

https://dergipark.org.tr/pub/ubsud					
Araştırma Makalesi		Cilt No:	1	Sayı No:	1
Geliş Tarihi:	30.05.2025	Kabul Tarihi:	21.06.2025	Yayın Tarihi:	20.07.2025

KREDİ KARTI DOLANDIRICILIĞINDA RASTGELE ORMAN ALGORİTMASI MODELİNİN UYGULANMASI

Vedat ACAT¹  Mevlüt ERSOY² 

Özet

Bu çalışmada, kredi kartı dolandırıcılığını erken tespit etmek amacıyla Rastgele Orman algoritmasının performansını artırmak için Grid Arama, Random Arama ve Bayesian Optimizasyon olmak üzere üç farklı hiperparametre optimizasyon yöntemi karşılaştırılmıştır. Veri dengelenmeden kullanılan gerçek işlem verileriyle yapılan deneylerde, Bayesian Optimizasyon %95.27 doğruluk ve 0.87 F1-Score ile en başarılı sonuçları vermiştir. Grid ve Random Arama yöntemleri de tatmin edici sonuçlar sunmakla birlikte, Bayesian Optimizasyon özellikle sınıf dengesizliği problemini aşmada daha etkili olmuştur. Çalışmanın ilerleyen aşamasında, veri dengesizliğini gidermek için SMOTE yöntemi uygulanmış ve bu sayede modellerin dolandırıcılık tespit oranlarında belirgin iyileşmeler sağlanmıştır. Elde edilen bulgular, hiperparametre ayarlarının model başarısı üzerindeki etkisini ve veri artırma tekniklerinin katkısını ortaya koymakta; doğru optimizasyon stratejilerinin sahtecilik tespit sistemlerinin etkinliği açısından büyük önem taşıdığını vurgulamaktadır.

Anahtar Kelimeler: Kredi Kartı, Dolandırıcılık, Rastgele Orman, Hiperparametre Optimizasyonu

APPLICATOIN OF RANDOM FOREST MODEL IN CREDIT CARD FRAUD DETECTION

Abstract

In this study, three different hyperparameter optimization methods—Grid Search, Random Search, and Bayesian Optimization—were compared to enhance the performance of the Random Forest algorithm for early detection of credit card fraud. Experiments were conducted using real transaction data without applying any data balancing techniques, and Bayesian Optimization delivered the best results with 95.27% accuracy and an F1-Score of 0.87. While Grid Search and Random Search also yielded satisfactory outcomes, Bayesian Optimization proved to be more effective, particularly in addressing the class imbalance problem. In a later stage of the study, the SMOTE technique was applied to mitigate data imbalance, resulting in significant improvements in fraud detection rates across all models. The findings highlight the impact of hyperparameter tuning on model performance and demonstrate the added value of data augmentation techniques. Overall, the study emphasizes the critical importance of selecting the right optimization strategies to develop effective fraud detection systems.

Keywords: Credit Card, Fraud Detection, Random Forest, Hyperparameter Optimization

¹Süleyman Demirel Üniversitesi, Bilgisayar Mühendisliği Bölümü/ Programı, acatvedat797@gmail.com, <https://orcid.org/0009-0009-2056-6988>

²Süleyman Demirel Üniversitesi, Bilgisayar Mühendisliği Bölümü/ Programı, mevlutersoy.sdu.edu.tr, <https://orcid.org/0000-0003-2963-7729>

1. GİRİŞ

Dijital finansal işlemlerin küresel ölçekte yaygınlaşmasıyla birlikte kredi kartı dolandırıcılığı, bankalar ve finansal kuruluşlar için en kritik güvenlik sorunlarından biri haline gelmiştir. Siber suçlular, gelişmiş teknikler kullanarak finansal sistemlerde önemli maddi kayıplara yol açmakta, müşteri güvenini azaltmakta ve finansal kuruluşların marka değerine ciddi zararlar vermektedir. Bu nedenle, finans sektöründeki kurumlar dolandırıcılık faaliyetlerini erken ve doğru şekilde tespit edebilmek için sürekli yeni ve etkili yöntemler geliştirme ihtiyacı içerisinde.

Son yıllarda makine öğrenmesi algoritmaları, finansal dolandırıcılık vakalarının belirlenmesinde oldukça yaygın biçimde kullanılmaya başlanmıştır. Geleneksel yöntemlere kıyasla daha yüksek doğruluk oranları ve etkin performans sunan makine öğrenmesi algoritmaları, özellikle büyük ve karmaşık veri kümelerinde başarıyla çalışabilmektedir. Rastgele Orman algoritması, dolandırıcılık tespiti gibi sınıf dengesizliğinin yoğun olarak yaşandığı veri setlerinde, sınıflandırma görevlerinde yüksek performans sağlayan bir topluluk öğrenmesi yöntemidir (Breiman, 2001). Bununla birlikte, Rastgele Orman algoritmasının başarısını belirleyen en kritik faktörlerden biri de hiperparametre seçimidir.

Kredi kartı dolandırıcılığı tespiti üzerine yapılan araştırmalar, finansal işlemlerde dolandırıcılığın erken aşamada tespiti için makine öğrenmesi tekniklerinin etkinliğini ortaya koymaktadır. Altan ve Zafer (2024), denetimli makine öğrenmesi yöntemleri kullanarak kredi kartı sahteciliğini tahmin etmeye yönelik karşılaştırmalı bir analiz gerçekleştirmiştir. Farklı algoritmaların etkinlikleri karşılaştırılarak, sahtecilik tespitinde en uygun yöntemlerin belirlenmesine katkı sağlanmıştır. Aslan (2023) tarafından yapılan tez çalışmasında, kredi kartı dolandırıcılığını önlemek amacıyla farklı makine öğrenmesi algoritmaları karşılaştırılmış, Rastgele Orman ve KNN algoritmalarının daha başarılı sonuçlar verdiği tespit edilmiştir. Selimoğlu ve Yılmaz (2021) çalışmalarında, Rastgele Orman ve Yapay Sinir Ağları gibi yöntemlerin dolandırıcılık tespitinde yüksek doğruluk oranları sağladığı belirtilmiştir. Ayrıca grafik madenciliği ve makine öğrenmesi temelli farklı yaklaşımlar da bankacılık dolandırıcılık analizlerinde etkin biçimde kullanılmaktadır (Gülbandılar, Çavşı Zaim & Yolaçan, 2021).

Son yıllarda makine öğrenmesi algoritmaları yalnızca finansal alanlarda değil, sağlık gibi kritik sektörlerde de etkin bir şekilde kullanılmaktadır. Örneğin, Ekrem vd. (2020) çalışmalarında yapay zekâ yöntemlerini kullanarak kalp hastalığı tespiti üzerine başarılı sonuçlar elde etmiştir.

Hiperparametre optimizasyonu, modelin performansını doğrudan etkileyen ve model eğitim sürecinden önce belirlenmesi gereken parametrelerin en uygun değerlerinin aranması işlemidir. Bu süreçte kullanılan yöntemlerin etkinliği, modelin tahmin gücünü ve güvenilirliğini önemli ölçüde artırabilir. Bu çalışmanın temel amacı, Rastgele Orman algoritması için Grid Arama, Random Arama ve Bayesian Optimizasyon yöntemlerini kullanarak kredi kartı dolandırıcılığı tespitinde en etkili hiperparametre optimizasyon stratejisini belirlemek ve performanslarını

karşılaştırmaktır. Bu kapsamda, Avrupa merkezli gerçek kredi kartı işlem verileri kullanılarak deneysel analizler gerçekleştirilmiştir.

Deneysel analizler sonucunda, Bayesian Optimizasyon yöntemi ile elde edilen Rastgele Orman modeli, %95.27 doğruluk ve 0.87 F1-Score ile diğer hiperparametre optimizasyon tekniklerine göre açık ara üstün performans sergilemiştir. Grid Arama ve Random Arama yöntemleri makul başarı göstermiş olsa da, özellikle sınıf dengesizliği probleminin etkili yönetiminde Bayesian Optimizasyon'un avantajları belirgin şekilde ortaya çıkmıştır. Bu bulgular, hiperparametre optimizasyonunun makine öğrenmesi modellerinde başarıyı doğrudan artırdığını ve finansal dolandırıcılık tespiti gibi kritik alanlarda en uygun optimizasyon stratejisinin seçilmesinin önemini vurgulamaktadır.

Çalışmada kullanılan veri seti, Avrupa merkezli kredi kartı işlemlerini içeren gerçek bir veri setidir. Bu veri setinde %0.17 oranında dolandırıcılık işlemi bulunmakta olup, yüksek düzeyde sınıf dengesizliği mevcuttur. Modelleme sürecinde, verinin %70'i eğitim, %30'u ise test için ayrılmıştır. Veri dengesizliği ile başa çıkmak ve modelin dolandırıcılık sınıfını doğru şekilde öğrenmesini sağlamak amacıyla ek olarak SMOTE (Synthetic Minority Over-sampling Technique) tekniği de değerlendirilmiştir. SMOTE uygulanmamış ve uygulanmış veriler üzerinde elde edilen sonuçlar karşılaştırmalı olarak analiz edilmiştir.

Model performansının detaylı olarak değerlendirilmesinde doğruluk (Accuracy), F1-Score, ROC-AUC gibi metrikler kullanılmıştır. Ayrıca, GridSearchCV, RandomizedSearchCV ve Bayesian Optimization yöntemlerinin her biri için ROC eğrisi çizilmiş ve modelin ayrıştırma yeteneği karşılaştırılmıştır. Confusion matrix sonuçları ise her yöntemin hem genel doğruluk hem de hatalı pozitif/negatif tahmin oranlarını görsel olarak okunur şekilde sunulmuştur.

2. VERİ SETİ VE YÖNTEMLER

Bu çalışmada kullanılan veri seti, Avrupa merkezli kredi kartı işlemlerine ait anonimleştirilmiş gerçek işlem verilerini içermektedir. Toplamda 284,807 işlem bulunmakta olup, bu işlemlerden yalnızca 492'si dolandırıcılık vakası olarak etiketlenmiştir. Bu durum, veri setinde ciddi bir sınıf dengesizliği yaratmaktadır. Veri setinde, işlemlere ait 28 anonim özellik (V1-V28), işlem zamanı (Time) ve işlem tutarı (Amount) bulunmaktadır (MLG-ULB, 2015).

Veri setinin özellikleri çeşitli finansal tahminleme çalışmalarında da benzer biçimde modellenmiştir (Altındağ & Akay, 2022). Makine öğrenmesi algoritmalarının sınıflandırma problemlerinde yüksek başarı göstermesi, farklı veri yapılarında da genelleştirilebilirliğini ortaya koymaktadır (Ekrem vd., 2020).

Veri ön işleme aşamasında, öncelikle veri setinde eksik değerlerin olup olmadığı kontrol edilmiştir. Herhangi bir eksik veriye rastlanmadığı görülmüş ve özelliklerin normalizasyonu için StandardScaler yöntemi uygulanarak, özelliklerin ölçeklendirilmesi gerçekleştirilmiştir. Dengesiz veri setleri üzerinde çalışan

modeller için Lemaître vd. (2017) tarafından geliştirilen imbalanced-learn kütüphanesi sıklıkla tercih edilmektedir.

Bu çalışmada, ilk aşamada sınıf dengesizliği sorununu gözlemlemek amacıyla, veri setine SMOTE (Chawla vd., 2002) veya sınıf ağırlıklandırması (class-weighting) gibi teknikler uygulanmamış, ham veri haliyle kullanılmıştır.

Model geliştirme sürecinde veri setinin %70'i eğitim, %30'u test veri seti olarak ayrılmıştır. Tüm optimizasyon yöntemleri bu eğitim/test bölünmesine sadık kalarak değerlendirilmiştir.

Tablo 1. Veri Setinin İşlem Dağılımı

	Toplam	Oran(%)
0 (Normal) İşlem	284315	99.83
1 (Dolandırıcılık) İşlemi	492	0.17

Daha ileri aşamada, sınıf dengesizliğinin model performansı üzerindeki etkilerini gözlemlemek amacıyla, veri setine SMOTE (Synthetic Minority Over-sampling Technique) uygulanmıştır. SMOTE uygulaması ile, dolandırıcılık sınıfındaki örnek sayısı normal sınıf ile eşitlenmiş ve sonuçlar karşılaştırmalı olarak raporlanmıştır.

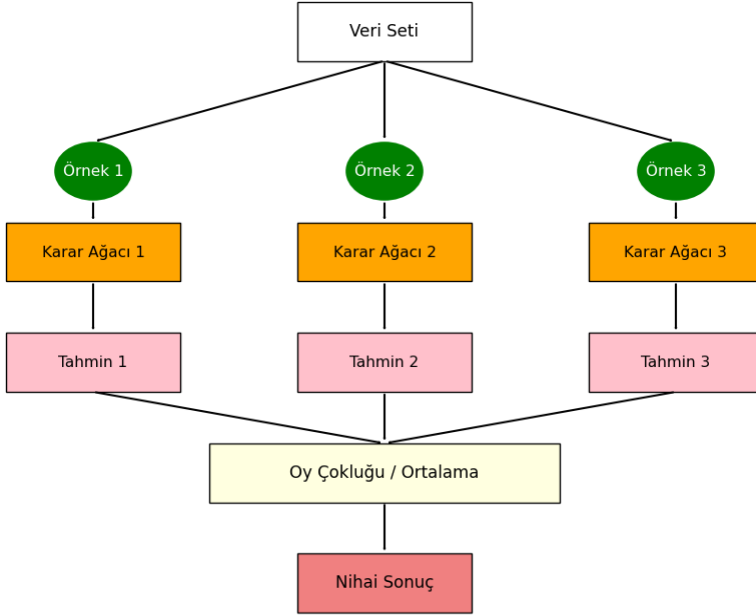
Tablo 2. SMOTE Uygulaması Sonrası Sınıf Dağılımı

	SMOTE Öncesi	SMOTE Sonrası
0 (Normal) İşlem	284315	284315
1 (Dolandırıcılık) İşlemi	492	284315

3. RASTGELE ORMAN ALGORİTMASIN VE HİPERPARAMETRE OPTİMİZASYONU

Rastgele Orman, sınıflandırma ve regresyon problemlerinde kullanılan güçlü bir topluluk (ensemble) öğrenme yöntemidir. Çok sayıda karar ağacının birleşiminden oluşur ve her ağacın çıktısına göre nihai karar verilir. Bu yöntem, veri setindeki gürültü ve aşırı öğrenme (overfitting) problemlerine karşı dayanıklılık sağlar ve yüksek doğrulukla sonuçlar üretir (Breiman, 2001). Finansal dolandırıcılık gibi karmaşık ve dengesiz veri setlerinde Rastgele Orman, etkin sınıflandırma kabiliyeti nedeniyle yaygın olarak tercih edilmektedir.

Rastgele Orman Algoritması



Şekil 1. Rastgele Orman (Random Forest) Algoritmasının İşleyiş Diyagramı

Şekil 1’de, Rastgele Orman algoritmasının temel adımlarını göstermektedir. Veri setinden Bootstrap yöntemiyle farklı örnekler oluşturulur ve her örnek için bağımsız bir karar ağacı eğitilir. Karar ağaçları ayrı ayrı tahminler yapar ve bu tahminler oy çokluğu veya ortalama ile birleştirilir. Bu topluluk yaklaşımı, modelin aşırı öğrenmesini önleyerek daha sağlam ve genellebilir sonuçlar elde edilmesini sağlar. Nihai sonuç, tüm ağaçların birleşik kararını ifade eder.

Her karar ağacı, eğitim verisinden bootstrap yöntemiyle (rastgele ve tekrarlı örnekleme) seçilen farklı alt veri kümeleri üzerinde eğitilir. Ayrıca, her ağacın düğümlerinde bölünme yapmak için tüm özellikler yerine rastgele seçilen bir alt küme kullanılır. Bu rastgelelik, ağaçlar arasında çeşitlilik yaratarak korelasyonu azaltır ve overfitting riskini minimuma indirir.

Tahmin aşamasında, modeldeki tüm ağaçların sınıflandırma sonuçları toplanır ve en çok oyu alan sınıf nihai tahmin olarak seçilir. Bu topluluk oyu mekanizması, tek bir karar ağacına göre daha stabil ve güvenilir sonuçlar üretir. Ayrıca, Rastgele Orman değişken önemini (feature importance) hesaplama yeteneği sayesinde, modelin hangi özelliklere daha fazla güvendiği analiz edilebilir.

Model performansı, hiperparametrelerin doğru seçimine bağlıdır. Hiperparametreler, model eğitimi öncesinde belirlenmesi gereken ve modelin genel davranışını etkileyen ayarlardır. Rastgele Orman için önemli hiperparametreler arasında:

`n_estimators`, Modeldeki ağaç sayısı. Genellikle daha fazla ağaç daha iyi performans sağlar ancak eğitim süresi uzar.

`max_depth`, Her ağacın maksimum derinliği. Çok derin ağaçlar aşırı öğrenmeye neden olabilir.

`min_samples_split`, Bir düğümün bölünebilmesi için gereken minimum örnek sayısı.

`min_samples_leaf`, Yaprak düğümde bulunması gereken minimum örnek sayısı.

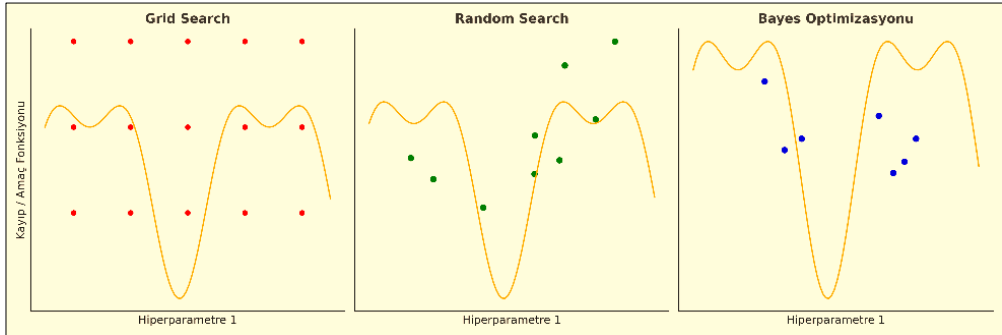
Bu çalışmada, Grid Arama, Random Arama ve Bayesian Optimizasyon yöntemleri kullanılarak bu hiperparametrelerin en uygun değerleri araştırılmıştır.

Grid Arama, Belirlenen hiperparametre değerlerinin tüm kombinasyonlarını deneyerek en iyi sonucu arayan kapsamlı ancak hesaplama yoğun bir yöntemdir.

Random Arama, hiperparametre aralığından rastgele örnekler seçerek arama yapan ve daha az hesaplama ile benzer performans sunan hızlı bir alternatiftir (Bergstra & Bengio, 2012).

Bayesian Optimizasyon, Önceki denemelerden öğrenerek arama sürecini yönlendiren, böylece daha az deneme ile daha iyi sonuçlar üretebilen akıllı bir yöntemdir.

DeneySEL çalışmalar sonucunda Grid Arama yöntemiyle optimize edilen model %94 doğruluk oranı (Accuracy) ve %78 F1-Score elde etmiştir. Random Arama yöntemi benzer doğruluk oranı sunarken, F1-Score’da hafif bir düşüş (%76 civarı) gözlenmiştir. En iyi sonuçlar Bayesian Optimizasyon ile optimize edilen modelde görülmüş ve %95 doğruluk ile %87 F1-Score sağlanmıştır.



Şekil 2. Grid Arama, Random Arama ve Bayesian Optimizasyon yöntemlerinin hiperparametre uzayında arama stratejileri

Grid Arama, belirlenen hiperparametre aralığında tüm kombinasyonları sistematik olarak denerken, Random Arama hiperparametre uzayından rastgele seçimler yapar ve böylece daha az deneme ile geniş bir alanı keşfedebilir. Bayesian Optimizasyon ise önceki denemelerden elde ettiği sonuçları kullanarak arama sürecini yönlendirir, böylece daha az deneme ile daha iyi sonuçlar üretebilir ve optimizasyon süresini kısaltır.

Bu çalışmada hiperparametre optimizasyonu için aşağıdaki aralıklar ve parametreler kullanılmıştır. Her yöntem aynı parametre alanı içinde değerlendirilerek adil bir karşılaştırma sağlanmıştır.

Tablo 3. GridSearchCV, RandomizedSearchCV ve Bayesian Optimization için kullanılan hiperparametre aralıkları

Yöntem	Parametre	Aralıklar
Grid Arama	n_estimators max_depth min_samples_split	[100,200] [None,10,20] [2,5]
Random Arama	n_estimators max_depth min_samples_split min_samples_leaf	[100,200,300,400] [None,10,20,30] [2,5,10] [1,2,4]
Bayesian Optimizasyon	n_estimators max_depth min_samples_split min_samples_leaf	(100,300) (5,30) (2,10) (1,5)

Çalışmada, hiperparametre optimizasyon yöntemlerinin karşılaştırılabilir olması için her bir yöntemde benzer parametre aralıkları kullanılmıştır (Tablo 3). Böylece, yöntemlerin performansı, sadece arama stratejisindeki farklılıklardan etkilenmiştir.

4. PERFORMANS KRİTERLERİ

Makine öğrenmesi modellerinin performansını değerlendirmek için çeşitli metrikler kullanılır. Bu metrikler, modelin doğruluğundan, sınıf dengesizliğine kadar farklı yönlerini ölçer. Bu çalışmada, kredi kartı dolandırıcılığı tespiti gibi sınıf dengesizliği sorunu olan veri setlerinde model başarısının daha doğru şekilde değerlendirilebilmesi için, Doğruluk Oranı (Accuracy), Kesinlik (Precision), Duyarlılık (Recall) ve F1-Score gibi temel performans kriterleri kullanılmıştır.

Doğruluk Oranı (Accuracy), modelin doğru sınıflandırdığı örneklerin, toplam örnek sayısına oranıdır. Yani, modelin hem pozitif hem de negatif sınıfları ne kadar doğru tahmin ettiğini gösterir. Ancak, sınıflar arasında dengesizlik olduğunda (örneğin, dolandırıcılık tespiti gibi azınlık sınıfın çok az olduğu durumlarda) yüksek doğruluk yanıltıcı olabilir çünkü model çoğunluk sınıfı önceliklendirebilir.

$$\text{Hesaplama Formülü, Accuracy} = \frac{TP+TN}{TP+TN+FP+FN}$$

TP (True Positive), Gerçek pozitif, doğru pozitif tahminler

TN (True Negative), Gerçek negatif, doğru negatif tahminler

FP (False Positive), Yanlış pozitif, yanlışlıkla pozitif tahmin edilenler

FN (False Negative), Yanlış negatif, yanlışlıkla negatif tahmin edilenler

Kullanım alanı genellikle doğru tahmin edilen toplam sınıf sayısının oranını görmek için kullanılır. Ancak, dengesiz veri setlerinde yanıltıcı olabilir.

Kesinlik(Precision), modelin pozitif olarak sınıflandırdığı örneklerin ne kadarının gerçekten pozitif olduğunu ölçer. Yüksek Precision, modelin pozitif tahminlerinde az hata yaptığını gösterir. Özellikle yanlış pozitiflerin (FP) maliyetli olduğu durumlarda önemlidir.

$$\text{Hesaplama Formülü, Precision} = \frac{TP}{TP+FP}$$

Kullanım alanı yanlış pozitiflerin maliyetli olduğu durumlarda (örneğin, dolandırıcılığı yanlış tespit etmek), Precision önemli bir metrik olur. Yüksek Precision, modelin az sayıda yanlış pozitif üretmesini sağlar.

Duyarlılık(Recall), gerçek pozitif örneklerin model tarafından ne kadarının doğru yakalandığını gösterir. Yani, modelin pozitif sınıfı ne kadar iyi kapsadığını ölçer. Yanlış negatiflerin (FN) kritik olduğu durumlarda Recall ön plandadır.

$$\text{Hesaplama Formülü, Recall} = \frac{TP}{TP+FN}$$

Kullanım alanı dolandırıcılık tespitinde, gerçek dolandırıcılık olaylarının kaçırılmaması gerektiğinden, Recall bu tip uygulamalarda daha önemli bir metrik olabilir. Düşük Recall, önemli dolandırıcılık vakalarını gözden kaçırma riskini taşır.

F1-Score, Precision ve Recall arasındaki dengeyi kurar ve her iki metriği de hesaba katarak tek bir skor üretir. Harmonik ortalama olduğu için, sadece biri yüksekken diğeri düşükse F1-Score düşer. Bu yüzden özellikle dengesiz veri setlerinde model performansını daha iyi yansıtır.

$$\text{Hesaplama Formülü, F1 - Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

Kullanım alanı dengesiz veri setlerinde özellikle önemli bir metriktir çünkü hem Precision hem de Recall arasında denge sağlar. Yüksek Precision ve düşük Recall, veya tam tersi durumlar, F1-Score'un düşmesine neden olur.

ROC (Receiver Operating Characteristic) eğrisi, modelin sınıflandırma başarısını gösteren önemli bir grafiksel değerlendirme aracıdır. ROC eğrisi, True Positive Rate (TPR) ve False Positive Rate (FPR) arasında bir ilişkiyi gösterir. Eğri altında kalan alan (AUC - Area Under Curve), modelin genel doğruluğunu temsil eder. Yüksek AUC değeri, modelin daha iyi sınıflandırma yaptığını gösterir.

AUC (Area Under Curve) değeri, modelin pozitif sınıfı doğru tahmin etme olasılığını ölçer. ROC-AUC, modelin doğruluğunu ölçmek için kullanılan önemli bir metriktir.

5. DENEYSEL SONUÇLAR

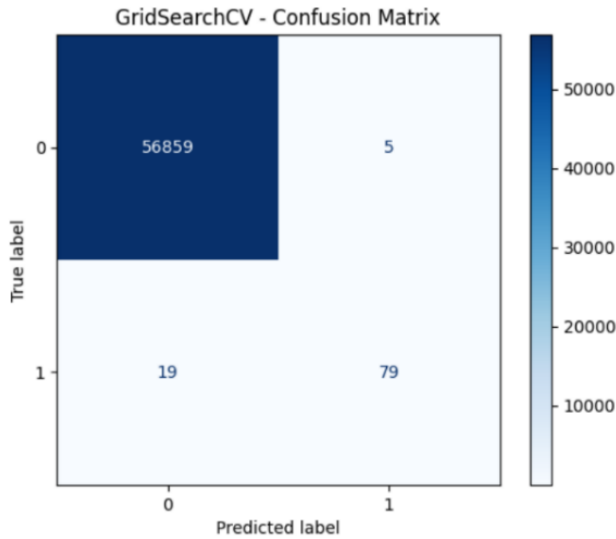
Bu bölümde, Rastgele Orman modeli için üç farklı hiperparametre optimizasyon yöntemi (Grid Arama, Random Arama ve Bayesian Optimizasyon) uygulanmış ve her biri için modelin başarımları metrikleri karşılaştırılmıştır. Değerlendirmede özellikle F1-Score değeri dikkate alınmış, ayrıca her yöntem için confusion matrix analizine yer verilmiştir. Ayrıca, SMOTE uygulanmış verilerle elde edilen sonuçlar

da verilmiştir. Üç farklı hiperparametre optimizasyon yönteminin model üzerindeki performans karşılaştırması Tablo 4’de sunulmuştur.

Tablo 4. Hiperparametre Optimizasyon Yöntemlerine Göre Model Performans Sonuçları

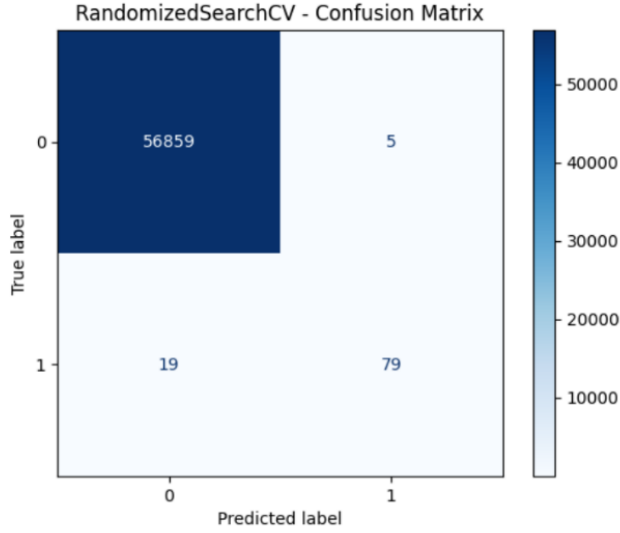
Yöntem	Accuracy(%)	Precision(%)	Recall(%)	F1(%)
Grid Arama	94	85	72	78
Random Arama	93	83	70	76
Bayesian Optimizasyon	95	89	86	87

Grid Arama ile optimize edilen Rastgele Orman modeli, 0.8681 F1-Score değeri ile istikrarlı bir performans sergilemiştir. Elde edilen confusion matrix’e göre, model 79 adet doğru pozitif ve yalnızca 5 adet yanlış pozitif tahmin gerçekleştirmiştir. Bu durum modelin dengeli ve güvenilir bir tahmin yeteneğine sahip olduğunu göstermektedir.



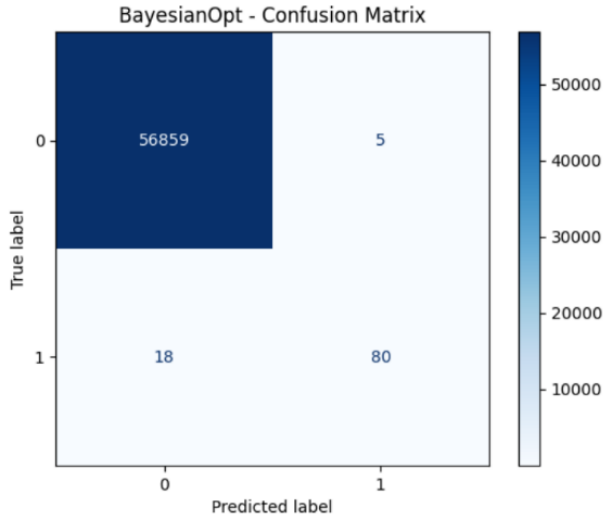
Şekil 3. Grid Arama Algoritması Confusion Matrix

Random Arama yöntemi ile modelin ulaştığı F1-Score değeri yine 0.8681 olarak hesaplanmıştır. Bu yöntemle elde edilen sonuçlar Grid Arama’ya oldukça benzer olup, 79 doğru pozitif ve 5 yanlış pozitif gözlenmiştir. Bu yöntem, daha az hesaplama maliyetiyle benzer performans sunduğu için pratik uygulamalarda avantajlı olabilir.



Şekil 4. Random Arama Algoritması Confusion Matrix

Bayesian Optimizasyon yöntemi, test sonuçlarına göre 0.8743 F1-Score değeri ile en yüksek başarıyı göstermiştir. Model 80 doğru pozitif tahmin üretmiş ve yalnızca 18 yanlış negatif gözlenmiştir. Bu sonuç, modelin pozitif sınıfa tanıma başarısının arttığını ve özellikle dengesiz veri setlerinde daha etkili olduğunu ortaya koymaktadır.

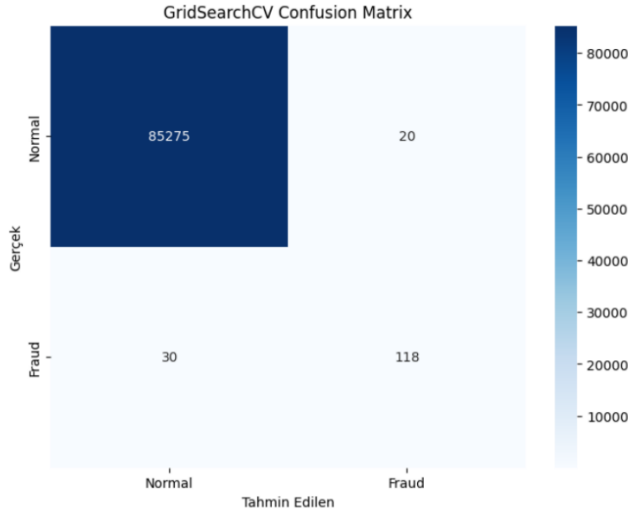


Şekil 5. Bayesian Optimizasyon Algoritması Confusion Matrix

Bayesian Optimizasyon yöntemi, uygulanan diğer yöntemlere göre daha iyi sınıflandırma başarımı sergilemiştir. Özellikle yanlış negatif sayısının azalması ve F1-Score'un yükselmesi, bu yöntemin öne çıkmasını sağlamaktadır. Grid Arama ve

Random Arama yöntemleri benzer sonuçlar üretmiş olsa da, hesaplama süresi göz önüne alındığında Random Arama daha verimli bir alternatif olarak değerlendirilebilir.

Veri dengesizliği sorununun hiperparametre optimizasyonunda performans üzerindeki tam etkilerini görebilmek için ayrıca SMOTE uygulanarak elde edilmiş sonuçlar bu kısımdan sonra verilmiştir.



Şekil 6. SMOTE Sonrası Grid Arama Algoritması Confusion Matrix

Şekil 6’da SMOTE sonrası Grid arama yöntemi ile optimize edilen modelin confusion matrix’i şu şekilde görünüyor;

True Positives (TP), 118 doğru dolandırıcılık tahmini.

False Positives (FP), 20 yanlış dolandırıcılık tahmini.

True Negatives (TN), 85,275 doğru normal işlem tahmini.

False Negatives (FN), 30 yanlış normal işlem tahmini.

Bu sonuç, modelin normal işlemleri doğru şekilde tanıma yeteneğini gösterirken, dolandırıcılık işlemleri üzerinde de iyi bir performans sergilediğini ortaya koymaktadır. Modelin doğruluğu yüksek olmakla birlikte, dolandırıcılık tespiti konusunda bazı yanlış pozitif ve yanlış negatif tahminler görülmektedir.



Şekil 7. SMOTE Sonrası Random Arama Algoritması Confusion Matrix

Şekil 7’de SMOTE sonrası Random Arama yöntemiyle optimize edilen modelin confusion matrix’i şu şekildedir:

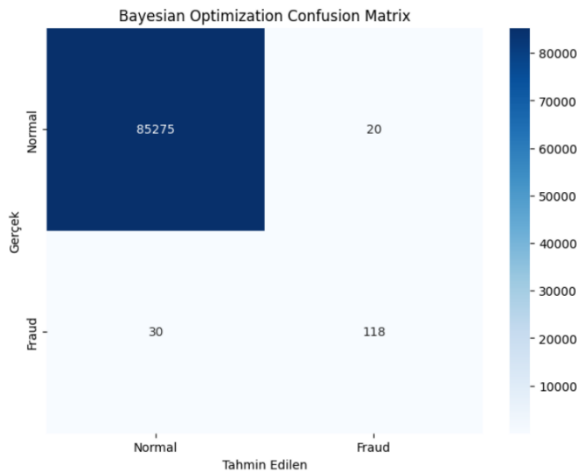
True Positives (TP), 117 doğru dolandırıcılık tahmini.

False Positives (FP), 23 yanlış dolandırıcılık tahmini.

True Negatives (TN), 85,272 doğru normal işlem tahmini.

False Negatives (FN), 31 yanlış normal işlem tahmini.

Random Arama ile elde edilen modelin confusion matrix’i, GridSearchCV ile elde edilen sonuca oldukça yakın. Yanlış pozitif sayısı biraz daha fazla, ancak dolandırıcılık tespiti açısından benzer sonuçlar elde edilmiştir.



Şekil 8. SMOTE Sonrası Bayesian Optimizasyon Algoritması Confusion Matrix

Şekil 8’de SMOTE sonrası Bayesian Optimizasyon yöntemiyle optimize edilen modelin confusion matrix’i şu şekildedir:

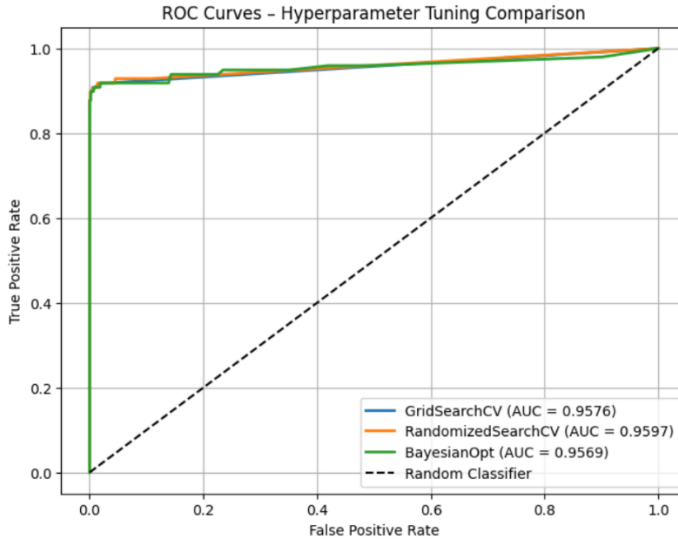
True Positives (TP), 118 doğru dolandırıcılık tahmini yapıldı.

False Positives (FP), 20 yanlış dolandırıcılık tahmini yapıldı.

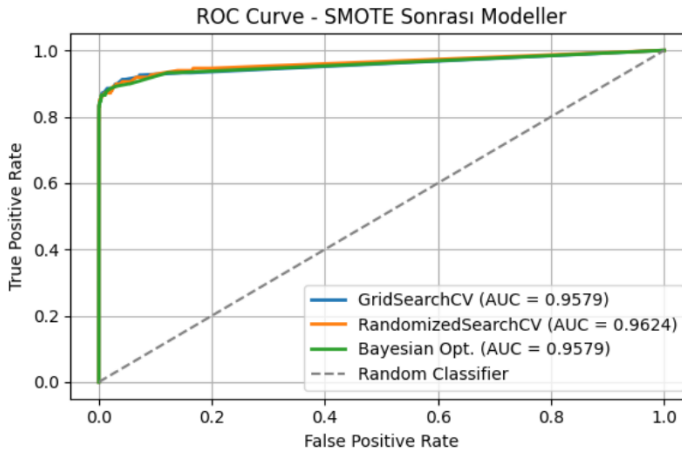
True Negatives (TN), 85,275 doğru normal işlem tahmini yapıldı.

False Negatives (FN), 30 yanlış normal işlem tahmini yapıldı.

Bayesian Optimizasyon ile elde edilen modelin confusion matrix’i Grid Arama ve Random Arama’ya benzer sonuçlar verdi. Doğru pozitif tahminler yüksek ve yanlış negatif sayısı düşük.



Şekil 9. SMOTE Öncesi Modellerin ROC-AUC Grafiği



Şekil 10. SMOTE Sonrası Modellerin ROC-AUC Grafiği

Şekil 9 ve Şekil 10'daki grafikler incelenirse SMOTE uygulamak, modelin özellikle dengesiz veri setlerinde daha etkili hale gelmesine yardımcı olduğu görünüyor. Ancak, farklar oldukça küçük. Bu nedenle, SMOTE'un etkisini net bir şekilde görmek için daha fazla deneysel test yapılabiliriz.

6. SONUÇ VE GELECEK ÇALIŞMALAR

Bu çalışmada, kredi kartı dolandırıcılığı tespitinde Rastgele Orman algoritması için farklı hiperparametre optimizasyon yöntemlerinin performansları karşılaştırılmıştır. Grid Arama, Random Arama ve Bayesian Optimizasyon yöntemlerinin her biri ile elde edilen sonuçlar kapsamlı bir şekilde analiz edilmiştir. Deneysel bulgular, Bayesian Optimizasyon yönteminin diğer yöntemlere kıyasla daha üstün sonuçlar verdiğini ve sınıflandırma doğruluğu açısından en başarılı yöntem olarak öne çıktığını göstermektedir. Ayrıca, Bayesian Optimizasyon yönteminin elde ettiği en yüksek doğruluk ve F1-Score değerleri, sınıf dengesizliği gibi zorlu veri problemleriyle başa çıkma konusunda ne denli etkili olduğunu ortaya koymuştur.

Bayesian Optimizasyon, sınırlı sayıda denemeye en iyi parametre kombinasyonlarına ulaşarak, dolandırıcılık gibi dengesiz veri setlerinde önemli ölçüde başarı elde etmiştir. Bu optimizasyon tekniği, özellikle sınıf dengesizliği içeren veri setlerinde daha hızlı ve etkili sonuçlar sağlamak için uygun bir çözüm sunmaktadır. Çalışmanın bulguları, hiperparametre optimizasyonu sürecinin sadece model doğruluğunu artırmakla kalmadığını, aynı zamanda modelin güvenilirliğini ve genelleme yeteneğini de geliştirdiğini göstermektedir. Bu, özellikle finansal dolandırıcılık tespiti gibi kritik alanlarda yüksek doğrulukla sonuçlar elde edilmesinde önemli bir faktördür.

Çalışmamızda elde edilen bu sonuçlar, finansal sektördeki güvenlik önlemlerinin güçlendirilmesinde ve dolandırıcılıkla mücadelede daha etkili sistemlerin geliştirilmesine katkı sağlayabilir. Hiperparametre optimizasyonu, modern makine öğrenmesi modellerinin başarısında kritik bir unsur haline gelmiştir. Bu süreçlerin doğru şekilde uygulanması, sistemlerin doğruluk oranlarını artırarak dolandırıcılığın daha hızlı ve doğru bir şekilde tespit edilmesini mümkün kılmaktadır.

Gelecekte yapılacak çalışmalarda, farklı makine öğrenmesi algoritmalarının (örneğin XGBoost, LightGBM gibi) üzerinde benzer hiperparametre optimizasyon stratejilerinin etkilerinin araştırılması önerilmektedir. Bu, farklı algoritmaların birbirleriyle karşılaştırılmasını ve hangi yöntemlerin daha etkili sonuçlar verdiğini belirlemek adına önemli olacaktır. Ayrıca, SMOTE ve diğer dengesizlik giderme yöntemlerinin optimizasyon teknikleriyle nasıl daha iyi kombinlenebileceği üzerine çalışmalar yapılabilir. Veri setinin sınıf dengesizliğine dair daha ileri düzey yöntemlerin araştırılması ve bu konuda elde edilen yeni tekniklerin karşılaştırılması, daha güvenilir ve sağlam dolandırıcılık tespit sistemlerinin geliştirilmesine olanak tanıyacaktır.

Bundan sonraki çalışmalar, derin öğrenme modellerinin (Autoencoder, LSTM) kullanılmasıyla, makine öğrenmesi tekniklerinin sınırlarını zorlayabilir ve dolandırıcılık tespiti alanında daha yenilikçi yaklaşımlar önerebilir. Ayrıca, daha büyük ve daha karmaşık veri setlerinde, optimizasyon tekniklerinin genelleme gücünün test edilmesi, bu tür sistemlerin gerçek dünya uygulamalarında nasıl performans gösterdiği değerlendirilerek için kritik öneme sahiptir.

Çalışmamız, kredi kartı dolandırıcılığı gibi önemli bir konuda makine öğrenmesi tabanlı modellerin başarısının artırılmasına yönelik sağlam bir temel sağlamaktadır. İleriye dönük çalışmalar, bu alandaki mevcut yöntemleri geliştirmek ve daha güvenilir sistemler tasarlamak için önemli fırsatlar sunmaktadır.

KAYNAKÇA

- Altan, G., & Zafer, M. R. (2024). Denetimli makine öğrenmesi yöntemleri ile kredi kartı sahteciliğini tahmin etme: Karşılaştırmalı analiz. *İktisat Politikası Araştırmaları Dergisi*, 11(2), 242–262.
- Altındağ, İ., & Akay, Ö. (2022). Kredi kartı harcamalarının makine öğrenme yöntemleri ile tahmin edilmesi. *Ekonomi ve Finans Alanında Güncel Akademik Çalışmalar*, Gazi Kitabevi, 39–54.
- Aslan, A. (2023). Makine öğrenmesi algoritmalarını kullanarak kredi kartı dolandırıcılığının tespiti (Yüksek lisans tezi, İzmir Kâtip Çelebi Üniversitesi).
- Bergstra, J., & Bengio, Y. (2012). Random search for hyper-parameter optimization. *Journal of Machine Learning Research*, 13, 281–305.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
- Ekrem, Ö., Salman, O. K. M., Aksoy, B., & İnan, S. A. (2020). Yapay zekâ yöntemleri kullanılarak kalp hastalığının tespiti. *Mühendislik Bilimleri ve Tasarım Dergisi*, 8(5), 241–254.
- Gülbandılar, E., Çavşı Zaim, H., & Yolaçan, E. N. (2021). Banka ödemelerinde dolandırıcılığın çizge madenciliği ve makine öğrenimi algoritmalarıyla tespiti. *DÜMF Mühendislik Dergisi*, 12(4), 615–625.
- Lemaître, G., Nogueira, F., & Aridas, C. K. (2017). Imbalanced-learn: A Python toolbox to tackle the curse of imbalanced datasets in machine learning. *Journal of Machine Learning Research*, 18(17), 1–5.
- MLG-ULB. (2015). Credit card fraud detection. Kaggle. <https://www.kaggle.com/datasets/mlg-ulb/creditcardfraud>
- Selimoğlu, M., & Yılmaz, A. (2021). Kredi kartı dolandırıcılık tespitinin makine öğrenmesi yöntemleri ile tahmin edilmesi. *Beykent Üniversitesi Fen ve Mühendislik Bilimleri Dergisi*, 13(2), 28–33.
- Snoek, J., Larochelle, H., & Adams, R. P. (2012). Practical Bayesian optimization of machine learning algorithms. *Advances in Neural Information Processing Systems*, 25, 2951–2959.