

Yüksek Mahkeme Kararlarının Sınıflandırılması için Veri Dengeleme ve Açıklanabilir Yapay Zekâ Tabanlı Karar Destek Sistemi

Mustafa Emirkan OKURSOY¹ , Tülin İNKAYA¹ 

¹Bursa Uludağ Üniversitesi, Mühendislik Fakültesi, Endüstri Mühendisliği Bölümü, Bursa, TÜRKİYE

Özet

Yapay zekâ (YZ) teknikleri hukuk alanında büyük bir potansiyele sahiptir. Hukukta yazılı metinler ağırlıklı olarak kullanıldığından, yapay zekâ tekniklerinin kullanımı hukuki metinlerin hızlı ve etkin bir şekilde işlenmesini sağlayarak karar verme süreçlerine yardımcı olabilir. Bunun yanında, yapay zekâ teknikleri hukuki süreçlerin iyileştirilmesinde ve davaların olası sonuçlarını tahmin etmede de kullanılabilir. Bu çalışmada, adli ve idari yargıdaki iki Yüksek Mahkeme olan Yargıtay ve Danıştay'ın kararları ele alınmıştır ve yapay zekâ tabanlı bir karar destek sistemi geliştirilmesi amaçlanmıştır. Yargıtay ve Danıştay kararları; hukuki birliğin sağlanması, bireylerin hak ve özgürlüklerinin korunması ve adaletin sağlanması açısından büyük önem taşımaktadır. Bu kapsamda, Ulusal Yargı Ağı Projesindeki Mevzuat ve İçtihat programından alınan her iki Yüksek Mahkemeye ait karar metinleri kullanılmış olup bu kararların %11'i onama, %89'u ise bozma kararıdır. Bu nedenle, ele alınan problem, dengesiz sınıflandırma problemi olarak tanımlanmıştır. Öncelikle veri ön işleme ve doğal dil işleme (NLP) teknikleri kullanılarak karar metinlerinden öznitelikler çıkarılmıştır. Sonrasında, verideki dengesizliği gidermek amacıyla üst örnekleme yöntemlerinden sentetik azınlık aşırı örnekleme tekniği (SMOTE) ve rastgele alt örnekleme uygulanmıştır. Son olarak, mahkeme kararlarının tahmin edilmesi amacıyla k-en yakın komşuluk, karar ağacı, destek vektör makinesi, rassal orman, LightGBM, XGBoost ve yapay sinir ağları ile sınıflandırma modelleri geliştirilerek modellerin performansları karşılaştırılmıştır. Elde edilen deneysel sonuçlar, önerilen karar destek sisteminin hukuk profesyonellerine fayda sağlama potansiyeli olduğunu göstermektedir.

Anahtar Kelimeler: Yapay Zekâ, Doğal Dil İşleme, Karar Destek Sistemi, Hukuk, Dengesiz Sınıflandırma

Data Balancing and Explainable Artificial Intelligence-Based Decision Support System for the Classification of Supreme Court Decisions

Abstract

Artificial intelligence (AI) techniques have great potential in the field of law. Since written texts are mainly used in law, the use of AI techniques can assist decision-making processes by ensuring fast and effective processing of legal texts. In addition, AI techniques can also be used to improve the legal processes and predict the possible outcomes of the cases. In this study, the decisions made by the Court of Cassation and the Council of State, two Supreme Courts in the judicial and administrative judiciary, are considered, and it is aimed to develop an AI-based decision support system. The decisions of the Court of Cassation and the Council of State are

Makale Bilgisi

Başvuru:
xx/xx/201x
Kabul:
xx/xx/201x

* İletişim e-posta: emirkanokursoy@uludag.edu.tr

crucial for ensuring legal unity, protecting the rights and freedoms of individuals and providing justice. In this context, the decisions of both Supreme Courts taken from the Legislation and Jurisprudence program of the National Judicial Network Project were used, and 11% of these decisions were approvals and 89% were reversals. Therefore, the problem considered is defined as an imbalanced classification problem. First, features were extracted from decision texts using data preprocessing and natural language processing (NLP) techniques. Then, Synthetic Minority Oversampling Technique (SMOTE) and random undersampling were applied to eliminate the imbalance in the data. Finally, in order to predict the court decisions, classification models with k-nearest neighbors, decision tree, support vector machine, random forest, LightGBM, XGBoost and artificial neural networks were developed, and the performances of the models were compared. The experimental results showed that the proposed decision support system has the potential to benefit legal professionals.

Keywords: Artificial Intelligence, Natural Language Processing, Decision Support System, Law, Imbalanced Classification

1 Giriş

Yapay zekâ, bilgisayarların insana benzer yetenekler sergileyebilmesini, diğer bir ifadeyle insan zekâsını taklit edebilmesini sağlayan; son yıllarda hızla gelişen bir teknoloji ve bilim alanıdır. Yapay zekâ, doğal dil işleme, makine öğrenmesi ve derin öğrenme gibi teknolojiler aracılığıyla karmaşık problemleri çözmeye, bilgiyi öğrenme ve karar verme yetileri kazanmıştır.

Yapay zekânın bir alt dalı olan doğal dil işleme (NLP) sistemleri; insanlar tarafından kullanılan doğal dillerin bilgisayarlarca anlaşılmasını sağlayan bir teknolojidir. Doğal dil işleme teknolojisi aracılığıyla yazılı metinler ile konuşmaların bilgisayar tarafından anlaşılabilirlik duygu analizi, makine çevirisi gibi işlevlerin yerine getirilmesi mümkündür.

Hukuk alanının temeli yazılı metin ve dokümanlar oluşturmaktadır. Söz konusu metin ve belgelerin incelenerek değerlendirilmesi hâkimler, savcılar ve avukatlar gibi hukuk profesyonellerinin iş yükünün önemli bir bölümünü oluşturmaktadır. Doğal dil işleme (NLP) teknikleri aracılığıyla bu belgelerin incelenerek sınıflandırılması, özetlenmesi ve belgeler üzerinde benzerlik analizleri yapılması mümkündür. Bu sayede hukuki kararların daha hızlı ve etkin bir şekilde verilme potansiyeli söz konusu olmaktadır.

Bu çalışmanın temel amacı, yapay zekâ yöntemlerini kullanarak hâkim, savcı ve avukatlar ile diğer hukuk profesyonelleri için bir karar destek sistemi geliştirmektir. Bunun için yüksek mahkemeler olan Yargıtay ve Danıştay kararlarının gerekçe kısımlarına bakarak karar sonucunu onama veya bozma olarak sınıflandıran bir yapay zekâ modeli kurulmuştur. Bunun yanında açıklanabilirlik

yöntemleri entegre edilerek karar alma süreçleri açıklanabilir kılınmaya çalışılmıştır. Hukuk gibi karar alma süreçlerinin büyük önem arz ettiği bir alanda, modelin belirli sonuçlara neden ve nasıl ulaştığının açıklanabilir olması, hukuk alanında yapay zekâ kullanımının içselleştirilmesi açısından büyük öneme sahiptir.

Yargıtay ve Danıştay ilamlarının sınıflandırılmasında açıklanabilir yapay zekâ yöntemlerinin kullanılması, gelecekteki benzer araştırma ve uygulamalar için referans teşkil edecektir.

Makalenin 2. Bölümü'nde literatürde yer alan çalışmalar incelenmiştir ve bu çalışmanın özgün yönü verilmiştir. 3. Bölüm'de çalışmanın gerçekleştirilmesinde kullanılan materyal ve yöntem açıklanmıştır. 4. Bölüm'de yapılan deneysel çalışmalar anlatılmıştır. 5. Bölüm'de ise yapılan deneysel çalışmaların sonuçları sunulurken açıklanabilirlik yöntemleriyle en iyi sonucu veren modelin kararları açıklanmıştır.

2 Literatür araştırması

Literatürde, yapay zekâ tekniklerinin Türk hukuk sistemine uygulandığı sınırlı sayıda çalışma bulunmaktadır. Örneğin; Arslan, Talaş ve Çubukçu'nun çalışmasında, bireylerin hukuki sorunlarına yönelik olarak hangi alanda hizmet veren bir hukuk danışmanına başvurmaları gerektiğini belirleyebilen yapay zekâ destekli bir web uygulaması geliştirilmiştir. Kullanıcılar tarafından girilen metinler çeşitli ön işleme aşamalarından geçirilerek kelime torbası (BOW) yöntemiyle sınıflandırılmıştır. Uygulama kullanıcıların verdiği metinlerdeki kelime frekanslarını değerlendirerek, kullanıcıları doğru alanda hizmet veren bir danışmana yönlendirmiştir. [1]

Turan, Küçüksille ve Kemaloğlu Alagöz tarafından yapılan bir çalışmada, Türkiye Cumhuriyeti Anayasa Mahkemesi tarafından verilen kararları sınıflandırılmıştır. Anayasa Mahkemesi'nin Kararlar Bilgi Bankası'nda yayınlanan veriler kullanılarak k-en yakın komşuluk (KNN), destek vektör makineleri (SVM), karar ağaçları (DT), rassal orman (RF) ve aşırı gradyan artırma (XGBoost) algoritmaları ile tahmin modelleri geliştirilmiştir. Makalede açıklanabilirlik tekniği olarak SHAP (shapley additive explanations) yönteminden yararlanılmıştır. [2]

Görentaş ve diğerleri tarafından yapılan çalışmada, makine öğrenmesi ve doğal dil işleme ile Uyuşmazlık Mahkemesi kararlarının sınıflandırılması ele alınmıştır. Veriler üzerinde ön işleme uygulanıp akabinde terim sıklığı-ters belge frekansı (TF-IDF) vektörizasyonu gerçekleştirilmiştir. Lojistik regresyon, destek vektör makineleri, karar ağaçları ve rassal orman gibi makine öğrenmesi algoritmalarıyla sınıflandırma yapılmıştır. [3]

Uçkan ve Görentaş tarafından yürütülen bir diğer çalışmada Uyuşmazlık Mahkemesi kararları kümelenmiştir. Önce veri üzerinde ön işleme yapıp TF-IDF vektörleri oluşturulmuştur. Sonrasında temsilciler kullanarak kümeleme (cure), k-ortalamlar (k-means), gürültülü uygulamaların yoğunluk tabanlı uzamsal kümelenmesi (DBscan), aglomeratif yuvalama (agnes), yakınlık yayılımı (affinity) ve dengeli iteratif azaltma ve hiyerarşiler kullanarak kümeleme (birch) kümeleme teknikleri uygulanmıştır. En iyi sonuç birch tekniği ile elde edilmiştir. [4]

Karabel ve Aydemir tarafından yürütülen çalışmada ise medeni usul hukuku alanında yapay zekâ ve çevrimiçi yargılama uygulamalarını ele alarak bunları usul ekonomisi ve mahkemeye erişim hakkı bakımından değerlendirmiştir. Yapay zekâ kullanımı Türk hukuk sistemindeki küçük uyuşmazlıkların yanı sıra diğer uyuşmazlık türleri bakımından da yargılamaların hızlandırılması perspektifinden ele alınarak ulusal yargı ağı projesi (UYAP) için uygun olabilecek değişiklikler önerilmiştir. [5]

Mumcuoğlu ve diğerleri tarafından Türkiye'deki yüksek mahkemelerin verecekleri kararların doğal dil işleme, makine öğrenmesi ve derin öğrenme yöntemleriyle tahmin edilmesine yönelik kapsamlı bir analiz yapılmıştır. Bu amaçla karar ağaçları, rassal orman, destek vektör makineleri, kapılı tekrarlayan birimler (GRU), uzun kısa süreli bellek (LSTM) ağları ve dikkat mekanizması ile güçlendirilmiş çift yönlü LSTM (BiLSTM) gibi hem makine öğrenmesi hem de derin öğrenme tabanlı yöntemler kullanılmıştır.

Çeşitli mahkeme türlerinde (Anayasa Mahkemesi, Yargıtay Hukuk Daireleri, Ceza Daireleri, Vergi Dava Daireleri vb.) modellerin başarısı incelenmiştir. Derin öğrenme modellerinin performansının yüksek olduğu, en yüksek doğruluk oranına Bölge İdare Mahkemesi Vergi Dava Dairelerinin verdiği kararlarla dikkat mekanizması eklenmiş BiLSTM modelinde ulaşıldığı belirtilmiştir. [6]

Yapay zekâ tekniklerini diğer ülkelerin hukuk sistemlerine uygulayan çalışmalar bulunmaktadır. Bunlardan biri Almuzaini ve Azmi tarafından Suudi Arabistan mahkemeleri için gerçekleştirilen çalışmadır. Adli metinlerin işaretlenmesi, davaların sınıflandırılması ve bilgiye erişim sağlanması gibi işlevlerin gerçekleştirilmesi için TaSbeeb adlı derin öğrenme tabanlı bir yargısal karar destek sistemi önerilmiştir. Adli metinlerin etiketlenmesi için yarı otomatik Ann-judicial tekniği, bilgiye erişimin iyileştirilmesi için ise Jud_RoBERTa adlı yargısal bir dil modeli geliştirilmiştir. Homojen ve heterojen derin öğrenme modelleri tasarlanarak dengesiz veri kümeleri üzerinde sınıflandırmalar yapılmış ve SMOTE ve rastgele alt örnekleme yöntemleri ile sınıflandırıcıların performansları değerlendirilmiştir. [7]

Dang ve diğerleri, büyük dil modellerinin (LLM) hukuki metinlerin işlenmesi üzerindeki etkisini kapsamlı bir şekilde incelemiştir. Çalışmada LLM'lerin hukuki metinlerin özetlenmesi, bilgi çıkarımı ve sınıflandırılması gibi çeşitli uygulamalardaki potansiyeline değinilmiş ve bu modellerin özellikle bağlamı anlama, uzun ve karmaşık yapıdaki cümleleri çözümleme gibi işlemler için kullanılabileceğini ele alınmıştır. [8]

Arriba-Pérez ve diğerleri tarafından gerçekleştirilen çalışmada, İspanyol mahkemelerinin kararları birden çok etiketle sınıflandırılmıştır. Oluşturulan modelde; makine öğrenmesi, derin hukuki muhakeme ve doğal dil işleme yöntemleri bir araya getirilmiştir. [9]

Aletras ve diğerleri, Avrupa İnsan Hakları Mahkemesi (AİHM) davalarının sonuçlarını tahmin etmek için metin madenciliği ve makine öğrenmesi yöntemlerini kullanmışlardır. Bahse konu çalışmada, dava metinleri n-gramlar ve modelleme teknikleriyle temsil edilmiş ve destek vektör makineleriyle ikili sınıflandırma problemi olarak modellenmiştir. Çalışma sonuçları mahkeme kararlarının %79 doğrulukla tahmin edilebildiğini ve Olgular bölümünün, kararlar üzerinde en yüksek etkiye sahip faktör olduğunu göstermiştir. Bu sonuçların kararların hukuki formalizm yerine çoğunlukla

olayların etkisiyle şekillendiğini öne süren hukuki realizm teorisiyle de uyumlu olduğu gözlemlenmiştir. [10]

Yapay zekâ teknikleri farklı alanlardaki metinlerin analiz edilmesinde ve örüntülerin çıkarılmasında yaygın olarak kullanılmaktadır. Bu alandaki çalışmaların bir bölümü dengesiz sınıflandırma problemini ele almaktadır. Örneğin; Budaya ve Suniantara'nın çalışmasında, Endonezya'nın Bali bölgesindeki halk sağlığı merkezleri ile ilgili Google yorumları üzerinde duygu analizi yapılmıştır. Sınıf dengesizliği SMOTE tekniğiyle giderilmiştir. TF-IDF yöntemi kullanılmış olup yüksek TF-IDF skoruna sahip kelimeler, pozitif ve negatif duyguların ayrıştırılmasında temel göstergeler olarak kullanılmıştır. Yorumlar olumlu ve olumsuz olarak ikili şekilde sınıflandırılmıştır. Sınıflandırma için lojistik regresyon, destek vektör makineleri, XGBoost, naive bayes ve rassal orman gibi makine öğrenmesi modelleri kullanılmıştır. Lojistik regresyon modelinin yüksek doğruluğu, yorumlanabilirliği ve ikili sınıflandırma problemlerine uygunluğu nedeniyle çoğu duygu analizi çalışması için uygun olduğunu, ancak karmaşık veri setlerinde XGBoost veya rassal orman gibi alternatif modellerin daha iyi sonuçlar verebileceği vurgulanmıştır. [11]

Khan, Towhid, Faruk ve diğerleri tarafından yürütülen çalışmada, Bangla haber veri seti üzerinde sınıf dengesizliğini gidermek amacıyla rastgele alt örnekleme ve sentetik azınlık aşırı örnekleme tekniği (SMOTE) uygulanmıştır. Sınıflandırma işlemlerinde lojistik regresyon, karar ağacı ve stokastik gradyan inişi gibi makine öğrenmesi modellerinin yanı sıra yapay sinir ağı (ANN), evrişimli sinir ağı (CNN) ve çift yönlü kodlayıcı dönüştürücüler (BERT) gibi derin öğrenme modelleri de kullanılmıştır. Modellerin performansları karşılaştırılarak en yüksek performansı BERT modelinin sergilediği belirlenmiştir. Ek olarak sınıflandırma modellerine yerel yorumlanabilir modelden bağımsız açıklamalar (LIME) ve SHAP teknikleri entegre edilerek model tahminlerinin daha iyi açıklanabilir hale getirilmesi sağlanmıştır. [12]

Rupapara ve diğerleri tarafından gerçekleştirilen çalışmada, sosyal medya platformlarında yer alan kaba ve saygısız (toksik) yorumların sınıflandırılması için bir sistem geliştirilmiştir. Özellikle dengesiz veri setlerinde toksik yorumların doğru tespiti için SMOTE ve rastgele alt örnekleme ile beraber lojistik regresyon ve destek vektör makinelerinin birleşiminden oluşan regresyon vektör oylama sınıflandırıcısı (RVVC) adlı yeni bir

model önerilmiştir. TF-IDF ve BOW yöntemleri ile metinlerden öznitelikler çıkarılmıştır. RVVC modelinin; veri dengelemenin SMOTE kullanılarak yapıldığı, vektörleştirme için TF-IDF uygulandığı durumda diğer makine öğrenmesi yöntemlerine kıyasla daha üstün bir performans sergilediği görülmüştür. [13]

Hukuk dışındaki alanlarda Türkçe dili için metin sınıflandırma ve doğal dil işleme tekniklerinin kullanıldığı çalışmalar da bulunmaktadır. Örnek olarak; Koçak ve diğerleri tarafından yapılan çalışmada, 500 adet Ar-Ge projesi ve ilgili makalelerden oluşan bir veri seti üzerinde doğal dil işleme ve derin öğrenme yöntemleri beraberce uygulanmıştır. Veri seti üzerinde ön işleme adımları uygulanarak metinler kelimelerine ayrılmış, tümü küçük harfe çevrilmiş ve noktalama işaretleri kaldırılmıştır. Sonrasında word2vec yöntemiyle kelime temsilleri üretilmiştir. CNN yöntemi ile sınıflandırma yapılmıştır. Çalışma sonucunda, uygulanan yöntemle Ar-Ge projelerinin etkin bir şekilde sınıflandırılabilceği gösterilmiştir. [14]

Toğaçar ve diğerleri tarafından yapılan bir araştırmada, sahte haberlerin tespitine yönelik olarak doğal dil işleme ve derin öğrenme yöntemleri birlikte kullanılmıştır. LSTM modelinin doğal dil işleme teknikleri ile birlikte uygulanmasının sahte haberlerin saptanmasında etkili olabileceği sonucuna varılmıştır. [15]

Sevli ve Kemaloğlu Alagöz, olağandışı olaylarla ilgili tweet'lerin gerçek veya gerçek dışı şeklinde sınıflandırılması için Google BERT modelini kullanmışlardır. Yapılan çalışmada 7613 adet tweet içeren, deprem, kötü hava şartları, kaza gibi durumları kapsayan bir veri seti kullanılmıştır. Veri setinin %80'lik kısmı eğitim, %20'lik kısmı ise test verisi olarak ayrıldıktan sonra BERT modeli ile eğitilmiştir. Sınıflandırma sonucunda %98,88 doğruluk oranına ulaşıldığı, bu yöntemle olağandışı gelişen durumlara ilişkin sosyal medyada yer alan bilgilerin hızlı ve etkili bir şekilde analiz edilebileceği belirtilmiştir. [16]

Öztürk, Durak ve Badıllı tarafından yapılan bir çalışmada, Twitter'dan elde edilen mesajlar incelenmiş ve bu mesajların halk sağlığı ile ilişkisi bulunup bulunmadığı, ilişkili olması halinde hangi hastalıklarla ilgili olduğu analiz edilmiştir. Veri toplama adımının ardından TF-IDF ve BOW yöntemleri ile öznitelikler çıkarılmıştır. Gözetimli ve gözetimsiz öğrenme algoritmaları ile kümeleme yapılarak gözetimli öğrenmenin gözetimsiz öğrenmeye kıyasla daha iyi performans gösterdiği

saptanmıştır. Sosyal medya verilerinin halk sağlığı alanında hastalıkların yayılımını izlemek ve epidemiyolojik analizler yapmak için uygun bir kaynak olabileceği gösterilmiştir. [17]

Bu çalışmanın literatürdeki çalışmalardan farkı Yargıtay ve Danıştay kararlarından oluşan bir veri setinde Türk hukuk sistemi için LIME ve yapay zekâ tabanlı karar destek sistemi geliştirilmesidir. Ayrıca, ele alınan veri setinde sınıflar arasında dengesizlik olduğu için doğruluk oranının artırılması amacıyla SMOTE ve rastgele alt örnekleme yöntemleri uygulanmıştır.

3 Materyal ve yöntem

Bu bölümde, çalışmada kullanılan veri seti tanıtılarak veri ön işleme adımları ve kullanılan yöntemler açıklanmıştır.

3.1 Veri kümesi

Bu çalışmada kullanılan veri kümesi mahkeme kararı verilerini içeren UYAP Mevzuat ve İçtihat Programından elde edilmiştir. Veri kümesi; Yargıtay Ceza ve Hukuk Daireleri, Ceza Genel Kurulu, Hukuk Genel Kurulu, Danıştay daireleri ve İdari Dava Daireleri Kuruluna ait 6425 adet mahkeme kararını içermektedir. Karar tarihleri 1975-2016 yılları arasındadır. Veri setindeki onama kararı sayısı 708, bozma kararı sayısı ise 5704'tür. Ayrıca veri kümesi içinde beş adet Anayasa Mahkemesi kararı mevcut olup bu kararlar silinmiştir. [18]

3.2 Veri ön işleme

Veri kümesi içerisinde yer alan mahkeme kararlarına veri ön işleme uygulanmıştır. Veri ön işleme aşamasında düz metin formatına getirilen mahkeme kararlarının gerekçe kısımları üzerinde etkisiz kelimeler (stop words) çıkartılmış ve kelimeler köklerine ayrılmıştır. Ardından TF-IDF vektörizasyonu ile sınıflar arası dengesizliğin giderilmesi için SMOTE ve rastgele alt örnekleme yöntemleri ayrı ayrı ve hibrit olarak uygulanmıştır. Bu bölümde TF-IDF vektörizasyonu ile veri kümesindeki dengesizliğin giderilmesi için kullanılan yöntemler açıklanmıştır.

3.2.1 TF-IDF

TF, belgelerdeki belirli kelimelerin sıklığını ifade etmektedir. TF'in matematiksel formülasyonu Denklem 1'de gösterilmiştir.

$$TF_{i,j} = \frac{n_{i,j}}{\sum_k n_{k,j}} \quad (1)$$

Burada $n_{i,j}$, t_i teriminin d_j belgesindeki frekansını, $\sum_k n_{k,j}$ ise d_j belgesindeki tüm terimlerin toplam frekansını ifade etmektedir. Kelimelerin TF değerleri yükseldikçe belge için önem düzeyi artmaktadır. Doküman frekansı (DF) ise, belirli bir kelimenin belge koleksiyonunda kaç farklı belgede bulunduğunu göstermektedir. Yüksek DF değeri olan kelimelere çoğu belgede rastlandığından bu kelimelerin önem seviyesinin düşük olduğu kabul edilmektedir. Buna bağlı olarak DF'nin tersi olan IDF, belge koleksiyonundaki anahtar kelimelerin öneminin ölçülmesi amacıyla kullanılmaktadır. IDF'in matematiksel formülasyonu Denklem 2'de gösterilmiştir.

$$IDF_{i,j} = \log \frac{|D|}{|d_j \in D: t_i \in d_j|} \quad (2)$$

Denklemde $|D|$ belge sayısını, $|d_j \in D: t_i \in d_j|$ belirli bir t_i teriminin bulunduğu belge sayısını ifade etmektedir. [19]

3.2.2 SMOTE

Veri kümesinde sınıflar arasındaki dengesizliği gidermek için kullanılan bir yöntem olan SMOTE, veri kümesinde azınlıkta olan sınıfa ait verilerden sentetik örnekler üreterek veri kümesini dengelemektedir. Denklem 3, SMOTE tekniğinin sentetik örnek üretimini göstermektedir.

$$s = x + u \cdot (x_R - x) \quad (3)$$

Burada, azınlık sınıfındaki her bir örnek için (x), öznitelik uzayındaki k en yakın komşu sayısı belirlenerek k en yakın komşulardan biri rastgele seçilir (x_R). Denklemdeki u katsayısı, 0 ile 1 arasında rastgele bir sayı olup üretilen sentetik örneğin (s), x ile x_R arasındaki doğru parçasında kalmasını sağlamaktadır. [20]

3.2.3 Rastgele alt örnekleme

Rastgele alt örnekleme (random undersampling) azınlık sınıfı dışındaki sınıfın gözlem sayısını azaltarak veri setindeki dengesizliği gidermeye çalışan bir yöntemdir. Çoğunluk sınıfından rastgele örnekler seçilir ve veri seti azaltılır. Böylece çoğunluk sınıfının azınlık sınıfına eşit olması veya azınlık sınıfının belli bir oranı kadar sayıda örnek içermesi sağlanır. [21]

3.3 Kullanılan yöntemler

Bu bölümde sınıflandırma işlemleri için kullanılan makine öğrenmesi yöntemleri açıklanmıştır.

3.3.1 K-en yakın komşuluk

KNN yöntemi veri kümesi üzerinde herhangi bir varsayımda bulunmayan, basit ve etkili bir denetimli öğrenme algoritmasıdır. Sınıf etiketi bulunan eğitim verileri üzerinden test verilerinin sınıf etiketlerinin tahmininde kullanılmaktadır. KNN algoritması, test verisindeki bir örneğin sınıfını belirlemek için eğitim veri kümesindeki k en yakın komşunun sınıfına bakar ve çoğunluk hangi sınıfta ise test verisini de o şekilde sınıflandırır. k en yakın komşuların belirlenmesinde kullanılan uzaklıklardan biri öklid uzaklığıdır. x ve y nesneleri arasındaki öklid uzaklığının hesaplanması Denklem 4'te gösterilmiştir:

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (4)$$

Burada, $d(x, y)$ iki veri noktası arasındaki Öklid uzaklığını; x_i ve y_i , sırasıyla x ve y noktalarının i 'inci öznitelikliğini ve n öznitelik sayısını göstermektedir. Algoritmada komşu sayısı (k değeri) çapraz doğrulama ile belirlenebilir. [22]

3.3.2 Karar ağacı

DT, verilerin özniteliklerine göre bölünerek ağaç yapısının oluşturulmasına dayalı bir sınıflandırma modelidir. Algoritma en iyi bölme öznitelikliğini seçer, bu öznitelik kök düğüm olarak belirlenir ve veri kümesi bu öznitelikliğin değerlerine göre dallara ayrılarak alt kümeler bölünür. Veri kümesinin olabildiğince saf alt kümeler bir başka deyişle aynı sınıftan örnekler içeren kümeler ayrılması amaçlanır. Her alt küme için bu işlem tekrarlanır. Alt kümelerin saf hale geldiği veya daha fazla bölme yapılamadığı noktada yaprak düğümler oluşturulur ve sınıf etiketleri atanır. Öznitelik seçiminde safsızlık ölçütleri olan entropi ve gini indeksi kullanılır. Entropi veri setindeki düzensizliği, gini indeksi ise heterojenliği ölçer. Entropi ve gini indekslerinin matematiksel formülasyonları Denklem 5 ve 6'da gösterilmiştir:

$$Entropi(S) = - \sum_{i=1}^c \frac{|S_i|}{|S|} \log \left(\frac{|S_i|}{|S|} \right) \quad (5)$$

$$Gini(S) = 1 - \sum_{i=1}^c \left(\frac{|S_i|}{|S|} \right)^2 \quad (6)$$

Burada, S veri kümesini, c sınıf sayısını, S_i i sınıfına ait örnekler kümesini ifade etmektedir. [23]

3.3.3 Destek vektör makineleri

SVM, veri kümesinde sınıflar arasındaki ayrımı en geniş marjinle sağlayan hiper düzlemi bulmayı hedefleyen bir makine öğrenmesi algoritmasıdır. Sınıfların lineer olarak ayrılabilirdiği veri kümelerinde sınıflar arasındaki marjini maksimize eden en uygun hiper düzlemi kullanarak verileri ayırır. Verilerin lineer olarak ayıramaması halinde daha yüksek boyutlu bir öznitelik uzayına geçiş için çekirdek fonksiyonları (kernel) kullanılır. Böylelikle orijinal uzayda lineer olmayan veriler yüksek boyutlu uzayda lineer olarak ayrılabilir hale gelir. Yaygın olarak doğrusal (linear), polinom (polynomial), gaussian radial basis function (RBF) ve sigmoid çekirdek fonksiyonları kullanılmaktadır. Verinin özniteliklerine göre uygun çekirdek fonksiyonu seçilerek modelin performansı optimize edilebilir. [24]

3.3.4 Rassal orman

RF, birden fazla karar ağacının bir araya gelmesiyle oluşan bir sınıflandırma algoritmasıdır. Algoritmayı oluşturan ağaçlar birbirinden bağımsız ve aynı dağılıma sahip rastgele vektörlerden oluşur. Her bir karar ağacı, sınıf etiketi atanmasında bir oy kullanır ve çoğunluğun oyuna göre nihai sınıf etiketi belirlenir. Rassal orman modeli bağımsız ve aynı dağılıma sahip rastgele vektörlerden ($\{h(x, \theta_k), k = 1, \dots\}$) oluşmaktadır. Bu yapıda x , girdi vektörünü, $h(x, \theta_k)$, k 'inci karar ağacını temsil etmektedir. Model K kez çalıştırıldıktan sonra sınıflandırıcılar dizisi ($h_1(x), h_2(x), \dots, h_K(x)$) ile birden fazla model elde edilmiş olur. Nihai sınıf etiketinin atanması için hesaplanan çoğunluk oyu, Denklem 7'de gösterilen formül ile elde edilir:

$$H(x) = \arg \max_Y \sum_{i=1}^K I(h_i(x) = Y) \quad (7)$$

Burada $H(x)$ sınıflandırma modellerinin kombinasyonunu ifade eder. (h_i) rassal orman içerisindeki i 'inci karar ağacı modelidir. Toplam karar ağacı sayısı K ile, indikatör fonksiyonu $I(\cdot)$ ile gösterilmiştir. Y ise çıktı değişkenini temsil etmektedir. [25]

3.3.5 LightGBM

LightGBM, Microsoft araştırmacıları tarafından geliştirilmiş kademeli güçlendirilmiş karar ağaçları (gradient boosting decision tree-GBDT) algoritmasının yüksek verime sahip bir uygulamasıdır.

GBDT algoritması, karar ağaçlarını kullanarak giriş uzayından (X_s) gradyan uzayına (G) bir fonksiyonun öğrenilmesini sağlar. Karar ağacı modeli, her düğümü bilgi kazancının en yüksek olduğu öznitelikte böler. Bir düğümün bölünmesiyle elde edilen bilgi kazancı varyans formülü ile ölçülmekte olup Denklem 8'de gösterilmiştir.

$$V_{j|O}(d) = \frac{1}{n_O} \left(\frac{\left(\sum_{\{x_i \in O: x_{ij} \leq d\}} g_i \right)^2}{n_{j|O}(d)} + \frac{\left(\sum_{\{x_i \in O: x_{ij} > d\}} g_i \right)^2}{n_{jr|O}(d)} \right) \quad (8)$$

Burada, O indisi karar ağacının sabit bir düğümünde bulunan eğitim verisini temsil etmektedir. Denklemde hesaplanan n_O düğümdeki veri örneği sayısını, $n_{j|O}(d)$ ve $n_{jr|O}(d)$ sırasıyla j özneliliğinin d noktasında bölünmesi sonrasında sağ ve sol çocukta veri örneklerinin sayılarını, g_i her bir veri noktası için, modelin hatasını azaltmak amacıyla kullanılan kayıp fonksiyonunun gradyan değerini, $V_{j|O}(d)$ ise bu bölünme sonucu oluşan varyans kazancını göstermektedir. LightGBM, kullandığı gradyan tabanlı tek taraflı örnekleme (GOSS) ve özel öznitelik demetleme (EFB) gibi teknikler sayesinde eğitim süresini önemli ölçüde kısaltarak hesaplama maliyetlerini azaltmaktadır. Dolayısıyla yüksek boyutlu veri setleri üzerinde etkili bir şekilde çalışabilmektedir. [26,27]

3.3.6 XGBoost

XGBoost, ağaç güçlendirme (tree boosting) tabanlı bir algoritma olup genellikle karar ağaçlarından oluşan bir dizi zayıf tahmin edicinin birbiri ardına eğitilip her birinin kendisinden önceki modelin hatalarını düzeltmeye çalışması ilkesine dayanır. Karmaşık ve büyük veri setlerinde tercih edilmektedir. Modelin en uygun yaprak ağırlığının (w_j^*) hesaplanması Denklem 9'da gösterilmiştir:

$$w_j^* = - \frac{\sum_{i \in I_j} g_i}{\sum_{i \in I_j} h_i + \lambda} \quad (9)$$

Burada $\sum_{I_j} g_i$ yaprak j 'deki tüm örnekler için birinci dereceden türevlerin toplamını, $\sum_{I_j} h_i$ ise yaprak j 'deki tüm örnekler için ikinci dereceden türevlerin toplamını temsil etmektedir. λ , modelin aşırı uyumunu (overfitting) engellemek için eklenen düzenleme parametresidir.

Denklem 10'da toplam kayıp fonksiyonu değerinin hesaplanması gösterilmiştir:

$$\widetilde{L^{(t)}}(q) = -\frac{1}{2} \sum_{j=1}^T \frac{\left(\sum_{i \in I_j} g_i \right)^2}{\sum_{i \in I_j} h_i + \lambda} + \gamma T \quad (10)$$

Burada T yaprak sayısını ifade etmekte olup γT modelin karmaşıklığını cezalandıran bir düzenleme parametresidir. [28, 29]

3.3.7 Yapay sinir ağırları

ANN, biyolojik sinir sistemlerinin çalışma şekline esinlenilerek oluşturulan matematiksel modellerdir. Bu modellerde girdi, gizli ve çıktı olmak üzere üç ana katman bulunmaktadır. İleri beslemeli yapay sinir ağına girdiler ilk hücreden n 'inci hücreye ($x_1, \dots, x_n$) doğru tek yönlü olarak iletilir. Aynı şekilde bir nöronun çıktısı (O) da tek yönlüdür. Girdilerin ağırlıklı toplamı bir (f) transfer fonksiyonu ile çıktıya dönüştürülmektedir. Denklem 11 sinir ağına çıktılarının üretilmesine ilişkin matematiksel formülasyonu göstermektedir:

$$O = f(net) = f \left(\sum_{j=1}^n w_j x_j \right) \quad (11)$$

Burada nöronun çıkış değeri (O) bir nörona gelen toplam giriş (net) bağlı olarak belirlenir. Formüldeki net değeri ağırlık vektörünün (w) transpozisi ile giriş vektörünün (x) skaler çarpımı yapılarak hesaplanır. Denklem 12'de net değerinin hesaplanması formüle edilmiştir:

$$net = w^T x = w_1 x_1 + \dots + w_n x_n \quad (12)$$

ANN modeli, hatanın geriye yayılımı (back propagation) gibi eğitim algoritmaları ve CNN gibi gelişmiş ağ yapıları kullanılarak çeşitli amaçlar için oluşturulabilir. [30,31]

3.4 Açıklanabilir yapay zekâ

Bu çalışmada sınıflandırma modellerinin verdiği kararları anlaşılır hale getirmek için açıklanabilir yapay zekâ yöntemlerinden LIME tekniği uygulanmış olup bu bölümde açıklanmıştır.

LIME, herhangi bir sınıflandırma modelinin tahminlerini, tahminin etrafındaki yerel bölgede modeli eğiterek açıklayan bir tekniktir. LIME veriyi pertürbe ederek belirli bir özniteliği açıklamak için ona yakın küçük değişikliklerle örnekler oluşturur. LIME modelden bağımsız olup modeli kara kutu olarak değerlendirir. Denklem 13'te LIME'in çalışma mekanizması formüle edilmiştir:

$$\xi(x) = \arg \min_{g \in G} L(f, g, \pi_x) + \Omega(g) \quad (13)$$

Buna göre LIME'in yapacağı açıklama $g \in G$ olarak tanımlanır. Burada G , doğrusal modeller, karar ağaçları veya düşen kural listeleri (falling rule lists) gibi olası yorumlanabilir modellerdir. Modelin veri kümesi $\{0,1\}^d$ boyutlarında olup, ikili (binary) yapıdadır ve g fonksiyonu, bu veri kümesindeki özniteliklerin varlığını ya da yokluğunu dikkate alır. $\xi(x)$, x örneğinin açıklaması olup, belirtilen fonksiyonun minimum değerini veren girdinin bulunması hedeflenir. x örneği alınarak bazı öznitelikleri rastgele değiştirilir ve pertürbe edilen örnekler oluşturulur. Bu örnekler x 'e olan yakınlıklarına göre ağırlıklandırılır, ardından $f(z)$ hesaplanarak etiketlenir. Son olarak, bu etiketlerle basit ve yorumlanabilir bir model (g) eğitilir. Sonuç olarak $L(f, g, \pi_x)$ fonksiyonu, x ve diğer örnekler arasındaki yakınlığı ölçen fonksiyon olan π_x 'i kullanarak, LIME'in açıklamasının modelin yaptığı sınıflandırmayı ne derece iyi temsil ettiğini gösterir. Denklemdeki $\Omega(g)$, LIME'in açıklama karmaşıklığını ölçer ve olabildiğince düşük tutulması hedeflenir. Bu şekilde karmaşık modelin davranışı yerel bazda basit ve anlaşılır bir modelle açıklanabilir. [32]

4 Deneysel çalışmalar

Bu bölümde, yapılan deneysel çalışmalar açıklanmıştır.

4.1 Deneysel koşullar

Deneysel çalışmalar Python programlama dilinde sklearn, BeautifulSoup, pandas, numpy, zemberek, nltk, imblearn ve lime kütüphaneleri kullanılarak yapılmıştır.

İlk olarak BeautifulSoup kütüphanesiyle veri kümesi içinde yer alan HTML formatındaki mahkeme kararları düz metin (plain text) formatına çevrilmiştir. Bu aşamadan sonra her bir mahkeme kararının onama veya bozma sınıfında olduğu tespit edilerek sınıf etiketi atanması için aşağıdaki işlemler yapılmıştır.

- Karar metinlerinin sonunda yer alan büyük harflerle "SONUÇ:" veya "HÜKÜM:" ile başlayan paragraf bulunmuştur.
- Bu paragraf baştan sona kadar alınıp içinde ONAMA, BOZMA veya DÜZELTİLEREK ONAMA kelimelerinin geçip geçmediğine bakılmıştır. Kararın farklı bir şekilde yazılması ihtimaline karşın paragrafta bu kelimeler aranmıştır.
- Bu şekilde her karar için sınıf etiketi belirlenmiştir.

Kararların gerekçe kısımlarında "onanması gerekir" veya "bozulması gerekir" gibi sınıf etiketleri bulunabileceğinden, kural tabanlı bir kod ile gerekçenin son cümleleri kontrol edilmiştir. Son cümlede "boz" veya "ona" kelimeleri (ve türevleri) yer alıyorsa, bu cümleler gerekçe metninden çıkarılmıştır.

Mahkeme kararlarının gerekçe kısımlarında bulunan tüm kelimeler küçük harfe çevrilmiştir. Sadece sayılardan oluşan kelimeler ile nltk Türkçe kütüphanesinde yer alan etkisiz kelimeleri (stop words) çıkarılmıştır. Bu aşamalardan sonra veri ön işleme adımlarına devam edilerek, Zemberek kütüphanesiyle metinlerdeki her kelime tespit edilerek köklerine ayrılmıştır. Kelimeler TF-IDF yöntemiyle vektörleştirilmiştir.

Veri kümesinin %80'i modelin eğitimi için kullanılmak üzere eğitim verisi, %20'si ise modelin performansını değerlendirmek amacıyla test verisi olarak ayrılmıştır.

Onama türü kararların sayısı ile bozma türü kararların sayısı arasındaki dengesizlik nedeniyle problem dengesiz sınıflandırma modeli olarak ele alınmıştır. Dengesizliğin giderilmesi için eğitim verilerine SMOTE ve rastgele alt örnekleme teknikleri ayrı ayrı ve hibrit olarak uygulanmıştır.

Veri ön işlemeden sonra tüm sınıflandırma modelleri için en iyi hiperparametreler belirlenmiştir. En iyi hiperparametrelerin belirlenmesinde Tablo 1'deki değer aralıklarında 5-katlı çapraz doğrulama yapılmıştır.

Dengesiz veri setlerindeki model performanslarının değerlendirilmesinde, doğruluk (accuracy), kesinlik (precision) ve duyarlılık (recall veya TPrate) ölçütlerinin yanı sıra, azınlık sınıfının doğru tahmin edilme oranını öne çıkaran metriklerin kullanılması önem arz etmektedir. [35] Bu bağlamda, geometrik ortalama (G-mean), pozitif ve negatif sınıflar arasındaki dengenin ölçülmesinde etkin bir metriktir. [35] Dengeli doğruluk indeksi (IBA) ise sınıflar arasındaki dengeyi dikkate alarak modelin genel performansını değerlendirmektedir. Ayrıca, F-ölçütü, duyarlılık ve kesinlik arasındaki dengeyi ortaya koyması nedeniyle, özellikle dengesiz veri setlerinde sınıflandırma performansını değerlendirmek için yaygın olarak tercih edilmektedir. [33, 34] Bu nedenlerle, bu çalışmada modellerin performansları doğruluk, kesinlik, duyarlılık, F-ölçütü, özgüllük, geometrik ortalama ve IBA metrikleriyle ölçülmüştür. Denklem 14-20'de performans ölçütlerinin formülleri verilmiştir. Performans ölçütlerinin hesaplanmasında Tablo

2'deki karmaşıklık matrisi kullanılmıştır. Burada, IBA değerinin hesaplanmasında kullanılan α , 0 ile 1 arasında değerler almaktadır. [33] Bu çalışmada, $\alpha = 0,1$ alınmıştır.

Tablo 1. Parametreler ve değer aralıkları

Yöntem	Parametreler ve Değer Aralıkları
KNN	K-en yakın komşu sayısı (k) $\in \{3, 5, 7\}$
DT	Maksimum ağaç derinliği (D) $\in \{3, 5, 7\}$
SVM	Düzenleştirme parametresi (C) $\in \{0,1, 1, 10\}$, kernel tipi (K) $\in \{\text{linear, rbf}\}$
RF,Light GBM, XGBoost	Karar ağacı sayısı $\in \{50, 100\}$, Maksimum ağaç derinliği $\in \{3, 5\}$
ANN	Gizli katmandaki nöron sayısı $\in \{50, 100\}$, Gizli katman için aktivasyon fonksiyonu $\in \{\text{tanh, relu}\}$, Ağırlık optimizasyonu için çözücü $\in \{\text{sgd, adam}\}$, Ceza terimi katsayısı $\in \{0,0001, 0,001\}$

Tablo 2. Karmaşıklık Matrisi

	Tahmin Pozitif	Tahmin Negatif
Gerçek Pozitif	Gerçek Pozitif (TP)	Yanlış Negatif (FN)
Gerçek Negatif	Yanlış Pozitif (FP)	Gerçek Negatif (TN)

$$\text{Doğruluk} = \frac{TP + TN}{TP + TN + FP + FN} \quad (14)$$

$$\text{Kesinlik} = \frac{TP}{TP + FP} \quad (15)$$

$$\text{Duyarlılık (TPRate)} = \frac{TP}{TP + FN} \quad (16)$$

$$\text{Özgüllük (TNrate)} = \frac{TN}{TN + FP} \quad (17)$$

$$F - \text{ölçütü} = 2 \cdot \frac{\text{Kesinlik} \cdot \text{Duyarlılık}}{\text{Kesinlik} + \text{Duyarlılık}} \quad (18)$$

$$\text{Geometrik Ortalama} = \sqrt{\text{TPRate} \cdot \text{TNrate}} \quad (19)$$

$$\text{IBA} = (1 + \alpha \cdot (\text{TPRate} - \text{TNrate})) \cdot \text{TPRate} \cdot \text{TNrate} \quad (20)$$

Sonuç olarak, geometrik ortalama, IBA ve F-ölçütü gibi metrikler, duyarlılık ve özgüllüğü değerlendirme sürecine dâhil eden ölçütlerdir. Duyarlılık, onama kararlarının doğru tespit edilme oranını ifade ederken, özgüllük bozma kararlarının doğru sınıflandırılma oranını gösterir. Model performanslarının değerlendirilmesinde bu metriklerle önem verilmesi, dengesiz veri setlerinde

model başarısının daha gerçekçi ve azınlık sınıfı açısından daha adil bir şekilde değerlendirilmesini sağlamaktadır.

4.2 Bulgular ve tartışma

Deneyisel çalışmalarda veri dengeleme yöntemleri şu şekilde incelenmiştir:

- Makine öğrenmesi modeli orijinal veri ile (veri dengeleme olmadan) çalıştırılmıştır.
- Eğitim verisinde azınlık (onama) sınıfındaki örnekler SMOTE yöntemiyle sırasıyla %100, %200, %300 ve %600 oranlarında artırılarak sınıflandırma modelleri geliştirilmiştir.
- Eğitim verisinde çoğunluk (bozma) sınıfındaki örnekler rastgele alt örnekleme (RA) yöntemiyle azaltılmıştır. Burada, x:y ifadesi azınlık sınıfı örnek sayısının çoğunluk (bozma) sınıfı örnek sayısına oranını göstermektedir. Örneğin, 1:2 ifadesinde çoğunluk sınıfı örnek sayısı, azınlık sınıfı örnek sayısının iki katıdır. Bu kapsamda, RA'da 1:2 ve 1:1,6 oranları incelenmiştir.
- Eğitim verisine SMOTE ve RA yöntemleri birlikte uygulanmıştır. İlk olarak, SMOTE ile eğitim verilerindeki azınlık (onama) sınıfı %140 oranında artırılmıştır. Ardından rastgele alt örnekleme uygulanarak azınlık sınıfı örnek sayısının çoğunluk (bozma) sınıfı örnek sayısına oranı 1:2 olacak şekilde düzenlenmiştir.

Tablo 3'te K-en yakın komşuluk, karar ağacı, destek vektör makinesi, rassal orman, LightGBM, XGBoost ve yapay sinir ağları modelleriyle yapılan sınıflandırma işlemlerinin test verileri için sonuçları gösterilmiştir. Her performans kriteri için en iyi değer * ile belirtilmiş, seçilen yöntem ise koyu ile işaretlenmiştir.

Tablo 3'e göre, veri dengeleme yöntemlerinin uygulanması, rassal orman ve LightGBM modellerinde geometrik ortalama kriterinde iyileşme sağlamıştır. Ancak, SMOTE'nin uygulanması KNN ve karar ağacı modellerinde F-ölçütü değerlerinin kötüleşmesine neden olmuştur. KNN modeli için hibrit yöntem dışındaki SMOTE ve rastgele alt örnekleme tekniklerinin ayrı ayrı uygulandığı durumlarda geometrik ortalama ve IBA değerlerinde iyileşme gözlenmiştir. Karar ağacı modelinde SMOTE uygulaması geometrik ortalama ve IBA değerlerini iyileştirmiştir. Destek vektör makinesi ve XGBoost modellerinde SMOTE'nin olumlu bir etkisi görülmemiştir.

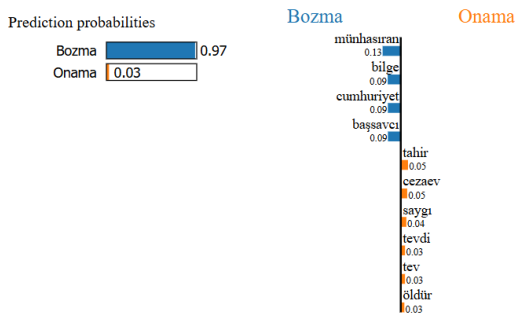
Tablo 3. Sınıflandırma Modellerinin Sonuçları

Model	Veri Dengeleme Yöntemi	Doğruluk	Kesinlik	Duyarlılık (TPrate)	Özgüllük (TNrate)	F-ölçütü	Geometrik ortalama	IBA
KNN	Yok	0,910	0,671	0,978	0,359	0,468	0,593	0,330
	SMOTE (%100)	0,533	0,175	0,492	0,866	0,291	0,653	0,442
	SMOTE (%200)	0,493	0,165	0,445	0,880	0,278	0,626	0,409
	SMOTE (%300)	0,479	0,163	0,426	0,901	0,277	0,620	0,402
	SMOTE (%600)	0,445	0,156	0,387	0,908	0,266	0,593	0,370
	RA (1:2)	0,772	0,275	0,788	0,648	0,387	0,714	0,503
	RA (1:1,6)	0,743	0,256	0,748	0,697	0,375	0,722	0,519
	SMOTE (%140)+RA (1:2)	0,313	0,134	0,234	0,951	0,235	0,472	0,238
DT	Yok	0,894	0,534	0,964	0,331	0,409	0,565	0,299
	SMOTE (%100)	0,882	0,443	0,957	0,275	0,339	0,513	0,245
	SMOTE (%200)	0,882	0,458	0,949	0,345	0,394	0,572	0,308
	SMOTE (%300)	0,828	0,315	0,872	0,472	0,377	0,641	0,395
	SMOTE (%600)	0,828	0,318	0,870	0,486	0,384	0,650	0,407
	RA (1:2)	0,834	0,354	0,862	0,606	0,447	0,723	0,509
	RA (1:1,6)	0,820	0,324	0,850	0,577	0,415	0,701	0,478
	SMOTE (%140)+RA (1:2)	0,835	0,339	0,874	0,521	0,411	0,675	0,439
SVM	Yok	0,924*	0,881	0,994	0,366	0,517	0,603	0,341
	SMOTE (%100)	0,924*	0,881	0,994	0,366	0,517	0,603	0,341
	SMOTE (%200)	0,924*	0,879	0,994	0,359	0,510	0,597	0,334
	SMOTE (%300)	0,923	0,877	0,994	0,352	0,503	0,592	0,327
	SMOTE (%600)	0,923	0,891	0,995	0,345	0,497	0,586	0,321
	RA (1:2)	0,876	0,457	0,906	0,634	0,531	0,758	0,559
	RA (1:1,6)	0,856	0,410	0,876	0,690	0,514	0,778*	0,594*
	SMOTE (%140) +RA (1:2)	0,921	0,759	0,983	0,423	0,543*	0,645	0,392
RF	Yok	0,891	1,000*	1,000*	0,014	0,028	0,119	0,013
	SMOTE (%100)	0,893	0,609	0,992	0,099	0,170	0,313	0,089
	SMOTE (%200)	0,893	0,558	0,983	0,169	0,259	0,408	0,153
	SMOTE (%300)	0,889	0,494	0,965	0,275	0,353	0,515	0,247
	SMOTE (%600)	0,860	0,388	0,909	0,465	0,423	0,650	0,404
	RA (1:2)	0,857	0,364	0,914	0,394	0,378	0,600	0,342
	RA (1:1,6)	0,838	0,328	0,887	0,444	0,377	0,627	0,376
	SMOTE (%140) +RA (1:2)	0,886	0,478	0,959	0,303	0,371	0,539	0,271
LightGBM	Yok	0,916	0,815	0,991	0,310	0,449	0,554	0,286
	SMOTE (%100)	0,910	0,680	0,979	0,359	0,470	0,593	0,330
	SMOTE (%200)	0,911	0,659	0,974	0,408	0,504	0,631	0,375
	SMOTE (%300)	0,902	0,580	0,963	0,408	0,479	0,627	0,372
	SMOTE (%600)	0,894	0,524	0,947	0,465	0,493	0,664	0,419
	RA (1:2)	0,853	0,395	0,882	0,620	0,482	0,739	0,532
	RA (1:1,6)	0,837	0,372	0,856	0,683	0,481	0,765	0,575
	SMOTE (%140) +RA (1:2)	0,893	0,517	0,939	0,528	0,523	0,704	0,475
XGBoost	Yok	0,377	0,123	0,329	0,760	0,212	0,500	0,261
	SMOTE (%100)	0,110	0,110	0,000	1,000*	0,199	0,000	0,000
	SMOTE (%200)	0,110	0,110	0,000	1,000*	0,199	0,000	0,000
	SMOTE (%300)	0,110	0,110	0,000	1,000*	0,199	0,000	0,000
	SMOTE (%600)	0,110	0,110	0,000	1,000*	0,199	0,000	0,000
	RA (1:2)	0,193	0,118	0,095	0,978	0,211	0,305	0,101
	RA (1:1,6)	0,384	0,146	0,314	0,943	0,253	0,544	0,315
	SMOTE (%140) +RA (1:2)	0,110	0,110	0,000	1,000*	0,199	0,000	0,000
ANN	Yok	0,914	0,670	0,973	0,444	0,534	0,657	0,409
	SMOTE (%100)	0,916	0,681	0,974	0,451	0,542	0,662	0,416
	SMOTE (%200)	0,901	0,568	0,958	0,444	0,498	0,652	0,403
	SMOTE (%300)	0,899	0,552	0,954	0,451	0,496	0,656	0,408
	SMOTE (%600)	0,899	0,548	0,951	0,479	0,511	0,675	0,434
	RA (1:2)	0,850	0,397	0,872	0,676	0,500	0,768	0,578
	RA (1:1,6)	0,814	0,339	0,826	0,718	0,460	0,770	0,587
	SMOTE (%140) +RA (1:2)	0,899	0,542	0,943	0,542	0,542	0,715	0,491

Destek vektör makinesinde sadece rastgele alt örnekleme tekniği uygulandığında geometrik ortalama ve IBA değerleri belirgin şekilde yükselmiştir. ANN modelinde SMOTE tekniği belirgin bir iyileşme sağlamazken, rastgele alt örnekleme tekniğinin uygulanması geometrik ortalama ve IBA değerlerini artırmıştır.

Onama ve bozma kararlarını doğru şekilde sınıflandırmak amacıyla öncelikle geometrik ortalama, F-ölçütü ve IBA metriklerinde en yüksek değerlere sahip modeller incelenmiştir. Uzman görüşü doğrultusunda, her iki sınıf etiketinde en az %70 başarı hedeflenmiştir. Bu nedenle, duyarlılık ve özgüllük değerleri en az 0,70 olan modeller değerlendirilmiştir. Rastgele alt örnekleme (1:1,6) ile oluşturulan ANN modeli, duyarlılık değeri 0,826 ve özgüllük değeri 0,718 ile bu kriterleri karşılamıştır. Ayrıca, bu modelin doğruluk değeri 0,814 olup, her iki sınıfı da başarıyla tespit etmektedir.

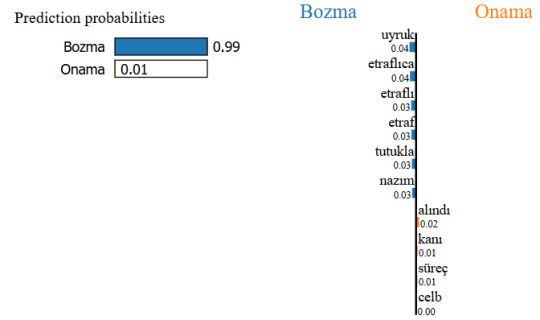
Seçilen modelin sınıflandırma kararları LIME tekniği kullanılarak açıklanmıştır. Doğru sınıflandırılmış birer onama ve bozma örneği rastgele seçilerek LIME çıktıları üretilmiştir. Şekil 1’de, doğru sınıflandırılmış bozma örneğinin LIME çıktısı gösterilmiştir. Buna göre, model tahminine en yüksek katkıda bulunan öznitelikler ve etkileri gösterilmektedir. ANN modeli bu örneği %97 olasılıkla bozma, %3 olasılıkla onama olarak tahmin etmiştir. Bozma sınıfı tahmininde etkili öznitelikler arasında “münhasıran”, “bilge” ve “cumhuriyet” kelimeleri %9 ila %13 oranında etkili olmuştur. Onama sınıfına katkıda bulunan özniteliklerin etkisi ise %0,03 ila %0,05 arasında değişmiştir.



Şekil 1. Bozma sınıfından doğru sınıflandırılmış rastgele bir örnek için LIME açıklaması

Şekil 2’de ise doğru sınıflandırılmış onama örneğinin LIME çıktısı gösterilmiştir. ANN modeli bu örneği %99 olasılıkla bozma, %1 olasılıkla onama olarak tahmin etmiştir. Onama tahmininde etkili öznitelikler arasında “alındı”, “kanı”, “celb” ve

“süreç” gibi kelimeler %2’ye kadar etki düzeyine sahiptir.



Şekil 2. Onama sınıfından doğru sınıflandırılmış rastgele bir örnek için LIME açıklaması

Hukuk alanında yapay zekâ tabanlı karar destek sistemlerinin geliştirilmesi ve bu sistemlerin açıklanabilirliğinin sağlanması hukuki süreçlerin iyileştirilmesinde katkı sağlayabilir. Böylece süreçler hızlandırılabilir ve süreçlerin şeffaflığı sağlanabilir.

5 Sonuç

Bu çalışmada, Yüksek Mahkeme kararlarını sınıflandırmak için yapay zekâ tabanlı bir karar destek sistemi önerilmiştir. Karar metinleri üzerinde veri ön işleme yapılmış ve dengesizlik sorununu gidermek için SMOTE ve rastgele alt örnekleme yöntemleri kullanılmıştır. En iyi performansı, rastgele alt örnekleme uygulanmış ANN modeli göstermiştir. LIME yöntemiyle doğru sınıflandırılmış örnekler için özniteliklerin model tahminine katkısı analiz edilmiştir.

Gelecek çalışmalarda diğer hukuk dallarında yapay zekâ tabanlı karar destek sistemleri geliştirilebilir. Ayrıca farklı veri ön işleme ve boyut azaltma yöntemlerinin uygulanması model performansını iyileştirebilir. Model kararlarının açıklanabilirliğini sağlamak için farklı açıklanabilirlik yöntemleri incelenebilir.

Teşekkür

Bu çalışma, TÜBİTAK tarafından desteklenmiştir (Proje No: 22AG001). Yazarlar ayrıca çalışmaya katkılarından dolayı Dr. Öğretim Üyesi Emine Gökçe KARABEL’e teşekkür ederler.

Kaynaklar

- [1] Arslan A, Talaş U, Çubukçu B. “LAWNAV: yapay zekâ destekli hukuk danışmanı yönlendirme uygulaması”. 10. International European Conference on Interdisciplinary Scientific Research, pp. 177–185, Zürich, Switzerland, 2024.
- [2] Turan T, Kucuksille EU, Kemaloğlu Alagöz N. “Prediction of Turkish Constitutional Court

- decisions with explainable artificial intelligence". *Bilge International Journal of Science and Technology Research*, 7(1), 2023.
- [3] Görentaş MB, Uçkan T, Bayram Arlı N. "Uyuşmazlık Mahkemesi Kararlarının Makine Öğrenmesi Yöntemleri ile Sınıflandırılması". *Yüzüncü Yıl Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, 28(3), 947-961, 2023.
- [4] Görentaş MB, Uçkan T. "Makine Öğrenmesi Yöntemleri Kullanılarak Mahkeme Kararlarının Kümelmesi". *Computer Science*, 8(2), 148-158, 2023.
- [5] Karabel EG, Aydemir D. "Medeni Usul Hukukunda Yargılamanın Hızlandırılması ve Adalet Erişim Hakkı Bakımından Çevrimiçi Yargılama ve Yapay Zekanın Kullanımı". *Marmara Üniversitesi Hukuk Fakültesi Hukuk Araştırmaları Dergisi*, 29(1), 530-573, 2023.
- [6] Mumcuoğlu E, Öztürk CE, Ozaktas HM, Koç A. "Natural language processing in law: Prediction of outcomes in the higher courts of Turkey". *Information Processing & Management*, 58(5), 102684, 2021.
- [7] Almuzaini HA, Azmi AM. "TaSbeeb: A judicial decision support system based on deep learning framework". *Journal of King Saud University Computer and Information Sciences*, 35(8), 101695, 2023.
- [8] Anh DH, Do DT, Tran V, Minh NL. "The impact of large language modeling on natural language processing in legal texts: A comprehensive survey". *15th International Conference on Knowledge and Systems Engineering (KSE)*, Hanoi, Vietnam, 1-7, 2023.
- [9] de Arriba-Pérez F, García-Méndez S, González-Castaño FJ, González-González J. "Explainable Machine Learning Multi-label Classification of Spanish Legal Judgements". *Journal of King Saud University - Computer and Information Sciences*, 34(10B), 10180-10192, 2022.
- [10] Aletras N, Tsarapatsanis D, Preotiuc-Pietro D, Lamos V. "Predicting judicial decisions of the European Court of Human Rights: A natural language processing perspective". *PeerJ Computer Science*, 2, e93, 2016.
- [11] Budaya IGBA, Suniantara IKP. "Comparison of sentiment analysis algorithms with SMOTE oversampling and TF-IDF implementation on Google Reviews for public health centers". *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 4(3), 1077-1086, 2024.
- [12] Khan MH, Towhid NA, Faruk KO, Mahmud JA, Mridha MF. "Strategies for enhancing the performance of news article classification in Bangla: Handling imbalance and interpretation". *Engineering Applications of Artificial Intelligence*, 125, 106688, 2023.
- [13] Rupapara V, Rustam F, Shahzad HF, Mehmood A, Ashraf I, Choi GS. "Impact of SMOTE on imbalanced text features for toxic comments classification using RVVC model". *IEEE Access*, 9, 78621-78633, 2021.
- [14] Kocak S, İç YT, Sert M, Dengiz B. "Ar-Ge projelerinin sınıflandırılması için doğal Türkçe dil işleme tabanlı yöntem". *Gazi Üniversitesi Mühendislik Mimarlık Fakültesi Dergisi*, 38(3), 1375-1388, 2023.
- [15] Toğaçar M, Eşidir KA, Ergen B. "Yapay zekâ tabanlı doğal dil işleme yaklaşımını kullanarak internet ortamında yayınlanmış sahte haberlerin tespiti". *Journal of Intelligent Systems: Theory and Applications*, 5(1), 1-8, 2022.
- [16] Seveli O, Kemalöğlu Alagöz N. "Olağandışı olaylar hakkındaki tweet'lerin gerçek ve gerçek dışı olarak Google BERT modeli ile sınıflandırılması". *Veri Bilimi Dergisi*, 4(1), 31-37, 2021.
- [17] Öztürk A, Durak Ü, Badıllı F. "Twitter verilerinden doğal dil işleme ve makine öğrenmesi ile hastalık tespiti". *KONJES*, 8(4), 839-852, 2020.
- [18] Adalet Bakanlığı. "UYAP Mevzuat ve İçtihat Programı" <https://mevzuat.adalet.gov.tr/> (11.07.2024)
- [19] Kim SW, Gil JM. "Research paper classification systems based on TF-IDF and LDA schemes". *Human-centric Computing and Information Sciences*, 9(1), 30, 2019.
- [20] Chawla NV, Bowyer KW, Hall LO, Kegelmeyer WP. "SMOTE: Synthetic Minority Over-sampling Technique". *Journal of Artificial Intelligence Research*, 16, 321-357, 2002.
- [21] Chawla NV. "Data Mining for Imbalanced Datasets: An Overview". In: Maimon O, Rokach L (eds) *Data Mining and Knowledge Discovery Handbook*. Springer, Boston, MA, 2005.
- [22] Taunk K, De S, Verma S, Swetapadma A. "A brief review of nearest neighbor algorithm for learning and classification". *2019 International Conference on Intelligent Computing and Control Systems (ICCS)*, pp. 1255-1260, Madurai, India, 2019.
- [23] Fürnkranz J. "Decision Tree". In: Sammut C, Webb GI (eds) *Encyclopedia of Machine Learning*. Springer, Boston, MA, 2011.
- [24] Cervantes J, Garcia-Lamont F, Rodríguez-Mazahua L, Lopez A. "A comprehensive survey on support vector machine classification: Applications, challenges and trends". *Neurocomputing*, 408, 189-215, 2020.
- [25] Liu Y, Wang Y, Zhang J. "New machine learning algorithm: Random forest". In: Liu B, Ma M, Chang J (eds) *Information computing and applications. ICICA 2012. Lecture Notes in Computer Science, Vol. 7473*, Springer, Berlin, Heidelberg, 2012.
- [26] Microsoft Corporation "LightGBM Dokümantasyonu" <https://lightgbm.readthedocs.io/> (26.06.2024)
- [27] Ke G, Meng Q, Finley T, Wang T, Chen W, Ma W, Ye Q, Liu T-Y. "LightGBM: A highly efficient gradient

- boosting decision tree". *Advances in Neural Information Processing Systems*, 30, 3146–3154, 2017.
- [28] Chen T, Guestrin C. "XGBoost: A scalable tree boosting system". In: *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 785–794, Association for Computing Machinery, 2016.
- [29] XGBoost Geliştiricileri, "XGBoost Dökümantasyonu" <https://xgboost.readthedocs.io/> (10.07.2024)
- [30] Abraham A. "Artificial Neural Networks". In: Sydenham PH, Thorn R (eds) *Handbook of Measuring System Design*, 2005.
- [31] Yang GR, Wang X-J. "Artificial neural networks for neuroscientists: A primer". *Neuron*, 107(6), 1013–1032, 2020.
- [32] Ribeiro MT, Singh S, Guestrin C. "Why Should I Trust You?" Explaining the Predictions of Any Classifier. In: *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 1135–1144, ACM, 2016.
- [33] García V, Mollineda RA, Sánchez JS. "Index of balanced accuracy: A performance measure for skewed class distributions". In: Araujo H, Mendonça AM, Pinho AJ, Torres MI (eds) *Pattern recognition and image analysis. IbPRIA 2009. Lecture Notes in Computer Science*, Vol. 5524, Springer, Berlin, Heidelberg, 2009.
- [34] He H, Garcia EA. "Learning from imbalanced data". *IEEE Transactions on Knowledge and Data Engineering*, 21(9), 1263–1284, 2009.
- [35] Kotsiantis S, Kanellopoulos D, Pintelas P. "Handling imbalanced datasets: A review". *GESTS International Transactions on Computer Science and Engineering*, 30(1), 25–36, 2006.