



Amazon veri setinde tüketici duygularını değerlendirmek için derin öğrenme yaklaşımı

A deep learning approach for assessing consumer sentiment in amazon dataset

Nazeeha Sayghn Khalid Khalid¹ , Serkan Savaş^{2*}

¹Fen Bilimleri Enstitüsü, Çankırı Karatekin Üniversitesi, Çankırı, Türkiye.

nazneftci18@gmail.com

²Bilgisayar Mühendisliği Bölümü, Mühendislik ve Doğa Bilimleri Fakültesi, Kırıkkale Üniversitesi, Kırıkkale, Türkiye.

serkansavas@kku.edu.tr

Geliş Tarihi/Received: 29.04.2024

Düzeltilme Tarihi/Revision: 12.02.2025

doi: 10.5505/pajes.2025.45753

Kabul Tarihi/Accepted: 20.03.2025

Araştırma Makalesi/Research Article

Öz

Bu çalışmada, Amazon ürün yorumları kullanılarak tüketici duyarlılık analizinde makine öğrenmesi tekniğinin etkinliği araştırılmıştır. Çalışmanın ana hedefi, metinsel yorumlarla ilgili yıldız puanlamaları arasındaki uyumu değerlendirmek ve bu uyumu tahmin etmek için makine öğrenimi modellerini kullanmaktır. Çalışmada, Destek Vektör Makinesi, Karar Ağacı ve K-En Yakın Komşu gibi makine öğrenmesi algoritmalarının yanı sıra Uzun Kısa Süreli Bellek gibi derin öğrenme algoritması da kullanılmıştır. Bu modellerin performansları karşılaştırılmış ve derin öğrenme modellerinde gizli katman sayısının doğruluk üzerindeki etkisi analiz edilmiştir. Çalışmanın bulguları, Uzun Kısa Süreli Bellek ağlarının tüketici yorumlarındaki doğal dilin karmaşıklıklarını ele almadaki etkinliğini göstermektedir. Uzun Kısa Süreli Bellek modeli, %98'lik doğruluk oranı ile test veri setinde en iyi performansı sergilemiştir. Buna karşın Karar Ağacı modeli, %77.8'lik doğruluk oranı ile en düşük performansı sergilemiştir. Bu sonuçlar, farklı makine öğrenmesi tekniklerinin duyarlılık analizindeki etkinliğine dair önemli fikirler sağlamaktadır. Ayrıca, bu sonuçlar hızla gelişen bu alanda gelecekteki araştırmalar için de önemli bir temel oluşturmaktadır.

Anahtar kelimeler: Tüketici duyarlılık analizi, Derin öğrenme, Uzun kısa süreli bellek, Amazon, Makine öğrenmesi.

Abstract

This study investigates the effectiveness of machine learning techniques in consumer sentiment analysis using Amazon product reviews. The main objective of the study is to use machine learning models to evaluate and predict the correspondence between textual reviews and the corresponding star ratings. In addition to classical machine learning algorithms such as Support Vector Machine, Decision Tree and K-Nearest Neighbor, deep learning algorithm such as Long Short-Term Memory is also used in the study. The performances of these models are compared and the impact of the number of hidden layers on the accuracy of deep learning models is analyzed. The findings of the study demonstrate the effectiveness of Long Short-Term Memory networks in handling the complexities of natural language in consumer reviews. The Long Short-Term Memory model performed best on the test dataset with an accuracy of 98%. In contrast, the Decision Tree model performed the worst with an accuracy of 77.8%. These results provide important insights into the effectiveness of different machine learning techniques in sensitivity analysis. Furthermore, these results provide an important foundation for future research in this rapidly evolving field.

Keywords: Consumer sentiment analysis, Deep learning, Long short-term memory, Amazon, Machine learning.

1 Giriş

E-ticaret ve modern lojistik, günümüzün internet çağının belirgin bir özelliği olarak hızla gelişmektedir. Wu ve diğerlerine göre, giderek daha fazla sayıda insan alışveriş tercihlerini çevrimiçi ürün satın almaya yönlendirmektedir [1]. Online alışveriş tüm yaş grupları arasında popüler olup, önemli bir kolaylık ve verimlilik sunmaktadır. Ancak, bu durum sanal bir ortamda ürünleri fiziksel olarak görememe ve ürün açıklamaları ile gerçek ürünler arasındaki tutarsızlıklar gibi sorunlara yol açabilir. Bu bağlamda ürün yorumlarının ürün değerlendirmelerinde temel bir referans olduğu belirtilmektedir [2]. E-ticaret siteleri, tüketicilerin geri bildirimlerini dikkate alarak alışveriş yapmalarını sağlar. İnternet, ürünler ve hizmetler hakkında önemli bir bilgi kaynağıdır ve milyonlarca geri bildirim üretir. Bu geri bildirimler, bir şirket hakkında detaylı görüşler sağlar ve perakendecilerin alıcıların ihtiyaçlarını göz önünde bulundurmasını kolaylaştırır. Bu da bir şirketi anlamak ve tüketicilerin tercihlerini değerlendirmek için kapsamlı bir

inceleme ve analiz süreci gerektirir. Amazon, dünyanın önde gelen çevrimiçi satıcı platformlarından biridir ve detaylı ürün yorumları barındırmaktadır. Bu yorumlar, ürünlerin kalitesi ve özellikleri hakkında değerli bilgiler sağlar. Le ve Mikolov'a göre Amazon'da alışveriş yapan tüketicilerin %88'i profesyonel tavsiyeleri kontrol etmektedir [3]. Bu yorumlar, duyarlılık analizi ve doğal dil işleme teknikleri kullanılarak değerlendirilebilir. Du ve dig. tarafından gerçekleştirilen çalışmaya göre, çevrimiçi alışveriş yapan insanların %91'i bir ürün veya hizmet satın almadan önce yorumları okur [4]. Askalidis ve Malthouse, Amazon platformlarındaki ürün yorumlarının kalitesinin önemli olduğunu ve müşterilerin genellikle ilk birkaç yoruma odaklandığını belirtmişlerdir [5]. Bu durum, müşterilerin yararlı bilgileri belirlemesini zorlaştırmaktadır. Bu nedenle, e-ticaret web sitelerindeki müşteri geri bildirimleri ve ürün yorumları, müşteri satın alma deneyimini geliştirmede ve işletmelerin ürün ve hizmetlerini iyileştirmede kritik bir rol oynamaktadır. İşletmeler, en karlı ürün ve hizmetlere odaklanarak, müşteri memnuniyetini ve sadakatini sağlamayı hedeflemektedir. Ayrıca, müşteri

*Yazışılan yazar/Corresponding author

deneyimini geliştirmeye ve olumlu görüşler teşvik etmeye odaklanırlar. Ek olarak, işletmeler ürün ve hizmetlerini geliştirmeye, iyileştirmeye ve müşteri memnuniyeti ile sadakatini artırmak için etkili stratejiler geliştirmeye odaklanmalıdır.

Bu çalışmada, Amazon ürün yorumları kullanılarak tüketici duyarlılık analizinde makine öğrenmesi ve derin öğrenme tekniklerinin etkinliği araştırılmıştır. Çalışmanın ana hedefi, metinsel yorumlarla ilgili yıldız puanlamaları arasındaki uyumu değerlendirmek ve bu uyumu tahmin etmek için makine öğrenimi ve derin öğrenme modellerini kullanmaktır. Çalışmada, Destek Vektör Makinesi (Support Vector Machine - SVM), Karar Ağacı (Decision Tree-DT) ve k-En Yakın Komşu (k-Nearest Neighbor-kNN) gibi makine öğrenmesi algoritmalarının yanı sıra Uzun Kısa Süreli Bellek (Long Short-Term Memory-LSTM) ağları gibi derin öğrenme algoritması da kullanılmıştır. Bu modellerin performansları karşılaştırmış ve derin öğrenme modellerinde gizli katman sayısının doğruluk üzerindeki etkisi analiz edilmiştir.

Çalışmanın organizasyonu şu şekilde yapılandırılmıştır: Bölüm 2 ilgili çalışmalara odaklanarak alandaki ana literatür ve gelişmelere genel bir bakış sunmaktadır. Bölüm 3, problem çözme için kullanılan metodolojileri açıklayarak, yeni yöntemlerin verimliliği ve etkinliğini vurgulamaktadır. Bölüm 4'te bu yöntemlerin belirlenen soruna uygulanmasının sonuçları sunulmuştur. Son bölümde sorunun kökeni, uygulanan metodolojiler ve bulguların potansiyel etkilerini anlamının önemi vurgulanmıştır.

2 İlgili çalışmalar

Birçok araştırma, duygu analizi alanındaki çeşitli teknikleri ve yöntemleri inceleyerek bu metodolojilerin özelliklerini, sınırlılıklarını ve etkinliklerini analiz etmektedir. Joseph [6], geniş bir Amazon veritabanı ve önceden eğitilmiş bir BERT modeli kullanarak ürün yorumlarını analiz etmek ve sınıflandırmak için bir makine öğrenimi modeli tanıtmıştır. Bu model, geri bildirim analizinde %96.3 doğruluk elde etmiştir. Shrestha ve Nasoz [7] ise Amazon yorumları ve puanları arasındaki uyumu değerlendirmek amacıyla duygu analizini kullanmışlardır. Derin öğrenme teknikleri ve tekrarlayan sinir ağı kullanarak Amazon yorumlarını vektör temsillerine dönüştüren bu çalışmanın modeli, ortalama olarak 0.5861 kesinlik puanı ve 0.4085 hatırlama oranı elde etmiştir. Gope ve diğ. [8] de, Amazon ürün puanları ve metin yorumları üzerine duygu analizi yapmak için bir makine öğrenimi metodolojisi geliştirmiş ve RNN-LSTM yaklaşımıyla %97.52'lik en yüksek doğruluk oranına ulaşmışlardır. Wedjdane ve diğ. tarafından gerçekleştirilen çalışmada, çevrimiçi yorumların akıllı telefon satışları üzerindeki etkisini inceleyen yazarlar, bu yorumları negatif/pozitif, çok negatif/çok pozitif ve nötr kategorilere ayırarak bilinçli tüketici ve üretici kararları için önemini vurgulamıştır [9]. Amazon yorum veri seti üzerinde AlZu'bi ve diğ. [10] tarafından yapılan çalışmada, kullanıcı oylama yüzdelarının pozitif veya negatif sınıflandırma için analiz edilmesinin etkisiz olduğu ve yorum sayısı ile alınan oylar arasında ters bir ilişki olduğu ortaya konulmuştur. Bu da daha etkili istatistiksel analiz ihtiyacını vurgulamıştır. Norinder U. ve Norinder P. tarafından gerçekleştirilen çalışmada, derin öğrenme ile doğal dil işlemenin birleştirilmesinin, on iki Amazon ürün yorumunda geçici duygunun doğru bir şekilde tahmin edilmesini sağladığı ve sınıf dengesizliklerini ele aldığı gösterilmiştir [11]. Paknejad ise çalışmasında, bilgisayar bilimlerindeki duygu analizinin hızlı büyümesine odaklanarak

fikir madenciliği, metin madenciliği ve duygu analizi tekniklerinin belirli ürünler hakkındaki bireylerin tutumları üzerindeki etkisine odaklanmıştır [12]. Hamdallah çalışmasında, Amazon yorumlarındaki müşteri deneyimlerini anlamak için duygu analizi yöntemlerini kullanarak şirketlere değerli fikirler sunmuştur [13]. Du ve diğ. tarafından gerçekleştirilen çalışmada, 142.8 milyon Amazon müşteri yorumu incelenmiş ve her bir yorumun yararlılığı ile önemsizliği belirlenmiştir [4]. Çalışmada özet başlıklar, ürün yorumları ve yararlılık göstergeleri analiz edilmiştir. Meenakshi ve diğ., çevrimiçi ürün yorum analizinin e-ticaretteki önemini vurgulamış ve gerçek hayattaki tüketici deneyimlerine dayanarak belirli ürünler veya hizmetler hakkında detaylı bilgi sağlama yeteneğine dikkat çekmiştir [14].

Duygu analizi, metin veya konuşmadaki duyguları tespit etmek için kullanılan bir tekniktir. Bu teknik, e-ticaret, sosyal medya ve pazarlama gibi çeşitli alanlarda kullanılmaktadır. Bu literatür özeti, duygu analizi alanında yapılan bazı önemli çalışmaları analiz etmektedir. Çalışmalar, duygu analizi için kullanılan çeşitli teknikleri, bu tekniklerin sınırlılıklarını ve sonuçları belirlemedeki etkinliklerini incelemektedir. Literatürdeki ana bulgular şöyle özetlenebilir:

- Duygu analizi için kullanılan yaygın teknikler arasında makine öğrenimi, doğal dil işleme ve duygu sözlükleri yer almaktadır,
- Makine öğrenimi, duygu analizi için en etkili tekniklerden biridir,
- Duygu analizi, çevrimiçi ürün yorumları, sosyal medya gönderileri ve konuşma kayıtları gibi çeşitli veri türlerinde kullanılabilir.

Bu kapsamda, literatürdeki çalışmalar yöntem, veri seti, model performansı ve değerlendirme ölçütleri açısından analiz edildiğinde bu çalışmalardan farklı olarak çalışmamızın sunduğu yenilikler şöyle detaylandırılabilir.

Özellikle bu çalışmada:

- Literatürde genellikle yalnızca belirli bir modelin performansı ele alınırken, çalışmamız Destek Vektör Makineleri, Karar Ağaçları, k-En Yakın Komşu ve Uzun Kısa Süreli Bellek gibi farklı yaklaşımları kapsamlı şekilde karşılaştırmaktadır,
- Literatürde bazı çalışmalar sadece tek bir modelin uygulanmasına odaklanırken, bu çalışmada hiperparametre optimizasyonları gerçekleştirilmiş ve farklı modellerin doğruluk, kesinlik, geri çağırma, F1 skoru ve AUC-ROC gibi çeşitli performans metrikleriyle detaylı analizi yapılmıştır,
- Literatürde genellikle daha küçük ölçekli veri setleri ele alınırken, bu çalışmada büyük ölçekli bir Amazon veri seti kullanılmış, böylece modelin genellenebilirliği artırılmıştır,
- Çalışmamız, LSTM modelinin geleneksel makine öğrenmesi yöntemlerine kıyasla daha yüksek doğruluk oranına ulaştığını göstererek derin öğrenme modellerinin duyarlılık analizi açısından avantajlarını vurgulamaktadır.

Gelecekteki araştırmalar, duygu analizi için yeni teknikler geliştirmeye ve bu tekniklerin sınırlılıklarını ele almaya odaklanmalıdır. Ayrıca, duygu analizinin farklı veri türlerinde

ve görevlerde etkinliğini değerlendirmek için daha fazla araştırmaya ihtiyaç vardır.

3 Metodoloji

Bu bölümde, veri toplama, analiz ve uygulanan ardışık adımlarla birlikte metodoloji açıklanmıştır. Başlangıçta Amazon web sitesinden kullanıcı yorumları ve yıldız puanlamaları çıkarılmıştır. Bu işlem Python programlama dili ile BeautifulSoup kütüphanesinin birlikte kullanılmasıyla (veri kazıma olarak bilinen bir teknik) gerçekleştirilmiştir [15]. Bunun ardından elde edilen veriler Excel formatında (.xlsx) kaydedilmiştir. Bu çalışma üç ana aşamada yürütülmüştür. İlk aşama, veri toplama sürecini ve verilerin birleştirilmesini içerir. İkinci aşamada, web sitelerinden elde edilen metinler, Amazon web sitesine özgü ön işlemeye tabi tutulur. Son adım, elde edilen sonuçların analiz edildiği aşamadır. Çalışmanın akış şeması Şekil 1'de gösterilmiştir.

3.1 Amazon ve Veri Seti Toplama

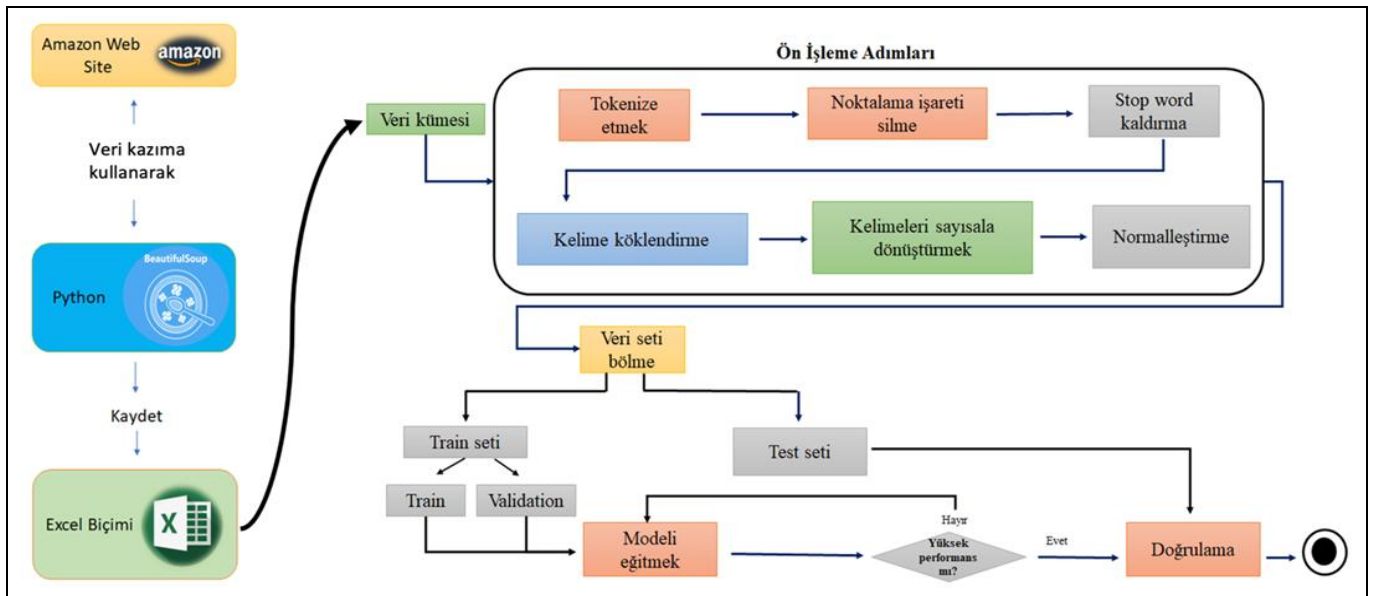
Amazon, kullanıcıların e-ticaret aktivitelerine dayanarak favori ürünlerine kolayca erişmelerini ve analiz etmelerini sağlayan bir platformdur. Platform, ürünlerin ve hizmetlerin popülaritesine ve müşteri memnuniyetine dayalı bir sıralama sistemini kullanır. Sistem, ürünleri puanlamalarına göre kategorize eder ve en yüksek ve en düşük puanlamaları sırasıyla en üst ve ilk beş olarak belirler. Platform ayrıca ürünlerin kalitesini analiz eder, en popüler ürünleri belirler ve buna göre kategorize eder. Platform, ürünlerin popülaritesinin arkasındaki nedenleri de analiz eder (örneğin hizmetin kalitesi ve ürünün kalitesi gibi). Bu analiz, en popüler ürünleri belirlemeye ve en iyi satıcıları tespit etmeye de yardımcı olur. Bu çalışmada, farklı kullanıcılardan çeşitli sayfalarda 15,000'den fazla yorum elde edilmiştir. Veri seti; Apple iPhone 11, OnePlus Nord 5G, Samsung Galaxy M31, Apple iPad 10.2 inç Wi-Fi 32GB, Apple Watch Series 5 GPS 40mm ve Sony Xperia dahil olmak üzere on farklı cihaz için yaklaşık 15,310 içeriği kapsamaktadır. Elde edilen veriler tek bir Excel dosyasında birleştirilmiştir. Bu yorumlarda bulunan yıldız puanlamaları ve açıklamalar daha derin bir analiz için çıkarılmış ve ardından bir CSV dosyasına aktarılmıştır.

Çalışmada, veri kazıma sürecini gerçekleştirmek için kullanılan araçlar arasında öne çıkan Python'da uygulanan BeautifulSoup kütüphanesidir. Amazon'un ürün puanlama kazıyıcı tablosundan türetilen bazı veri alanları şunları içermektedir: Ürün adı, İnceleme başlığı, İnceleme içeriği/Metin, Puanlama, İnceleme yayınlama tarihi, Yazar adı ve URL.

Şekil 1'de görüldüğü gibi toplanan veri seti yapısında, sürecin ilk aşaması verileri içe aktarma, inceleme ve düzenlemedir. Bu işlemin bir parçası olarak, boş veya puanlaması olmayan satırların ortadan kaldırılması da gerekmektedir. Veri toplama yöntemindeki dengesizlik, toplanan puanlamalar ile gerçek tüketici yorumları arasındaki bir farklılığı gösteren dikkate değer bir gözlemdir. 1'den 5'e kadar değişen puanlamalar, müşteri memnuniyetinin objektif bir ölçüsü olarak yorumlanır. Burada 1 (bir) en düşük memnuniyet seviyesini ve 5 (beş) en yüksek memnuniyeti ifade ederken, 2 (iki), 3 (üç) ve 4 (dört) ise artan memnuniyeti belirtir. Analiz için, beş farklı grup oluşturan bir kümeleme yöntemi kullanılmış ve her grup, en az memnun olandan (1) en çok memnun olana (5) kadar farklı memnuniyet derecelerini temsil edilmiştir. Bu kümeleme tekniği, müşteri yorumlarındaki ürünler hakkındaki kelime benzerliklerinin analizine dayanmaktadır.

3.2 Ön işleme

Bu aşamada yorumlardan yeni satır karakterleri (\n) kaldırılmıştır. "3.0 out of 5" gibi kelimeler ve rakamlar içeren, emoji vb. karakterler yorum verilerinden çıkarılmıştır. Bu anlayışa uygun olarak, puanlamanın tam sayısı olarak ifade edilmesini sağlamak için değişiklikler yapılmıştır. Birden fazla .csv dosyasından müşteri inceleme verilerini okuduktan ve birleştirdikten sonra, özelliklerin veri türleri Pandas kütüphanesinin info() fonksiyonu kullanılarak incelenmiştir. Ayrıca, puanlama verilerinin veri türü gözden geçirilerek bu sütunda da düzenlemeler yapılmıştır. Bazı puanlama değerlerinin kayan noktalı biçimde (örn., '3.0') bulunduğu tespit edilmiştir. Puanlama özelliği tam sayı türüne dönüştürülmüştür.



Şekil 1. Veri seti ve sınıflandırma akış şeması.

Figure 1. Data set and classification flow chart.

Derin öğrenme veya makine öğrenimi algoritmalarının tahminlere dayalı olarak işlev görmesi için, metinsel verileri doğrudan işlemek mümkün değildir. Bu nedenle, tüm metinsel özellikler, sayısal değerleri kabul edebilen bir formata dönüştürülmelidir. Bu dönüşüm, metinleri sabit uzunlukta sayısal vektörlere dönüştürerek gerçekleştirilir. Bu vektör, metinlerin metinsel kodlamasını temsil eder. Bu dönüşüm için, Sklearn kütüphanesinin CountVectorizer fonksiyonu kullanılarak vektörler oluşturulur [16]. Bu yöntem, metinleri parçalayarak kelime dağılımını temsil eden sayısal bir dizi oluşturur. Metin yorumları sabit uzunlukta sayısal vektörlere dönüştürülür ve ardından terim frekansı-ters belge frekansı (TF-IDF) olarak bilinen bir dönüşüme tabi tutulur. Bu dönüşüm, temel kelimelere dayalı bir derleme oluşturma ve algoritma eğitimi için veri madenciliği yapma yöntemidir [17],[18].

3.2.1 Tokenleştirme

Amazon web sitesinden elde edilen ön işlenmiş verilerden sonra, duygu analizi üzerine odaklanan analiz işlemleri gerçekleştirilmiştir. Veriler, pandas, numpy ve matplotlib gibi kütüphaneler kullanılarak Python programlama dili ile işlenmiştir. Veri yorumlarını ön işleme için doğal dil işleme (Natural Language Processing - NLP) teknikleri kullanılmıştır. Burada cümle tokenleştirme işlemi gerçekleştirilmiştir. Bu yöntem, kullanıcı yorumlarını analiz için ayrı cümlelere ayırmakta ve metinde anlamlı bileşenleri belirlemekte kullanılır. Python'daki NLTK kütüphanesinden sent_tokenize() fonksiyonu, metni cümlelere ayırmak için kullanılmıştır. Böylece her cümle için içeriği ve duygusu hakkında daha detaylı bir anlayış sağlanmıştır. Şekil 2, cümleyi tokenlere ayırmak için kullanılan Tokenleştirme şemasını göstermektedir.

Cümle tokenleştirmenin ötesinde, yorumların ön işlenmesi, noktalama işaretlerinin kaldırılması, köklemeyi uygulama ve çeşitli kelime dönüşümleri de bu aşamada gerçekleştirilmiştir. Noktalama işaretlerini kaldırmak, özellikle duygu analizinde

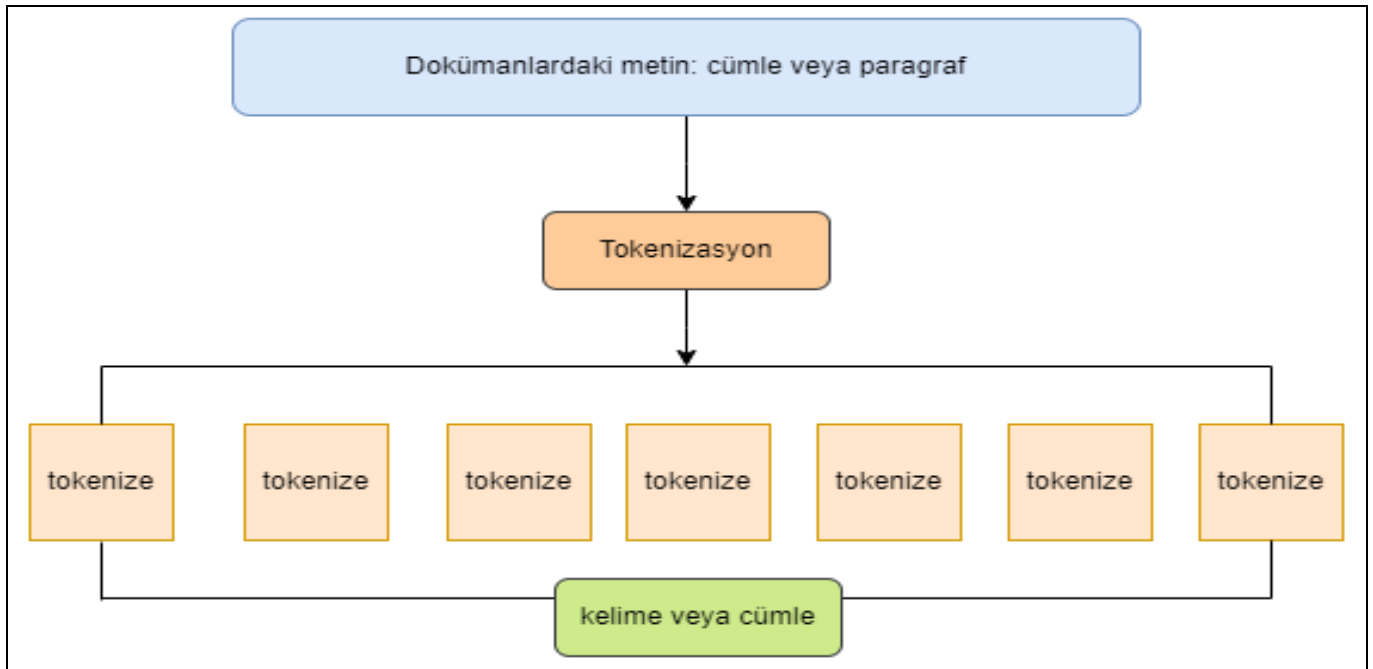
düzgün bir metin oluşturmada önemli bir adımdır. Çünkü bu sembollerin genellikle anlamlı bir katkı sağlamazlar. Köklemek, kelimeleri temel formlarına dönüştürerek veri içinde tutarlılık sağlamada kritik bir rol oynar. Bu süreç, aynı kelimenin farklı formlarını tek bir varlık olarak ele almayı sağlayarak, metinsel analizin bütünlüğünü korumaya yardımcı olur.

3.2.2 Noktalama işaretlerinin kaldırılması

Bu yöntem, yorumlardaki nokta, virgül, soru işareti ve ünlem gibi tüm noktalama işaretlerini ortadan kaldırarak metni standartlaştırmak ve gereksiz karakterlerin analizin doğruluğu üzerindeki olası etkisini en aza indirmek için kullanılmıştır. Dolayısıyla NLP bağlamında bu teknik, kritik bir öneme sahiptir. Noktalama işaretlerinin kaldırılması, yorumları daha temiz ve daha tutarlı bir forma sokar ve analizin sonraki adımlarını büyük ölçüde kolaylaştırır. Noktalama işaretlerini kaldırmak için çeşitli Python kütüphaneleri kullanılmıştır. Şekil 3, noktalama işaretlerinin nasıl kaldırıldığını göstermektedir.

3.2.3 WordNetLemmatizer nesnesi

NLP içeren yapay zekâ projelerinde WordNetLemmatizer nesnesinin kullanımı, metin kalitesini ve sonraki analizlerin doğruluğunu artırmak için önemli bir bileşen olarak kabul edilmektedir. Gereksiz karakterlerin ve noktalama işaretlerinin kaldırılması gibi teknikler uygulandıktan sonra, WordNetLemmatizer nesnesi, metin analizi için uygun olan orijinal formata (kök kelimelerine) dönüştürmek için kullanılır. Amazon'daki kullanıcı yorumlarından türetilen metin verilerini analiz ederken, WordNetLemmatizer nesnesini kullanmak, temel kelimelerin çıkarılmasına olanak tanır. Sonuç olarak, kullanıcı yorumlarından elde edilen verilerin analizi daha verimli ve doğru bir şekilde gerçekleştirilebilir. Bu bilgiler, ürün tasarımını iyileştirmek veya müşteri hizmetlerini geliştirmek gibi alanlarda kullanılabilir.



Şekil 2. Cümleyi tokenlere ayırmak için tokenleştirme şeması.
Figure 2. Tokenization scheme for splitting the sentence into tokens.

	A	B	C	D
1	Text with Punctuation			Remove Punctuation
2	"Apple"			Apple
3	(Pear). 5			Pear 5
4	{{Orange}}			Orange
5	Lemon;;; :::			Lemon
6	Lychee!			Lychee
7	<Blueberry>			Blueberry
8	Dash-test			Dashtest
9	TEST~!#\$%^&*()_+{} '";<>?.,			TEST

Şekil 3. Noktalama işaretlerinin kaldırılması.

Figure 3. Removing punctuation marks.

3.2.4 CountVectorizer

Scikit-learn kütüphanesi tarafından sağlanan CountVectorizer fonksiyonu, betiklerdeki kelime temsillerini dönüştürmek için kullanılır. CountVectorizer() ile oluşturulan vector_count nesnesi, metni her kelimenin temsil sayısını içeren bir vektöre dönüştürür. Özetle, bu kod metinleri bir diziye dönüştürerek, her metindeki her kelimenin oluşumunu temsil eden vektörler oluşturur. Bu vektörler daha sonra words_feature dizisinde saklanır.

3.2.5 Normalizasyon

Çeşitli stillerde ve formatlarda çok sayıda yorumu analiz etmek, verileri karşılaştırılabilir ve analiz edilebilir bir formata dönüştürmeyi gerektirir. Normalizasyon, verileri standart bir forma dönüştürmek için kullanılır. İlk aşamada, yorumlardan gereksiz harfler ve semboller çıkarılır. Ardından, metinler belirli bir harf ve rakam kombinasyonuna dönüştürülür. Daha sonra, metinler sayılara dönüştürülür ve belirli bir aralık belirlenir. Bu, yorumlar arasında analiz ve karşılaştırma doğruluğunu artırmaya yardımcı olur. Genel olarak, normalizasyon verileri standartlaştırmada yardımcı olur ve bunları karşılaştırma ve kullanım için daha uygun hale getirir.

3.3 Veri bölme

Modelin geliştirme ve değerlendirme aşamaları için, veri seti üç ana segmente sistematik bir şekilde kategorize edilmiştir: eğitim, doğrulama ve test. Veri kümesinde ön işleme adımları uygulandıktan sonra çalışma deneyleri 8906 veri üzerinde gerçekleştirilmiştir. Bu verilerden %40 oranında test verisi seçilerek, eğitim sürecinde modellerin görmeyeceği verilere ayrılmıştır. Veri setini %60 eğitim, %40 test olarak ayırma kararı, modelin genelleme yeteneğini artırmak ve aşırı öğrenmeyi önlemek için belirlenmiştir. Makine öğrenmesi ve derin öğrenme literatüründe, tipik olarak %70-30, %80-20 veya %60-40 gibi farklı veri bölme stratejileri kullanıldığı görülmektedir [19].

Özellikle geniş veri setlerinde, test verisinin büyük bir kısmının (%40) ayrılması modelin gerçek dünya verileri üzerindeki genelleme kapasitesini daha iyi değerlendirmek açısından önemli bir avantaj sağlamaktadır. Paknejad çalışmasında, Amazon ürün yorumları üzerine yapılan duygu analizi çalışmalarında %60-40 veri bölme stratejisinin sıklıkla tercih edildiği ve bu bölme oranının model doğruluğunu artırmada etkili olduğu belirtilmiştir [12]. Diğer yaygın bölme stratejileriyle karşılaştırıldığında ise %70-30 veya %80-20 bölme stratejileri genellikle daha küçük veri kümeleri için önerilmektedir [20]. Ancak büyük ölçekli veri setlerinde, test

kümesinin daha geniş tutulması önerilir. %60-40 bölme, büyük veri setlerinde test verisini daha geniş tutarak modelin genellenabilirliğini daha iyi değerlendirmeye olanak tanır. Bu nedenle, çalışmamızda kullanılan %60 eğitim ve %40 test bölme stratejisi hem literatürde desteklenen bir yöntemdir hem de modelin gerçek dünya verileri üzerindeki performansını daha iyi değerlendirmek için tercih edilmiştir. Kalan %60 oranındaki veriden %80'i eğitim, %20'si ise doğrulama için kullanılmıştır.

3.4 Çalışmada kullanılan modeller

Bu bölümde çalışmada kullanılan makine öğrenimi ve derin öğrenme modelleri açıklanmış, modellerin oluşturulmasında kullanılan parametreler belirtilmiştir. Parametreler, deneme yanılma yöntemine göre modellerin en iyi performans gösterdiği parametreler belirlenene kadar optimize edilmiştir.

3.4.1 Destek vektör makinesi

Destek Vektör Makinesi, sınıflandırma ve regresyonla ilgili sorunları çözmek için kullanılan güçlü bir makine öğrenimi algoritmasıdır. SVM'nin ana özellikleri arasında, veriden maksimum marjlı sapmayı belirleyen hiperparametre seçimi ve modeli daha doğru hale getirmek için çekirdek kullanılması bulunur. Ayrıca, küçük ve büyük veri kümelerinde veri sınıflandırma ve çoklu setlerde veri sınıflandırma yeteneği gibi geniş bir özellik yelpazesi sunar. SVM, finans, biyomühendislik, tıp ve cerrahi gibi çeşitli alanlarda kullanılır. Optimal hiperparametreleri modelin performansını büyük ölçüde etkiler [21],[22].

SVM modelinin performansını artırmak için CalibratedClassifierCV metodu kullanılmıştır. Bu metod, Scikit-Learn kütüphanesinde bulunan, olasılık tahminlerini daha güvenilir hale getirmek için kullanılan bir tekniktir. SVM gibi bazı sınıflandırıcılar, varsayılan olarak olasılık tahminlerini doğrudan sağlamaz. CalibratedClassifierCV, bu tahminleri kalibre ederek, daha doğru ve güvenilir olasılık değerleri elde edilmesini sağlar. Çalışmada ayrıca, çekirdek olarak kernel='rbf' (Radial Basis Function) seçilmiştir. Seçilen rbf çekirdeği, doğrusal olmayan ayrımların modellenmesi için en yaygın kullanılan çekirdek fonksiyonudur. Regularizasyon Parametresi ise C=1.0 olarak belirlenmiştir. C değeri, modelin hata toleransını belirler. Düşük C değerleri ile model daha genelleştirilebilir olur, ancak bazı hatalara izin verir. Yüksek C değerlerinde ise model hataları daha fazla cezalandırır, ancak aşırı öğrenme riski artar.

3.4.2 Karar ağacı

Karar Ağacı, önde gelen makine öğrenimi algoritmalarından bir diğeridir. Karar Ağacının ana özellikleri arasında büyük ve çeşitli veri setlerini analiz etme, müşterileri segmentlere ayırma, kredi riskini analiz etme ve pazarın çeşitli yönlerini analiz etme yeteneği bulunur [23], [24]. Veri analizi, sınıflandırma ve regresyon gibi konuları ele alır. Modelin performansı, veri kalitesi, modelin güvenilirliği, aşırı uyum riski, farklı tüketici türlerinin sayısı, modelin karmaşıklığı gibi kriterler üzerine kuruludur. Çalışmada Karar Ağacı uygulaması gerçekleştirilirken kullanılan parametreler şöyle özetlenebilir.

DT modeli, "entropy" kriteri ile oluşturulmuş ve maksimum derinliği (max_depth=3) olarak sınırlandırılmıştır. Criterion parametresi dallanma işlemi sırasında hangi ölçütün kullanılacağını belirler ve genellikle kullanılan iki temel seçenek vardır (gini ve entropy). Çalışmada her düğümde bilgi

kazancını maksimize eden bölme tercih edilmiştir. Literatürde, gini ve entropy arasında büyük performans farkları gözlenmemekle birlikte, "entropy" genellikle daha hesaplama yoğun ancak teorik olarak biraz daha iyi sonuç verebilen bir kriterdir. Maksimum derinlik parametresi ağacın derinliğini sınırlandırır. Derinlik arttıkça model daha karmaşık hale gelir ve aşırı öğrenme riski artar. Maksimum derinliğin küçük tutulması, modelin genelleme yeteneğini artırabilir ancak çok küçük bir değer de yetersiz öğrenmeye yol açabilir. Çalışmada kullanılan max_depth=3 değeri ile ağacın 3 seviye derinlikte büyüebilmesi sağlanmıştır. Diğer parametrelerden bölme yöntemi 'best' (en iyi), düğümlerin bölünebilmesi için minimum örnek sayısı ise 2 olarak belirlenmiştir. Her bölmede kullanılacak maksimum özellik sayısına bir sınırlandırma getirilmemiştir.

3.4.3 K-en yakın komşu

K-En Yakın Komşu, makine öğrenimi uygulamalarında kullanılan temel ve etkili bir öğrenme yöntemidir. Öklid, Manhattan, Minkowski vb. gibi çeşitli uzaklık ölçütleri kullanılarak elde edilen mesafe hesaplarıyla birlikte kullanılabilir. kNN algoritmalarının özellikleri arasında basitlik, doğruluk, parametrik optimizasyon ve performans bulunur. Hem sınıflandırma hem de regresyon problemlerinde kullanılabilirler [25]. Çalışmada kNN algoritması için en önemli iki parametreden biri olan komşu sayısı (k) 3 olarak belirlenmiş, uzaklık ölçütü olarak ise Öklid uzaklığı kullanılmıştır.

3.4.4 Uzun kısa süreli bellek

Uzun kısa süreli bellek, eğitimde kullanılan derin öğrenme modelidir ve öğrenmeyle ilgili sorunları analiz etmek ve çözmek için NLP yaklaşımını kullanır. Öğrenme sürecine odaklanarak, öğrenme sürecinin altında yatan desenleri belirler ve elde edilen bilgileri depolar. LSTM'nin ana özellikleri arasında, öğrenme sürecinin bilgi akışını kontrol eden gizli kapılar bulunur. Ayrıca, öğrenme sürecinin dinamik bileşenlerini kullanarak daha karmaşık öğrenme sonuçlarına olanak tanır. LSTM, öğrenme sürecindeki sorunları bozan ve modelin performansını zayıflatan kaybolan gradyan problemi ile de ilgilenir. LSTM, çeşitli öğrenme problemleriyle başa çıkabilen güçlü bir öğrenme modelidir [26],[27]. Çalışmada oluşturulan LSTM modelinde 82 nöronlu bir katman kullanılmıştır. Giriş boyutu veriye uyarlanmış ve çıkış katmanını 3 sınıflı bir katman olarak belirlenmiştir. Çıkış katmanında Softmax aktivasyon fonksiyonu kullanılmıştır. Çalışmada batch size olarak 32 değeri kullanılmıştır. Çok sınıflı sınıflandırma problemlerinde en yaygın kullanılan kayıp fonksiyonu olan loss='categorical_crossentropy' fonksiyonu tercih edilmiştir. Optimizer parametresinde 'adam' algoritması kullanılmıştır.

3.5 Model performansı

Makine öğrenimi ve derin öğrenme modelleri, doğruluk ve verimliliklerini değerlendirmek için çeşitli performans metrikleri kullanılarak değerlendirilir. Bu metrikler arasında doğruluk, kesinlik, geri çağırma, F1 puanı ve AUC_ROC bulunur. Doğruluk (Denklem (1)), doğru tahmin edilen gözlemlerin toplam gözlem sayısına oranını ölçerek bir modelin tahmin yeteneğinin bütünsel bir ölçüsünü sağlar. Kesinlik (Denklem (2)), tahmin edilen pozitif örneklerin oranını gerçek pozitiflere karşılık gelenlerle karşılaştırarak, pozitif vakaları sınıflandırmada modelin kesinliğini ayırt eder. Geri çağırma (Denklem (3)), bir modelin gerçek pozitifleri doğru bir şekilde tanımlama yeteneğini değerlendirirken, F1 puanı (Denklem

(4)) kesinlik ve geri çağırma uyumlaştırarak tahmin performansının dengeli bir değerlendirmesini sağlar. Son olarak, AUC_ROC metriği, bir modelin pozitif ve negatif sınıflar arasında ayırt etme yeteneğini gösterir.

$$\text{Doğruluk} = \frac{DP + DN}{DP + DN + YP + YN} \quad (1)$$

$$\text{Kesinlik} = \frac{DP}{DP + YP} \quad (2)$$

$$\text{Geri Çağırma} = \frac{DP}{DP + YN} \quad (3)$$

$$\text{F1 Puanı} = \frac{2 \times \text{Kesinlik} \times \text{Geri Çağırma}}{\text{Kesinlik} + \text{Geri Çağırma}} \quad (4)$$

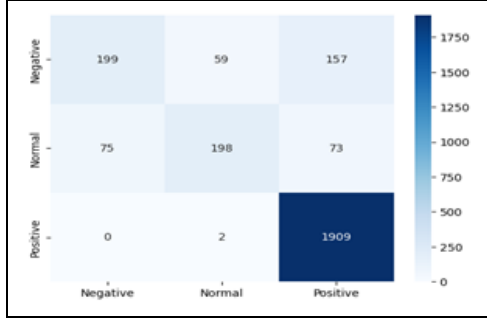
Denklemlerde doğru pozitif değerler için DP, yanlış pozitif değerler için YP, doğru negatif değerler için DN ve yanlış negatif değerler için YN kısaltması kullanılmıştır. Bu metrikler bir sınıflandırma algoritmasının performansını kapsamlı bir şekilde değerlendirerek çeşitli uygulamalardaki etkinliği konusunda bilinçli karar verme sürecini kolaylaştırır [28],[29].

4 Sonuçlar ve tartışma

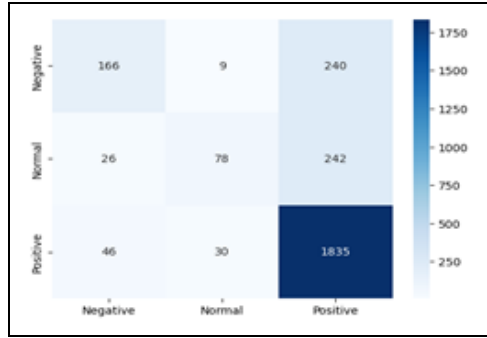
Bu bölüm, işlenmiş yorum verilerine sınıflandırma görevleri için SVM, DT, kNN ve LSTM algoritmalarının uygulanmasından elde edilen sonuçları sunmaktadır. Bu yorumlar, orijinal veri setinden türetilen ve işlem sonrası sonuçları temsil etmektedir. Farklı kategoriler arasındaki olası ilişkileri değerlendirmek için, önerilen tüm algoritmalar (SVM, DT, kNN ve LSTM) için bir karışıklık matrisi hesaplanmıştır ve Şekil 4'te gösterilmiştir.

Karışıklık matrisi, her sınıf için tahminlerin ayrıntılı bir dökümünü sunarak, doğru ve yanlış tahminler arasındaki oranı vurgular. Matristeki değerler, doğru pozitifler, yanlış pozitifler ve yanlış negatifler göstererek sınıflar arasındaki ilişkileri ve tahmin hatalarını anlamamıza yardımcı olur. SVM için olumlu yorumlarda doğru pozitif sayısı 1909, yanlış normal sayısı 2 ve yanlış negatif oranı 0'dır. Normal ve olumsuz yorumlar için sayılar, doğru ve yanlış sınıflandırmaların bir karışımını gösterir ve SVM modelinin iyi performans gösterdiği alanları ve geliştirilmesi gereken yerleri vurgular. Buna karşılık, DT modeli, yorum kategorileri arasında doğru pozitifler, yanlış pozitifler ve yanlış negatifler için farklı oranlar göstermiştir. Olumlu yorumlar için bu sayılar sırasıyla 1835, 30 ve 46'dır. Normal ve olumsuz yorumlardaki performans, modelin doğruluğundaki ve sınıflar arasındaki ayrımı yapma yeteneğindeki varyasyonu gösterir. Benzer şekilde, kNN modelinin karışıklık matrisi, olumlu yorumlar için doğru pozitif sayısı 1880, yanlış normal oranı 19 ve yanlış negatif oranı 12 olan sınıflandırma yeteneklerini ortaya koymuştur. Normal ve olumsuz yorumlar için sayılar yine, modelin çeşitli yorum türlerini sınıflandırmadaki güçlü ve zayıf yönlerine dair bilgiler sağlamaktadır. Son olarak, LSTM modeli, özellikle olumlu yorumlarda doğru pozitif sayısı 2534 olan, sıfır yanlış normal ve yanlış negatif oranları ile dikkat çekici bir doğruluk sergilemiştir. Normal ve olumsuz yorumlardaki performans bu eğilimi sürdürerek, LSTM modelinin farklı duygulara sahip yorumları doğru bir şekilde sınıflandırmada iyi bir seçenek olduğunu kanıtlamıştır. Özetle, her model duygu sınıflandırmada benzersiz güçlü yönler ve iyileştirme alanları sunmaktadır. LSTM modeli, tüm yorum türleri arasında yüksek doğrulukla öne çıkmıştır, diğer modeller ise çeşitli derecelerde etkililik gösterir. Bu bilgiler, bu algoritmaları daha da

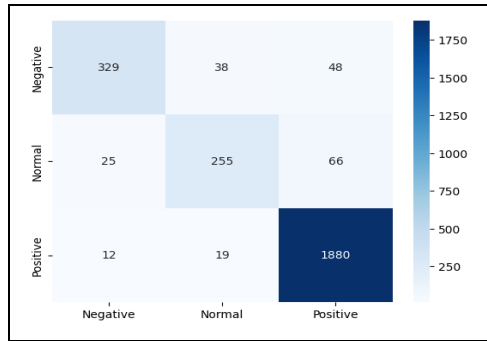
geliştirmek ve duygu analizinde uygulamalarını iyileştirmek için hayati öneme sahiptir.



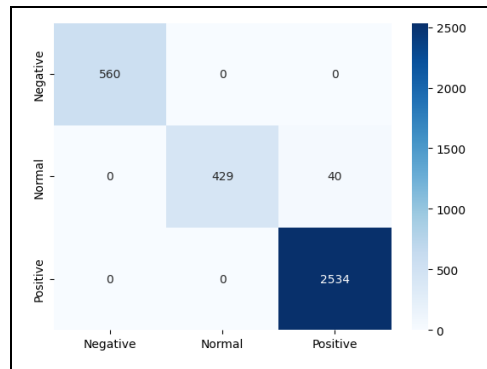
(a)



(b)



(c)



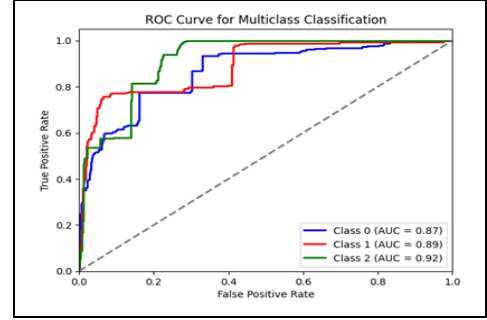
(d)

Şekil 4. Önerilen algoritmalar için karışıklık matrisi. (a): SVM. (b): DT. (c): kNN. (d): LSTM.

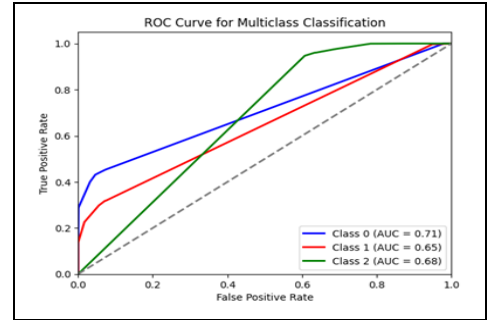
Figure 4. Confusion matrix for the proposed algorithms. (a): SVM. (b): DT. (c): kNN. (d): LSTM.

Ayrıca, çeşitli veri türleri arasında sınıflandırma modellerinin performansını değerlendirmede kritik bir ölçüm olan AUC-ROC

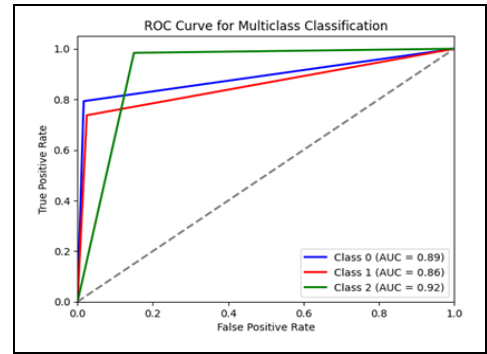
endeksi, bu çalışmada kullanılan farklı modellerin etkinliğini doğrulamak için hesaplanarak Şekil 5'te gösterilmiştir.



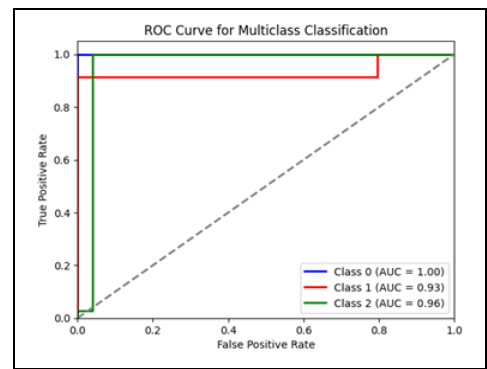
(a)



(b)



(c)



(d)

Şekil 5. Önerilen algoritmalar için ROC eğrisi. (a): SVM. (b): DT. (c): kNN. (d): LSTM.

Figure 5. ROC curve for the proposed algorithms. (a): SVM. (b): DT. (c): kNN. (d): LSTM.

Şekil 5(a)'da görüldüğü gibi SVM modeli, ROC eğrisinin yukarı ve sol yönlü eğilimiyle yüksek bir doğru pozitif oranını

vurgulamış olup bu eğitim modelin pozitif ve negatif sınıflar arasında ayırma yapma yeteneğini göstermektedir. SVM modelinin 0.87, 0.89 ve 0.92 olan AUC değerleri, modelin sınıflandırmadaki iyi yeteneğini öne çıkarmıştır. Karar Ağacı modeli söz konusu olduğunda, ROC eğrisinin hareketi benzer şekilde yüksek bir doğru pozitif oranını yansıtmıştır. Bu hareket, pozitif ve negatif sınıflar arasında kolay bir ayırma sağlamaktadır. DT modelinin 0.71, 0.65 ve 0.68 olan AUC değerleri, SVM'den daha düşük olmasına rağmen, makul bir ayırma derecesini göstermektedir. KNN modelinin performansı, ROC eğrisinin sol üst köşeye yaklaşmasıyla yüksek bir doğru pozitif oranını işaret etmektedir. Bu modelin 0.89, 0.86 ve 0.92 olan AUC değerleri, farklı sınıfları doğru bir şekilde sınıflandırmada etkinliğini vurgular. En dikkat çekici olarak, LSTM modelinin ROC eğrisi neredeyse sol üst köşeye ulaşmıştır. Bu da yüksek bir doğru pozitif oranını ima eder. Bu model, 1.00, 0.93 ve 0.96 olan AUC değerleriyle diğerlerinden daha başarılı performans sergileyerek sınıflar arasında ayırma yapma yeteneğinde üstünlüğünü göstermiştir. Son olarak, her modelin benzersiz güçlü ve zayıf yönleri vardır. Çalışmada modeller için yüksek genel doğruluk oranları elde edilmiştir: SVM için %86, DT için %77 ve kNN için %92. Bu oranlar, araştırmada kullanılan makine öğrenmesi algoritmalarından elde edilen yüksek doğruluk seviyesini göstermektedir. Ancak LSTM ile ulaşılan oran, makine öğrenmesi algoritmalarından daha başarılıdır. Bu çalışmadan elde edilen LSTM sonucu, modelin genel doğruluğunun etkileyici bir şekilde %98 olduğunu göstermiştir. Algoritmaların performansı Tablo 1'de gösterilmiştir.

Tablo 1. Sonuç karşılaştırması.

Table 1. Results comparison.

Yöntem	Doğruluk	Keskinlik	Geri Çağırma	F1 Puanı	AUC_ROC
LSTM	0.98	0.95	0.98	0.98	0.98
DT	0.778	0.76	0.78	0.74	0.68
kNN	0.922	0.92	0.92	0.92	0.878
SVM	0.85	0.85	0.86	0.85	0.841

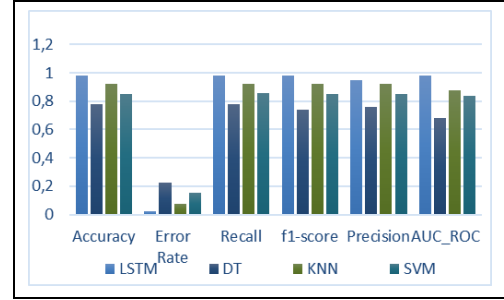
Bu sonuçlar, Amazon yorumlarını sınıflandırmaya yönelik farklı makine öğrenimi modellerinin karşılaştırmalı analizi sunmaktadır. Elde edilen bulgular arasında dikkat çeken en önemli sonuç, LSTM modelinin %98 doğruluk oranına ulaşarak en başarılı model olarak öne çıkmasıdır. Bu üstün performans, modelin ROC-AUC eğrisi altında en geniş alanı kapsayabilme yeteneğine dayanmaktadır; dolayısıyla veri sınıflandırma doğruluğu açısından yüksek bir başarı sergilediği görülmektedir.

Buna karşın DT modeli, %77.8 doğruluk oranı ile en düşük performansı göstermiştir. Ancak, DT modelinin görece düşük doğruluk oranına rağmen, değişkenler arasındaki nedensel ilişkileri analiz etme ve modelin yorumlanabilirliğini artırma açısından önemli bir katkı sunduğu gözlemlenmiştir. Bu özellik, özellikle veri içerisindeki örüntüleri anlamının kritik olduğu uygulamalarda Karar Ağacı modelini değerli bir araç haline getirmektedir.

Öte yandan kNN modeli, %92.2 doğruluk oranı ile dikkate değer bir performans sergilemiştir. Modelin basitliği ve hesaplama verimliliği, hızlı ve doğrudan çözümler gerektiren uygulamalar için avantaj sağlamaktadır. Benzer şekilde, doğrusal ve doğrusal olmayan ayırıştırma tekniklerini kullanan SVM modeli, %85 doğruluk oranına ulaşarak kabul edilebilir bir sınıflandırma başarısı göstermiştir. Bu sonuç, özellikle kesin ve

verimli veri sınıflandırmanın önemli olduğu senaryolarda SVM modelinin etkili bir seçenek olabileceğini göstermektedir.

Bu yüksek doğruluk seviyesi, farklı kategorilerdeki verileri sınıflandırmada seçilen modellerin verimliliğini ve güvenilirliğini vurgulayarak bu modellerin pratik uygulamalardaki potansiyelini ortaya koymuştur. Şekil 6'da sağlanan grafiksel temsil, çalışmanın bulgularını karşılaştırmalı göstermekte ve her model arasındaki farkları vurgulayarak daha fazla analiz için yararlı bir referans olarak hizmet etmektedir.



Şekil 6. Karşılaştırmalı sonuçlar.

Figure 6. Comparative results.

Genel olarak, en uygun modelin seçimi, uygulamanın gereksinimlerine ve veri kümesinin özelliklerine bağlı olarak değişiklik göstermektedir. Eğer temel hedef en yüksek doğruluk oranına ulaşmak ise, LSTM modeli en etkili yöntem olarak öne çıkmaktadır. Çalışmanın bulguları, makine öğrenimi modellerinin belirli koşullar altında iyi performans gösterebildiğini, ancak daha yüksek doğruluk oranlarına ulaşabilmek için ek çaba ve optimizasyon gerektirdiğini ortaya koymaktadır. Bununla birlikte, derin öğrenme yaklaşımlarının, özellikle LSTM algoritmasının, klasik makine öğrenimi yöntemlerine kıyasla daha yüksek doğruluk oranları sunduğu tespit edilmiştir. Modelin performansı, veri setinin genişletilmesi ve eğitim sürecinin optimize edilmesi ile daha da artırılabilir.

Bu alanda daha önce yapılan çalışmalarla karşılaştırıldığında, çalışmamız bazı önemli farklılıklar ve avantajlar sunmaktadır. Örneğin, Joseph çalışmasında, önceden eğitilmiş BERT modeli ile Amazon yorumları üzerinde duyarlılık analizi gerçekleştirilmiş ve %96.3 doğruluk oranı elde edilmiştir [6]. Benzer şekilde, Gope ve diğ., RNN-LSTM modelini kullanarak %97.52 doğruluk oranına ulaşmıştır [8]. Çalışmamızda ise LSTM modelinin optimize edilmesi ve geniş bir veri seti kullanılması sayesinde %98 doğruluk oranına ulaşılmıştır. Bu durum, modelin doğal dilin karmaşıklığını başarılı bir şekilde ele alabildiğini göstermektedir.

Makine öğrenmesi tabanlı yaklaşımlar açısından, Hamdallah çalışmasında SVM ve Naive Bayes modelleri kullanılarak duyarlılık analizi gerçekleştirilmiş, ancak doğruluk oranları %85'in üzerine çıkamamıştır [13]. Bu çalışmada ise SVM modeli %85 doğruluk oranı elde ederken, kNN modeli %92.2 ve LSTM modeli %98 doğruluk oranına ulaşmıştır. Bu sonuçlar, derin öğrenme tabanlı yöntemlerin özellikle büyük veri setlerinde daha yüksek doğruluk sağladığını göstermektedir.

Bunun yanı sıra, literatürdeki bazı çalışmalar daha küçük ölçekli veri setleriyle gerçekleştirilmiş ve farklı ön işleme teknikleri kullanılmıştır. Örneğin, AlZu'bi ve diğ., Amazon yorumlarını pozitif ve negatif olarak sınıflandırmış ancak kullanıcı oylama yüzdelерinin analizinde bazı sınırlamalarla

karşılaşmıştır [10]. Çalışmamız, farklı derecelerde duygu analizine (olumlu, nötr, olumsuz) odaklanarak daha ayrıntılı bir değerlendirme sunmaktadır.

5 Sonuç

Bu makale, müşteri sadakatini tahmin etmede Müşteri İlişkileri Yönetiminde duygu analizinin önemini vurgulamaktadır. Alışveriş alışkanlıklarına odaklanan geleneksel yöntemler genellikle müşteri sadakati sürücülerinin tam spektrumunu, örneğin memnuniyet ve ürün deneyimini yakalamada başarısız olur. Bu makale, çevrimiçi bir ortamdan çeşitli müşteri yorumlarını analiz etmek için makine öğrenimi ve derin öğrenime dayalı ileri NLP teknikleri kullanmaktadır. Çalışmada, veri toplama için Python kullanılmıştır ve Amazon'daki on elektronik cihazdan müşteri yorumlarını sınıflandırmak için kNN, SVM, DT ve LSTM gibi çeşitli algoritmalar kullanmıştır. Bulgular, bu ileri analitik tekniklerin duyguları doğru bir şekilde sınıflandırmadaki ve müşteri memnuniyeti ve sadakatini artırmadaki potansiyelini göstermektedir. Elde edilen bulgular, derin öğrenme yöntemlerinin geleneksel makine öğrenimi modellerine kıyasla daha yüksek doğruluk oranlarına ulaştığını ve tüketici yorumlarındaki duygu analizinin etkin bir şekilde gerçekleştirilebildiğini göstermektedir. Özellikle LSTM modeli, doğal dilin karmaşıklıklarını başarıyla işleyerek müşteri memnuniyeti ve sadakatini değerlendirme süreçlerinde güçlü bir araç olarak öne çıkmıştır. Bu çalışmadaki veri seti, duygu tanıma ve spam tespiti alanında daha ileri keşifler için değerli bir kaynaktır.

6 Conclusion

This paper highlights the importance of sentiment analysis in Customer Relationship Management in predicting customer loyalty. Traditional methods focusing on shopping habits often fail to capture the full spectrum of customer loyalty drivers, e.g. satisfaction and product experience. This paper uses advanced NLP techniques based on machine learning and deep learning to analyze various customer reviews from an online environment. The study uses Python for data collection and employs various algorithms such as kNN, SVM, DT and LSTM to classify customer reviews from ten electronic devices on Amazon. The findings demonstrate the potential of these advanced analytical techniques in accurately classifying sentiments and improving customer satisfaction and loyalty. The findings show that deep learning methods achieve higher accuracy rates compared to traditional machine learning models and can effectively perform sentiment analysis on consumer reviews. In particular, the LSTM model successfully handled the complexities of natural language and emerged as a powerful tool for evaluating customer satisfaction and loyalty. The dataset in this study is a valuable resource for further exploration in the field of sentiment recognition and spam detection.

7 Yazar katkı beyanı

Gerçekleştirilen çalışmada Nazeeha Sayghn Khalid Khalid, fikrin oluşması, tasarımı, yapılıması, literatür taraması, veri elde etme, analiz ve sonuçların yorumlanması, çalışmanın yazılması ve düzenlenmesi konularında katkı sağlamıştır. Serkan Savaş elde edilen sonuçların incelenmesi, sonuçların değerlendirilmesi, çalışmanın yazılması, kontrolü ve düzenlenmesi konularında katkı sağlamıştır.

8 Etik kurul onayı ve çıkar çatışması beyanı

"Hazırlanan makalede etik kurul izni alınmasına gerek yoktur". "Hazırlanan makalede herhangi bir kişi/kurum ile çıkar çatışması bulunmamaktadır".

9 Kaynaklar

- [1] Wu F, Shi Z, Dong Z, Pang C, Zhang B. "Sentiment analysis of online product reviews based on SenBERT-CNN". *International Conference on Machine Learning and Cybernetics (ICMLC)*, Adelaide, Australia, 02 December 2020.
- [2] Soleymani M, Garcia D, Jou B, Schuller B, Chang SF, Pantic M. "A survey of multimodal sentiment analysis". *Image Vis Comput*, 65, 3-14, 2017.
- [3] Le VQ, Mikolov T. "Distributed Representations of Sentences and Documents". *arXiv*, 2018. <https://arxiv.org/pdf/1405.4053.pdf>.
- [4] Du J, Rong J, Michalska S, Wang H, Zhang Y. "Feature selection for helpfulness prediction of online product reviews: An empirical study". *PLoS One*, 14(12), e0226902, 2019.
- [5] Askalidis G, Malthouse EC. "The value of online customer reviews". *10th ACM Conference on Recommender Systems*, Boston, MA, USA, 15-19 September 2016.
- [6] Joseph RPS. Amazon Reviews Sentiment Analysis: A Reinforcement Learning Approach. MSc Thesis. Griffith College, Dublin, Ireland, 2020.
- [7] Shrestha N, Nasoz F. "Deep learning sentiment analysis of amazon.com reviews and ratings". *International Journal on Soft Computing, Artificial Intelligence and Applications (IJSCAI)*, 8(1), 1-15, 2019.
- [8] Gope JC, Tabassum T, Mabrur MM, Yu K, Arifuzzaman M. "Sentiment analysis of amazon product reviews using machine learning and deep learning models". *International Conference on Advancement in Electrical and Electronic Engineering, ICAEEE 2022*, Gazipur, Bangladesh, 24-26 February 2022.
- [9] Wedjdane N, Khaled R, Okba K. "Better decision making with sentiment analysis of amazon reviews". *International Conference on Information Systems and Advanced Technologies (ICISAT)*, Tebessa, Algeria, 27-28 December 2021.
- [10] AlZu'bi S, Alsmadiv A, AlQatawneh S, Al-Ayyoub M, Hawashin B, Jararweh Y. "A brief analysis of amazon online reviews". *Sixth International Conference on Social Networks Analysis, Management and Security (SNAMS)*, Granada, Spain, 22-25 October 2019.
- [11] Norinder U, Norinder P. "Predicting amazon customer reviews with deep confidence using deep learning and conformal prediction". *Journal of Management Analytics*, 9(1), 1-16, 2022.
- [12] Paknejad S. "Sentiment classification on Amazon reviews using machine learning approaches". KTH Royal Institute of Technology, Stockholm, Sweeden, Degree Project in Technology, DD142X, 2018.
- [13] Hamdallah AR. Amazon Reviews Using Sentiment Analysis. MSc Thesis. Rochester Institute of Technology, Dubai, 2021.
- [14] Meenakshi, Banerjee A, Intwala N, Sawant V. *Sentiment Analysis of Amazon Mobile Reviews*. Editors: Tuba M, Akashe S, Joshi A. ICT Systems and Sustainability, 43-52, Singapore, Springer Singapore, 2020.

- [15] Richardson L, "Beautiful Soup Documentation". <https://www.crummy.com/software/BeautifulSoup/bs4/doc/> (10.02.2025).
- [16] Athanasiou V, Maragoudakis M, Kermanidis KL, Makris C, Mylonas P, Sioutas S. "A novel, gradient boosting framework for sentiment analysis in languages where NLP resources are not plentiful: a case study for modern greek". *Algorithms*, 10(1), 34, 2017.
- [17] Gündüz AB, Yavuz AG. "Comparative analysis of word vectoring methods for malware detection". *29th Signal Processing and Communications Applications Conference (SIU)*, İstanbul, Türkiye, 09-11 June 2021.
- [18] Salur MU, Aydın İ, Jamous M. "An ensemble approach for aspect term extraction in Turkish texts". *Pamukkale Üniversitesi Mühendislik Bilimleri Dergisi*, 28(5), 769-776, 2022.
- [19] Hastie T, Tibshirani R, Friedman J. *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. 2nd ed. New York, USA, Springer, 2009.
- [20] Brownlee J. *Machine Learning Mastery with Python: Understand Your Data, Create Accurate Models, and Work Projects End-to-End*. v1.20 ed. Melbourne, Australia, Independently Published, 2021.
- [21] Karakış R. *Destek Vektör Makinesi*. Editors: Savaş S, Buyrukoğlu S. Teori ve Uygulamada Makine Öğrenmesi, 93-118, Ankara, Türkiye, Nobel Akademik Yayıncılık Eğitim Danışmanlık Tic. Ltd. Şti., 2022.
- [22] Güler O. "Turbofan motorlarının kestirimci bakımında makine öğrenimi algoritmaları performanslarının karşılaştırılması". *Niğde Ömer Halisdemir Üniversitesi Mühendislik Bilimleri Dergisi*, 13(1), 99-106, 2024.
- [23] Tuğal İ. *Karar Ağaçları*. Editors: Savaş S, Buyrukoğlu S. Teori ve Uygulamada Makine Öğrenmesi, 119-148, Ankara, Türkiye, Nobel Akademik Yayıncılık Eğitim Danışmanlık Tic. Ltd. Şti., 2022.
- [24] Yavuz ÖÇ, Calp MH, Erkengel HC. "Prediction of breast cancer using machine learning algorithms on different datasets". *Ingeniería Solidaria*, 19(1), 1-32, 2023.
- [25] Bütüner R, Calp MH. *K-En Yakın Komşu Algoritması*. Editors: Savaş S, Buyrukoğlu S. Teori ve Uygulamada Makine Öğrenmesi, 223-250, Ankara, Türkiye, Nobel Akademik Yayıncılık Eğitim Danışmanlık Tic. Ltd. Şti., 2022.
- [26] Bilen B, Horasan F. "LSTM Network based Sentiment analysis for customer reviews". *Politeknik Dergisi*, 25(3), 959-966, 2022.
- [27] Zeybel Peköz A, İnkaya T. "Derin öğrenme ile talep tahmini: Bir üçüncü parti lojistik firması için COVID-19 döneminde vaka analizi". *Pamukkale Üniversitesi Mühendislik Bilimleri Dergisi*, 29(4), 331-339, 2023.
- [28] Horasan F, Yurttakal AH, Gündüz S. "A novel model based collaborative filtering recommender system via truncated ULV decomposition". *Journal of King Saud University-Computer and Information Sciences*, 35(8), 101724, 2023.
- [29] Ayan E. "Using a convolutional neural network as feature extractor for different machine learning classifiers to diagnose pneumonia". *Sakarya University Journal of Computer and Information Sciences*, 5(1), 48-61, 2022.