



K-medoids Kümeleme, Boosting ve BERT algoritmalarıyla Kredi Kartlarındaki Şüpheli Davranışların Tespit Edilmesi

Detection of Suspicious Behaviors in Credit Cards Using K-medoids Clustering, Boosting, and BERT Algorithms

Merve Pınar^{1*}, Ayşe Berna Altinel²

¹ Marmara Üniversitesi, Teknoloji Fakültesi, Bilgisayar Mühendisliği Bölümü, merve.pinar@marmara.edu.tr
ORCID: <https://orcid.org/0000-0003-3041-6958>

² Marmara Üniversitesi, Teknoloji Fakültesi, Bilgisayar Mühendisliği Bölümü, berna.altinel@marmara.edu.tr
ORCID: <https://orcid.org/0000-0001-5544-0925>

MAKALE BİLGİLERİ

Makale Geçmişi:

Geliş 20 Şubat 2025
Revizyon 18 Nisan 2025
Kabul 24 Nisan 2025
Online 30 Haziran 2025

Anahtar Kelimeler:

*Şüpheli Davranış Tespiti,
Dolandırıcılık Tespiti,
Boosting,
BERT,
K-medoids Kümeleme*

ÖZ

Finansal dolandırıcılıkların giderek artan karmaşıklığı karşısında, kredi kartı işlemlerinde şüpheli davranışların erken ve doğru şekilde tespiti, güvenlik sistemleri açısından kritik bir gereksinim haline gelmiştir. Bu çalışmada, K-medoids tabanlı kümeleme ile denetimli sınıflandırma algoritmalarının entegrasyonundan oluşan hibrit bir model önerilmiştir. Model, veri kümesini alt gruplara ayırarak sınıflandırıcıların daha homojen örüntüler üzerinde eğitilmesini sağlamak ve genel sınıflandırma başarımını artırmayı hedeflemektedir. Çalışma kapsamında, Avrupa'da gerçekleştirilmiş 284.807 işlem den oluşun, PCA dönüşümleriyle anonimleştirilmiş bir kredi kartı veri seti kullanılmıştır. Eğitim verisi oranları %30, %40 ve %80 olarak belirlenmiş ve bu oranların model başarımı üzerindeki etkisi sistematik olarak incelenmiştir. K-medoids kümelendiği veri üzerinde Lojistik Regresyon, Destek Vektör Makineleri, Yapay Sinir Ağları, Rastgele Orman, XGBoost, Gradyan Artırma, Hard/Soft Voting Ensemble ve BERT tabanlı modeller (BERT, XML-RoBERTa, DistilBERT, Electra) çalıştırılmıştır. Ayrıca, filtre tabanlı Bilgi Kazancı, Fisher Skor ve Ki-kare özellik seçimi tekniklerinin algoritma performansları üzerindeki etkisi detaylı şekilde analiz edilmiştir. Deneyisel sonuçlar, özellikle Gradyan Artırma algoritmasının %98.26 F1 skoruna ulaşarak en yüksek başarıyı elde ettiğini göstermiştir. Bilgi Kazancı ve Fisher Skor yöntemleri sınıflar arası ayrımı daha etkili biçimde modelleyerek sınıflandırma performansını artırırken, Ki-kare yöntemi, veri setinin sürekli ve dönüşümlü doğası nedeniyle daha düşük doğruluk oranlarına ulaşmıştır. Elde edilen bulgular, kümeleme destekli hibrit sınıflandırma yaklaşımlarının, dengesiz dağılımlı yüksek boyutlu veri setleri üzerinde yüksek doğruluk ve genellenebilirlik sunduğunu ortaya koymakta ve sahtekarlık tespitine yönelik karar destek sistemlerinin geliştirilmesinde etkin bir çerçeve sunmaktadır.

ARTICLE INFO

Article history:

Received 20 February 2025
Received in revised form 18 April 2025
Accepted 24 April 2025
Available online 30 June 2025

Keywords:

*Suspicious Behavior Detection,
Fraud Detection,
Boosting,
BERT,
K-medoids Clustering*

DOI: 10.24012/dumf.1643370

* Sorumlu Yazar

ABSTRACT

The increasing complexity of financial fraud necessitates the early and accurate detection of suspicious behaviors in credit card transactions for robust financial security systems. In this study, a hybrid model integrating K-medoids-based unsupervised clustering with various supervised classification algorithms is proposed. The model partitions the dataset into subgroups to enable classifiers to learn from more homogeneous patterns, thereby enhancing the overall classification performance. The experimental setup utilizes a publicly available credit card dataset comprising 284,807 anonymized transactions collected from Europe, where feature anonymization is achieved through Principal Component Analysis (PCA). The impact of different training-to-test splits (30%, 40%, and 80%) on model performance is systematically evaluated. On the clustered data, a diverse set of classification models is applied, including Logistic Regression, Support Vector Machines, Artificial Neural Networks, Random Forest, XGBoost, Gradient Boosting, Hard and Soft Voting Ensembles, as well as BERT-based transformer models (BERT, XML-RoBERTa, DistilBERT, Electra). In addition, three supervised filter-based feature selection methods—Information Gain, Fisher Score, and Chi-Squared—are employed to assess their effects on classification performance. The experimental results indicate that the Gradient Boosting algorithm achieves the highest F1 score, reaching up to 98.26%, especially when combined with Fisher Score-based feature selection. Both Information Gain and Fisher Score techniques significantly enhance the classification performance by capturing inter-class discriminative power. However, the Chi-Squared method exhibits comparatively lower effectiveness due to its incompatibility with the transformed and continuous nature of the dataset. Overall, the findings demonstrate that clustering-assisted hybrid classification architectures provide superior accuracy and generalizability in high-dimensional, imbalanced datasets and offer a promising framework for developing intelligent fraud detection systems.

Giriş

Finansal kayıpların önlenmesi ve güvenlik açıklarının minimize edilmesi, kredi kartı dolandırıcılığının erken ve doğru tespitini zorunlu hale getirmiştir. Dijital ekonominin hızla büyümesiyle birlikte dolandırıcılık vakalarının artışı, kredi kartı işlemlerinde şüpheli davranışların belirlenmesini kritik bir araştırma alanına dönüştürmüştür. Geleneksel kural tabanlı yöntemler, yeni nesil dolandırıcılık tekniklerine karşı yetersiz kalırken, Makine Öğrenmesi (MÖ) (Machine Learning (ML)) tabanlı yaklaşımlar, büyük ölçekli işlem verilerinden örüntüleri öğrenerek dolandırıcılık faaliyetlerini daha etkin bir şekilde tespit edebilme potansiyeline sahiptir.

MÖ tekniklerinin dolandırıcılık tespit sistemlerine entegrasyonu, dolandırıcılık faaliyetlerinin belirlenmesinde doğruluğu ve operasyonel verimliliği artırabilir. Kredi kartı işlem veri kümelerinin dengesiz dağılımı ve dolandırıcıların sürekli gelişen taktikleri, dolandırıcılık tespit sistemleri için önemli zorluklar oluşturmaktadır. Çeşitli araştırmalarda, bu tür zorlukların üstesinden gelmek ve dolandırıcılığı daha etkin bir şekilde tespit edebilmek amacıyla hem denetimli hem de denetimsiz öğrenme yöntemlerini içeren farklı MÖ yaklaşımları incelenmiştir. Bu çalışmalar, özellikle dengesiz veri kümeleri üzerinde algoritmaların başarımını artırmaya yönelik stratejiler geliştirerek, dolandırıcılık faaliyetlerinin daha yüksek doğrulukla belirlenmesini sağlamaya odaklanmıştır.

MÖ teknikleri genel olarak denetimli ve denetimsiz öğrenme yöntemleri olmak üzere iki ana kategoriye ayrılmaktadır. Denetimli öğrenme, geçmiş verilere dayalı olarak geçerli ve hileli işlemleri ayırt edebilen modellerin, etiketli veri kümeleri üzerinde eğitilmesini içerir. Karar Ağaçları (KA) (Decision Trees (DT)), Rastgele Orman (RO) (Random Forest (RF)) ve Destek Vektör Makineleri (DVM) (Support Vector Machines (SVM)) gibi yaygın algoritmalar, dolandırıcılık tespiti görevlerinde yüksek doğruluk oranları elde etme konusunda etkili olduğu gösterilen yöntemler arasındadır [1]. Özellikle Yapay Sinir Ağları (YSA) (Artificial Neural Networks (ANN)) tabanlı modellerin, kredi kartı dolandırıcılığı tespitinde önemli başarılar sağladığı ve bu alandaki denetimli öğrenme yaklaşımlarının yüksek potansiyele sahip olduğu belirtilmektedir [2].

Denetimsiz öğrenme yöntemleri, etiketli verinin sınırlı olduğu durumlarda tercih edilmekte olup, genellikle işlem verileri içerisindeki anormal ya da sıra dışı kalıpları belirlemeye odaklanmaktadır. K-ortalama kümeleme, Lokal Aykırılık Faktörü ve İzolasyon Ormanı (İO) (Isolation Forest (IF)) gibi yöntemler, kredi kartı işlemlerinde olağandışı davranışları tespit etmek amacıyla kullanılmaktadır [3]. Bu yöntemler, dolandırıcılık işlemlerini istatistiksel olarak aykırı noktalar şeklinde ele alarak tespit sürecinin etkinliğini artırmayı hedeflemektedir [4]. Denetimli ve denetimsiz öğrenme tekniklerinin bir arada kullanılması, dolandırıcılık tespitinde daha yüksek başarı oranlarına ulaşılmasını

sağlayarak, farklı yöntemlerin eksikliklerini tamamlayıcı bir yaklaşımla gidermesine olanak tanımaktadır.

Mohsen, Nassreddine ve Massoud, kredi kartı dolandırıcılığının tespitinde MÖ tekniklerinin etkinliğini değerlendirmiştir. Çalışmada, Lojistik Regresyon (LR) (Logistic Regression), YSA ve DVM gibi sınıflandırma tekniklerinin yanı sıra, topluluk öğrenme yöntemlerinden RO algoritması kullanılmıştır. Farklı MÖ yöntemlerinin doğruluk oranları karşılaştırılarak, kredi kartı dolandırıcılığının tespiti için en uygun modelin belirlenmesi amaçlanmıştır [5].

Oketola ve arkadaşları, çevrimiçi işlemlerin artmasıyla birlikte kredi kartı dolandırıcılığının nasıl evrildiğini incelemiş ve bu bağlamda denetimli ve denetimsiz makine öğrenmesi yöntemlerini değerlendirmiştir. Çalışmada, LR, KA, RO, K-En Yakın Komşu (K-EYK) (K-Nearest Neighbors (KNN)), DVM ve Naive Bayes (NB) gibi sınıflandırma tekniklerinin yanı sıra, K-ortalama kümeleme yöntemi kullanılmıştır. Ayrıca, hibrit yöntemler de incelenmiş ve modellerin performansları doğruluk, kesinlik, duyarlılık, F1 skor, ROC ve AUC gibi değerlendirme ölçütleri açısından karşılaştırılmıştır [6].

Muameleci, Avrupa'daki bir aylık kredi kartı işlemlerinden oluşan bir veri kümesinde anormal tespiti amacıyla Çok Değişkenli Genelleştirilmiş Pareto Dağılımı (ÇDGPD) (Multivariate Generalized Pareto Distribution (MGPD)) yönteminin performansını incelemiştir. Çalışmada, ÇDGPD' nin performansı, denetimli bir makine öğrenme algoritması olan Tam Bağlantılı Konvolüsyonel Sinir Ağı (TBKSA) (Fully Connected Convolutional Neural Networks (FFCNN)) ile denetimsiz algoritmalar olan İO ve DVM ile karşılaştırılmıştır. Sonuçlar, ÇDGPD' nin İO ve DVM'ye kıyasla bazı durumlarda daha iyi performans gösterdiğini, ancak en yüksek başarı oranının TBKSA modeli tarafından elde edildiğini ortaya koymuştur [7].

Awoyemi ve arkadaşları, yüksek derecede dengesiz kredi kartı dolandırıcılığı verisi üzerinde NB, K-EYK ve LR algoritmalarının performansını incelemiştir. Çalışmada, veri dengesizliğini gidermek amacıyla alttan ve üstten örnekleme tekniklerinin birlikte kullanıldığı hibrit yöntemler uygulanmıştır. Elde edilen sonuçlar, K-EYK algoritmasının diğer yöntemlere kıyasla en iyi performansı gösterdiğini ortaya koymuştur [8].

Hussein ve arkadaşları, bulanık kaba en yakın komşu yöntemi ile lojistik regresyonu birleştiren hibrit bir model önermiş ve topluluk yöntemlerinin tespit doğruluğunu artırma potansiyelini ortaya koymuştur [9]. Benzer şekilde, Setiawan ve arkadaşları, LR ile KA yöntemlerini karşılaştırmış ve her iki tekniğin de etkili olduğunu, ancak performanslarının veri kümesinin özelliklerine ve mevcut dolandırıcılık desenlerine bağlı olarak değişkenlik gösterdiğini belirlemiştir [10].

Alarfağ ve arkadaşları, yaptıkları çalışmada, KA, RO, DVM, LR ve Konvolüsyonel Sinir Ağı (KSA) (Convolutional Neural Network (CNN)) gibi yöntemleri değerlendirilmiş, KSA modelinin doğruluk %99.9, F1 skoru %85.71, kesinlik %93 ve AUC %98 açısından üstün performans gösterdiği bulunmuştur. Sonuçlar, derin öğrenme modellerinin dolandırıcılık tespitinde daha başarılı olduğunu ve katman sayısının artırılmasının doğruluğu iyileştirdiğini göstermektedir [11].

Varmedja ve arkadaşları, kredi kartı dolandırıcılığının tespitinde farklı MÖ algoritmalarının performansını karşılaştırmıştır. Çalışmada, LR, RO, NB ve Çok Katmanlı Algılayıcı (ÇKA) (Multilayer Perceptron (MLP)) algoritmaları değerlendirilmiş, dengesiz veri yapısını dengelemek için Sentetik Azınlık Aşırı Örnekleme Tekniği (SAAÖT) (Synthetic Minority Over-sampling Technique (SMOTE)) kullanılmıştır. Deney sonuçlarına göre, RO algoritması %99.96 doğruluk, %96.38 kesinlik ve %81.63 duyarlılık ile en iyi performansı göstermiştir. Ayrıca, uygun veri ön işleme ve örnekleme teknikleri kullanıldığında geleneksel MÖ algoritmalarının, derin öğrenme yöntemleriyle benzer doğruluk seviyelerine ulaşabileceği ortaya konmuştur

Yapılan çalışmalar, kredi kartı dolandırıcılığının tespitinde farklı MÖ ve derin öğrenme yöntemlerinin etkinliğini ortaya koymuş ve veri dengesizliği gibi sorunların giderilmesiyle model performansının iyileştirilebileceğini göstermiştir. Bu doğrultuda, mevcut yaklaşımların yanı sıra, yeni yöntemlerin entegrasyonu ve hibrit modellerin kullanımı dolandırıcılık tespitinin doğruluğunu artırma potansiyeline sahiptir. Bu çalışmada, kredi kartı dolandırıcılığının tespitine yönelik çeşitli MÖ teknikleri kapsamlı bir şekilde incelenmiş ve özellikle K-medoids kümeleme yöntemi ile Boosting ve BERT tabanlı sınıflandırma algoritmalarının entegrasyonu analiz edilmiştir. Çalışma kapsamında, denetimli ve denetimsiz öğrenme yöntemlerinin hibrit bir yapıda birleştirilmesiyle dolandırıcılık tespitinde daha yüksek başarı oranına sahip bir model geliştirilmiştir. Farklı eğitim-test oranlarında gerçekleştirilen deneysel analizler sonucunda, LR, YSA, DVM, RO, XGBoost, Gradyan Artırma (GA) (Gradient Boosting (GB)), Hard ve Soft Voting Topluluk yöntemleri, Minkowski-K-EYK ve BERT (Bidirectional Encoder Representations from Transformers) gibi algoritmaların performansı değerlendirilmiştir. Algoritmaların başarımı, doğruluk ve F1 skoru metrikleri üzerinden karşılaştırılmıştır. K-medoids kümeleme, Boosting tabanlı modeller ve BERT ile entegre edilerek hibrit olarak geliştirilen bu yöntemde F1 skoru ve doğrulukla değerlendirilen modeller arasında GA algoritması tüm senaryolarda en yüksek F1 skorunu elde ederek en başarılı yöntem olmuştur.

Elde edilen sonuçlar, dolandırıcılık tespiti alanında daha yüksek doğruluk oranlarına ulaşmak için hibrit yaklaşımların etkili bir çözüm sunduğunu göstermektedir.

Bu doğrultuda geliştirilen yöntemlerin, gerçek zamanlı finansal güvenlik sistemlerinde uygulanabilirliği artırarak, kredi kartı dolandırıcılığıyla mücadelede önemli katkılar sağlaması hedeflenmektedir.

Materyal ve Metot

Çalışmada, Avrupa'daki kredi kartı işlemlerinden oluşan, anonimleştirilmiş 284.807 kayıt içeren bir veri seti kullanılmıştır. Kümeleme ve sınıflandırma teknikleri birlikte kullanılarak hibrit bir yaklaşım benimsenmiş ve modelin başarımını artırmaya yönelik çeşitli optimizasyon yöntemleri uygulanmıştır. Öncelikle, K-medoids kümeleme algoritması ile veri setindeki benzer işlem grupları belirlenmiş, ardından sınıflandırma algoritmaları kullanılarak dolandırıcılık tespit süreci gerçekleştirilmiştir. Veri setine uygulanan ön işleme adımlarının ardından, K-medoids kümeleme yöntemi ile veri gruplandırılmış; her bir küme üzerinde çeşitli sınıflandırma algoritmaları çalıştırılmış ve elde edilen sonuçlar doğruluk ve F1 skoru gibi metriklerle değerlendirilmiştir.

Özellik Çıkarımı Yöntemleri

Özellik mühendisliği, doğru özelliklerin seçilmesi yoluyla modelin performansını önemli ölçüde artırabileceği, karmaşıklığını azaltabileceği ve yorumlanmasını kolaylaştırabileceği için, makine öğrenimi sürecinde kritik bir adımdır. Bu bağlamda, özellik seçimi teknikleri orijinal özellikleri dönüştürmek yerine, bu özelliklerin en anlamlı alt kümesini seçmeye odaklanır. Seçilen bu alt küme, genellikle hem modelin hesaplama süresini ve karmaşıklığını azaltır hem de modelin tahmin gücünü ve sınıflandırma başarımını artırabilir. Başka bir deyişle, özellik seçimi; en alakalı ve öngörü gücü yüksek özellikleri belirleyerek, modelin yeni ve görülmemiş veriler üzerindeki doğruluğunu iyileştirmeye katkı sağlayabilir [12].

Bilgi Kazancı

Bilgi Kazancı (BK) (Information-Gain (IG)) tekniği, entropiyi kullanarak özellik mühendisliği yapmaktadır. Entropi, sistemdeki düzensizliği ölçen bir kavramdır. Dolayısıyla BK özellik çıkarım tekniği de veri kümesindeki özellikler için her bir kategori bazında bilgi kazanım oranını ölçerek özellik mühendisliği yapmaktadır ve denklemi (1) şöyledir [13]:

$$\text{Bilgi Kazancı} = \text{Entropi (Sınıf)} - \text{Entropi (Sınıf / Özellik)} \quad (1)$$

Fisher Skor

Fisher Skor (FS), FS fonksiyonu aracılığıyla, veri kümesindeki sınıflar arasında en yüksek ayrım gücüne sahip öznelilikleri seçmeyi amaçlayan bir özellik seçimi yöntemidir. FS özellik çıkarımının hesaplama yaptığı formül Denklem 2'de verilmektedir [14]:

$$F(x_i) = \frac{|\mu_i^k - \mu_i^j|}{\alpha_i^k - \alpha_i^j} \quad (2)$$

Burada, k sembolü veri kümesindeki sınıflardan birini ve j sembolü de veri kümesindeki diğer sınıfı; α_i^+ ve α_i^- değerleri de bu sınıfları standart sapma değerlerini ve μ_i^k ve μ_i^j değerleri yine bu sınıfların aritmetik ortalamalarını göstermektedir. FS özellik çıkarım yöntemi, özelliklerin sıralamasını fisher puanına göre azalan sırada döndürür ve daha sonrasında özellikler istenilen sayıda seçilebilir.

Ki-Kare

Ki-kare (Chi-Squared) özellik çıkarımı tekniği oldukça yaygın olarak kullanılan ve istatistiksel olarak anlamlılık hesaplaması yapan bir özellik seçim yöntemidir. Bu özellik çıkarımı tekniği, genellikle birbirinden bağımsız iki değer arasında niteliksel olarak ne ölçüde bir bağlantı olup olmadığını ölçen Ki-kare isimli bir teste dayanmaktadır [15]. Her özellik için Ki-kareyi hesaplar ve en iyi Ki-kare puanlarına sahip istenilen sayıda özelliği seçer. Ki-karenin gerçekçi bir hesaplama yapabilmesi için değişkenler kategorik olmalı ve bağımsız olarak örneklenmelidir.

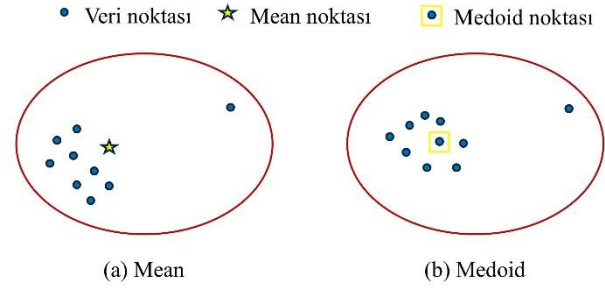
Bu çalışmamızda denetimli özellik seçimi yöntemlerinden BK, FS ve Ki-kare teknikleri ile denetimsiz özellik seçimi yöntemlerinden PCA tekniği kullanıldı. Özellik seçimi yöntemleri veri kümesi üzerinde kullanıldıktan sonra ayrı ayrı, K-medoids+LR, K-medoids+DVM, K-medoids+TSA, K-medoids+BERT, K-medoids+Minkowski-K-EYK, K-medoids+RO, K-medoids+HV-Ensemble, K-medoids+SV-Ensemble, K-medoids+XGBoost ve K-medoids+GA algoritmaları çalıştırılarak deneyler gerçekleştirilmiştir.

Kümeleme Algoritmaları

Kümeleme, bir nesne veya örnek setini gruplama süreci olarak tanımlanabilir. Bu sürecin temel amacı, aynı küme içerisindeki nesnelerin birbirine benzer özellikler göstermesi, ancak farklı kümelerde yer alan nesnelerden belirgin şekilde ayrışmasıdır [16]. Kümeleme yöntemleri genel olarak hiyerarşik ve hiyerarşik olmayan yaklaşımlar olmak üzere iki ana gruba ayrılmaktadır. Hiyerarşik olmayan kümeleme algoritmaları arasında en yaygın kullanılanlar, K-ortalama (K-means) ve medoid merkezli kümeleme (K-medoids) algoritmalarıdır. (Şekil 1).

K-ortalama algoritması, küme sayısı kadar merkez noktası belirleyerek, kümedeki örneklerin koordinatlarını kullanarak yinelemeli bir şekilde bu merkezleri güncellemekte ve her nesneyi en yakın merkez noktasına atamaktadır. Buna karşılık, K-medoids algoritması, kümedeki nesnelerin koordinatlarının ortalamasını almak yerine, kümenin ortalamasına en yakın olan bir nesneyi merkez olarak seçmekte ve bu nesneyi kümenin temsili merkezi olarak kabul etmektedir [17]. Bu iki algoritma arasındaki temel farklar dikkate alındığında, K-ortalama algoritmasının düşük hesaplama karmaşıklığına sahip

olması ve kolay uygulanabilirliği avantaj sağlarken, K-medoids algoritmasının aykırı değerler karşısında daha dayanıklı olduğu ve bazı veri kümelerinde daha yüksek performans gösterebildiği belirtilmektedir [18].



Şekil 1. K-ortalama ve K-medoids kümeleme yöntemleri arasındaki farkın grafiksel gösterimi [19]

Sınıflandırma Algoritmaları

Lojistik Regresyon

LR, bağımlı değişkenin kategorik olduğu durumlar için kullanılan istatistiksel bir sınıflandırma algoritmasıdır. Özellikle ikili sınıflandırma problemlerinde yaygın olarak tercih edilen bu yöntem, bir olayın gerçekleşme olasılığını tahmin etmek için sigmoid fonksiyonundan yararlanmaktadır. Model, giriş değişkenleri ile hedef değişken arasındaki ilişkiyi doğrusal bir fonksiyon üzerinden kurmakta ve çıktı değerlerini 0 ile 1 arasında bir olasılık şeklinde üretmektedir [20]. Hesaplama verimliliği ve yorumlanabilirliği sayesinde istatistiksel analizlerde ve MÖ uygulamalarında tercih edilmekte olup, bankacılıkta kredi değerlendirme süreçlerinde de yaygın olarak kullanılmaktadır.

Destek Vektör Makineleri

DVM, öncelikli olarak sınıflandırma problemlerinde kullanılan bir MÖ algoritmasıdır. Temel prensibi, veri örneklerini ayrıştırarak sınıflandırmak ve sınıflar arasındaki en iyi ayrımı sağlayan hiper düzlemi belirlemektir. Bu ayrım çizgisi marjin olarak adlandırılmakta olup, sınırın her iki tarafında yer alan ve en yakın veri noktalarını temsil eden destek vektörleri ile belirlenmektedir. DVM, doğrusal olmayan veri setlerinde çekirdek fonksiyonları kullanarak veriyi daha yüksek boyutlu uzaylara dönüştürmekte ve böylece doğrusal olmayan veri dağılımlarında da etkili bir sınıflandırma gerçekleştirebilmektedir. Ancak, düşük boyutlu veri setlerinde başarılı sonuçlar vermesine karşın, büyük ölçekli veri setlerinde yüksek hesaplama maliyeti nedeniyle işlem süresi önemli ölçüde artabilmektedir.

Tekrarlayan Sinir Ağları

Tekrarlayan Sinir Ağları (TSA) (Recurrent Neural Network (RNN)), özellikle sıralı verilerin işlenmesine yönelik geliştirilen güçlü bir sinir ağı mimarisi olarak öne çıkmaktadır. Sahip olduğu geri besleme mekanizması, geçmiş bilginin depolanmasına ve mevcut hesaplamalara dahil edilmesine olanak tanıyarak, doğal dil işleme (DDİ) (Natural Language Processing (NLP)), konuşma tanıma ve zaman serisi tahmini gibi zamana bağlı veri içeren görevlerde yüksek performans sergilemesini sağlamaktadır [21].

$$h_t = \tanh(W_{hh}h_{t-1} + W_{xh}x_t) \quad (3)$$

Denklem 3'te W ağırlık, h gizli katman, W_{hh} bir önceki gizli katman ağırlığı, W_{hx} mevcut gizli katman ağırlığı ve $\tanh$ aktivasyon fonksiyonunu temsil etmektedir.

BERT

BERT, Google tarafından geliştirilen ve DDİ amacıyla kullanılan bir MÖ modelidir. BERT, kelimeleri yalnızca bağımsız olarak değerlendirmek yerine, cümleyi bütünsel olarak analiz ederek bağlam içindeki anlamlarını belirlemektedir. Cümle içindeki belirli kelimelere odaklanmaksızın tüm cümleyi dikkate alarak yorumlama yapar ve elde ettiği anlam doğrultusunda arama sonuçlarını oluşturur. Bu nedenle, cümle yapısının düzeni ve kelimelerin uyumu, modelin beklenen doğrulukla çalışabilmesi açısından önemlidir [22].

BERT, kelimeleri bireysel veya gruplar halinde incelemek yerine, cümlelerin tamamını ele alıp yorumlamakta ve çıkarımlarını bu bütüncül analiz doğrultusunda yapmaktadır. Bu nedenle, BERT'in girdi olarak aldığı cümlelerin düzeni, kelimelerin sıralaması ve birbirleriyle olan uyumu gibi anlamı değiştirebilecek her türlü dilbilgisel unsur, algoritmanın performansını doğrudan etkileyebilmektedir. Tıpkı Long Short-Term Memory (LSTM) ağı gibi, bir cümledeki her kelimenin hem önceki hem de sonraki kelimelerle olan ilişkisini anlamlandırabilen çift yönlü bir yaklaşım kullanmaktadır. Bu özelliği sayesinde, doğal dilin karmaşıklığı ve bağlamsal zenginliği yüksek ortamlarda etkili bir şekilde çalışarak başarılı sonuçlar üretebilmektedir [23],[24]. Bu çalışmada, kredi kartı kullanımındaki şüpheli davranışları tespit etmek amacıyla BERT'in yanı sıra, önceden eğitilmiş modellerinden biri olan XLM-R gibi alternatif modeller de deneylerimizde değerlendirilmiştir [25]-[28].

Minkowski Mesafesi k-En Yakın Komşular

K-EYK, sınıflandırma ve regresyon için kullanılan parametrik olmayan bir tekniktir. Bir veri noktasını Minkowski mesafesine göre en yakın komşularının çoğuna göre sınıflandırır. Genel olarak Öklid ve Manhattan mesafesini genelleştirir ve şu şekilde tanımlanır:

$$f(x) = (\sum_{i=1}^n (x_i - y_i)^p)^{\frac{1}{p}} \quad (4)$$

Burada Denklem 4'te p metriğin türünü belirler. Bu esneklik K-EYK algoritmasının farklı veri dağılımlarına ve uygulama gereksinimlerine uyum sağlamasına olanak tanır [29].

Rastgele Orman

RO, sınıflandırma ve regresyon için kullanılan bir topluluk öğrenme tekniğidir. Eğitim sırasında birden fazla karar ağacı oluşturmakta ve sınıflandırmada sınıfların dış modlarını regresyonda ise tek tek ağaçların ortalama tahminlerini kullanmaktadır. Bu her bir ağacın eğitim verilerinden rastgele bir alt küme ve rastgele bir özellik alt kümesi tarafından oluşturulmasını sağlayarak modeli aşırı öğrenme ve genelleme sorunlarını azaltmada sağlam hale getirmektedir [30].

Topluluk Algoritmaları

Hard Voting Ensemble Algoritması

Hard Voting Ensemble (HV-Ensemble) yöntemi doğruluğu artırmak için çeşitli MÖ algoritmalarından yapılan tahminleri birleştirir. Topluluktaki her sınıflandırıcı çoğunluk oyu ile nihai çıktıyı tahmin etmeye çalışır. Bu yaklaşım özellikle bireysel modeller verilerin farklı bölümlerinde iyi performans gösterdiğinde ve dolayısıyla daha sağlam bir genel tahmin sağlamak için güçlü yönlerinden yararlandığında işe yaramaktadır [31].

Soft Voting Ensemble Algoritması

Soft Voting Ensemble (SV-Ensemble), hard voting yöntemine benzemekte ancak sınıf etiketleri yerine tahmin edilen olasılıkları kullanmaktadır. Topluluktaki her sınıflandırıcıdan her sınıf için bir olasılık tahmini sunması beklenmektedir. Nihai tahmin bu olasılıkların ortalamasıdır ve en yüksek ortalama olasılığa sahip sınıf seçilmektedir. Bu genellikle hard voting yönteminden daha iyi çalışmakta çünkü her sınıflandırıcının tahmininin güvenini dikkate almaktadır [32].

Extreme Gradient Boosting

Extreme Gradient Boosting (XGBoost) verimliliği, esnekliği ve taşınabilirliği ile bilinen güçlü bir MÖ algoritmasıdır. Tahminler yapmak için GA yöntemini kullanmakta ve genellikle sıralama, sınıflandırma ve regresyon problemlerinde üstün performans sergilemektedir.

Yüksek doğruluğu hızlı hesaplama süresi ve ölçeklenebilirliği nedeniyle yarışmalarda ve endüstriyel uygulamalarda sıklıkla tercih edilir. Algoritma büyük veri setlerinde ve karmaşık modellerde etkili olup optimizasyon ve düzenleme teknikleriyle aşırı öğrenmeyi azaltır ve genel performansı artırır. Paralel hesaplama yetenekleri sayesinde eğitim sürecini hızlandırır ve büyük veri kümeleriyle çalışırken avantaj sağlar. XGBoost ardışık olarak bir dizi zayıf öğrenici (genellikle karar ağaçları) inşa eder ve her yeni ağaç önceki ağaçların hatalarını düzeltmeye çalışır. Algoritma belirli bir kayıp fonksiyonunu minimize ederek çalışır. XGBoost fonksiyonu, tüm sınıflandırıcılar (j tree'ler) için bir düzenleme fonksiyonunun ve tamamında hesaplanan bir kayıp fonksiyonunun ürünüdür. Denklem 5, j tree ile tahmin üreten matematiksel yolun formülüdür.

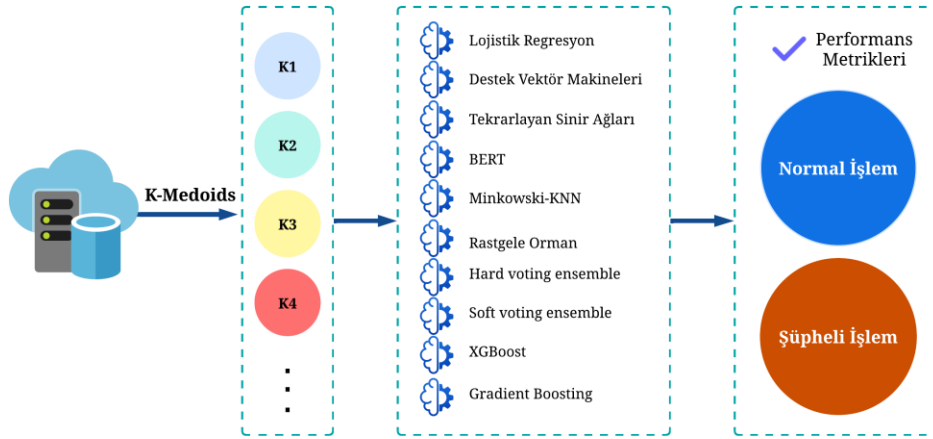
$$obj(0) = \sum_i^n l(y - y^{\wedge}) + \sum_{j=1}^J \Omega(f_j) \quad (5)$$

Öğrenmenin ardından yerleşik ağaçları veya algoritmaları kullanarak yeni veri noktaları Denklem 6 çözülerek bulunabilir. [33].

$$F2(x) = H0(x) + \eta(H1(x)) + \eta(H2(x)) \quad (6)$$

Gradyan Artırma

GA ise makine öğrenmesinin öne çıkan algoritmalarından biridir ve çoğunlukla sınıflandırma ve regresyon faaliyetleri için kullanılır.



Şekil 2. Önerilen yöntemin akış şeması

GA, zayıf öğrencileri bir araya getirerek güçlü bir tahmin modeli oluşturan bir topluluk öğrenme tekniğidir. Ardışık olarak zayıf tahmin ediciler ekleyerek ve bu tahmincilerin hatalarını düzelterek çalışır. Her bir zayıf tahmin edici önceki tahmincilerin hatalarını düzeltmek için bir hata fonksiyonu kullanır. GA algoritması karmaşık yapıları ve yüksek boyutlu veri setlerini etkili bir şekilde modelleyebilir, ancak öğrenme süreci daha uzun sürebilir ve aşırı uydurma riski taşır. Bu nedenle GA'nın hiper parametrelerinin dikkatlice ayarlanması önemlidir [32].

Hibrit Yaklaşım: Gözetimli ve Gözetimsiz Algoritmaların Bir Arada Kullanılması

Hem akademik çalışmalarda hem de sektörel problemlerin çözümünde denetimsiz kümeleme algoritmaları yaygın olarak kullanılmaktadır. K-ortalama algoritması, K-medoids algoritması, aşağıdan yukarı şeklinde bir yaklaşım olan aglomeratif (agglomerative) kümeleme algoritması, yukarıdan aşağı şeklinde bir yaklaşım olan bölünmüş kümeleme algoritması, BIRCH algoritması, OPTICS algoritması, Spektral Kümeleme algoritması, DBSCAN algoritması, DENCLUE algoritması en çok kullanılan kümeleme algoritmaları arasındadır [33].

Kredi kartları üzerindeki şüpheli davranış tespiti konulu bu çalışmada hem kümeleme hem de sınıflandırma algoritmalarını bir araya getiren hibrit bir yaklaşım tasarlanmıştır. Bu çalışmada, daha önce ağlarda nüfuz tespiti için kullanılan hibrit yaklaşım modeli esas alınarak, kredi kartlarındaki şüpheli davranışların tespitine yönelik entegrasyon sağlanmıştır [34]. Toplamda 284.807 işlem içeren veri seti, öncelikle K-medoids kümeleme algoritması kullanılarak, genel ve varsayılan parametreler doğrultusunda kümeler ayrıştırılmıştır. Daha sonra, belirlenen kümeler üzerinde LR, DVM, TSA, BERT, Minkowski-K-EYK, RO, HV-Ensemble, XGBoost ve GA sınıflandırma algoritmaları uygulanmıştır. Son değerlendirme aşamasında, alt kümelerden elde edilen sınıflandırma sonuçlarına ait karışıklık matrisinde yer alan Doğru Pozitif (DP), Yanlış Pozitif (YP), Doğru Negatif (DN) ve Yanlış Negatif (YN) değerleri toplanarak, tüm test seti için sınıflandırma başarısını temsil eden toplam bir değer elde edilmiştir.

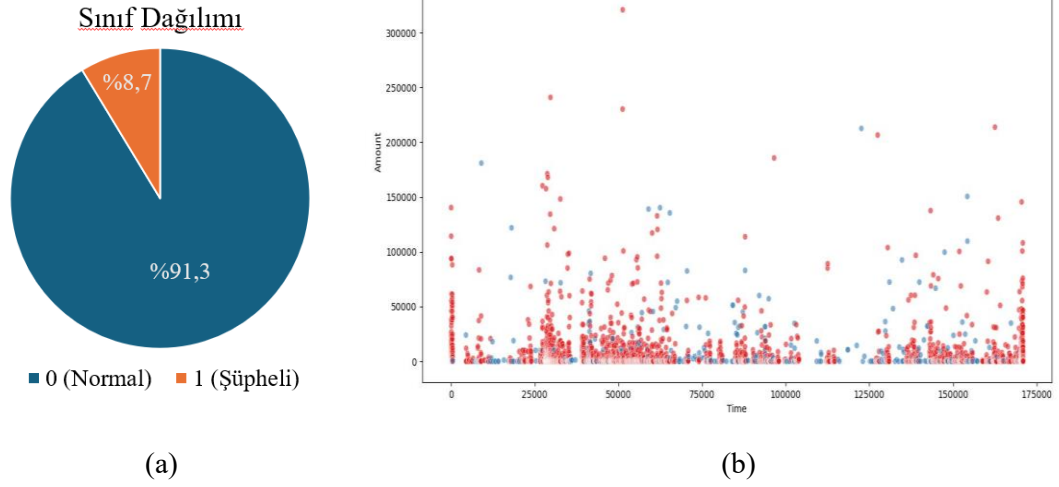
İlk olarak, veri seti ön işleme aşamasına tabi tutulmaktadır. Ardından, metin verilerinin uygun gösterimleri oluşturularak dönüştürme işlemi gerçekleştirilmektedir. Daha sonra, dönüştürülmüş veri seti üzerinde kümeleme işlemi uygulanmakta ve her kümede ayrı ayrı sınıflandırma yapılmaktadır. Son aşamada ise, her kümenin sınıflandırma sonuçları ve performans ölçütleri detaylı olarak değerlendirilerek modelin genel başarımı analiz edilmektedir (Şekil 2).

DeneySEL Sonuçlar

Bu bölümde, kullanılan MÖ algoritmalarının performansları detaylı bir şekilde analiz edilmiştir. Çalışmada kullanılan veri seti, Avrupa'daki kredi kartı işlemlerinden oluşmakta ve PCA dönüşümleri ile anonimleştirilmiştir. Algoritmaların doğruluk ve F1 skoru performans metriklerine göre karşılaştırılması yapılmıştır. Sonuçlar, özellikle dengesiz veri setlerinde doğru model seçiminin dolandırıcılık tespiti üzerinde önemli bir etkisi olduğunu göstermektedir. XGBoost algoritmasının üstün performans sergilediği gözlemlenmiş, bu da doğru model seçiminin dolandırıcılığı önleme çabalarına önemli katkılar sağlayabileceğini ortaya koymuştur.

Deneyler kapsamında BERT, XML-RoBERTa, DistilBERT ve Electra modelleri değerlendirilmiş olup, elde edilen sonuçlar karşılaştırmalı olarak analiz edilmiştir. Yapılan deneysel analizler sonucunda, geleneksel BERT modelinin diğer varyantlara kıyasla en yüksek doğruluk ve F1 skoru sağladığı gözlemlenmiştir. Bu nedenle, makalede sunulan nihai sonuçlar, BERT modeli üzerinden raporlanmıştır.

Veri kümesi, algoritmaların performansını daha iyi değerlendirebilmek ve literatürdeki çalışmaların sonuçlarıyla karşılaştırma yapabilmek amacıyla farklı oranlarda eğitim ve test kümelerine ayrılmıştır. Tüm algoritmalar, veri kümesinin %80 eğitim - %20 test, %40 eğitim - %60 test ve %30 eğitim - %70 test oranlarında ayrı ayrı çalıştırılmış; elde edilen sonuçlar sırasıyla Tablo 2, Tablo 3, Tablo 4 ve Tablo 5'te raporlanmıştır. Gerçek dünyadaki pek çok problemde olduğu gibi, kredi kartı dolandırıcılığıyla ilgili problemlerde de etiketli veriye ulaşmak ve bu veriyi oluşturmak hem zor hem de



Şekil 3. Veri setindeki sınıf dağılımı(a) ve zamana-miktara göre sınıf dağılımı(b)

maliyetlidir. Bu nedenle, özellikle %30 eğitim - %70 test gibi düşük eğitim oranlarında da algoritmalar çalıştırılarak performansları değerlendirilmiştir.

Veri Seti

Bu çalışmada, Avrupa'da gerçekleştirilen kredi kartı işlemlerinden oluşan ve sahtekarlık tespiti amacıyla kullanılan bir veri seti analiz edilmiştir. Erişimi herkese açık olan bu veri setinin içeriği şu şekildedir:

- Veri seti toplamda 284.807 işlemden oluşmakta olup, bu işlemlerin yalnızca %0.17'si dolandırıcılık olarak belirlenmiştir (Şekil 3).
- İşlemler, 2 günlük bir süre boyunca gerçekleştirilen Avrupa banka kartı sahiplerine ait verileri içerir. Bu süre boyunca, toplamda 492 sahte işlem tespit edilmiştir.
- Veri seti, 30 özellikten oluşmakta olup, PCA yöntemi ile anonimleştirilmiştir. Bu dönüşüm, özelliklerin gizliliğini korumak amacıyla gerçekleştirilmiş olup, orijinal veriler hakkında doğrudan bilgi sağlamamaktadır.
- "Time" özelliği işlemin zamanını, "Amount" özelliği işlemin miktarını, "Class" özelliği ise işlemin sahtekarlık olup olmadığını gösterir. "Class" özelliği, 0 ve 1 değerlerini alır; 0, sahte olmayan işlemleri ve 1, sahte işlemleri temsil eder [35]. Veri seti, iki farklı sınıf içermektedir. Sınıf 0, dolandırıcılık içermeyen işlemleri temsil etmekte olup, gerçek ve herhangi bir şüpheli davranış barındırmayan finansal işlemleri kapsamaktadır. Sınıf 1 ise, dolandırıcılık içeren işlemleri ifade etmekte ve şüpheli olarak rapor edilen, etik olmayan veya yasa dışı faaliyetleri işaret eden işlemleri içermektedir. Veri setindeki zamana ve işlem miktarına göre sınıf dağılımı, aşağıdaki grafikte gösterilmektedir.

Veri setinin dengesiz sınıf dağılımı göz önünde bulundurulduğunda, dolandırıcılık tespiti için

geliştirilecek MÖ modellerinde sınıf dengesizliği yönetimi kritik bir faktör olarak öne çıkmaktadır. Bu veri seti, dolandırıcılık tespiti ve güvenlik analizleri için önemli bir kaynak oluştururken, PCA dönüşümü sayesinde hem veri gizliliği korunmuş hem de işlemlerin derinlemesine analizi kolaylaştırılmıştır.

Sonuçlar

Eğitim oranlarına göre bakıldığında tüm algoritmalar %80 eğitim oranında en yüksek doğruluk skorlarına ulaşmaktadır. Bu durum daha fazla veri ile eğitimin algoritmaların performansını önemli ölçüde artırdığını göstermektedir.

Tablo 1. Veri seti özellikleri

Sınıf Tipi	Açıklama
Time	Her iki kart işleminin arasında geçen zamanı belirtir.
PCA Dönüşümleri	Veri setinde 30 tanesi mevcuttur. Gizlilik esasına uygun olması için yapılmış olan dönüşümlerdir.
Amount	Yapılan işlemin miktarını belirtir.
Class	Yapılan işlemin sahte olup olmadığını belirtir. Burada "1" sahte olduğunu, "0" ise olmadığını belirtir.

Eğitim ve test oranlarına göre çeşitli makine öğrenme algoritmalarının performansını değerlendirdiğimizde, K-medoids+GA algoritmasının tüm veri bölme senaryolarında en yüksek başarı oranlarını elde ettiği görülmektedir. %30 eğitim %70 test oranında %95.78, %40 eğitim %60 test oranında %95.98 ve %80 eğitim %20 test oranında %96.9 başarı oranlarına ulaşmıştır. Bu, K-medoids+GA algoritmasının yüksek doğruluk ve

genelleme kabiliyeti ile en iyi performans gösteren model olduğunu göstermektedir (Şekil 4).

İkinci en yüksek başarıyı K-medoids+XGBoost algoritması göstermektedir. Sırasıyla %30 eğitim %70 test oranında %94.56, %40 eğitim %60 test oranında %94.97 ve %80 eğitim %20 test oranında %95.45 doğruluğa ulaşmıştır.

Üçüncü sırada K-medoids+ SV-Ensemble algoritması yer almaktadır. Bu algoritmanın performansları %30 eğitim %70 test oranında %93.04, %40 eğitim %60 test oranında %93.43 ve %80 eğitim %20 test oranında %93.86 olarak kaydedilmiştir.

Tablo 2’de, medoid merkezli kümeleme yöntemi ile desteklenen farklı sınıflandırma algoritmalarının F1 skorları karşılaştırmalı olarak verilmiştir. Eğitim verisinin oranı arttıkça, tüm modellerin performansında genel bir iyileşme gözlemlenmiştir. En yüksek F1 skoruna, %80 eğitim oranıyla Gradient Boosting modeli ulaşmıştır. XGBoost ve Soft Voting Ensemble algoritmaları da başarılı sonuçlar elde etmiş, özellikle %95’in üzerinde F1 skorlarına ulaşarak dikkat çekmiştir.

Tablo 2. K-medoids kümeleme yöntemi ile kullanılan algoritmaların F1 skor değerleri

K-medoids+ Algoritma	%30 eğitim %70 test	%40 eğitim %60 test	%80 eğitim %20 test
K-medoids+LR	88.1	89.32	90.48
K- medoids+DVM	91.32	91.67	92.14
K- medoids+TSA	91.68	91.99	92.75
K- medoids+BERT	92.2	92.69	92.78
K- medoids+ Minkowski-K-EYK	85.1	86.37	86.56
K- medoids+RO	87.34	88.07	88.96
K- medoids+ HV-Ensemble algoritması	92.64	92.78	92.95
K- medoids+ SV-Ensemble algoritması	93.04	93.43	93.86
K- medoids+XGBoost	94.56	94.97	95.45
K- medoids+GA	95.78	95.98	96.9

Medoid merkezli kümeleme yöntemi ile entegre edilen sınıflandırma algoritmaları değerlendirildiğinde, en

yüksek performansı GA algoritması göstermiştir. Bunu sırasıyla XGBoost ve SV-Ensemble yöntemleri takip etmektedir. HV-Ensemble ve BERT modelleri de tatmin edici sonuçlar vermiştir. TSA, DVM ve LR algoritmalarının performansı orta düzeyde kalırken, RO ve Minkowski-K-EYK algoritmaları diğer yöntemlere kıyasla daha düşük başarı oranları sergilemiştir. Bu sonuçlar, kullanılan yöntemlerin veri seti üzerindeki başarımlarını açıkça ortaya koymaktadır.

Ayrıca, özellik seçimi tekniklerinden BK, FS ve Ki-kare testi yöntemleri ayrı ayrı tüm algoritmalar ve tüm eğitim/test yüzdeleri üzerinde uygulanmış ve deney sonuçları Tablo 3, Tablo 4 ve Tablo 5’te raporlanmıştır. Tablo 3 BK tekniği ve K-medoids kümeleme yöntemi ile kullanılan algoritmaların F1 skor değerlerini göstermektedir. Tablo 3’e göre filtrelemeye dayalı denetimli bir özellik seçimi tekniği olan BK kullanımı ile nerdeyse tüm algoritmalar ve eğitim/test yüzdeleri bazında sınıflandırma performansının arttığı gözlenmektedir.

Tablo 3. Bilgi Kazancı tekniği ve K-medoids kümeleme yöntemi ile kullanılan algoritmaların F1 skor değerleri

K-medoids+ Algoritma	%30 eğitim %70 test	%40 eğitim %60 test	%80 eğitim %20 test
K-medoids+LR	89.42	90.66	91.81
K-medoids+DVM	92.61	93.01	93.52
K-medoids+TSA	93.02	93.17	94.12
K-medoids+BERT	93.29	93.89	94.18
K-medoids+ Minkowski-K-EYK	86.32	87.59	87.74
K-medoids+RO	88.46	89.29	90.21
K-medoids+ HV-Ensemble algoritması	93.95	94.08	94.25
K-medoids+ SV-Ensemble algoritması	94.34	94.74	95.17
K-medoids+XGBoost	95.67	96.14	96.60
K-medoids+GA	96.93	97.12	98.04

Benzer bir durum yine bir başka denetimli özellik seçimi tekniği olan ve literatürde yine BK tekniği gibi sıklıkla kullanıldığını gözlemlediğimiz FS tekniğinde de görülmektedir. Tablo 4 FS tekniği ve K-medoids

kümeleme yöntemi ile kullanılan algoritmaların F1 skor değerlerini göstermektedir.

Tablo 4. Fisher Skor tekniği ve K-medoids kümeleme yöntemi ile kullanılan algoritmaların F1 skor değerleri

K-medoids+ Algoritma	%30 eğitim %70 test	%40 eğitim %60 test	%80 eğitim %20 test
K-medoids+LR	89.30	90.57	91.75
K-medoids+DVM	91.85	92.92	93.43
K-medoids+TSA	92.96	93.28	94.05
K-medoids+BERT	93.49	93.79	93.78
K-medoids+ Minkowski-K-EYK	86.19	87.74	87.65
K-medoids+RO	88.54	89.28	90.17
K-medoids+ HV-Ensemble algoritması	93.84	94.08	94.25
K-medoids+ SV-Ensemble algoritması	94.34	94.74	95.17
K-medoids+XGBoost	95.88	96.27	96.79
K-medoids+GA	97.12	97.40	98.26

Tablo 5 ise Ki-kare tekniği ve K-medoids kümeleme yöntemi ile kullanılan algoritmaların F1 skor değerlerini göstermektedir. BK ve FS özellik seçimi yöntemlerinin aksine, Ki-kare özellik seçimi yöntemiyle yapılan deneylerdeki sınıflandırma başarısının PCA tekniği kullanılarak oluşturulmuş olan deney durumlarının sınıflandırma başarısına göre daha düşük olduğu gözlenmektedir.

BK yöntemi, entropi temelli hesaplama yaparak bilgi açısından en değerli özellikleri ön plana çıkarırken; FS yöntemi, sınıflar arasındaki farkı artırıp sınıf içindeki benzerliği azaltarak en ayırt edici özellikleri seçmeye çalışır. Bu nedenle her iki yöntem de veri kümesindeki önemli ve anlamlı öznitelikleri başarıyla belirleyebilmekte ve sınıflandırma performansını artırmaktadır. Öte yandan, Ki-kare yöntemi istatistiksel olarak iki değişken arasındaki ilişkiyi ölçer ve etkili çalışabilmesi için özellikle kategorik verilerle ve bağımsız örneklerle çalışması gerekir. Ancak bu çalışmada kullanılan veri seti, PCA dönüşümüyle sayısal ve sürekli hale getirilmiş olduğu için Ki-kare yönteminin sınıflandırma performansına katkısı sınırlı kalmıştır. Bu nedenle, Ki-kare ile yapılan özellik seçimi sonrası elde edilen başarı oranları, BK ve FS yöntemleriyle elde edilen

sonuçların gerisinde kalmıştır. Bu durum, kullanılan özellik seçimi yöntemlerinin veri tipi ve yapısıyla ne kadar uyumlu olduğunun sonuçlar üzerinde doğrudan etkili olduğunu göstermektedir.

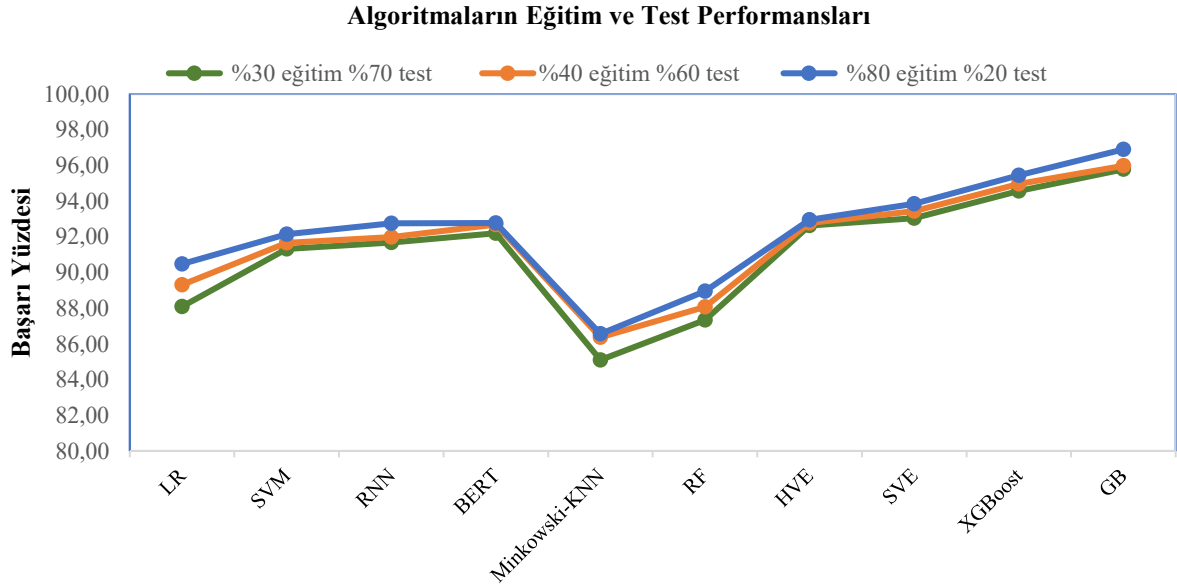
Tablo 5. Ki-kare tekniği ve K-medoids kümeleme yöntemi ile kullanılan algoritmaların F1 skor değerleri

K-medoids+ Algoritma	%30 eğitim %70 test	%40 eğitim %60 test	%80 eğitim %20 test
K-medoids+LR	87.31	88.52	89.67
K-medoids+DVM	90.48	90.84	91.31
K-medoids+TSA	90.85	91.12	91.75
K-medoids+BERT	91.37	91.86	91.94
K-medoids+ Minkowski-K-EYK	84.33	85.57	85.78
K-medoids+RO	86.55	87.38	88.16
K-medoids+ HV-Ensemble algoritması	91.81	91.94	92.23
K-medoids+ SV-Ensemble algoritması	92.20	92.59	93.02
K-medoids+XGBoost	93.71	94.13	94.59
K-medoids+GA	94.92	95.10	96.13

Bu çalışmada, kredi kartı dolandırıcılığının tespiti amacıyla, kümeleme yöntemi ile çeşitli sınıflandırma algoritmalarının entegrasyonunu içeren hibrit bir yaklaşım önerilmektedir. Veri kümesi üzerinde farklı algoritmalar ve çeşitli eğitim-test oranları kullanılarak elde edilen deneysel sonuçlar hem ticari hem de akademik ortamlar açısından oldukça motive edici ve umut vericidir. Önerilen yaklaşım, geliştirilebilecek bir kullanıcı arayüzü aracılığıyla platformdan bağımsız hale getirilebilir ve gerçek zamanlı çalışan bir sisteme entegre edilebilir.

Tartışma ve Sonuç

Bu çalışmada, suçluluk analizi açısından literatüre katkı sağlamak amacıyla farklı MÖ yöntemlerinin doğruluk oranları karşılaştırılarak kredi kartı dolandırıcılığı tespiti için etkili bir model oluşturulması hedeflenmiştir. Sınıflandırma teknikleri olarak LR, KA, RO, K-EYK, DVM yöntemleri değerlendirilmiş, kümeleme tekniği olarak ise K-medoids yöntemi incelenmiştir. Ayrıca, birden fazla yöntemin birleştirildiği hibrit modeller de analiz edilmiştir. Eğitilen modellerin performansları, doğruluk ve F1 skoru açısından karşılaştırılarak, kredi



Şekil 4. Algoritmaların Eğitim ve Test Performansları

kartı dolandırıcılığını tespit etmede etkin bir model geliştirilmesi amaçlanmıştır.

Kredi kartı dolandırıcılığını tespit eden tek ve mutlak üstünlüğe sahip bir algoritma bulunmamaktadır. Her veri madenciliği ve MÖ tekniği, kendine özgü güçlü ve zayıf yönleri sahiptir; bazı yöntemler yüksek doğruluk oranlarına ulaşırken, eğitim süresi açısından daha yavaş olabilmekte ya da tam tersi bir durum söz konusu olabilmektedir.

Çeşitli algoritmalar kullanılarak kredi kartı işlem veri seti üzerinde yapılan analizler sonucunda, modellerin yeni bir işlemin dolandırıcılık olup olmadığını belirleme süreçlerinde farklı sınıflandırma performansları sergilediği gözlemlenmiştir.

Bu çalışmada kredi kartı dolandırıcılığının tespiti için medoid merkezli kümeleme yöntemi ile çeşitli sınıflandırma algoritmalarının entegrasyonunu içeren hibrit bir yaklaşım önerilmektedir. Farklı eğitim ve test oranlarında gerçekleştirilen deneyler sonucunda, en yüksek başarıyı GA algoritmasının elde ettiği, bunu sırasıyla XGBoost ve SV-Ensemble algoritmalarının takip ettiği belirlenmiştir. Elde edilen sonuçlar, sınıflandırma algoritmalarının başarımının kullanılan veri setinin özelliklerine, uygulanan ön işleme tekniklerine ve veri dağılımına duyarlı olduğunu göstermektedir.

Algoritmaların F1 skoru ve doğruluk değerleri %85'in üzerinde gerçekleşmiş olup, bu durum tüm modellerin veri kümesi üzerinde makul düzeyde başarı sağladığını ortaya koymaktadır. Ayrıca, eğitim verisi oranı arttıkça performansın da yükseldiği, bu durumun ise modellerin öğrenme kapasitesi ile doğrudan ilişkili olduğu gözlemlenmiştir. En iyi sonuçlara ulaşan GA algoritması %96,9 F1 skoru ile öne çıkarken, klasik yöntemlerden LR ve RO daha düşük başarımla sergilemiştir.

Elde edilen bulgular, literatürde yer alan diğer çalışmalarla karşılaştırıldığında önemli benzerlikler ve farklılıklar göstermektedir. Örneğin, Awoyemi ve arkadaşları [36], K-EYK algoritmasının yüksek başarı sağladığını bildirirken, bu çalışmada Minkowski-K-EYK yöntemi diğer yöntemlere göre daha düşük performans sergilemiştir. Benzer şekilde, Muameleci [37], TBKSA modelinin sahtekarlık tespitinde üstün performans gösterdiğini ifade etmiştir; ancak bu çalışmada klasik makine öğrenimi algoritmaları arasında GA algoritması, %96,9 F1 skoru ile en başarılı model olmuştur. Alarfaj ve arkadaşları [38], derin öğrenme tabanlı modellerin yüksek doğruluk sağladığını belirtmiş; bu çalışmada ise BERT modeli başarılı olmakla birlikte, topluluk öğrenme tabanlı yöntemlerin (özellikle GA ve XGBoost) daha yüksek performans sunduğu görülmüştür. Bu sonuçlar, model başarımının kullanılan veri seti, veri dengesizliği, ön işleme yöntemleri ve algoritma parametrelerine göre değişebileceğini göstermekte; bu yönüyle çalışma, mevcut literatüre önemli bir katkı sağlamaktadır.

Bu çalışmada uygulanan özellik seçimi tekniklerinin sınıflandırma performansı üzerindeki etkisi açık biçimde gözlemlenmiştir. BK yöntemi ile en yüksek F1 skoru %98.04'e (GA algoritması ile) ulaşırken, FS yöntemiyle bu değer %98.26'ya yükselmiştir. Buna karşın Ki-kare yöntemi ile elde edilen en yüksek F1 skoru %96.13'te kalmıştır. Bu sonuçlar, BK ve FS yöntemlerinin sınıflar arası ayrımı daha iyi temsil ettiğini ve sınıflandırma başarımını artırmada daha etkili olduğunu göstermektedir. Ki-kare'nin daha düşük performansı ise veri yapısının sürekli ve PCA dönüşümlü olmasıyla ilişkilendirilmiştir.

BERT tabanlı modelin genel olarak yüksek doğruluk sağlamasına rağmen, çalışmamızda özellikle topluluk öğrenme yöntemleri (örneğin GA ve XGBoost) ile

karşılaştırıldığında daha düşük F1 skoru elde ettiği gözlemlenmiştir. Bu durum, kullanılan veri setinin özellikleri ve BERT'in doğasına bağlı olarak açıklanabilir. Çalışmada kullanılan veri seti, PCA ile dönüştürülmüş sayısal ve bağlamsal bilgi içermeyen anonimleştirilmiş özelliklerden oluşmaktadır. BERT gibi bağlam tabanlı dil modelleri, metin verilerinde anlamsal ilişkileri modelleyerek üstün performans göstermektedir; ancak bu çalışmada kullanılan yapay olarak oluşturulmuş, metinsel olmayan, bağlamdan yoksun veri yapısı, BERT'in derin bağlam modelleme yeteneklerinden tam anlamıyla faydalanmasına engel olmuştur. Buna karşın, ağaç tabanlı ensemble algoritmalar, sayısal verilerdeki örüntüleri daha etkin bir şekilde modelleyebilmiş ve bu nedenle daha yüksek performans sergilemiştir. Bu bağlamda, BERT'in potansiyelinin daha çok metinsel ya da bağlamsal veri içeren görevlerde ortaya çıktığı ve bu tür yapısal veri setleriyle sınırlı kaldığı söylenebilir.

Genel olarak değerlendirildiğinde, geliştirilen hibrit yaklaşımın, kredi kartı işlemlerinde şüpheli davranışları tespit etmede başarılı sonuçlar sunduğu görülmektedir. Bu yaklaşım, işlem verilerini kümelere ayırarak modellerin her alt gruba özel öğrenme yapmasına olanak tanımış ve dolandırıcılık tespitindeki doğruluğu artırmıştır. Özellikle topluluk temelli algoritmalar ile derin öğrenme modellerinin birlikte kullanımı hem geleneksel yöntemlerin yorumlanabilirliğini hem de derin öğrenmenin bağlamsal analiz gücünü birleştirerek daha güçlü sonuçlar üretmiştir.

Bu çalışmada geliştirilen sistem, bir grup kredi kartı işlem kümesindeki kullanıcı davranışlarını analiz ederek şüpheli işlemleri tespit etmeyi amaçlamaktadır. Önerilen yaklaşım, kümeleme ve çeşitli sınıflandırma algoritmalarının birleşiminden oluşan hibrit bir mimariye sahiptir. Bu yapı, sınıflandırma temelli herhangi bir veri kümesine entegre edilerek etkin bir şekilde sınıflandırma yapabilen esnek bir sistem sunmaktadır. Ayrıca, sistemde kullanılan özellik seçimi, kümeleme ve sınıflandırma algoritmaları ihtiyaçlara göre değiştirilebilir veya geliştirilebilir. Kısacası, bu yaklaşım sınıflandırmaya dayalı farklı problem alanlarına da kolaylıkla uyarlanabilecek esneklikte tasarlanmıştır.

Önerilen yöntemin yüksek sınıflandırma performansı sergilemesi hem sektörel uygulamalar hem de akademik literatür açısından önemli bir katkı olarak değerlendirilmektedir. Ancak, dünya genelinde dolandırıcılık vakalarının artması, bu alandaki sistemlerin daha karmaşık ve dinamik yapılandırılmasını gerekli kılmaktadır. Özellikle kredi kartı dolandırıcılığı tespit sistemlerinin gerçek zamanlı ve dinamik bir şekilde çalışabilmesi büyük fayda sağlayacaktır. Bu bağlamda, çalışmada sunulan yöntemin gerçek zamanlı olarak çalışmaması, en önemli kısıtlarından biri olarak öne çıkmaktadır. Ayrıca, önerilen yaklaşımın farklı işletim sistemleri ve yazılım kütüphaneleri üzerinde test

edilmemiş olması, platforma bağımlılığı artırmakta ve bu durum da bir diğer önemli kısıt olarak değerlendirilmektedir. Bunlara ek olarak, kredi kartlarının yoğun olarak kullanıldığı internet tabanlı bankacılık, ödeme sistemleri ve e-ticaret uygulamalarıyla entegrasyon sağlayabilecek bir kullanıcı arayüzünün bulunmaması, yöntemin uygulanabilirliğini sınırlayan bir başka kısıt olarak görülebilir.

Gelecekte, önerilen sistemin gerçek zamanlı sistemlerle entegre edilerek sahtekarlıkların anlık olarak tespit edilmesi ve önlenmesi hedeflenmektedir. Ayrıca, farklı türde işlem verileriyle genişletilmiş veri kümeleri ve zaman serisi özelliklerinin dikkate alındığı dinamik modeller ile başarımın daha da artırılacağı düşünülmektedir.

Etik kurul onayı ve çıkar çatışması beyanı

Hazırlanan makalede etik kurul izni alınmasına gerek yoktur.

Hazırlanan makalede herhangi bir kişi/kurum ile çıkar çatışması bulunmamaktadır.

Bu çalışma Marmara Üniversitesi Bilimsel Araştırma Projeleri Koordinasyon Birimi tarafından 11558 nolu proje kapsamında desteklenmiştir.

Kaynaklar

- [1] J. K. Afriyie et al., "A supervised machine learning algorithm for detecting and predicting fraud in credit card transactions," *Decision Analytics Journal*, vol. 6, p. 100163, 2023.
- [2] S. Singh, M. Kashyap, and N. Tantubay, "Comparative Analysis of ANN, RNN, and GRU for Credit Card Fraud Detection," 2025.
- [3] S. Ounacer et al., "Using Isolation Forest in anomaly detection: the case of credit card transactions," *Periodicals of Engineering and Natural Sciences (PEN)*, vol. 6, no. 2, pp. 394–400, 2018.
- [4] S. Jiang et al., "Credit card fraud detection based on unsupervised attentional anomaly detection network," *Systems*, vol. 11, no. 6, p. 305, 2023.
- [5] O. R. Mohsen, G. Nassreddine, and M. Massoud, "Credit Card Fraud Detector Based on Machine Learning Techniques," *Journal of Computer Science and Technology Studies*, vol. 5, no. 2, pp. 16–30, 2023.
- [6] A. G. Oketola, T. Gbadebo-Ogunmefun, and A. Agbeja, "A Review of Credit Card Fraud Detection Using Machine Learning Algorithms," 2023.
- [7] K. Muameleci, "Anomaly Detection in Credit Card Transactions using Multivariate Generalized Pareto Distribution," 2022.
- [8] J. O. Awoyemi, A. O. Adetunmbi, and S. A. Oluwadare, "Credit card fraud detection using machine learning techniques: A comparative analysis," in *Proc. Int. Conf. Computing Networking and Informatics (ICCNI)*, 2017, pp. 1–9.

- [9] N. S. Alfaiz and S. M. Fati, "Enhanced credit card fraud detection model using machine learning," *Electronics*, vol. 11, no. 4, p. 662, 2022.
- [10] R. Setiawan et al., "Fraud Detection in Credit Card Transactions Using HDBSCAN, UMAP and SMOTE Methods," *Int. J. Sci. Technol. Manage.*, vol. 4, no. 5, pp. 1333–1339, 2023.
- [11] F. K. Alarfaj et al., "Credit card fraud detection using state-of-the-art machine learning and deep learning algorithms," *IEEE Access*, vol. 10, pp. 39700–39715, 2022.
- [12] D. Theng and K. K. Bhoyar, "Feature selection techniques for machine learning: a survey of more than two decades of research," *Knowl. Inf. Syst.*, vol. 66, pp. 1575–1637, 2024. doi:10.1007/s10115-023-02010-5.
- [13] M. A. Hall, Correlation-based feature selection for machine learning, Ph.D. dissertation, Univ. Waikato, 1999.
- [14] Q. Gu, Z. Li, and H. Han, "Generalized Fisher Score for Feature Selection," *arXiv preprint, arXiv:1202.3725*, 2012.
- [15] H. Budak, "Özellik Seçim Yöntemleri ve Yeni Bir Yaklaşım," *Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, vol. 22, pp. 21–31, 2018.
- [16] J. Han, "Spatial clustering methods in data mining: A survey," in *Geographic Data Mining and Knowledge Discovery*, pp. 188–217, 2001.
- [17] N. K. Kaur, U. Kaur, and D. Singh, "K-Medoid clustering algorithm-a review," *Int. J. Comput. Appl. Technol.*, vol. 1, no. 1, pp. 42–45, 2014.
- [18] T. S. Madhulatha, "Comparison between k-means and k-medoids clustering algorithms," in *Proc. Int. Conf. Adv. Comput. Inf. Technol.*, Springer, Berlin, 2011.
- [19] A. Entezami, H. Sarmadi, and B. S. Razavi, "An innovative hybrid strategy for structural health monitoring by modal flexibility and clustering methods," *J. Civil Struct. Health Monit.*, vol. 10, no. 5, pp. 845–859, 2020.
- [20] M. P. LaValley, "Logistic regression," *Circulation*, vol. 117, no. 18, pp. 2395–2399, 2008.
- [21] A. Sherstinsky, "Fundamentals of recurrent neural network (RNN) and long short-term memory (LSTM) network," *Physica D: Nonlinear Phenomena*, vol. 404, p. 132306, 2020.
- [22] M. V. Koroteev, "BERT: a review of applications in natural language processing and understanding," *arXiv preprint, arXiv:2103.11943*, 2021.
- [23] A. Kara, "Global solar irradiance time series prediction using long short-term memory network," *Gazi Univ. J. Sci. Part C*, vol. 4, no. 7, pp. 882–892, 2019.
- [24] M. Lan et al., "Supervised and traditional term weighting methods for automatic text categorization," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 31, no. 4, pp. 721–735, 2008.
- [25] K. Clark, "Electra: Pre-training text encoders as discriminators rather than generators," *arXiv preprint, arXiv:2003.10555*, 2020.
- [26] A. Conneau, "Unsupervised cross-lingual representation learning at scale," *arXiv preprint, arXiv:1911.02116*, 2019.
- [27] V. Sanh, "DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter," *arXiv preprint, arXiv:1910.01108*, 2019.
- [28] S. Schweter, "BERTurk-BERT models for Turkish," *Zenodo*, 2020. doi:10.5281/zenodo.3770924.
- [29] P. Müller, "Flexible K nearest neighbors classifier: derivation and application for ion-mobility spectrometry-based indoor localization," in *Proc. Int. Conf. Indoor Positioning and Indoor Navigation (IPIN)*, IEEE, 2023.
- [30] S. Xuan et al., "Random forest for credit card fraud detection," in *Proc. IEEE Int. Conf. Networking, Sensing and Control (ICNSC)*, 2018.
- [31] P. Tomar, S. Shrivastava, and U. Thakar, "Ensemble learning based credit card fraud detection system," in *Proc. Conf. Information and Communication Technology (CICT)*, IEEE, 2021.
- [32] M. A. Mim, N. Majadi, and P. Mazumder, "A soft voting ensemble learning approach for credit card fraud detection," *Heliyon*, vol. 10, no. 3, 2024.
- [33] A. E. Ezugwu et al., "A comprehensive survey of clustering algorithms: State-of-the-art machine learning applications, taxonomy, challenges, and future research prospects," *Eng. Appl. Artif. Intell.*, vol. 110, p. 104743, 2022.
- [34] S. K. Çalışkan and İ. Soğukpınar, "Kxknn: K-means ve k en yakın komşu yöntemleri ile ağlarda nüfuz tespiti," *EMO Yayınları*, pp. 120–124, 2008.
- [35] A. Dal Pozzolo et al., "Credit card fraud detection: a realistic modeling and a novel learning strategy," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 29, no. 8, pp. 3784–3797, 2017.
- [36] J. O. Awoyemi, A. O. Adetunmbi, and S. A. Oluwadare, "Credit card fraud detection using machine learning techniques: A comparative analysis," in *Proc. Int. Conf. Computing Networking and Informatics (ICCNI)*, IEEE, 2017, pp. 1–9.
- [37] K. Muameleci, "Anomaly Detection in Credit Card Transactions using Multivariate Generalized Pareto Distribution," 2022.
- [38] F. K. Alarfaj et al., "Credit card fraud detection using state-of-the-art machine learning and deep learning algorithms," *IEEE Access*, vol. 10, pp. 39700–39715, 2022.