	Nümerik	-1.65	38.11	-290	-7.00	0.0	8.00	280
X ₁₀	Nümerik	-2.56	37.54	-302	-8.00	0.0	8.00	251
X ₁₁	Nümerik	-3.52	37.34	-276	-8.00	-1.0	5.00	262
X ₁₂	Nümerik	-4.78	37.47	-306	-11.00	-2.0	7.00	283
X ₁₃	Nümerik	-2.16	38.14	-288	-10.00	0.0	10.00	296
X ₁₄	Nümerik	-2.91	36.64	-290	-8.00	0.0	9.00	291
X ₁₅	Nümerik	-4.36	36.24	-323	-9.00	-2.0	5.00	251
X ₁₆	Nümerik	-4.11	35.93	-317	-6.00	-1.0	4.00	270

İlk olarak, çoklu doğrusal bağlantının varlığını araştırmak üzere korelasyon matrisi ve buna karşılık gelen ısı haritası incelenmiştir. Veri setinde yer alan özniteliklere ait ısı matrisi Şekil 5 ile sunulmuştur. Öznitelikler arasındaki ilişkiler incelendiğinde ikili korelasyonların çok yüksek olduğu görülmektedir. Isı haritasında renklerin ağırlıklı kırmızıya yaklaşması katsayıların 0.70–0.91 aralığında yer alması, belirgin bir çoklu doğrusal bağlantı yapısına işaret etmektedir. Özellikle ana köşegenin dışında yer alan birbirine yakın değişken çiftlerinde görülen 0.85–0.91 düzeyindeki korelasyonlar, bu değişkenlerin büyük ölçüde aynı bilgiyi taşıdığını göstermektedir. Öte yandan birkaç çok düşük hatta hafif negatif korelasyon da gözlenmektedir. Bu durum özniteliklerin diğerlerinden görece bağımsız ya da tamamlayıcı bilgi taşıdığını işaret edebilmektedir.



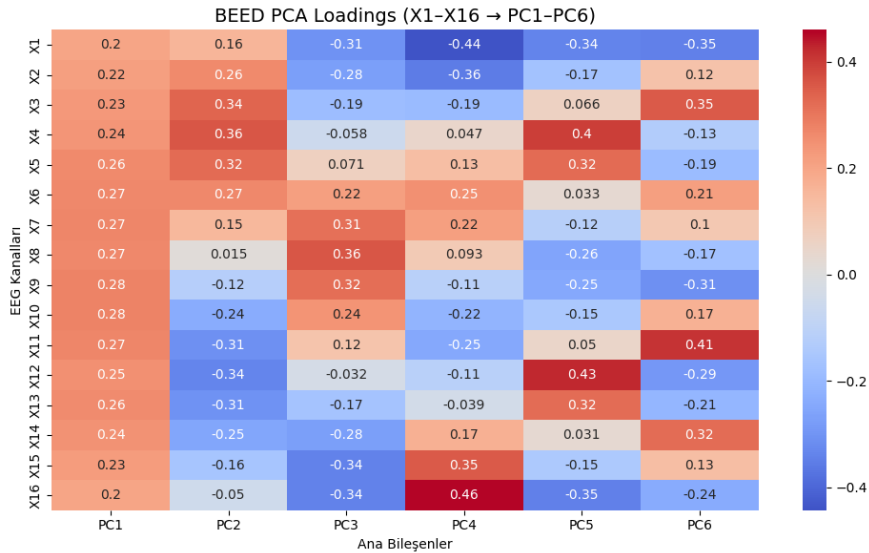
Şekil 5: Korelasyon Isı Matrisi

Varyans şişirme faktörü (VIF) değerleri hesaplanmıştır. Montgomery vd. (2021)'e göre $5 \leq VIF < 10$ orta düzey, $VIF \geq 10$ ise ciddi çoklu doğrusal bağlantıya işaret etmektedir. Tablo 5'teki bulgulara göre X_1 (7.82) ve X_{16} (7.89) orta düzeyde, diğer öznitelikler ise ciddi çoklu doğrusal bağlantı göstermektedir. Modeldeki değişkenlerin önemli bir kısmı arasında yüksek düzeyde çoklu doğrusal bağlantı bulunduğu saptanmıştır.

Tablo 5: VIF Değerleri

X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8
7.82	19.35	20.28	14.22	10.22	13.18	16.83	15.80
X_9	X_{10}	X_{11}	X_{12}	X_{13}	X_{14}	X_{15}	X_{16}
15.60	19.08	17.09	10.44	12.80	13.87	18.19	7.89

Lojistik regresyon gibi doğrusal (GLM) modellerde yüksek korelasyon katsayı tahminlerini dengesizleştirmekte, standart hataları büyütmekte; bu da yorumlanabilirliğini ve güvenilirliğini düşürmektedir. Marjine dayalı SVM ve ağaç tabanlı XGBoost bu duruma görece daha dayanıklı olsa da yüksek ölçüde benzer öznitelikler model karmaşıklığını artırıp aşırı uyum riskini tetikleyebilmektedir. Bu sebeple boyut indirgeme amacıyla PCA uygulanmıştır. PCA öncesinde öznitelikler standartlaştırılmış, kümülatif açıklanan varyans eşiği %95 olarak alınmış ve toplam 6 temel bileşen ($PC_1 - PC_6$) elde edilmiştir. Böylece 16 kanal 6 bileşene indirgenmiştir. Şekil 6'da bu bileşenler üzerindeki yük değerleri görselleştirilmiştir. Pozitif (kırmızı tonlu) ve negatif (mavi tonlu) değerler, ilgili kanalın o bileşen üzerindeki yönlü katkısını göstermektedir. Şekilden görüldüğü üzere PC_1 bileşeni tüm kanallar boyunca pozitif yükler taşımakta ve sinyallerin ortak varyansının büyük kısmını temsil etmektedir. Bu durum EEG kanalları arasında güçlü korelasyon bulunduğunu ve PCA'nın bu ortak varyansı ilk birkaç bileşende yoğunlaştırdığını göstermektedir. Öte yandan diğer bileşenlerde yük değerlerinin işaret değiştirmesi, belirli kanalların (özellikle X_8-X_{12} aralığında) özgül varyanslarının yakalandığını ve bu kanalların nöbet sınıflandırması açısından ayırt edici bilgi taşıdığını düşündürmektedir. Özellikle PC_4 bileşeninde X_{16} kanalına ait yüksek pozitif yük (0.46), bu kanalın kendi sinyal örüntüsüyle diğerlerinden farklılaştığını ortaya koymaktadır. Negatif yüklerin dağılımı ise (örneğin PC_3 ve PC_5 'te) bazı kanalların birbirine zıt varyans yönelimleri sergilediğini göstermekte, bu da PCA'nın kanallar arasındaki çoklu doğrusal bağlantıyı etkili biçimde çözümlediğini desteklemektedir.



Şekil 6: Temel Bileşen Yükleri

Bileşenlerin varyans açıklama oranları Tablo 6 ile sunulmuştur. Buna göre genel olarak, ilk üç temel bileşenin verideki ortak varyansı özetlediği, sonraki bileşenlerin ise daha çok kanallara özgü varyans bileşenlerini temsil ettiği söylenebilir. Bu sonuç, PCA'nın veri boyutunu azaltırken bilgi kaybını minimize ettiğini ve modelleme sürecinde hesaplama yükünü azaltmak için uygun bir ön-işleme adımı olarak işlev gördüğünü göstermektedir.

Tablo 6: Varyans Katkısı

Açıklanan Varyans Oranı	Temel Bileşenler						Toplam
	PC ₁	PC ₂	PC ₃	PC ₄	PC ₅	PC ₆	
	0.367	0.248	0.237	0.057	0.028	0.016	

Sonraki aşamada, hem ham öznitelik uzayına (PCA'sız) hem de PCA ile elde edilen düşük boyutlu temel bileşenlere makine öğrenmesi yöntemleri uygulanmış; hesaplanan performans ölçütleri kayda geçirilmiştir. Hiperparametreler ızgara araması (grid search) ile eniyilenmiştir. Tüm deneyler, veri sızıntısını önlemek için ölçekleme/PCA dahil ön işleme adımlarının yalnızca eğitim verisi üzerinde öğrenildiği PCA'lı ve PCA'sız iki ayrı işlem hattında (pipeline) yan yana değerlendirilmiştir. Test paylarına göre (%30–%20–%10) PCA'lı/PCA'sız LR–SVM–XGBoost performans ve süre karşılaştırması, ayrılmış test kümesi için Tablo 7'de, on katlı çapraz doğrulama (CV10) için Tablo 8'de, eğitim kümesi için Ek 1'de sunulmuştur.

Tablo 7: BEED Test Kümesi Performans Karşılaştırması

Test	Ölçüt	Algoritma					
		LR	PCA - LR	SVM	PCA - SVM	XGBoost	PCA - XGBoost
%30	Doğruluk	0.46	0.46	0.82	0.72	0.97	0.84
	Kesinlik	0.50	0.49	0.84	0.74	0.97	0.84
	Duyarlılık	0.46	0.46	0.82	0.72	0.97	0.84
	F1 Ölçütü	0.47	0.47	0.81	0.71	0.97	0.84
	Süre	1sn	1sn	4dk 5sn	4dk 26sn	1dk 17sn	1 dk 17sn
%20	Doğruluk	0.46	0.44	0.82	0.72	0.97	0.86
	Kesinlik	0.50	0.47	0.84	0.74	0.97	0.86
	Duyarlılık	0.46	0.44	0.82	0.72	0.97	0.86
	F1 Ölçütü	0.47	0.45	0.82	0.71	0.97	0.86
	Süre	1sn	2sn	5dk 7sn	5dk 40sn	1dk 22sn	1dk 24sn
%10	Doğruluk	0.45	0.45	0.82	0.71	0.98	0.85
	Kesinlik	0.49	0.48	0.84	0.72	0.98	0.85
	Duyarlılık	0.45	0.45	0.82	0.71	0.98	0.85
	F1 Ölçütü	0.46	0.45	0.82	0.69	0.98	0.85
	Süre	2sn	2sn	7dk 22sn	7dk 55sn	1dk 27sn	1dk 26sn

Tablo 7'ye göre PCA'sız XGBoost tüm test oranlarında doğruluk/kesinlik/duyarlılık/F1 açısından yaklaşık 0.97–0.98 aralığıyla ve en kısa süre (1:17–1:27) ile açık ara en iyi performansı sergilemektedir. PCA etkisi değerlendirildiğinde, SVM'de belirgin bir düşüş görülmektedir (örneğin, F1: 0.81–0.82 → 0.69–0.71). XGBoost'ta PCA eklenmesi mutlak 0.11–0.13 puanlık kayba yol açmaktadır (0.97–0.98 → 0.86–0.85). Ayrıca XGBoost'a PCA eklendiğinde performansın düştüğü, aşırı öğrenme riskinin arttığı görülmektedir (EK 1). Lojistik regresyon zaten düşük bir taban performansına sahiptir (0.45–0.50); bu nedenle PCA eklenmesi anlamlı bir iyileşme sağlamamıştır.

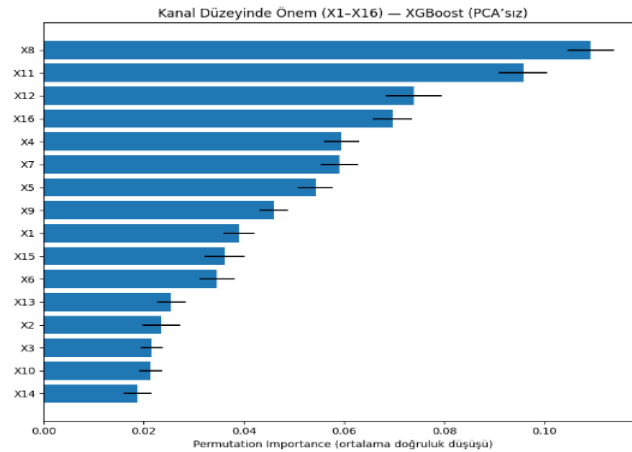
Lojistik regresyon en hızlıdır; ancak düşük doğruluk nedeniyle pratikte yetersizdir ve yalnızca süre karşılaştırmalarında referans olarak kullanılabilir. SVM en yavaş modeldir (4–8 dk); PCA'lı SVM süreyi daha da uzatırken performansı düşürmekte, dolayısıyla zayıf bir hız–başarım dengesi sunmaktadır. PCA'nın zaman etkisi değerlendirildiğinde, 16 öznitelikten 6 bileşene indirgemeye karşın dönüşümün ek hesaplama maliyeti ve küçük boyut nedeniyle SVM ve XGBoost'ta çalışma süreleri hızlanmamış; tersine çoğu durumda aynı kalmış ya da bir miktar uzamıştır. Ayrıca eğitim seti büyüdükçe (test oranı %30'dan %10'a düştükçe) tüm modellerin çalışma süreleri artmaktadır. Elde edilen sonuçlar, PCA'nın uyguladığı boyut indirgeme nedeniyle oluşan bilgi kaybının modellerin tahmin performansını olumsuz etkilediğini göstermektedir.

Tablo 8: On Katlı Çapraz Doğrulama Sonuçları

Test	Ölçüt	CV10	
		SVM	XGBoost
%30	Doğruluk	0.8204 ± 0.0165	0.9623 ± 0.0076
	Kesinlik	0.8448 ± 0.0172	0.9627 ± 0.0076
	Duyarlılık	0.8204 ± 0.0165	0.9623 ± 0.0076
	F1 Ölçütü	0.8157 ± 0.0172	0.9624 ± 0.0076
%20	Doğruluk	0.8264 ± 0.0100	0.9677 ± 0.0062
	Kesinlik	0.8501 ± 0.0091	0.9682 ± 0.0062
	Duyarlılık	0.8264 ± 0.0100	0.9677 ± 0.0062
	F1 Ölçütü	0.8222 ± 0.0104	0.9677 ± 0.0062
%10	Doğruluk	0.8246 ± 0.0123	0.9729 ± 0.0055
	Kesinlik	0.8463 ± 0.0159	0.9731 ± 0.0054
	Duyarlılık	0.8246 ± 0.0123	0.9729 ± 0.0055
	F1 Ölçütü	0.8209 ± 0.0127	0.9729 ± 0.0055

Tablo 8 incelendiğinde, CV10 sonuçları SVM için yalnızca marjinal bir artış göstermektedir (F1 yaklaşık 0.81'den 0.85'e). XGBoost'ta CV10 sonuçları test performansı ile uyumludur (0.96–0.97), bu da aşırı uyum belirtisi olmadığını düşündürmektedir. Test ve CV puanları arasındaki farkın küçük olması (<0.05) genel olarak aşırı uyum riskinin düşük olduğuna işaret etmektedir. Ayrılmış test doğruluğu ile on katlı çapraz doğruluğun 0,97 düzeyinde ve birbirine çok yakın olması, modelin hem mevcut veri üzerinde hem de genel durumda güçlü bir genelleme performansı sergilediğine işaret eder. Eğitim skoru ile Test/CV skorları arasındaki yaklaşık 0,03'lük fark genellenmenin pratik olarak iyi olduğunu göstermektedir. Bu bulgular, XGBoost'un veri içindeki genel örüntüleri etkili biçimde öğrendiğini ve bu örüntüleri yeni örneklerle taşıırken tahmin doğruluğunu koruduğunu ortaya koymaktadır.

XGBoost gibi yüksek kapasiteli ağaç modelleri, kanallar arası benzerliklerin yüksek ve öznelilik korelasyonlarının belirgin olduğu veri kümelerinde eğitim verisine çok hızlı ve neredeyse tam uyum sağlayabilir. Bu nedenle, eğitim performans ölçütlerinin 1'e yakın ya da 1 çıkması literatürde gözlenen bir durumdur; buna karşın erken durdurma, subsample/colsample ayarları ve regülerizasyon gibi mekanizmalar gerekli görüldüğünde aşırı uyumu sınırlamak için kullanılabilir. Test oranına duyarlılık değerlendirildiğinde; %30 → %20 → %10 test dilimlerinde görece performans sıralaması değişmemektedir: XGBoost (PCA'sız) > SVM (PCA'sız) > Lojistik regresyon. Özetle, inceleme kapsamında BEED veri kümesinde (PCA'sız) doğrudan öznelilik uzayında eğitilen XGBoost, hem doğruluk/F1 hem de çalışma süresi ve uygulama kolaylığı açısından en güçlü aday olarak öne çıkmaktadır.



Şekil 7: XGBoost Algoritmasına Göre Önemli Kanallar

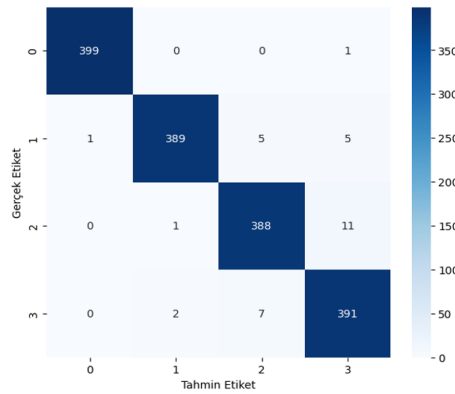
Modelde yer alan özneliliklerin görece önem düzeylerini incelemek amacıyla değişken önem analizi gerçekleştirilmiştir. Bu amaçla kullanılan permutation importance (öznelilik önemliliği) yöntemi, her bir özneliliğin model performansına katkısını ölçmek için ilgili özneliliğin değerlerini rastgele permüte ederek (karıştırarak) modelin doğruluk kaybını hesaplamaktadır. Dolayısıyla bir değişkenin permütasyonu sonucunda model performansında büyük bir düşüş gözleniyorsa, o değişkenin model açısından yüksek öneme sahip olduğu kabul edilmektedir (Breiman, 2001). Buna göre Şekil 7'de XGBoost(PCA'sız) en önemli öznelilikleri X₈, X₁₁, X₁₂,... olacak şekilde sıralamış; modelin genel performansı doğruluk=0.9793 ve macro-F1=0.9794 olarak raporlanmıştır.

XGBoost hiperparametreleri RandomizedSearchCV ile doğruluk kriterini maksimize edecek şekilde ayarlanmıştır. Parametre değerleri satır örnekleme oranı (subsample) 0.7, L2 regularizasyon (reg_lambda) 1.5, L1 regularizasyon (reg_alpha) 0.1, toplam ağaç sayısı (n_estimators) 400, derinlik (max_depth) 7, öğrenme oranı (learning_rate/eta) 0.2, sütun örnekleme oranı (colsample_bytree) 0,9 olarak optimize edilmiştir. Diğer hiperparametreler için varsayılan değerler kullanılmıştır. Bu parametrelerle eğitilen modelde CV10 için doğruluk=0.9808 ve macro-F1 =0.9808 elde edilmiş; ayrılmış test skoru ile yüksek uyum gözlenmiş ve belirgin bir aşırı uyum belirtisine rastlanmamıştır.

Nöbet bazlı ayırmalı oluşturulan test kümesi üzerinde (PCA'sız) XGBoost için sınıflandırma raporu Tablo 9'da, karmaşıklık matrisi ise Şekil 8'de sunulmuştur. Sınıflandırma raporuna göre sınıf bazında F1 değerleri 0.97–1.00 aralığındadır; karmaşıklık matrisi ise hataların ağırlıkla 1–2–3 sınıfları arasında yoğunlaştığını göstermektedir. 0 (sağlıklı) sınıfı neredeyse hatasız sınıflanmış; en belirgin karışıklık 2 (fokal) ile 3 (nöbet aktivitesi) sınıfları arasında gözlenmiştir.

Tablo 9: Sınıflandırma Raporu

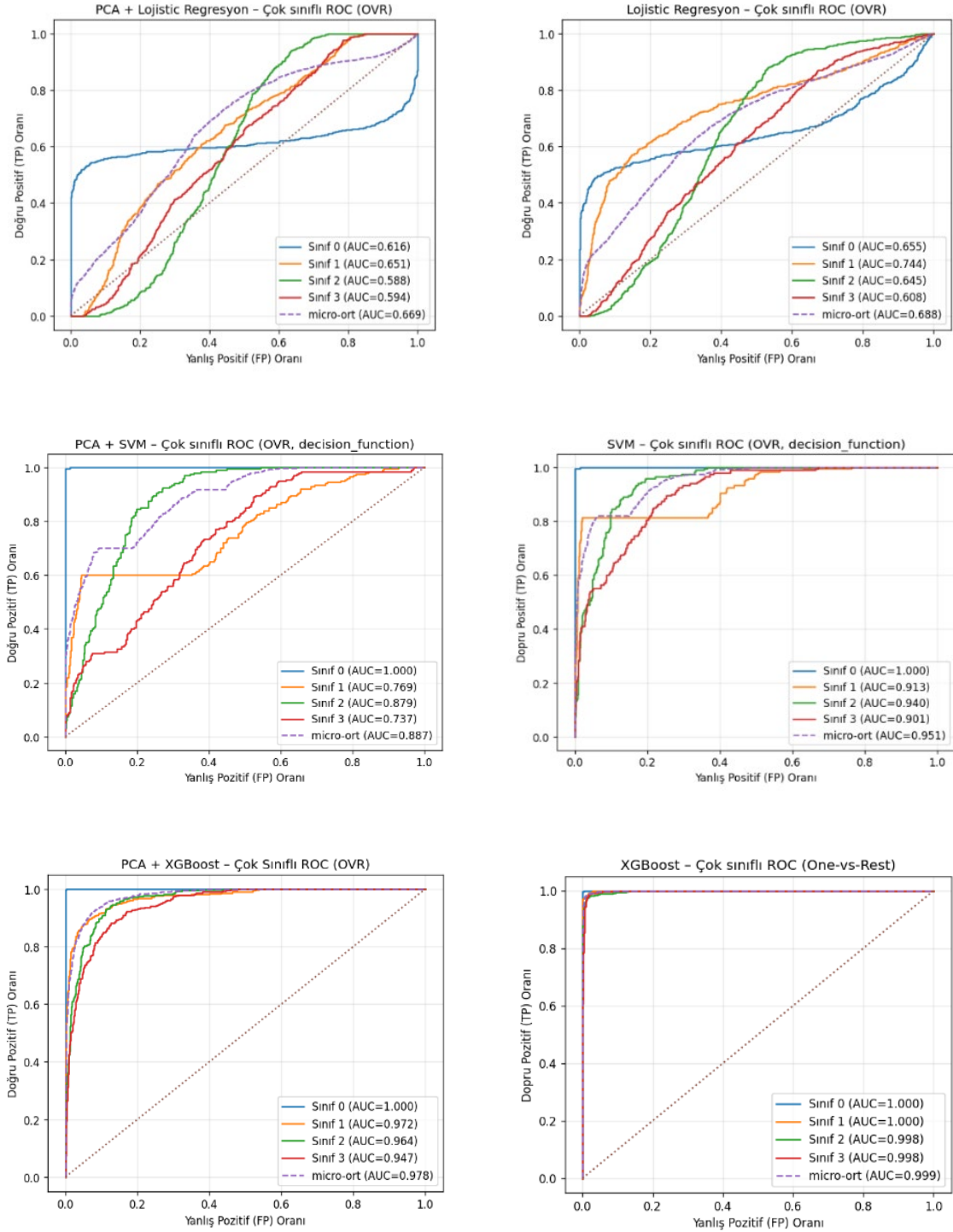
Sınıf	Kesinlik	Duyarlılık	F1	Veri
Sağlıklı (0)	1.00	1.00	1.00	400
Jeneralize nöbet (1)	0.99	0.97	0.98	400
Fokal nöbet (2)	0.97	0.97	0.97	400
Nöbet aktiviteleri (3)	0.96	0.98	0.97	400



Şekil 8: Karmaşıklık Matrisi (Gerçek x Tahmin)

Şekil 9'da sunulan ROC-AUC sonuçlarına göre PCA'sız XGBoost, micro-AUC = 0.999 ile açık ara en yüksek performansı göstermektedir. XGBoost (PCA'sız) AUC değerleri 1.000/1.000/0.998/0.998 olup sınıfların neredeyse tamamen ayrıldığını göstermektedir. Bu bulgu, önceki doğruluk/macro-F1 = 0.97 düzeylerinin neden tutarlı biçimde yüksek olduğunu açıklamaktadır. Sınıflar dengeli olduğundan, macro-AUC (sınıf-başına ortalama) ile micro-AUC (örnek-ağırlıklı ortalama) doğal olarak birbirine çok yakın çıkmaktadır; nitekim $(1+1+0.998+0.998)/4=0.999$ olup raporlanan micro-AUC ile uyumludur.

SVM (PCA'sız) çok güçlü micro-AUC = 0.951 değerine sahip olup sınıf AUC değerleri 1.000/0.913/0.940/0.901 şeklindedir. Lojistik regresyon zayıf micro-AUC = 0.688 değerine sahip olup sınıf başına 0.59–0.74 sonuçları problemin doğrusal olmadığını göstermektedir. PCA eklendiğinde tüm modellerde AUC değerleri düşmektedir. Etki lojistik regresyonda sınırlı iken, SVM'de belirgin bir azalma görülmektedir. Bunun olası nedeni, SVM'nin ham öznelilik uzayındaki ayrımı iyi yakalaması ve 0.95 açıklanan varyans eşliğinin sınıf ayırıcı bazı eksenleri bastırmasıdır. XGBoost'ta gözlenen hafif düşüş ise ağaç tabanlı modellerin PCA'ya tipik olarak ihtiyaç duymadığını düşündürmektedir. Sonuç olarak, PCA varyansı koruyan ancak sınıf ayırıcılığı garanti etmeyen bir dönüşüm olup bu veri kümesinde sınıf ayırıcı bilginin bir kısmını kaybettiği için AUC değerini düşürmektedir.



Şekil 9: BEED için ROC-AUC Eğrileri

6. SONUÇ

Çalışmada, dengeli BEED veri kümesi üzerinde PCA'lı ve PCA'sız iki ayrı işlem hattında denetimli makine öğrenmesi sınıflandırıcıları olan Lojistik regresyon, SVM ve XGBoost karşılaştırılmıştır. Test payı %30, %20, %10 olacak şekilde üç senaryo çalıştırılmış; ayrıca on katlı çapraz doğrulama raporlanmıştır. Tüm test paylarında PCA'sız XGBoost en yüksek başarıyı vermiştir. SVM (PCA'sız) ikinci sırada, lojistik regresyon ise tüm senaryolarda 0.50 ve altında kalarak düşük bir baz performansındadır. PCA eklendiğinde hem SVM hem de XGBoost performansında düşüş ile mutlak kayıp gözlenmiştir. Lojistik regresyon performansı zaten düşük olduğundan PCA anlamlı bir artış sağlamamıştır. XGBoost'un çapraz doğruluğunun %97'ye ulaşması ayrılmış test skoruyla uyumludur. Belirgin bir

aşırı uyum göstergesi bulunmamaktadır. SVM’de çağraz doğrulama sadece marjinal bir iyileşme sağlamaktadır. Farkların küçük olması modelin genelleme yetisinin iyi olduğuna işaret etmektedir. XGBoost’a PCA eklendiğinde performans düşmüştür; aşırı öğrenme riski artmıştır. Karmaşıklık matrisi hataların çoğunlukla 1 (jeneralize nöbet), 2 (fokal nöbet), 3 (nöbet aktivitesi) sınıfları arasında yoğunlaştığını; 0 (sağlıklı) sınıfının ise neredeyse hatasız ayrılabilirdiğini göstermektedir. En belirgin karışıklık ise fokal nöbet ile nöbet aktivitesi arasında görülmektedir. PCA’sız XGBoost, en iyi micro-AUC değerine sahip olup sınıflar dengeli olduğundan macro-AUC değeri de çok yakındır. PCA eklendiğinde Lojistik regresyon/SVM/XGBoost’un ROC-AUC değerleri düşmektedir. Bu durum, PCA’nın yüksek varyans açıklama yönlerini korurken sınıf ayırıcı doğrultuların bir kısmını bastırabildiğini göstermektedir. Etki, SVM’de daha belirgindir; XGBoost’ta düşüş sınırlıdır ve bu sonuç ağaç tabanlı modellerin PCA’ya ihtiyaç duymadığını teyit etmektedir. PCA’nın 16 öznelikten 6 bileşene indirgemeye karşın zaman kazancı yoktur. Dönüşümün ek maliyeti ve küçük boyut nedeniyle sürelerin aynı kaldığı ya da hafifçe arttığı gözlenmiştir. Performanslar test payına duyarlı olmayıp eğitim miktarı arttıkça çalışma süresi artmış ancak performans sıralaması değişmemiştir. BEED üzerinde (PCA’sız) XGBoost, hem doğruluk/F1 hem de çalışma süresi ve uygulama basitliği açısından en güçlü adaydır. PCA, bu veri kümesinde ayırıcı bilgiyi kısmen kaybettirdiği için tüm modellerde başarısı düşürmektedir. Ayrıca zaman kazandırmamaktadır. Bu bulgular, yüksek korelasyonlu EEG öznelikleri için ağaç tabanlı yaklaşımların boyut indirgemesiz işlem hatlarının daha avantajlı olabileceğini göstermektedir.

Gelecek çalışmalarda SVM ve XGBoost için PCA temelli boyut indirgeme yaklaşımının yanı sıra, öznelik seçimini doğrudan hedef değişkenle ilişkilendiren yöntemlerin kullanılması model başarımını, hesaplama verimliliğini ve yorumlanabilirliği artırabilir. Bu kapsamda, Minimum Redundancy Maximum Relevance (mRMR), Karşılıklı Bilgi (Mutual Information), ReliefF, Boruta, ve L1-tabanlı (LASSO) regularizasyon gibi yöntemlerin uygulanması, EEG öznelikleri arasında bilgilendirici alt kümelerin belirlenmesine katkı sağlayabilir. Lojistik regresyon için doğrusal olmayan genişletmeler (polinom/etkileşim terimleri, kernel-lojistik regresyon) ve gözetimli indirgeme yöntemleri (Doğrusal Ayırt Edici Analiz-LDA ve Kısmi En Küçük Kareler Ayırt Edici Analizi-PLS-DA) değerlendirilebilir. Ayrıca daha farklı makine öğrenmesi sınıflandırıcılarıyla ve derin öğrenme yöntemleriyle performans karşılaştırmaları yapılabilir.

Beyan ve Açıklamalar (Disclosure Statements)

1. Bu çalışmanın yazarı, araştırma ve yayın etiği ilkelerine uyduğunu kabul etmektedirler (The author of this article confirm that their work complies with the principles of research and publication ethics).
2. Yazar tarafından herhangi bir çıkar çatışması beyan edilmemiştir (No potential conflict of interest was reported by the author).
3. Bu çalışma, intihal tarama programı kullanılarak intihal taramasından geçirilmiştir (This article was screened for potential plagiarism using a plagiarism screening program).

KAYNAKÇA

- Abdi, H., ve Williams, L. J. (2010). Principal Component Analysis. *Wiley Interdisciplinary Reviews: Computational Statistics*, 2(4), 433–459. <https://doi.org/10.1002/wics.101>.
- Alpar, R. (2013). *Çok Değişkenli İstatistiksel Yöntemler*, Detay Yayıncılık: Ankara, Türkiye.
- Berrich, Y. ve Guennoun, Z. (2025). EEG-Based Epilepsy Detection Using CNN-SVM and DNN-SVM with Feature Dimensionality Reduction By PCA. *Sci Rep* 15, 14313 (2025). <https://doi.org/10.1038/s41598-025-95831-z>.
- Bishop, C. M. (2006). *Pattern Recognition and Machine Learning*, Springer.
- Bradley, A. P. (1997). The Use of the Area Under the ROC Curve in the Evaluation of Machine Learning Algorithms. *Pattern Recognition*, 30(7), 1145–1159. [https://doi.org/10.1016/S0031-3203\(96\)00142-2](https://doi.org/10.1016/S0031-3203(96)00142-2)
- Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>.
- Carvajal-Dossman, J.P., Guio, L., García-Orjuela, D. vd. (2025). Retraining and Evaluation of Machine Learning and Deep Learning Models for Seizure Classification from EEG Data. *Sci Rep* 15, 15345 (2025). <https://doi.org/10.1038/s41598-025-98389-y>.
- Chan, J.-L., Leow, S., Bea, K., Cheng, W., Phoong, S., Hong, Z.-W. ve Chen, Y. L. (2022). Mitigating the Multicollinearity Problem and its Machine Learning Approach: A Review, *Mathematics*, 10(8), 1283. <https://doi.org/10.3390/math10081283>.

- Chen, T., ve Guestrin, C. (2016). XGBoost: A Scalable Tree Boosting System, *In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794). <https://dl.acm.org/doi/pdf/10.1145/2939672.2939785>.
- Chen, T., ve He, T. (2020). XGBoost: Extreme Gradient Boosting. R Package Version 1.0.0.2.
- Cortes, C., ve Vapnik, V. (1995). Support-Vector Networks. *Machine Learning*, 20(3), 273–297. <https://doi.org/10.1007/BF00994018>.
- Diler, S. ve Demir, Y. (2024). Çoklu Doğrusal Bağlantı Olması Durumunda Veri Madenciliği Algoritmaları Performanslarının Karşılaştırılması, *Nicel Bilimler Dergisi*, 6(1), 40-67. doi: 10.51541/nicel.1371834.
- Fawcett, T. (2006). An introduction to ROC analysis. *Pattern Recognition Letters*, 27(8), 861–874. <https://people.inf.elte.hu/kiss/13dwhdm/roc.pdf>
- Friedman, J., Hastie, T., ve Tibshirani, R. (2001). The Elements of Statistical Learning: Data Mining, *Inference, and Prediction*. Springer. <https://tinyurl.com/4cpndh32>.
- Friedman, J. H. (2001). Greedy Function Approximation: A Gradient Boosting Machine, *Annals of Statistics*, 29(5), 1189–1232.
- Geron, A. (2019). *Hands-On Machine Learning with Scikit-Learn, Keras, and TensorFlow, Concepts, Tools, and Techniques to Build Intelligent Systems*, 2nd Edition. O'Reilly Media, Inc., Canada. <https://www.oreilly.com/catalog/errata.csp?isbn=9781492032649>.
- Gupta, A. ve Goel, S. (2023). Deep Learning Empowered Weather Image Classification for Accurate Analysis. *IEEE 2nd International Conference on Industrial Electronics: Developments & Applications (ICIDEA)*, Imphal, India, 2023, pp. 567-572, doi: 10.1109/ICIDEA59866.2023.10295216.
- Han, J., Kamber, M., ve Pei, J. (2012). *Data Mining: Concepts and Techniques* (3rd ed.). Morgan Kaufmann. <https://homes.di.unimi.it/ceselli/IM/2012-13/slides/02-KnowYourData.pdf>.
- Hanley, J. A., ve McNeil, B. J. (1982). The Meaning and Use of the Area Under a Receiver Operating Characteristic (ROC) Curve. *Radiology*, 143(1), 29–36. <https://doi.org/10.1148/radiology.143.1.7063747>.
- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2019). *Multivariate Data Analysis*, (8th ed.). Cengage.
- Hastie, T., Tibshirani, R., ve Friedman, J. (2009). *The Elements of Statistical Learning*, (2nd ed.). Springer.
- Hosmer, D. W., Lemeshov, S. ve Sturdivant, R. X. (2013). *Applied Logistic Regression*, (Third Edition). John Wiley & Sons, Inc: New Jersey, USA.
- Islam, M. S., Thapa, K., ve Yang, S.-H. (2022). Epileptic-Net: An Improved Epileptic Seizure Detection System Using Dense Convolutional Block with Attention Network from EEG. *Sensors*, 22(3), 728. <https://doi.org/10.3390/s22030728>.
- Jolliffe, I. T. (2002). *Principal Component Analysis* (2nd ed.). Springer Series in Statistics. Springer-Verlag, New York.
- Kecman, V. (2005). *Support Vector Machines: Theory and Applications*. In L.Wang (Eds). *Support Vector Machines- An Introduction* (pp. 1–47). Springer, Berlin Heidelberg. <https://doi.org/10.1007/b95439>.
- Kendirkıran, G. ve Doğan, S. (2024). Makine Öğrenimi Teknikleri ile Kredi Risk Tahmininde Yeniden Örneklem Yöntemlerinin Karşılaştırılması, *Söke İşletme Fakültesi Dergisi*, Yıl: 2024, Cilt: 1, Sayı: 2, ss.48-60 <https://dergipark.org.tr/tr/download/article-file/4364826>.
- Kıvrak, O. (2025). Araç Fiyat Tahmininde Makine Öğrenmesi Algoritmalarının Karşılaştırılması ve Performans Analizi, *İktisadi İdari ve Siyasal Araştırmalar Dergisi*, Cilt: 10 Sayı: 27, 454-474, 29.06.2025 <https://doi.org/10.25204/iktisad.1494020>.
- Kong, G., Ma, S., Zhao, W., Wang, H., Fu, Q., & Wang, J. (2024). A Novel Method for Optimizing Epilepsy Detection Features Through Multi-Domain Feature Fusion and Selection, *Frontiers in Computational Neuroscience*, 18, 1416838. doi:10.3389/fncom.2024.1416838.
- Lewis, N. D. (2017), *Machine Learning Made Easy with R: An Intuitive Step by Step Blueprint for Beginners, CreateSpace Independent Publishing Platform*: Carolina, USA.
- Li, Q., Cao, W, Zhang, A. (2025). Multi-Stream Feature Fusion of Vision Transformer and CNN for Precise Epileptic Seizure Detection from EEG Signals. *J Transl Med*, Aug 6;23(1):871. doi: 10.1186/s12967-025-06862-z. PMID: 40770757; PMCID: PMC12329966.
- Mason, C. H. ve Perreault, W. D. (1991), Collinearity, Power, and Interpretation of Multiple Regression Analysis, *Journal of Marketing Research*, 28(3), 268–280. <https://doi.org/10.2307/3172863>.

- Moguerza, J., ve Muñoz, A. (2006). Support Vector Machines with Applications, *Statistical Science*, 21, 322- 336. <https://doi.org/10.1214/088342306000000493>.
- Montgomery, D. C., Peck, E. A. ve Vining, G. G. (2021). *Introduction to Linear Regression Analysis*. John Wiley ve Sons.
- Murphy, K. P. (2012). *Machine Learning: A Probabilistic Perspective*. MIT Press.
- Najmüsseher ve Banu, P. K. N. (2024). BEED: Bangalore EEG Epilepsy Dataset. [Dataset]. UCI Machine Learning Repository. <https://doi.org/10.24432/C5K33B>. <https://archive.ics.uci.edu/dataset>.
- Najmüsseher ve Banu, P. K. N. (2025). Feature Engineering for Epileptic Seizure Classification Using SeqBoostNet. *International Journal of Computing and Digital Systems*, VOL. 17, NO. 1, 1–15. <http://dx.doi.org/10.12785/ijcds/1571020131>.
- Rahman, M. M., Ghasemi, Y., Suley, E., Zhou, Y., Wang, S. ve Rogers, J. (2021). Machine Learning Based Computer Aided Diagnosis of Breast Cancer Utilizing Anthropometric and Clinical Features, *IRBM*, 42(4), 215-226.
- Powers, D. M. W. (2011). Evaluation: From Precision, Recall and F-measure to ROC, Informedness, *Markedness and Correlation*. *Journal of Machine Learning Technologies*, 2(1), 37–63. <https://arxiv.org/abs/2010.16061>.
- Provost, F., ve Fawcett, T. (2013). *Data Science for Business*. O'Reilly Media. <https://www.oreilly.com/library/view/data-science-for/9781449374273/>.
- Saito, T., ve Rehmsmeier, M. (2015). The Precision-Recall Plot is More Informative Than the ROC Plot When Evaluating Binary Classifiers on Imbalanced Datasets. *PLOS ONE*, 10(3), e0118432. <https://doi.org/10.1371/journal.pone.0118432>.
- Shalev-Shwartz, S., ve Ben-David, S. (2014). *Understanding Machine Learning: From Theory to Algorithms*. Cambridge University Press.
- Shlens, J. (2014). A Tutorial on Principal Component Analysis. arXiv preprint arXiv:1404.1100.
- Smola, A. J. ve Schölkopf, B. (2004). A Tutorial on Support Vector Regression. *Statistics and Computing*, 14, 199-222. <https://doi.org/10.1023/B:STCO.0000035301.49549.88>.
- Wei, L. ve Mooney, C. (2020). Epileptic Seizure Detection in Clinical EEGs Using an XGboost-based Method, *IEEE SPMB*, v1.0:1-6. https://isip.piconepress.com/conferences/ieee_spmb/2020/papers/I02_03.pdf?utm.
- Wu, J., Zhou, T., Li, T. (2020). Detecting Epileptic Seizures in EEG Signals with Complementary Ensemble Empirical Mode Decomposition and Extreme Gradient Boosting. *Entropy (Basel)*, 24;22(2):140. doi: 10.3390/e22020140.

EK 1: BEED Eğitim Performansı Karşılaştırması

Eğitim	Ölçüt	Algoritma					
		LR	PCA - LR	SVM	PCA - SVM	XGBoost	PCA - XGBoost
%70	Doğruluk	0.49	0.46	0.83	0.72	1	1
	Kesinlik	0.52	0.49	0.85	0.74	1	1
	Duyarlılık	0.48	0.46	0.83	0.72	1	1
	F1 Ölçütü	0.498	0.47	0.83	0.70	1	1
%80	Doğruluk	0.48	0.46	0.83	0.72	1	1
	Kesinlik	0.52	0.49	0.86	0.75	1	1
	Duyarlılık	0.48	0.46	0.83	0.72	1	1
	F1 Ölçütü	0.49	0.46	0.83	0.71	1	1
%90	Doğruluk	0.48	0.46	0.84	0.73	1	1
	Kesinlik	0.52	0.49	0.86	0.75	1	1
	Duyarlılık	0.48	0.46	0.84	0.73	1	1
	F1 Ölçütü	0.79	0.47	0.83	0.71	1	1