

Türkiye'de Emsal Kararların Makine Öğrenmesi Algoritmaları ile Sınıflandırılması

Süleyman Kürşat DEMİR^{1*}, Emrah AYDEMİR², Yasin SÖNMEZ³

¹Gazi University, Vocational School of Technical Sciences, Department of Design, Ankara, Turkey

Makale Bilgisi

Araştırma makalesi
Başvuru: 31/03/2025
Düzeltilme: 04/06/2025
Kabul: 02/07/2025

Anahtar Kelimeler

Yapay Zekâ
Doğal Dil İşleme
Makine Öğrenimi
Emsal Kararlar
Gradient Boosting

Article Info

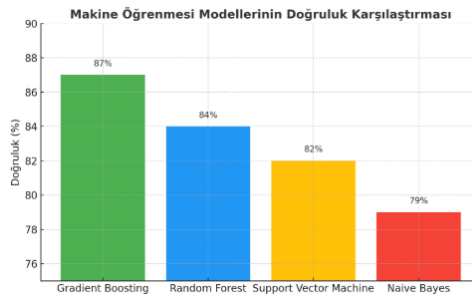
Research article
Received: 31/03/2025
Revision: 04/06/2025
Accepted: 02/07/2025

Keywords

Artificial Intelligence
Natural Language
Processing
Machine Learning
Precedent Decisions
Gradient Boosting

Graphical/Tabular Abstract (Grafik Özet)

Bu çalışma Türkiye'deki emsal kararların makine öğrenmesi ile sınıflandırılmasını ele alıyor. Gradient Boosting %87 doğrulukla en iyi sonucu verirken, hukukta yapay zekanın potansiyelini ortaya koydu. / This study examines classification of precedent decisions in Turkey using machine learning. Gradient Boosting achieved 87% accuracy, demonstrating AI's potential in the legal field



Şekil A: Modellerin karşılaştırılması / Figure A: Comparison of models

Önemli noktalar (Highlights)

- Gradient Boosting algoritması üçlü sınıflandırmada %83, ikili sınıflandırmada %87 doğruluk oranı ile en başarılı sonuçları verdi. / Gradient Boosting achieved the highest accuracy with 83% in multi-class and 87% in binary classification
- Hükümler arasında en sık kullanılan üç sınıf (Reddi, Kabulü, İadesi) üzerine odaklanıldı. / Focus was on the three most frequent judgment classes: Rejection, Acceptance, Return.
- Türk hukuk dilinin karmaşıklığı nedeniyle makine öğrenmesi modellerinin özelleştirilmesi gerekliliği vurgulandı. / The need for customization of machine learning models due to the complexity of Turkish legal language was emphasized.
- Çalışma, hukuk alanında yapay zekâ kullanımının potansiyelini ve emsal karar analizinde makine öğrenmesinin etkinliğini ortaya koymaktadır. / The study highlights AI's potential in law and effectiveness of machine learning in precedent decision analysis.

Amaç (Aim): Türk yargı sistemindeki emsal kararların yüksek doğruluklu makine öğrenmesi modeli ile sınıflandırılmasını geliştirerek, hukuk araştırmaları ve karar alma süreçlerinin etkinliğini artırmak. / To develop a high-accuracy machine learning model for classifying precedent decisions in the Turkish judiciary system, enhancing the efficiency of legal research and decision-making processes.

Özgünlük (Originality): Bu çalışma, UYAP sisteminden elde edilen büyük ölçekli Türk emsal kararları veri setine makine öğrenimi algoritmalarını uygulayarak, hukuki hükümlerin kapsamlı analiz ve sınıflandırılmasını özgün bir şekilde sunmaktadır. / This study uniquely applies machine learning algorithms to a large-scale dataset of Turkish precedent decisions from the UYAP system, providing comprehensive analysis and classification of legal judgments.

Bulgular (Results): Gradient Boosting algoritması, çoklu sınıflandırmada %83, ikili sınıflandırma senaryolarında ise %87'ye varan doğruluk oranlarıyla diğer modelleri geride bırakmıştır. / Gradient Boosting algorithm achieved the highest accuracy with 83% in multi-class classification and up to 87% in binary classification scenarios, outperforming other models.

Conclusion (Sonuç): Makine öğrenmesi modelleri, özellikle Gradient Boosting, hukuki karar alma süreçlerini destekleyip hızlandırma potansiyelini güçlü bir şekilde göstermektedir; ancak transformatör modeller Türk hukuk metinleri için daha fazla uyarılma gerektirmektedir. / Machine learning models, especially Gradient Boosting, demonstrate strong potential to support and accelerate legal decision-making processes, although transformer models require further adaptation for Turkish legal texts.



Türkiye'de Emsal Kararların Makine Öğrenmesi Algoritmaları ile Sınıflandırılması

Süleyman Kürşat DEMİR^{1*} , Emrah AYDEMİR² , Yasin SÖNMEZ³

¹Gazi University, Vocational School of Technical Sciences, Department of Design, Ankara, Turkey

Makale Bilgisi

Araştırma makalesi
Başvuru: 31/03/2025
Düzeltilme: 04/06/2025
Kabul: 02/07/2025

Anahtar Kelimeler

Yapay Zekâ
Doğal Dil İşleme
Makine Öğrenimi
Emsal Kararlar
Gradient Boosting

Öz

Son yıllarda yapay zekâ (YZ), büyük veri ve veri analitiği alanlarındaki hızlı gelişmeler, birçok sektörde iş süreçlerini kolaylaştırarak iş yükünü azaltmıştır. Ancak hukuk alanında, YZ'nin dünya çapında önemli bir dönüşüm yarattığına dair belirgin bir gelişme henüz gözlemlenmemiştir. Bu durum, hukukta yapay zekânın belirli ilerlemeler kaydettiğini, fakat daha büyük potansiyel barındırdığını göstermektedir.

Bu çalışmada Türkiye'deki emsal mahkeme kararları analiz edilmiştir. UYAP sistemi üzerinden indirilen veri setleri metin ve hüküm olarak sınıflandırılmıştır. En sık tekrar eden üç hüküm olan Reddi, Kabulü ve İadesi kararları üzerine odaklanılmıştır. Doğal Dil İşleme (NLP) teknikleri ve çeşitli makine öğrenimi algoritmaları (Gradient Boosting, Yapay Sinir Ağları, Karar Ağaçları vb.) kullanılarak sınıflandırma ve tahmin işlemleri yapılmıştır. Sonuçlar, doğruluk ve F1 skoru gibi metriklerle değerlendirilmiş; Gradient Boosting algoritması %83 doğruluk oranı ile en iyi performansı göstermiştir.

Çalışmanın devamında modelin performansını daha detaylı incelemek için veri seti ikili sınıflandırma görevlerine ayrılmıştır: "Kabulü-Reddi", "Reddi-İadesi" ve "Kabulü-İadesi". Bu senaryolarda da Gradient Boosting %87 doğrulukla en başarılı sonuçları vermiştir. Diğer algoritmalar da benzer başarı göstermiştir.

Bu çalışma, yapay zekânın hukuk alanındaki potansiyelini ortaya koymakla kalmayıp, emsal kararların sistematik analizine ve hukukta veriye dayalı karar alma süreçlerinin gelişimine önemli katkı sağlamaktadır.

Classification of Precedent Decisions in Türkiye with Machine Learning Algorithms

Article Info

Research article
Received: 31/03/2025
Revision: 04/06/2025
Accepted: 02/07/2025

Keywords

Artificial Intelligence
Natural Language
Processing
Machine Learning
Precedent Decisions
Gradient Boosting

Abstract

In recent years, rapid advancements in artificial intelligence (AI), big data, and data analytics have facilitated business processes and reduced workloads in many sectors. However, no significant global transformation driven by AI has yet been observed in the legal field. This indicates that while AI has made some progress in law, it still holds considerable untapped potential.

This study analyzes precedent court decisions in Turkey. The dataset, obtained from the UYAP system, was classified into text and judgment categories. The focus was on the three most frequently occurring judgments: Rejection, Acceptance, and Return. Classification and prediction tasks were performed using Natural Language Processing (NLP) techniques and various machine learning algorithms such as Gradient Boosting, Artificial Neural Networks, and Decision Trees. The results were evaluated using metrics like accuracy and F1 score, with the Gradient Boosting algorithm showing the best performance at an accuracy rate of 83%.

To further examine the model's performance, the dataset was divided into binary classification tasks: "Acceptance-Rejection," "Rejection-Return," and "Acceptance-Return." In these scenarios, Gradient Boosting again achieved the highest accuracy at 87%, with other algorithms also showing similar success.

This study not only reveals the potential of artificial intelligence in the legal domain but also contributes significantly to the systematic analysis of precedent decisions and the advancement of data-driven decision-making processes in law.

1. GİRİŞ (INTRODUCTION)

Türk yargı sistemi, hiyerarşik bir düzen içinde işler ve yargılama süreci ilk derece mahkemelerinde başlayarak istinaf (Bölge Adliye Mahkemeleri) ve temyiz (Yargıtay ve Danıştay) aşamalarıyla devam eder. Anayasa Mahkemesi, yasaların anayasaya uygunluğunu denetlerken bireysel başvuruları da değerlendirir; Sayıştay, kamu kurumlarının mali denetimini gerçekleştirir; Uyuşmazlık Mahkemesi ise yargı organları arasındaki görev ve yetki çatışmalarını çözüme kavuşturur. İç hukuk yollarının tükenmesi durumunda bireyler Avrupa İnsan Hakları Mahkemesi'ne başvurma hakkına sahiptir. [1] Tüm bu süreçlerde verilen kararlar kamuya açık olarak yazılı şekilde sunulur ve avukatlar bu kararları, benzer davalarda müvekkilleri lehine içtihat olarak kullanabilir. Bu nedenle üst mahkeme kararları hem alt mahkemeler hem de taraf vekilleri için bağlayıcı ve yol gösterici niteliktedir. Emsal kararlar, geçmişte benzer uyuşmazlıklarda verilmiş kararlar olup avukatlar tarafından detaylı incelenerek savunmalarda etkili şekilde kullanılmalıdır. Bu bağlamda, Adalet Bakanlığı tarafından geliştirilen Ulusal Yargı Ağı Bilişim Sistemi (UYAP), yargı süreçlerinin dijitalleşmesini sağlamış, adli ve idari yargı birimlerinin otomasyonunu mümkün kılmıştır. UYAP sayesinde emsal karar araştırmaları daha erişilebilir hale gelmiş olsa da, mevcut sistemler genellikle anahtar kelimeye dayalı çalıştığı için, çok sayıda karar arasında en uygun olanı bulmak hâlâ zaman alan bir süreçtir. Dava konuları karmaşıklaştıkça araştırılması gereken içtihat sayısı artmakta, bu da avukatlar açısından yoğun bir mesleki çaba gerektirmektedir. Sonuç olarak, Türk yargı sistemi teknolojik altyapı ve üst mahkeme kararlarının rehberliğiyle adaletin etkinliğini artırmayı, taraf vekillerinin savunma süreçlerini daha verimli ve hızlı hale getirmeyi amaçlamaktadır.

Son yıllarda büyük veri ve yapay zekâ (YZ) teknolojisinin hızla gelişmesi, birçok sektörde geniş bir uygulama yelpazesi sunmaktadır. Eğitim, savunma, sağlık, ticaret ve mühendislik gibi [2], [3] alanlarda çeşitlilik gösteren YZ uygulamaları, hukuk sektöründe de büyük bir potansiyel taşımaktadır. Hukuk alanında, doğal dil işleme (NLP) teknolojilerinin etkisiyle, dava kararı tahmini, belge analizi, hukuki yardım ve sözleşme oluşturma gibi konularda YZ çözümleri geliştirilmiştir. Bu uygulamalar, Hukuki Yargı Tahmini (HYT) adı altında, dava kararlarını hızla ve duygusallıktan uzak bir şekilde sonuçlandırmayı hedeflemektedir. Özellikle Hindistan ve Brezilya gibi ülkelerdeki dava yoğunluğu göz önüne

alındığında, bu tür uygulamalar, hukuk profesyonelleri için büyük faydalar sunmaktadır. [4], [5] Türkiye'de ise Adalet Bakanlığı'nın mahkeme kararlarını dijitalleştirme kararı, YZ ve hukuk entegrasyonunun önünü açmaktadır. Ancak Türk hukuk veri setlerinin gelişimi henüz erken aşamalarda olup, bu alandaki çalışmaların hızlandırılması gerekmektedir. Bu bağlamda, yeni bir Türk hukuk veri seti oluşturulması ve Türk hukukuna yönelik NLP uygulamalarına katkı sağlanması amaçlanmaktadır. Yapay zeka ve makine öğrenimi teknolojilerinin hukuk alanındaki yükselişi, karar alma süreçlerinde insan kaynaklı önyargıları azaltarak adaletin daha hassas bir şekilde tesis edilmesine katkıda bulunma potansiyeli taşımaktadır.[6] Ancak, bu potansiyel aynı zamanda, dilin kültürel bağlamlarından ve çeviri süreçlerinden kaynaklanan anlamsal sorunlar gibi bazı zorlukları da beraberinde getirmektedir. Bu nedenle, yapay zeka sistemlerinin adalet sistemine entegrasyonunda, yüksek kaliteli veri setlerinin kullanılması ve doğal dil işleme tekniklerinin (N-gramlar, TF-IDF gibi) etkin kullanımı büyük önem arz etmektedir. Günümüzde birçok ülke, kamusal bilgilerin dijital erişilebilirliğini artırarak yasal veri analizinin önünü açmakta, bu da yapay zeka uygulamalarının hukuk alanında daha fazla kullanılmasını teşvik etmektedir.[7] Hukuki metinlerin analizi için geliştirilen yapay zeka modelleri, özellikle doğal dil işleme teknikleri ile birlikte, dava sonuçlarını tahmin etme konusunda umut vaat eden sonuçlar ortaya koymaktadır.[8] Ancak, yapay zeka uygulamalarının eğitim verilerindeki olası önyargılar, karar alma süreçlerindeki şeffaflık sorunları ve etik kaygılar, bu teknolojilerin dikkatli ve sorumlu bir şekilde kullanılması gerektiğini göstermektedir. [6] Bu bağlamda, hukuki metinler üzerinde yapılan analizler, hem karar verme süreçlerinde insanlara destek olacak uygulamaların geliştirilmesine hem de yapay zeka tabanlı sanal yargıç gibi çözümlere temel oluşturabilecek niteliktedir. Bu çalışmalar, hukuk sisteminin yapay zeka ve makine öğrenmesini daha etkin ve adil hale gelmesi yönündeki tartışmaları derinleştirmeyi amaçlamaktadır.

1.1. Türkiye'de Yapay Zeka ve Hukuk Alanındaki Gelişmeler (Developments in Artificial Intelligence and Law in Türkiye)

Türkiye'de hukuk alanında yapay zekâ uygulamaları giderek artan bir ilgi görmektedir. Özellikle son yıllarda, hukuki süreçleri daha verimli hale getirmek, bilgiye erişimi kolaylaştırmak ve karar alma süreçlerini desteklemek amacıyla çeşitli yapay zekâ araçları ve yöntemleri kullanılmaya

başlanmıştır. Bu alandaki gelişmeler, hem yargı sistemini hem de hukuki araştırmaları derinden etkileme potansiyeline sahiptir.

Türkiye’de hukuki bilgiye erişim, özellikle Yargıtay, Danıştay ve Anayasa Mahkemesi gibi yüksek yargı organlarının kararlarına erişim açısından önemli adımlar atılmıştır. Bu mahkemelerin kararlarına, özel olarak geliştirilmiş arama motorları aracılığıyla ulaşılabilir. Benzer şekilde, Bölge Adliye Mahkemeleri ve bazı yerel mahkemelerin kararlarına da arama motorları vasıtasıyla erişim imkanı sunulmaktadır. Ayrıca, kanun yollarına bozma kararları gibi olağanüstü kanun yollarına ait bilgilere de yine özel arama motorları sayesinde ulaşılabilir. Bu durum, hukukçuların ve araştırmacıların hukuki bilgiye daha hızlı ve etkin bir şekilde erişmelerini sağlamaktadır. [9]

Yapay zekâ teknolojileri, hukuki metinlerin analizinde de önemli bir rol oynamaktadır. Özellikle doğal dil işleme (NLP) teknikleri, hukuki metinlerden otomatik olarak kişi, yasal düzenleme, mahkeme ve kuruluş ismi gibi önemli bilgileri çıkarabilme kabiliyetine sahiptir. Bu tür çalışmalar, büyük hukuki veri tabanlarından (örneğin, Legalbank) elde edilen metinler üzerinde derin öğrenme yöntemleri kullanılarak başarılı sonuçlar vermiştir. Hukuki metinlerdeki önemli bilgilerin otomatik olarak tespit edilmesi, hem hukuki araştırmaları hızlandırmakta hem de hukukçuların bilgiye erişimini kolaylaştırmaktadır. [13]

Yapay zekânın hukuk alanındaki en dikkat çekici uygulamalarından biri de karar tahminidir. Özellikle dava dosyalarının içeriği kullanılarak, dava sonucunda verilecek hukuki kararların otomatik olarak tahmin edilmesi, hukuki süreçlerin verimliliğini artırma potansiyeline sahiptir. Türkiye’de de mahkeme kararlarının doğal dil işleme teknikleri kullanılarak işlenmesi ve otomatik karar tahminine yönelik çalışmalar az da olsa bulunmaktadır. Bu çalışmalar, gelecekte mahkeme süreçlerinin daha öngörülebilir ve şeffaf hale gelmesine katkıda bulunabilir.

1.2. Literatür (Literature)

Yapay Zeka (YZ) ve Makine Öğrenimi (ML) teknolojilerinin hukuk alanına entegrasyonu, son yıllarda büyük bir ivme kazanmış ve hukuk profesyonellerinin çalışma yöntemlerini dönüştürme potansiyeli taşımaktadır. Bu durum, aynı zamanda yeni araştırma alanlarının ortaya çıkmasına da zemin hazırlamıştır [5], [10]. Literatürde YZ’nin; hukuki araştırma süreçlerinden

belge analizine, dava sonuçlarının tahmininden hukuki risk değerlendirmesine, otomatik belge oluşturmada kişiselleştirilmiş hukuk danışmanlığı hizmetlerine, hukuki doğruluk ve etik kontrollerden otomatik uyumsuzluk çözümüne kadar geniş bir uygulama alanına sahip olduğu görülmektedir [11], [12], [13]. Ayrıca, hukuk güvenliği, veri koruma, suçlu profillemeye, suç öngörüsü ve suç analitiği gibi alanlarda YZ’nin artan kullanımı, bu teknolojilerin hukuk sistemleri üzerindeki etkisini daha da belirgin hale getirmiştir [14].

Ancak, hukuk dilinin karmaşık ve özgün yapısı ile her dile özgü dil özellikleri, YZ ve hukuk çalışmalarında disiplinler arası bir iş birliğini zorunlu kılmaktadır. Hukuki metinlerin doğru anlaşılması ve yorumlanabilmesi için dilbilim, Doğal Dil İşleme (DDİ) ve hukuk alanlarında uzmanlaşmış bilim insanlarının ortak çalışması gereklidir. Bu disiplinler arası yaklaşım, YZ, büyük veri analitiği, makine öğrenimi ve DDİ gibi teknolojilerin hukuk alanında daha etkili ve verimli çözümler sunmasını sağlamaktadır.

1.3. Metin Sınıflandırma ve Hüküm Tahmini Çalışmaları (Text Classification and Verdict Prediction Studies)

Metin sınıflandırma, DDİ’nin temel unsurlarından biridir ve hukuki metinlerin belirli kategorilere ayrılmasında kritik bir rol oynar. Örneğin, bir e-postanın spam olup olmadığının belirlenmesinden, haber makalelerinin kategorize edilmesine veya sosyal medya paylaşımlarının duygu analizine kadar geniş bir uygulama alanı bulunmaktadır. Hukuki metinlerde bu teknik, aile içi şiddet gibi hassas konularda veri kümelerinin sınıflandırılması ve tanımlanması yoluyla hukuk profesyonellerine destek sağlayabilecek sistemlerin geliştirilmesine de olanak tanır.

Hüküm tahmini ise dava tanımlarına dayanarak, hukuki terimler, önceden verilmiş hükümler, para cezaları ve suçlar gibi unsurları kullanarak mahkeme kararlarını tahmin etmeyi amaçlar. Örneğin, Çin’de hukuki davalar üzerinde yapılan çalışmalar, derin öğrenme modellerinin, özellikle çoklu füzyon yaklaşımlarıyla birlikte dava tanımlarını analiz etmede ve hüküm sonuçlarını tahmin etmede umut vadeden sonuçlar verdiğini göstermiştir. Bu çalışmalarda TextCNN, TextRNN, Wide & TextCNN ve TextDenseNet gibi çeşitli modeller karşılaştırılmış ve TextDenseNet’in tahmin doğruluğu açısından öne çıktığı belirlenmiştir [15]. Ayrıca, hukuki metinlerin analizinde TF-IDF ve BiLSTM kombinasyonunun, Word2vec, SVM ve LR gibi yöntemlerle

birleştirilerek yüksek doğruluk oranları sağladığı tespit edilmiştir. Özellikle Suudi Arabistan'daki velayet davalarında bu yaklaşımlar etkili bir karar tahmin sistemi geliştirilmesine olanak tanımıştır [16].

Hibrit sinir ağı modelleri (örneğin, CNN+LSTM), mahkeme kararlarını tahmin etmede önceki çalışmalara göre daha yüksek doğruluk, hassasiyet, geri çağırma ve F1 puanları sağlamıştır. Avrupa İnsan Hakları Mahkemesi (AİHM) kararlarının analizinde Konvolüsyonel Sinir Ağları (CNN), AİHM'nin belirli maddelerinde (örneğin, 3, 5, 6, 8, 13, 2, 10, 11 ve 14) Destek Vektör Makineleri (SVM) modellerine kıyasla %7 oranında daha yüksek bir performans sergileyerek %82'ye varan eğitim doğruluğu sağlamıştır [17].

ABD'de yapılan bir çalışmada ise Yüksek Mahkeme dava sonuçlarının doğal dil işleme (NLP) ve makine öğrenimi teknikleri kullanılarak tahmin edilmesi araştırılmıştır. Bu model, iki hedef sınıf arasında ayrım yapabileceğini göstererek 0.324 F1 puanı ve 0.68 AUC değeri elde etmiştir. Ancak, bu modellerin yüksek mahkeme seviyesinde kullanılabilmesi için daha fazla araştırmaya ihtiyaç olduğu belirtilmiştir [8]. Brezilya mahkemelerinden elde edilen 4.043 dava üzerinde yapılan bir çalışmada ise, çeşitli sınıflandırıcılar ve derin öğrenme modelleri ile %80.2 F1-skor performansı gösteren bir prototip geliştirilmiştir. Bu çalışma, Brezilya mahkemelerinin karar sonuçlarını tahmin etmeye yönelik bilinen ilk yaklaşımlardan biri olarak değerlendirilmektedir [4].

Türkiye'deki hukuki emsal kararlar üzerinde yapılan bir çalışmada ise Uyuşmazlık Mahkemesi'nin olumsuz görev uyuşmazlığı davaları, adli ve idari olmak üzere iki sınıfa ayrılarak tahmin edilmiştir. Bu çalışmada lojistik regresyon, destek vektör makineleri (SVM), karar ağaçları ve rastgele orman algoritmalarının performansı değerlendirilmiş ve destek vektör makinelerinin %87 doğruluk oranıyla öne çıktığı görülmüştür [18].

Son olarak, Türkiye'de yapılan ve bu çalışmaya benzer nitelikteki bir araştırmada, dava

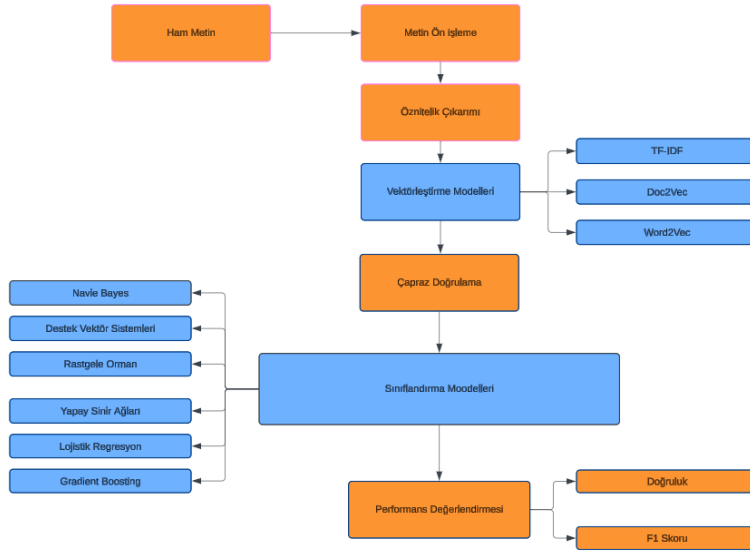
metinlerinden yola çıkarak Türk Anayasa Mahkemesi ve İstinaf Mahkemelerinin kararlarını tahmin etme üzerine çalışılmıştır. Bu tahmin sürecinde çeşitli Makine Öğrenmesi ve Derin Öğrenme algoritmaları kullanılarak yüksek doğruluk ve F1-skoru değerleri elde edilmiştir. [19]

1.4. Bu Çalışmanın Katkısı (Contribution of This Study)

Türkiye'deki emsal kararların mevcut arama sistemlerinin sınırlılıkları ve artan karar hacmi nedeniyle yaşanan zorluklara çözüm sunmayı amaçlayan bu çalışma, makine öğrenimi tekniklerinin Türk yargı sistemindeki emsal kararların nihai hükümlerini (Reddi, Kabulü, İadesi) otomatik olarak sınıflandırma potansiyelini araştırmayı ve yüksek doğruluklu bir model geliştirmeyi hedeflemektedir. Ve elde edilecek makine öğrenimi modelleri, emsal karar veritabanlarının daha akıllı ve verimli bir şekilde organize edilmesine, dolayısıyla hukuk profesyonellerinin emsal karar araştırmalarını kolaylaştırmasına nasıl bir katkı sağlayabilir? Gibi sorulara cevap vermeye çalışmaktadır.

Bu doğrultuda, makine öğrenimi algoritmalarının üç ana hükmü ne derece doğrulukla sınıflandırabileceği ve ikili sınıflandırmanın doğruluk oranını artırıp artırmadığı gibi temel araştırma sorularına yanıt aranmıştır. Büyük bir emsal karar veri kümesi üzerinde yapılan eğitimler sonucunda, üç sınıflı sınıflandırmada %70'in üzerinde, ikili sınıflandırmalarda ise %80'in üzerinde doğruluk oranlarına ulaşılarak önemli bir başarı elde edilmiştir. Elde edilen bu sonuçlar, makine öğreniminin karar metinleri analizindeki etkinliğini kanıtlamakla kalmayıp, emsal kararların daha sistematik değerlendirilmesi ve hukuk profesyonellerinin emsal araştırmalarını kolaylaştırması için güçlü bir temel oluşturarak, adaletin daha hızlı ve etkin tecellisine önemli bir katkı sağlamaktadır.

2. MATERYAL VE METOD (MATERIALS AND METHODS)



Şekil 1. Çalışmanın akış şeması (Flow chart of the study)

Günümüz bilgi çağında, büyük veri setlerinden anlamlı bilgiler çıkarmak kritik öneme sahiptir. Makine öğrenmesi yöntemleriyle yapılan çıkarımsal metin özetleme, bu ihtiyaca cevap veren önemli bir alandır. Çıkarımsal özetleme, orijinal metinden temel bilgileri seçerek kısa ve öz bir özet oluşturmayı hedefler, böylece metnin boyutu küçülürken önemli bilgiler korunmuş olur.[20] Doğal Dil İşleme (DDİ), metin analitiği, metin ön işleme, makine öğrenimi ve istatistik gibi farklı alanları bir araya getirerek, yapılandırılmış veya yapılandırılmamış metin verilerinden anlamlı bilgiler elde etmeyi ve bu verileri makinelerin anlayabileceği bir forma dönüştürmeyi amaçlayan bir disiplindir. Bilgisayarların insan dilini anlamasını, analiz etmesini ve işlemesini sağlayan DDİ, metin verilerini yorumlama, çıkarımlar yapma ve yanıtlar üretme gibi karmaşık dil tabanlı görevleri yerine getirme yeteneği sunar. Metin sınıflandırma, veri madenciliğinin bir alt alanı olarak, belgelerin önceden tanımlanmış kategorilere atanması sürecini ifade eder. Bu süreç, belgedeki verilere bakarak, kitabın türünün belirlenmesinden e-postaların spam olarak işaretlenmesine, yazarın üslubunun tespitinden konuşmaların tanınmasına kadar geniş bir uygulama yelpazesi sunar. Metin sınıflandırma uygulamaları, belge erişimi, e-posta sınıflandırması, haber organizasyonu gibi birçok alanda kullanıldığı gibi; bilgi alma, bilgi çıkarımı, duyarlılık analizi, öneri sistemleri, bilgi yönetimi, metin özetleme gibi çeşitli görevleri de içerir.[4] Metin sınıflandırma sürecinde, denetimli makine öğrenimi teknikleri kullanılarak belgelerin hangi kategoriye ait olduğunu, o kategoriye özgü

kelimelere veya terimlere bakarak belirlemek amaçlanır. Bu süreç, veri toplama, metin ön işleme, eğitim ve test verilerine ayırma, farklı sınıflandırma teknikleri uygulama ve performans değerlendirmesi gibi aşamalardan oluşur.[21],[22]

Öznitelik çıkarımı, metin verilerinden anlamlı ve temsilci özelliklerin çıkarılma sürecidir. Metinler, genellikle karmaşık ve yüksek boyutlu verilerdir ve doğrudan kullanılamayacakları için, öznitelik çıkarımı adımı, metin verilerini daha düşük boyutlu ve daha anlamlı bir temsile dönüştürerek, makine öğrenimi veya diğer analiz yöntemlerine girdi olarak kullanılabilecek veri özelliklerini elde etmeyi amaçlar.

2.1. Vektörleştirme Yöntemleri (Vectorization Methods)

Word2Vec: Tahmin tabanlı (prediction-based) kelime temsil yöntemi olup, temelinde yapay sinir ağı ile iki farklı model kullanarak kelimelerin eğitilmesi amaçlanıp geliştirilmiştir. [4]

- Word2Vec'in kullandığı iki model CBOW (Continuous Bag of Words) ve Skip-Gram Model'dir.
- Continuous Bag of Words: CBOW modelinde pencere boyutu merkezinde olmayan kelimeler girdi olarak alınıp, merkezinde olan kelimeler çıktı olarak tahmin edilmeye çalışılmaktadır.
- Skip Gram: Skip Gram modelinde pencere boyutu merkezinde olan kelimeler girdi olarak alınıp, mikeside olmayan kelimeler

çıktı olarak tahmin edilmeye çalışılmaktadır.

Doc2Vec: Paragraf Vektörü olarak da adlandırılır ve Doğal Dil İşleme'de belgelerin vektörler olarak gösterilmesini sağlayan bilinen bir tekniktir. Bu teknik, kelimeleri sayısal vektörler olarak gösterme yaklaşımı olan Word2Vec'e bir uzantı olarak tanıtılmıştır. Word2Vec kelime yerleştirmelerini öğrenmek için kullanılırken, Doc2Vec belge yerleştirmelerini öğrenmek için kullanılır Doc2Vec, belgelerin dağıtılmış gösterimini öğrenen sinir ağı tabanlı yaklaşımdır. Her belgeyi yüksek boyutlu bir uzayda sabit uzunlukta bir vektöre eşleyen gözetimsiz bir öğrenme tekniğidir. Vektörler, benzer belgelerin vektör uzayındaki yakın noktalara eşlendiği şekilde öğrenilir. Bu, belgeleri vektör gösterimlerine göre karşılaştırmamızı ve belge sınıflandırması, kümeleme ve benzerlik analizi gibi görevleri gerçekleştirmemizi sağlar [23]

TF-IDF: Bir metindeki veya dizideki kelimelerin metinle ne kadar ilgili olduğunun hesaplanması yöntemidir. TF-IDF Makine Öğrenmesi, Metin Madenciliği, Doğal Dil İşlemi gibi uygulamalarda sıkça kullanılır. Bu yöntem 2 ölçümden oluşur:[23]

Bir belgede kelimenin kaç kez görüldüğü (TF),

Bir dizi belgede (farklı belgelerde) kelimenin ters belge sıklığı (IDF).

Terimler:

t — term (kelime)

d — document (belge)

N —count of corpus (belge sayısı)

corpus —belgeler dizisi

Terim Frekansı: Bir belgede belli bir kelimenin sayısını temsil eder.

$$t f(t, d) = \frac{d \text{ belgesindeki } t \text{ kelimesinin sayısı}}{d \text{ belgesindeki toplam kelime sayısı}}$$

Belge Frekansı: Terim frekansına benzer bir şekilde birkelimenin tüm dizi belgelerdeki anlamını test eder.

$$d f(t) = t \text{ kelimesinin geçtiği belge sayısı}$$

Ters Belge Frekansı: IDF yani Ters Belge Frekansı bir terimin (t) bilgilendiriciliğini (anlamlılığını) ölçen değerdir, Belge Frekansı'nın tersidir.

$$idf(t) = \log_e \frac{\text{toplam belge sayısı}(N) + 1}{t \text{ kelimesinin geçtiği belge sayısı} + 1}$$

2.2. Makine Öğrenme Algoritmaları (Machine Learning Algorithms)

Navie Bayes: Naive Bayes sınıflandırıcıları, olasılıkları bulmak için Bayes Teoremi'ne dayanan sınıflandırma görevleri için kullanılan denetlenen makine öğrenme algoritmalarıdır, Naive Bayes sınıflandırıcısının arkasındaki ana fikir, verilerin özelliklerine göre farklı sınıfların olasılıklarına göre verileri sınıflandırmak için Bayes Teoremi'ni kullanmaktır. [24]

Destek Vektör Makinesi: Destek Vektör Makinesi (SVM), hem doğrusal hem de doğrusal olmayan sınıflandırma, regresyon ve aykırı değer algılama görevleri için yaygın olarak kullanılan güçlü bir makine öğrenme algoritmasıdır. SVM'ler son derece uyarlanabilir ve bu da onları metin sınıflandırması, görüntü sınıflandırması ve anormallik algılama gibi çeşitli uygulamalar için uygun hale getirir. [24]

Rastgele Orman: Rastgele Orman algoritması, Makine Öğrenmesinde tahminler yapmak için güçlü bir ağaç öğrenme tekniğidir ve daha sonra tahmin yapmak için tüm ağaçların oylamasını yapılır. Sınıflandırma ve regresyon görevi için yaygın olarak kullanılırlar. Her ağacı eğitmek için veri kümesinin farklı rastgele parçaları alınır ve ardından sonuçları ortalamasını alarak birleştirilir. Bu yaklaşım tahminlerin doğruluğunu artırmaya yardımcı olur. Random Forest, topluluk öğrenimine dayanır. [25]

Yapay Sinir Ağları: Yapay Sinir Ağı (YSA), beyinden esinlenen bir bilgi işleme paradigmasıdır. YSA'lar, insanlar gibi örneklerle öğrenir. Bir YSA, bir öğrenme süreci aracılığıyla desen tanıma veya veri sınıflandırması gibi belirli bir uygulama için yapılandırılır. Öğrenme büyük ölçüde nöronlar arasında var olan sinaptik bağlantılarda ayarlamalar içerir. Yapay Sinir Ağları (YSA), insan beyninin yapısı ve işlevinden esinlenen bir tür makine öğrenme modelidir. Bilgileri işleyen ve ileten birbirine bağlı "nöron" katmanlarından oluşurlar. [24]

Lojistik Regresyon: Lojistik regresyon, amacın bir örneğin belirli bir sınıfa ait olup olmadığını tahmin etmek olduğu sınıflandırma görevleri için kullanılan bir gözetimli makine öğrenme alg oritmasıdır. Lojistik regresyon, iki veri faktörü arasındaki

ilişkiyi analiz eden bir istatistiksel algoritmadır. [26]

Gradient Boosting: Gradyan artırma, genellikle karar ağaçları olan zayıf modellerin bir topluluğunu ardışık olarak oluşturur. Her yeni model, önceki modelin hatalarını düzeltmeye odaklanarak genel doğruluğu iyileştirir. Gradyan artırma, verilerdeki aykırı değerlere ve gürültüye karşı hassastır ve özellik ölçekleme ve eksik değer ataması gibi ön işleme adımları gerektirir. Her yinelemede, GBM, kayıp fonksiyonunun tahminlerine göre negatif gradyanını hesaplar. Bu gradyan, öğrenme yönünü temsil eder ve yeni ağacı bu artıklara daha iyi uyacak şekilde yönlendirir. Öğrenme oranı, bu gradyan boyunca atılan adım boyutunu kontrol ederek her bir ağacın etkisini etkiler. [27]

2.3. Transformatörler (Transformers)

BERT: Google'ın, belirli bir bağlama dayalı olarak kullanıcıların sorgularındaki dilsel varyasyonları anlamasını sağlayan bir yapay zekadır. Algoritma, kısaca bir metindeki kelimeler arasındaki bağlamsal ilişkiyi anlamak için bir mekanizma olan sözde bir dönüştürücü kullanır.

RoBERTa: BERT'i temel alır ve önemli hiperparametreleri değiştirir, sonraki cümle ön eğitim hedefini kaldırır ve çok daha büyük mini gruplar ve öğrenme oranları ile eğitim verir. [28]

Yukarıda bahsedilen makine öğrenme algoritmaları, vektörleştirme yöntemleri ve transformatör tabanlı sistemler; Visual Studio Code ortamında, Python programlama dili kullanılarak bir dizüstü bilgisayar üzerinde uygulandı. Bu süreçte, ham metin verisi ile başlatılan işlem adımları sırasıyla metin ön işleme (küçük harfe çevirme, noktalama temizliği, kök/gövde bulma vb.), özellik çıkarımı ve

vektörleştirme (TF-IDF, Word2Vec, Doc2Vec) yöntemleriyle gerçekleştirilmiştir. Sayısallaştırılmış metinler, çapraz doğrulama yöntemi ile farklı veri setlerinde test edilerek; Naive Bayes, Destek Vektör Makineleri, Rastgele Orman, Yapay Sinir Ağları, Lojistik Regresyon ve Gradient Boosting gibi sınıflandırma algoritmaları kullanılarak eğitilmiştir. Son olarak, her bir modelin başarımı doğruluk ve F1 skoru gibi performans metrikleri ile ölçülmüş ve karşılaştırmalı değerlendirmeler bu kriterler temelinde yapılmıştır ve elde edilen sonuçlar kayıt altına alınarak yöntem aşaması tamamlanmıştır.

3. RESULTS (BULGULAR)

Bu çalışmada, 27 Ekim 2024 tarihinde UYAP Emsal Karar veri tabanından web scraping yöntemiyle 530.000 adet karar indirilmiştir. İndirilen bu veriler; Daire, Esas No, Karar No, Dava Konusu, Hüküm ve Metin sütunlarını içeren bir Pandas DataFrame'ine dönüştürülmüştür. -Tablo 1 de gösterilmiştir- Veri temizleme sürecinde boş verilerin kaldırılması sonucunda toplam karar sayısı 280.000'e düşmüştür. Bu düşüşün nedenleri arasında, dava konusu sütununda kararların %5'inde veri eksikliği bulunması (bazı mahkeme kararlarında dava konusunun belirtilmemesi), hüküm sütununda kararların %10'unda eksiklik olması (bazı kararların hüküm bölümüne sahip olmaması), hükmün kesinleşmemiş olması ve her sütun içinde oluşan bazı boş veriler yer almaktadır.

Tablo 1. DataFrame'e dönüştürülmüş veriler (Data converted to DataFrame)

Daire	Esas No	Karar No	Dava Konusu	Hüküm	Metin
Kayseri Bölge Adli..	2020/494	2021/399	İtirazın İptali	Reddine	Mahkememizde..
Sakarya Bölge Adli..	2019/205	2020/700	Ticari Şirket	Kabulüne	Mahkememizde..
İstanbul Bölge Adli..	2022/424	2023/555	Menfi Tespit	İadesine	Mahkememizde..
Ankara Bölge Adli..	2023/456	2019/123	Diğer	Reddine	Mahkememizde..

Hüküm sütununda oluşan temel hüküm sınıfları ve frekansları aşağıdaki gibidir:

Tablo 2. Hüküm sınıfları, en yüksek frekansa sahip temel hükümler dikkate alınarak oluşturulmuştur. Detay hükümlerin sınıflandırılması DİĞER sınıfına topluca atanılmıştır (The provision classes were created by taking into account the main provisions with the highest frequency. The classification of the detailed provisions was assigned collectively to the OTHER class.)

Hüküm Sınıfı	Yüzde (%)	Adet (Sayı)
Reddine	30%	99.000
Kabulüne	24%	79.000
İadesine	15%	50.000
İptaline	0,01%	2.150
Gönderilmesine	1,50%	5.400
Görevsizliğine	0,01%	1.800
Kaldırılmasına	0,001%	1.100
Karar Verilmesine Yer Olmadığına	1,20%	4.200
Açılmamış Sayılmasına	0,30%	12.500
Geri Çevrilmesine	0,001	428
Diğer	5%	25.000

Çalışmada elde edilen hukuk kararlarını sınıflandırmaya Çalışmanın bu aşamasında, veri seti makine öğrenmesi modelleri için uygun hale getirmek amacıyla bir dizi aşamalar gerçekleştirildi. Öncelikle, veri seti yer alan her bir sınıf için rastgele 2000 adet veri seçimi yapılarak, dengeli bir analiz zemini oluşturma hedeflendi. Bu seçilen veriler üzerinde, metinleri sayısal verilere dönüştürmek ve makine öğrenmesi algoritmalarının anlayabileceği bir format sunmak amacıyla çeşitli vektörleştirme teknikleri uygulandı. Bu bağlamda, TF-IDF (Term Frequency-Inverse Document Frequency) tekniği ile kelime önemlerini belirlerken, Doc2Vec (Belge Temsili) ve Word2Vec (Kelime Temsili) modelleri ile de metinlerin anlamsal bütünlüğünü koruyarak vektör uzayında temsil edildi. Veri vektörleştirme aşamasının ardından, bu dönüştürülmüş veriler üzerinde altı farklı makine öğrenmesi algoritmasını deneyerek sınıflandırma performansı değerlendirmesi yapıldı. Kullanılan algoritmalar arasında Naive Bayes, Support Vector Machines (SVC), Random Forest, Neural Network, Logistic Regression ve Gradient Boosting yer almaktadır. Bu algoritmalar aracılığıyla hukuk kararlarının sınıflandırılmasında en başarılı sonuçları elde etme amaçlanıldı.

yönelik izlenilen veri işleme adımlarını ve kullanılan yöntem ve teknikleri detaylandırılmıştır. İlk olarak, "REDDİ, KABULÜ ve İADESİ" sütunlarından oluşturulan bir CSV dosyasını,

"Metin" ve "Sınıf" olmak üzere iki sütun içeren bir DataFrame yapısına dönüştürüldü. -Tablo 3'te gösterilmiştir- Ardından, veri ön işleme aşamasına geçerek metinleri analizimize uygun hale getirildi. Bu kapsamda, metinlerde bulunan sayıları ve noktalama işaretlerini temizlendi. Kelime analizinde anlamlı sonuçlar elde edebilmek için, frekans değeri 5'in altında olan kelimeleri çalışmamızdan çıkarıldı. Ayrıca, metinlerin anlamını doğrudan etkilemeyen "stopword" olarak adlandırılan sık kullanılan kelimeler de filtrelendi. Metin uzunluğunu standardize etmek amacıyla, 1000 karakterden kısa olan metinleri veri setimizden çıkarıldı. Son olarak, sınıflandırma sürecini kolaylaştırmak için, her bir kararı ait olduğu sınıfa göre yeniden etiketlendi: 'REDDİ' kararlarını 1, 'KABULÜ' kararlarını 2 ve 'İADESİ' kararlarını 3 olacak şekilde. Bu ön işlemler, verileri sınıflandırma algoritmaları için hazır hale getirmeye olanak tanımıştır

Tablo 3. Label Encoding edilmiş DataFrame
(Label Encoded DataFrame)

Filterelenmiş Metin	Sınıf
dava karar mahkememizde görülen taz...	3
dava karar mahkememizde görülen me...	2
dava karar mahkememizde görülen itira...	1
dava karar mahkememizde görülen ala...	2
dava karar mahkememizde görülen taz...	1

Çalışmanın bu aşamasında, veri seti makine öğrenmesi modelleri için uygun hale getirmek amacıyla bir dizi aşamalar gerçekleştirildi. Öncelikle, veri seti yer alan her bir sınıf için rastgele 2000 adet veri seçimi yapılarak, dengeli bir analiz zemini oluşturma hedeflendi. Bu seçilen veriler üzerinde, metinleri sayısal verilere dönüştürmek ve makine öğrenmesi algoritmalarının anlayabileceği bir format sunmak amacıyla çeşitli vektörleştirme teknikleri uygulandı. Bu bağlamda, TF-IDF (Term Frequency-Inverse Document Frequency) tekniği ile kelime önemlerini belirlerken, Doc2Vec (Belge Temsili) ve Word2Vec (Kelime Temsili) modelleri ile de metinlerin anlamsal bütünlüğünü koruyarak vektör uzayında temsil edildi. Veri vektörleştirme aşamasının ardından, bu dönüştürülmüş veriler üzerinde altı farklı makine öğrenmesi algoritmasını deneyerek sınıflandırma performansı değerlendirmesi yapıldı. Kullanılan algoritmalar arasında Naive Bayes, Support Vector Machines (SVC), Random Forest, Neural Network, Logistic Regression ve Gradient Boosting yer almaktadır. Bu algoritmalar aracılığıyla hukuk kararlarının sınıflandırılmasında en başarılı sonuçları elde etme amaçlanıldı.

Çalışmanın derinlemesine analiz aşamasında, ilk değerlendirme sonucunda en yüksek performansı sergileyen Gradient Boosting algoritmasının yanı sıra, Random Forest ve Neural Network algoritmalarını da ileri analizler için seçildi. Bu üç güçlü algoritmanın potansiyelini daha da artırmak amacıyla, her sınıftan 5000 rastgele örneklem

Tablo 4. Her sınıf için 2000 tane verinin kullanıldığı doğruluk değerleri (Accuracy values using 2000 data for each class)

Algoritma	TF-IDF Doğruluk Değeri	Word2Vec Doğruluk Değeri	Doc2Vec Doğruluk Değeri
Navie Bayes	0.59	0.49	0.42
Destek Vektör Sistemleri	0.71	0.54	0.46
Rastgele Orman	0.70	0.57	0.50
Yapay Sinir Ağları	0.70	0.59	0.48
Lojistik Regresyon	0.71	0.59	0.48
Gradient Boosting	0.73	-	-

üzerinde dönüştürücü (transformatör) mimarilere sahip, günümüz doğal dil işleme alanında iyi işler çıkartan BERT ve RoBERTa modelleri uygulandı. Bu modeller sayesinde, hukuk metinlerinin karmaşık dil yapılarını daha iyi anlamayı ve daha yüksek sınıflandırma başarısı elde etme hedeflenildi, aşağıda-Tablo 5- modellerin f1-Skorları verilmiştir.

Tablo 5. Bu aşamadaki düşük performans, transformatör modellerinin bu veri seti için uygun olmadığını göstermiştir (The poor performance at this stage indicated that the transformer models were not suitable for this data set)

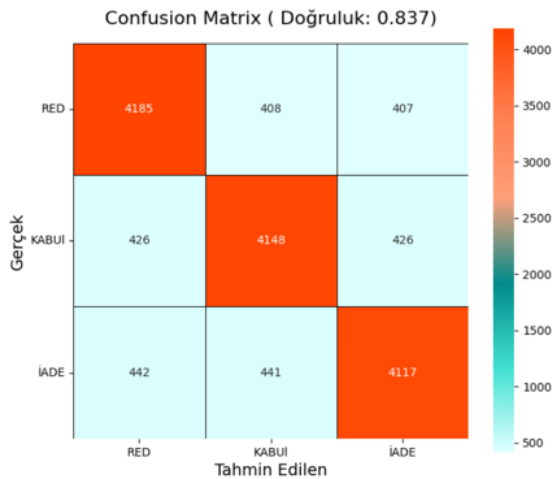
Algoritma	F1-Skoru
Rastgele Orman	0.16
Yapay Sinir Ağları	0.16
Gradient Boosting	0.37

Çalışmada transformatör modelleri beklenen başarıyı gösterememiştir. Bu başarısızlığının nedenleri Türk davalarının karmaşıklığı, uzunluğu ve modellerin Türkçe dilinde yeterli düzeyde eğitilememesi düşünülebilir. Beklenen başarıyı elde

edilememesinin ardından, her sınıftan 5000 tane rastgele örneklem üzerinde seçilen üç algoritmayı (Gradient Boosting, Random Forest ve Neural Network) daha da geliştirmek için bir dizi performans iyileştirici yöntemler uygulandı. Öncelikle, Türkçe dilinin özelliklerine uygun olarak kelime köklerine inme işlemini gerçekleştirmek için TurkishStemmer'ı kullanıldı. Ardından, metin verilerini sayısal formata dönüştürmek ve kelime önemlerini belirlemek için TF-IDF vektörleştirme tekniğini uygulandı. Modelin genelleme yeteneğini artırmak için çapraz doğrulama (cross-validation) yöntemini kullanılarak, modelin farklı veri alt kümelerinde ne kadar iyi performans gösterdiğini ölçüldü. Son olarak, her algoritma için en uygun parametreleri bulmak amacıyla Grid Search algoritmasının kullanılarak, model performansı optimize etme hedeflendi. Bu iyileştirme sürecinin sonunda, aşağıdaki değerleri elde edildi.

Tablo 6. 5000 tane veri ile geliştirilmiş modelin, f1-skoru değerleri (F1-score values of the model developed with 5000 data)

Algoritma	F1-Skoru	Hassasiyet	Geri Çağırma
Rastgele Orman	0.76	0.75	0.77
Yapay Sinir Ağları	0.72	0.73	0.72
Gradient Boosting	0.83	0.82	0.84



Şekil 2. Gradient Boosting algoritmasının (0.83) Confusion matrix'i (Gradient Boosting algoritmasının (0.83) Confusion matrix'i) Yukarıdaki confusion matrix, Gradient Boosting algoritmasının sınıflandırma performansını özetlemektedir. Doğru sınıflandırmalar incelendiğinde, RED sınıfı için 4185, KABUL sınıfı

için 4148 ve İADE sınıfı için 4117 doğru tahmin yapılmıştır. Yanlış sınıflandırmalarda ise RED yerine KABUL (408) veya İADE (407), KABUL yerine RED (426) veya İADE (442), İADE yerine RED (442) veya KABUL (441) tahmin edilmiştir. Bu sonuc kapsamında Gradient Boosting algoritmasının seçilen en iyi hiperparametreleri şunlardır: learning_rate: 0.15, max_depth: 7, n_estimators: 270 Genel olarak model, %83 doğruluk oranı ile iyi bir performans sergilemiş olsa da bazı sınıflar arasında karışıklık yaşandığı görülmektedir.

Tablo 6 deki değerlerden yola çıkarak, modelin performansını farklı sınıflandırma senaryolarında değerlendirmek amacıyla, ikili sınıflandırma görevlerine odaklanıldı. Bu kapsamda, veri seti "KABULÜ-REDDİ", "REDDİ-İADESİ" ve "KABULÜ-İADESİ" olmak üzere üç farklı iki sınıflı veri setine ayrıldı. Böylece, modelin farklı kararlar arasındaki ayrımı ne kadar başarılı bir şekilde yapabildiğini daha detaylı inceleme amaçlandı. Bu üç ikili sınıflandırma senaryosunda, Random Forest, Neural Network ve Gradient Boosting algoritmalarını kullanarak modeller tekrar eğitildi ve performanslarını değerlendirildi. Elde edilen doğruluk değerleri, -Tablo 7 belirtilmiştir-modelin farklı senaryolardaki etkinliğini daha net bir şekilde anlaşılması sağlandı.

Tablo 7. Modellerin üçlü sınıflandırmadaki doğruluk değerleri (Accuracy values of the models in triple classification)

Veri Seti	Rastgele Orman	Yapay Sinir Ağları	Gradient Boosting
Kabulü-Reddi	0.79	0.80	0.86
Reddi-İadesi	0.83	0.82	0.87
Kabulü-İadesi	0.83	0.82	0.87

4. SONUÇLAR (CONCLUSIONS)

Bu çalışmada, Türk yargı sisteminde mahkeme kararlarının çevrimiçi ortamlarda erişilebilir hale gelmesiyle ortaya çıkan bilgi yoğunluğunu yönetmek ve hukuk profesyonellerinin karar alma süreçlerini hızlandırmak amacıyla, UYAP veri tabanından elde edilen hukuki karar metinleri üzerinde doğal dil işleme (DDİ) ve makine öğrenmesi yöntemleri kullanılarak bir emsal karar tahmin modeli geliştirilmiştir. Çalışma, 530.000 adet yapılandırılmamış veri ile başlanmış ve

çalışma kapsamı için bu veri seti 280.000 adet ile sınırlandırılmıştır. Veri setindeki yanlışlık ve tutarsızlık gibi etkileri en aza indirmek ve sonuçların daha geniş kapsamda geçerli olmasını sağlamak için çeşitli temizleme ve ön işleme adımları uygulanmıştır. Elde edilen veri seti ile transformatör modellerinin beklendiği doğrultuda başarılı sonuçlar vermemesi üzerine, Gradient Boosting, Random Forest ve Neural Network algoritmaları denenmiştir. Bu algoritmalar, TurkishStemmer ile kelime köklerine indirme ve TF-IDF ile vektörleştirme teknikleri kullanılarak uygulanmıştır. Algoritmaların performansı, çapraz doğrulama ve Grid Search ile parametre optimizasyonu ile değerlendirilmiştir. Elde edilen sonuçlar, Gradient Boosting algoritmasının %83 doğruluk oranıyla en başarılı performansı sergilediğini göstermiştir. İki sınıflı model değerlendirmesinde ise, "KABULU-REDDİ", "REDDİ-İADESİ" ve "KABULU-İADESİ" veri setleri üzerinde yapılan testlerde, Gradient Boosting algoritması yine en yüksek doğruluk değerlerine ulaşarak (sırasıyla 0.86, 0.87 ve 0.87) diğer algoritmaları geride bırakmıştır.

Bu sonuçlar, makine öğrenmesi modellerinin hukuk alanındaki karar alma süreçlerini destekleme ve hızlandırma potansiyelini açıkça ortaya koymaktadır. Özellikle Gradient Boosting algoritmasının yüksek doğruluk oranları, bu algoritmanın hukuki karar metinlerinin sınıflandırılması ve emsal karar tahmininde etkili bir araç olabileceğini göstermektedir. Ancak, transformatör modellerinin bu çalışmada başarısız olması, hukuki metinlerin kendine özgü yapısı ve karmaşıklığı nedeniyle, bu tür modellerin daha fazla özelleştirme ve optimizasyona ihtiyaç duyabileceğini işaret etmektedir. Ayrıca, farklı veri setleri ve algoritmalar kullanılarak elde edilen sonuçlardaki küçük farklılıklar, modelin başarısının veri kalitesine ve kullanılan tekniklere duyarlı olduğunu göstermektedir.

Çalışmanın sonuçları, hukuk profesyonellerinin içtihat ve emsal karar araştırmalarını kolaylaştırarak, karar verme süreçlerini daha etkin hale getirebilecek bir araç sunmaktadır. Ancak, bu tür modellerin pratik uygulamaları için, modelin şeffaflığı, yorumlanabilirliği ve etik sorunları gibi konuların da göz önünde bulundurulması gerekmektedir. Gelecek çalışmalarda, farklı veri setleri, öznetelik çıkarma teknikleri ve makine öğrenmesi algoritmaları kullanılarak modelin

performansı daha da artırılabilir. Ayrıca, modelin insan uzmanlar tarafından denetlenmesi ve yorumlanması, hukuki karar alma süreçlerinin daha güvenilir ve adil olmasına katkı sağlayabilir. Bu bağlamda, makine öğrenmesi ve yapay zeka teknolojilerinin hukuk alanındaki potansiyelini tam olarak değerlendirebilmek için daha fazla araştırmaya ve işbirliğine ihtiyaç vardır.

ETİK STANDARTLARIN BEYANI (DECLARATION OF ETHICAL STANDARDS)

Bu makalenin yazarı çalışmalarında kullandıkları materyal ve yöntemlerin etik kurul izni ve/veya yasal-özel bir izin gerektirmediğini beyan ederler

The author of this article declares that the materials and methods they use in their work do not require ethical committee approval and/or legal-specific permission.

YAZARLARIN KATKILARI (AUTHORS' CONTRIBUTIONS)

Süleyman Kürşat DEMİR: Yazılım geliştirme, görselleştirme, saha/araştırma çalışmaları, kaynak sağlama, veri düzenleme ve yazımın ilk taslağını hazırlama süreçlerinde görev almıştır.

He participated in software development, visualization, field/research studies, sourcing, data editing, and preparation of the first draft of the manuscript.

Emrah AYDEMİR: Danışmanlık, kavramsal çerçevenin oluşturulması, yöntem geliştirme ve yazımın gözden geçirilip düzenlenmesi aşamalarına katkı sağlamıştır.

He contributed to the consultancy, conceptual framework development, method development, and manuscript review and editing stages.

Yasin Sönmez: Araştırmanın metodolojik tasarımına önemli katkılar sunmuş, makalenin yazımının gözden geçirilmesi ve düzenlenmesi aşamalarına katkıda bulunmuştur.

He made significant contributions to the methodological design of the research and contributed to the review and editing stages of the article.

ÇIKAR ÇATIŞMASI (CONFLICT OF INTEREST)

Bu çalışmada herhangi bir çıkar çatışması yoktur.

There is no conflict of interest in this study.

KAYNAKLAR (REFERENCES)

- [1] B. Kılıç and Y. Öner, "Yargıtay Kararlarının Suç Türlerine Göre Makine Öğrenmesi Yöntemleri İle Sınıflandırılması," *Veri Bilim Dergisi*, vol. 4, no. 3, pp. 61-71, 2021, doi: 10.51203/veri.1011206.
- [2] E. Mohammed, E. Mustapha, and A. Mourad, "Using Machine Learning to Predict Public Prosecution Judges Decisions in Moroccan Courts," *Procedia Computer Science*, vol. 220, pp. 998-1002, 2023, doi: 10.1016/j.procs.2023.08.199.
- [3] J. Morison and T. McInerney, "When should a computer decide? Judicial decision-making in the age of automation, algorithms and generative artificial intelligence," 2024. [Online]. Available: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4723280
- [4] A. Lage-Freitas et al., "Predicting Brazilian Court Decisions," *PeerJ Computer Science*, vol. 8, e904, 2022, doi:10.7717/peerj-cs.904.
- [5] R. Sil, A. Alpana, and A. Roy, "Machine Learning Approach for Automated Legal Text Classification," *International Journal of Computer Information Systems and Industrial Management Applications*, vol. 13, no. 1, pp. 242-251, 2021.
- [6] K. Javed and J. Li, "Artificial intelligence in judicial adjudication: Semantic biasness classification and identification in legal judgement (SBCILJ)," *PLoS One*, vol. 19, no. 5, e0301841, 2024, doi: 10.1371/journal.pone.0301841.
- [7] N. A. K. Rosili, N. H. Zakaria, R. Hassan, S. Kasim, F. Z. Che Rose, and T. Sutikno, "A systematic literature review of machine learning methods in predicting court decisions," *IAES International Journal of Artificial Intelligence*, vol. 9, no. 4, pp. 638-651, 2020, doi: 10.11591/ijai.v9.i4.pp638-651.
- [8] K. Lockard, R. Slater, and B. Sucrese, "Using NLP to model U.S. Supreme Court Cases," *SMU Data Science Review*, vol. 7, no. 1, Article 4, 2023. [Online]. Available: <https://scholar.smu.edu/datasciencereview/vol7/iss1/4>
- [9] D. Küçük and F. Can, "Hukuki metinlerin otomatik işlenmesinde yapay zekâ teknolojilerinin kullanımı," *Bilişim Hukuku Dergisi*, vol. 2024, no. 1, pp. 29-53, 2024, doi: 10.59752/bhd.1451000.
- [10] M. Medvedeva, M. Vols, and M. Wieling, "Using machine learning to predict decisions of the European Court of Human Rights," *Artificial Intelligence and Law*, vol. 28, no. 3, pp. 237-266, 2020, doi: 10.1007/s10506-019-09255-y.
- [11] S. Abbara et al., "ALJP: An Arabic legal judgment prediction in personal status cases using machine learning models," *arXiv preprint arXiv:2309.00238*, 2023.
- [12] D. Alghazzawi et al., "Efficient Prediction of Court Judgments Using an LSTM+CNN Neural Network Model with an Optimal Feature Set," *Mathematics*, vol. 10, no. 5, 683, 2022, doi: 10.3390/math10050683.
- [13] T. Turan, E. U. Küçüksille, and N. K. Alagöz, "Prediction of Turkish Constitutional Court Decisions with Explainable Artificial Intelligence," *Bilge International Journal of Science and Technology Research*, vol. 7, no. 2, pp. 128-141, 2023, doi: 10.30516/bilgesci.1317525.
- [14] R. Sil, A. Alpana, and A. Roy, "Machine Learning Approach for Automated Legal Text Classification," *International Journal of Computer Information Systems and Industrial Management Applications*, vol. 13, no. 1, pp. 242-251, 2021.
- [15] Y. Wen and P. Ti, "A Study of Legal Judgment Prediction Based on Deep Learning Multi-Fusion Models—Data from China," *SAGE Open*, vol. 14, no. 2, pp. 1-10, 2024, doi: 10.1177/21582440241257682.
- [16] M. B. Görentaş, T. Uçkan, and N. B. Arlı, "Classification of Decisions of the Court of Jurisdictional Disputes of Türkiye Using Machine Learning Methods," *Yüzüncü Yıl Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, vol. 28, no. 3, pp. 947-961, 2023, doi: 10.53433/yyufbed.1292275.
- [17] N. Aletras, D. Tsarapatsanis, D. Preotiuc-Pietro, and V. Lampos, "Predicting judicial decisions of the European Court of Human Rights: a natural language processing perspective," *PeerJ Computer Science*, vol. 2, e93, 2016, doi: 10.7717/peerj-cs.93.
- [18] M. Bilgin, "Kelime Vektörü Yöntemlerinin Model Oluşturma Sürelerinin Karşılaştırılması," *Gazi Üniversitesi Bilimsel Teknik Dergisi*, vol. 29, no. 4, pp. 695-705, Apr. 2019, doi: 10.17671/gazibtd.472226.

- [19] E. Mumcuoğlu, C. E. Öztürk, H. M. Ozaktas, and A. Koç, "Natural language processing in law: Prediction of outcomes in the higher courts of Turkey," *Information Processing and Management*, vol. 58, no. 6, 102684, 2021, doi: 10.1016/j.ipm.2021.102684.
- [20] T. Uçkan and K. Karabulut, "The Effectiveness of Machine Learning Algorithms in Extractive Text Summarization: A Comparative Analysis of K-Means, Random Forest, GBM, Logistic Regression, and SVM," *Dicle Üniversitesi Mühendislik Fakültesi Dergisi*, vol. 7, no. 2, pp. 77–91, 2024, doi: 10.57244/dfbd.1538959.
- [21] J. Cui, X. Shen, and S. Wen, "A Survey on Legal Judgment Prediction: Datasets, Metrics, Models and Challenges," *IEEE Access*, vol. 11, pp. 111598-111626, 2023, doi: 10.1109/ACCESS.2023.3317083.
- [22] I. Habernal et al., "Mining legal arguments in court decisions," *Artificial Intelligence and Law*, 2023, doi:10.1007/s10506-023-09936-8.
- [23] Kınık, D., & Güran, A. (2021). TF-IDF and Doc2Vec Based Turkish Text Classification System Performance Enhancement with Sequential Word Group Detection. *European Journal of Science and Technology* (21), 323-332. <https://doi.org/10.31590/ejosat.774144>
- [24] Harman, G. (2021). Diabetes Mellitus Prediction Using Support Vector Machines and Naive Bayes Classification Algorithms. *European Journal of Science and Technology*, (32), 7-13. DOI: 10.31590/ejosat.1041186
- [25] Parmar, A., Katariya, R., and Patel, V. (2018). A Review on Random Forest: An Ensemble Classifier. In *International Conference on Intelligent Data Communication Technologies and Internet of Things*, pp. 758-763. Cham: Springer.
- [26] Çokluk, Ö. (2010). Logistic Regression Analysis: Concept and Application. *Educational Sciences: Theory & Practice*, 10(3), 1357-1407.
- [27] Natekin, A., & Knoll, A. (2013). Gradient Boosting Machines: A Tutorial. *Frontiers in Neurorobotics*, 7, 21.
- [28] Rothman, D. (2021). Transformers for Natural Language Processing: Build Innovative Deep Neural Network Architectures for NLP with Python, PyTorch, TensorFlow, BERT, RoBERTa, and More. Packt Publishing Ltd. Bert