

**YAŞAM BOYU DUYGU ANALİZİ İLE MÜŞTERİ ÜRÜN
YORUMLARININ SINIFLANDIRILMASI**

Figen KAS BAYRAKDAROĞLU

**Haziran 2025
DENİZLİ**

**YAŞAM BOYU DUYGU ANALİZİ İLE MÜŞTERİ ÜRÜN
YORUMLARININ SINIFLANDIRILMASI**

**Pamukkale Üniversitesi
Sosyal Bilimler Enstitüsü
Doktora Tezi
İşletme Ana Bilim Dalı
Genel İşletme Doktora Programı**

Figen KAS BAYRAKDAROĞLU

Danışman: Prof. Dr. Esra AYTAÇ ADALI

**Haziran 2025
DENİZLİ**

Bu tezin tasarımı, hazırlanması, yürütülmesi, araştırmalarının yapılması ve bulgularının analizlerinde bilimsel etiğe ve akademik kurallara özenle riayet edildiğini; bu çalışmanın doğrudan birincil ürünü olmayan bulguların, verilerin ve materyallerin bilimsel etiğe uygun olarak kaynak gösterildiğini ve alıntı yapılan çalışmalara atıfta bulunulduğunu beyan ederim.

Figen KAS BAYRAKDAROĞLU

ÖN SÖZ

Öncelikle tüm doktora sürecim boyunca ve bu tezin hazırlanma sürecinde bana değerli bilgileriyle yol gösteren, beni motive eden ve her zaman destek olan saygıdeğer danışman hocam Prof. Dr. Esra AYTAÇ ADALI'ya teşekkür ederim. Saygıdeğer tez izleme komite üyeleri Prof. Dr. Muhsin ÖZDEMİR ve Prof. Dr. Ayşegül TUŞ hocalarıma ve saygıdeğer tez savunma jüri üyeleri Doç. Dr. Algin OKURSOY ve Dr. Öğr. Üyesi Gülin Zeynep ÖZTAŞ hocalarıma teşekkürlerimi sunarım.

Hayatıma, varlığıyla ışık katan biricik kızım Doğay'a ve beni her konuda destekleyen, her zaman arkamda duran ve bana güç veren canım eşim Halil İbrahim'e sonsuz sabırları ve sevgileri için teşekkür ederim.

Beni bu günlere getiren ve her daim destekçim olan canım annem Fatma KAS ve canım babam Kamil KAS'a teşekkür ederim.

Sizler olmasanız başaramazdım.

ÖZET

YAŞAM BOYU DUYGU ANALİZİ İLE MÜŞTERİ ÜRÜN YORUMLARININ SINIFLANDIRILMASI

Kas Bayrakdaroğlu, Figen
Doktora Tezi
İşletme Anabilim Dalı
Genel İşletme Doktora Programı
Danışman: Prof. Dr. Esra Aytaç Adalı

Haziran 2025, viii + 95 sayfa

Bu çalışmada, farklı alanlardan toplanan müşteri yorumlarının yaşam boyu öğrenme yaklaşımı kapsamında duygu sınıflandırması yapılmıştır. Bu sınıflandırma için bir online alışveriş sitesinden rastgele seçilmiş farklı ürünlerden toplanan müşteri yorumları kullanılmıştır. Her bir ürün, yaşam boyu öğrenme yaklaşımı için farklı bir alan olarak ele alınmıştır. Müşteri yorumları, ön işleme sürecinden geçirilmiş ve kelime vektörlerinin oluşturulmasında Countvectorizer ve TF-IDF yöntemleri kullanılmıştır. Yaşam boyu duygu sınıflandırması kapsamında yaşam boyu naive Bayes ve yaşam boyu lojistik regresyon sınıflandırma yöntemleri kullanılmış ve Python programlama dili ile kodlanmıştır. Bu iki yönteme ek olarak literatürden esinlenilerek toplama dayalı duygu sınıflandırması modeli oluşturulmuştur. Toplama dayalı duygu sınıflandırması, kelime sıklıklarının toplamına dayalı bir yaşam boyu öğrenme sürecini esas almaktadır. Kullanılan tüm yöntemlerin performansları karşılaştırılmış ve en iyi sonucu veren yöntemin, toplama dayalı duygu sınıflandırması modeli olduğu görülmüştür. Elde edilen tüm sonuçlar değerlendirildiğinde ele alınan problem, hem yaşam boyu duygu sınıflandırması algoritmalarının performanslarının test edilmesi ve karşılaştırılması açısından hem de e-ticaret işlemleri için önemli ve günceldir. Ayrıca yaşam boyu öğrenme kapsamında ilk kez Türkçe veri kümesi kullanılarak yapılan bu çalışmanın, sonraki çalışmalar için bir rehber niteliğinde olduğu söylenebilir.

Anahtar Kelimeler: Yaşam Boyu Öğrenme, Duygu Analizi, Metin Madenciliği, Doğal Dil İşleme, Toplama Dayalı Sınıflandırma, Naive Bayes, Lojistik Regresyon.

ABSTRACT

CUSTOMER PRODUCT REVIEWS CLASSIFICATION WITH LIFELONG SENTIMENT ANALYSIS

Kas Bayrakdaroğlu, Figen
Doctoral Thesis
Business Administration Department
Ph.D. in Business Administration
Adviser of Thesis: Prof. Dr. Esra Aytaç Adalı

June 2025, viii+ 95 pages

In this study, sentiment classification of customer comments collected from different areas was performed within the scope of lifelong learning approach. Customer comments collected from different products randomly selected from an online shopping site were used for this classification. Each product was considered as a different area for lifelong learning approach. Customer comments were preprocessed and Countvectorizer and TF-IDF methods were used to form word vectors. Lifelong naive Bayes and lifelong logistic regression classification methods were used within the scope of lifelong sentiment classification, and coded with Python programming language. In addition to these two methods, an aggregation-based sentiment classification model was presented inspired by the literature. Aggregation-based sentiment classification is based on a lifelong learning process based on the sum of word frequencies. The performances of all methods used were compared, and it was observed that the method that gave the best result was the aggregation-based sentiment classification model. When all the obtained results were evaluated, the problem addressed is important and up-to-date both in terms of testing and comparing the performances of lifelong sentiment classification algorithms and for e-commerce transactions. In addition, it can be said that this study, which was conducted for the first time using Turkish dataset within the scope of lifelong learning, is a guide for future studies.

Keywords: Lifelong Learning, Sentiment Analysis, Text Mining, Natural Language Processing, Aggregation-Based Classification, Naive Bayes, Logistic Regression.

İÇİNDEKİLER

ÖN SÖZ	i
ÖZET	ii
ABSTRACT.....	iii
İÇİNDEKİLER	iv
ŞEKİLLER DİZİNİ.....	vi
TABLolar DİZİNİ	vii
SİMGE VE KISALTMALAR DİZİNİ	viii
GİRİŞ	1

BİRİNCİ BÖLÜM

DUYGU ANALİZİ

1.1. Metin Madenciliği	5
1.1.1. Metin Ön İşleme Süreci.....	8
1.1.2. Metin Belgesinin Gösterimi	10
1.1.3. Modelin Performansının Değerlendirilmesi	12
1.2. Doğal Dil İşleme	14
1.3.1. Sözlüğe Dayalı (Sözlük Temelli) Yaklaşımlar.....	18
1.3.2. Makine Öğrenmesi Tabanlı Yöntemler	18
1.3.2.1. Naive Bayes yöntemi	20
1.3.2.2. Lojistik regresyon yöntemi	20
1.4. Türkçe’yi Konu Alan Duygu Analizi Çalışmalarına İlişkin Literatür Taraması	22

İKİNCİ BÖLÜM

YAŞAM BOYU ÖĞRENME YAKLAŞIMI

2.1. Yaşam Boyu Öğrenme	34
2.2. Yaşam Boyu Öğrenme Süreci	38
2.3. Yaşam Boyu Denetimli Öğrenme	41
2.3.1. Yaşam Boyu Naive Bayes Sınıflandırma Yöntemi	42
2.3.2. Yaşam Boyu Lojistik Regresyon Sınıflandırma Yöntemi	44
2.3.3. Toplama Dayalı Yaşam Boyu Duygu Sınıflandırması	46
2.4. Yaşam Boyu Öğrenme Yaklaşımına ve Duygu Analizine İlişkin Literatür Taraması	50

ÜÇÜNCÜ BÖLÜM
MÜŞTERİ ÜRÜN YORUMLARININ
YAŞAM BOYU DUYGU ANALİZİ İLE SINIFLANDIRILMASINA
İLİŞKİN BİR UYGULAMA

3.1. Veri Kümesi	55
3.2. Veri Ön İşleme	56
3.3. Model Eğitimi İçin Verilerin Hazırlanması	58
3.4. Metin Vektörlerinin Oluşturulması	58
3.5. Araştırma Modeli	59
3.6. Uygulama Çıktıları.....	59
3.7.1. Alan Sıralamasının Model Performansına Etkisi	63
3.7.2. Doğrulama Yöntemlerinin Model Performansına Etkisi.....	66
SONUÇ	70
KAYNAKLAR	73
EKLER.....	82

ŞEKİLLER DİZİNİ

	Sayfa
Şekil 1. Metin Madenciliği Adımları	5
Şekil 2. Metin Madenciliğinin İlişkili Olduğu Disiplinler	6
Şekil 3. Duygu Analizi Yöntemleri.....	17
Şekil 4. Lojistik Fonksiyon Grafiği.....	21
Şekil 5. Klasik Makine Öğrenmesi Paradigması.....	35
Şekil 6. Yaşam Boyu Öğrenme Sistem Yapısı.....	39
Şekil 7. Yaşam Boyu Naive Bayes Sınıflandırması	44
Şekil 8. Yaşam Boyu Lojistik Regresyon Sınıflandırması	45
Şekil 9. Toplama Dayalı Modele Ait Algoritma.....	49
Şekil 10. Alan Sıralamasının Model Performansına Etkisi.....	64
Şekil 11. İzole Modelin ve Toplama Dayalı Modelin Doğruluğu (Alan Sıralaması 1) ..	65

TABLOLAR DİZİNİ

	Sayfa
Tablo 1. Karışıklık Matrisi.....	13
Tablo 2. Türkçe Dilini Ele Alan Duygu Analizi Çalışmalarının Bir Özeti.....	30
Tablo 3. Yaşam Boyu Duygu Sınıflandırmasını Ele Alan Çalışmaların Bir Özeti.....	53
Tablo 4. Yorum Veri Yapısı Örneği	56
Tablo 5. Alan Sıralaması 1.....	60
Tablo 6. Naive Bayes ve Lojistik Regresyon Sınıflandırma Yöntemlerinin Sonuçları (Alan Sıralaması 1)	60
Tablo 7. Toplama Dayalı Yaşam Boyu Duygu Sınıflandırması Model Doğruluğu (Alan Sıralaması 1)	61
Tablo 8. Toplama Dayalı Model Bilgi Tabanı Örneği	62
Tablo 9. Naive Bayes ve Lojistik Regresyon Modellerinin Doğrulama Sonuçları.....	68
Tablo 10. Toplama Dayalı Duygu Sınıflandırması Model Doğrulama Sonuçları	69

SİMGE VE KISALTMALAR DİZİNİ

BERT	Bidirectional Encoder Representations from Transformers
CountVec	CountVectorizer
ÇD	Çapraz Doğrulama
FN	Yanlış Negatif (False Negative)
FP	Yanlış Pozitif (False Positive)
IDF	Ters Belge Frekansı (Inverse Document Frequency)
LLSA	Lifelong Sentiment Analysis
LR	Lojistik Regresyon
NB	Naive Bayes
NLTK	Doğal Dil Araç Takımı (Natural Language Toolkit)
TF	Terim Frekansı
TF-IDF	Terim Frekansı-Ters Doküman Frekansı (Term Frequency-Inverse Document Frequency)
TN	Gerçek Negatif (True Negative)
TP	Gerçek Pozitif (True Positive)

GİRİŞ

Büyük veri çağında internet, cep telefonları, giyilebilir cihazlar, bilgisayarlar ve kayıt ekipmanlarından gelen verilerin miktarındaki ve karmaşıklığındaki hızlı artış; insan davranışlarını, söylemlerini ve etkileşimlerini incelemek için yeni fırsatlar sunmaktadır. Başka bir deyişle, büyük verinin yükselişi ve veri madeninde insanların kolaylıkla göremediği bilgi zenginliği, araştırma sürecinde bir değişimi beraberinde getirmiştir. Bu anlamda veri madeninde gömülü ya da saklı olan bilgi, veri madenciliği yöntemlerinin ya da araçlarının ilginç, anlamlı ve sağlam kalıpları keşfetmesini gerektirmektedir (Shu ve Ye, 2023: 1). Bu açıklamalardan yola çıkarak veri madenciliği, büyük veri kümeleri içindeki ilginç desenleri, modelleri ve diğer tipteki bilgileri ortaya çıkarma süreci olarak tanımlanabilir (Han vd., 2022: 1). Veri madenciliği teknikleri, farklı türdeki araştırma problemlerini uygulamak ve çözmek için kullanılır. Veri madenciliğine ilişkin alanları; metin madenciliği, internet madenciliği, görüntü madenciliği, sıralı desen madenciliği, mekânsal madencilik, tıbbi madencilik, multimedya madenciliği, yapı madenciliği ve grafik madenciliği olarak özetlemek mümkündür (Vijayarani vd., 2015: 7).

Geleneksel veri madenciliği teknikleri, büyük veri içindeki bilgiyi çıkarmak için zaman ve çaba gerektirdiğinden metinsel verileri ele alma becerisine sahip değildir. Bu durum, “metin madenciliği” kavramını ortaya çıkarmıştır. Metin madenciliği, metinsel verilerden bilgiyi keşfetmek için ilginç ve önemli kalıpları çıkarmaya yönelik bir süreç olup bilgi erişimi, veri madenciliği, makine öğrenimi, istatistik ve hesaplamalı dil bilimine dayanan çok disiplinli bir alandır (Talib vd., 2016: 414). Diğer yandan metin madenciliği süreci, veri madenciliği ile aynıdır. Ancak veri madenciliği araçları, yapılandırılmış verileri işlemek için tasarlanmıştır. Metin madenciliği ise e-postalar, HTML dosyaları ve tam metin belgeleri vb. gibi yapılandırılmamış veya yarı yapılandırılmış veri kümelerini işleyebilir (Vijayarani vd., 2015: 7).

Metin madenciliğinin bir alt alanı olan duygu analizi, yazılı metinler aracılığıyla ifade edilen duyguların ve hislerin otomatik olarak tanımlanmasına ve kategorize edilmesine odaklanmaktadır. Duygu analizi; sosyal medyanın katlanarak büyümesi, kamuoyu görüşlerinin ve duyguların artması gibi sebepler ile oldukça önemli bir araç haline gelmiştir (Tan vd., 2023: 16-17). Duygu analizi, yapılandırılmamış veya yarı yapılandırılmış metinlerdeki duyguları olumlu, olumsuz veya nötr olarak yorumlanmasına ve duyguların ne kadar güçlü veya zayıf olduğunun hesaplanmasına yardımcı olur. Günümüzde duygu analizi; iş, finans, siyaset, eğitim, sağlık, acil durum

yönetimi, marka pazarlaması ve ticari değerleme gibi çeşitli alanlarda yaygın olarak kullanılmakta ve araştırmacıların, siyasetçilerin ve iş insanlarının daha etkin kararlar almasına yardımcı olmaktadır (Rodríguez-Ibáñez vd. 2023: 1).

Bu çalışma kapsamında veri madenciliği yöntemlerinden duygu analizi kullanılarak farklı alanlarda yapılan müşteri yorumlarının sınıflandırılması hedeflenmiştir. Duygu analizi çalışmaları genellikle tek bir alana ait veri kümesi kullanılarak yapılmaktadır (Chen ve Liu, 2018: 3-4). Günümüzde hızla artan veri sayısına bağlı olarak ihtiyaç duyulan analiz miktarı göz önünde bulundurulduğunda, tek bir alan için model oluşturmak hem zaman hem de para kaybıdır. Bu aşamada yaşam boyu öğrenme yaklaşımı literatürde yer bulmuştur. Yaşam boyu öğrenme, insanın öğrenme sürecini ele almaktadır. İnsan, yaşam süresi boyunca bilgileri öğrenir ve yeni bir durumla karşılaştığında eski öğrendiği deneyimlerinden faydalananak yeni duruma tepki verir. Yaşam boyu öğrenme insan öğrenmesine benzer şekilde, ardı ardına gelen farklı alanları sırayla öğrenerek yeni gelen alanı, eski alanlardan elde ettiği bilgiyi kullanarak sınıflandırmaya çalışır. Bilgi aktarımı sayesinde her bir alan için denetimli öğrenmede gerekli olan etiketli eğitim veri kümesi ve denetimsiz öğrenmede gerekli olan büyük miktardaki eğitim kümesi gerekliliği ortadan kalkmaktadır (Chen ve Liu, 2018: 1).

Yapılan bu tez çalışmasında, bir online alışveriş sitesinden toplanan müşteri yorumları, bir yaşam boyu öğrenme yaklaşımı altında değerlendirilmiştir. Bu anlamda bu tez çalışmasının temel amacı; makine öğrenmesi yöntemleri aracılığıyla farklı ve birden fazla alandan toplanmış müşteri yorumlarının duygu sınıflandırmasını yapmaktır. Bu temel amaca ek olarak sınıflandırma sürecinde, her bir alan için ayrı sınıflandırma yapmak yerine sürekli öğrenmeyi temel alan bir sınıflandırma modeli kullanarak yaşam boyu duygu sınıflandırması amaçlanmaktadır. Bahsedilen tüm amaçlar gözetilerek farklı alanlardan toplanan müşteri yorumları, 3 farklı duygu analizi yöntemi (yaşam boyu naive Bayes yöntemi, yaşam boyu lojistik regresyon yöntemi ve toplama dayalı yaşam boyu duygu sınıflandırması modeli) ile işlenmiştir. Tüm yöntemlerden elde edilen sonuçlar verilmiş ve yöntemlerin performansları karşılaştırılarak analiz edilmiştir. Ayrıca uygulama bölümünde yaşam boyu duygu sınıflandırması modelinin performansını etkilediği düşünülen iki faktör (alan sıralaması ve doğrulama yöntemleri) detaylı bir biçimde incelenmiş ve elde edilen sonuçlar tartışılmıştır.

Tüm açılardan değerlendirildiğinde yapılan bu çalışmada teorik bilgilere ek olarak uygulamalı bir analiz sunulmuştur. Ele alınan problem, hem yaşam boyu duygu sınıflandırması algoritmalarının performanslarının test edilmesi ve karşılaştırılması

açısından hem de e-ticaret işlemleri için önemli ve günceldir. Diğer yandan bu çalışma, yaşam boyu duygu sınıflandırması kapsamında Türkçe veri kümesi kullanılarak yapılması bakımından bir ilk olma niteliği taşımaktadır. Bu nedenle çalışmanın, Türkçe veri kümesi ya da farklı dillerde oluşturulmuş veri kümeleri kullanılarak yapılacak çalışmalara ışık tutacağı düşünülmektedir.

Bu çalışma, giriş ve sonuç bölümleri hariç üç bölümden oluşmaktadır. *Birinci bölümde* metin madenciliği kavramına değinilmiş ve bu kavram, detaylı bir şekilde açıklanmıştır. Aynı bölümde bu çalışmanın ana konusu olan duygu analizi ve doğal dil işlemeden bahsedilmiştir. Son olarak Türkçe'nin temel alındığı literatür, detaylı bir şekilde taranmıştır. *İkinci bölümde* öncelikle yaşam boyu öğrenme kavramına giriş yapılmış, bu kavrama ilişkin tanımlamalar verilmiş ve yaşam boyu öğrenme ile duygu analizini birlikte ele alan çalışmaların bir özeti verilmiştir. Bu bölümde ayrıca uygulama bölümünde kullanılan sınıflandırma yöntemleri detaylı bir şekilde açıklanmıştır. *Üçüncü bölümde* uygulamaya yer verilmiştir. Sonuç bölümünde ise uygulama sonuçları tartışılmış, çalışmanın kısıtları belirtilmiş ve sonraki çalışmalar için önerilerde bulunulmuştur.

BİRİNCİ BÖLÜM

DUYGU ANALİZİ

Verilerin boyutu, her geçen gün üstel oranlarda artmaktadır. Neredeyse tüm kurum, kuruluş ve iş sektörleri verilerini elektronik olarak saklamaktadır. Bu büyük veri hacminden değerli bilgi çıkarmak için uygun kalıpları ve eğilimleri belirlemek zorlayıcı bir iştir (Talib vd., 2016: 414). Bu anlamda basit bir tanım ile büyük veri yığınlarının belirli işlemlerden geçerek anlamlı hale getirilmesi olarak ifade edilebilen *veri madenciliği*, günümüzde sağlık, eğitim, mühendislik gibi birçok alanda kullanılmaktadır (Artsın, 2020: 344). Veri madenciliği ile eldeki büyük miktardaki veriden gizli kalmış, belirsiz halde bulunan, önceden bilinmeyen ancak potansiyel olarak yararlı bilgi ortaya konmaya çalışılmaktadır (Terzi vd., 2011: 30). Bu işlemler; sınıflandırma, kümeleme ve birliktelik kuralları analizi olmak üzere üç ana başlık altında toplanabilen yöntemler aracılığı ile yapılmaktadır (Yıldırım ve Birant, 2018: 2).

Diğer yandan her şeyin sözcüklerle ifade edildiği ve internet kullanımının arttığı dünyada; forumlar, bloglar, sosyal medya siteleri ve haber siteleri vb. alanlarda da metin verisi hızla artmaktadır. Bu nedenle kişi ve kuruluşlar kararlarında, bu alanlardaki metin verisini kullanmayı tercih etmektedir. Örneğin; kişi bir ürünü satın almak istediğinde, ürün hakkında internetteki halka açık forumlarda çok sayıda kullanıcı incelemesi ve tartışması olduğundan farklı görüşlere kolayca ulaşabilir. Bir kuruluş için kamuya açık bu tür bilgilerin bolluğu nedeniyle, kamuoyunu toplamak için anketler, kamuoyu yoklamaları ve odak grupları yürütmek artık gerekli olmayabilir. Bununla beraber tüm bu bilgilerin toplanması, seçilmesi ve analiz edilmesi giderek zorlaşan bir durum olup bu işlemleri yapabilen sistemlere ihtiyaç duyulmaktadır. Bu açığın kapatılmasına ilişkin hem akademik hem de endüstriyel anlamda duygu analizi ile ilgili çok sayıda çalışma yapılmaktadır.

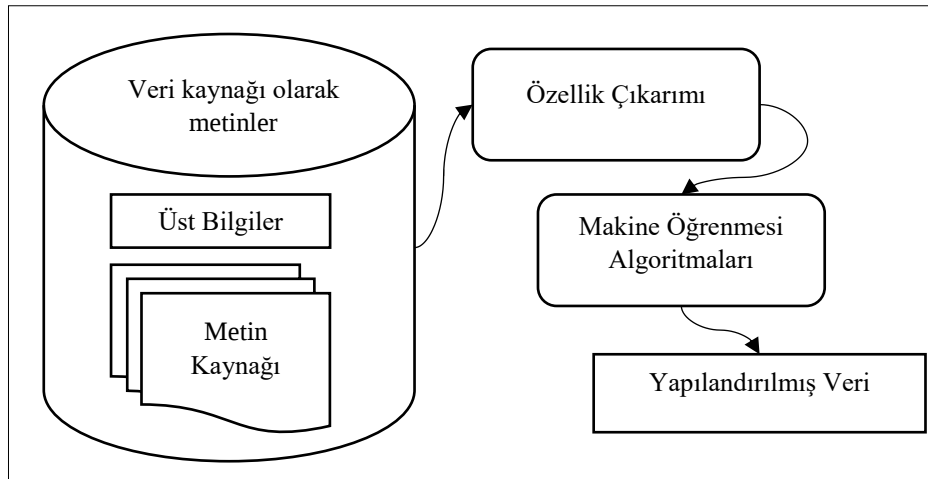
Metin verisi kullanılarak yapılacak çalışmalarda dikkat edilmesi gereken önemli bir nokta, kullanılan dilin özellikleridir. Yapısal anlamda her dil, birbirinden farklılık göstermektedir. Bu farklılıkların analiz edilmesi, metin verisini oluşturan dilin özelliklerinin anlaşılmasını sağlayacak sistemlerin başarı oranını olumlu yönde etkileyecektir. Bu amaçla metin madenciliği çalışmalarında kullanılan doğal dil işleme, her dilin kendi özelliklerinin dikkate alınarak anlaşılması ve incelenmesi çalışmasıdır.

1.1. Metin Madenciliği

Sayısal veri, veri denildiğinde ilk akla gelen veri tipidir. Bunun aksine metin verisi, daha çok üretilmekte ve depolanmaktadır. Kitaplar, makaleler, haberler, blog yazıları, sosyal medya gönderileri vb. veriler metin verisini oluşturmaktadır. Metin verisi, onu oluşturan dile ve dilin özelliklerine bağlı olarak şekillenir. Metin içerisinde bulunan cümlelerin yapısını, kelimeleri, kelime ek ve kök ilişkisini kullanılan dilin özellikleri belirlemektedir. Bu açıdan metin madenciliği, dil bilgisi ve doğal dil işleme ile yakından ilgilidir.

Metin madenciliği geniş anlamda, bir kullanıcının bir dizi analiz aracı kullanarak bir belge koleksiyonuyla zaman içinde etkileşime girdiği, bilgi yoğun bir süreç olarak tanımlanabilir. Metin madenciliği, veri madenciliğine benzer bir şekilde ilginç örüntülerin tanımlanması ve araştırılması yoluyla veri kaynaklarından faydalı bilgiler çıkarmaya çalışır. Bununla birlikte veri kaynakları, belge koleksiyonlarıdır ve belgelerin içinde yer alan yapılandırılmamış veya yarı yapılandırılmış metinsel verilerde ilginç modeller bulunur (Feldman ve Sanger, 2007: 14).

Yapılandırılmamış veya yarı yapılandırılmış veriler, belirli bir formatı ve veri modeli olmayan verilerdir. Metin verileri, görüntü verileri ve video verileri, yapılandırılmamış veri örneklerinden bazılarıdır. Bu tip verilerin, kuruluşların çoğu için değerli bilgilerin % 80'ini temsil ettiği tahmin edilmektedir (Gupta ve Gupta, 2019: 1).

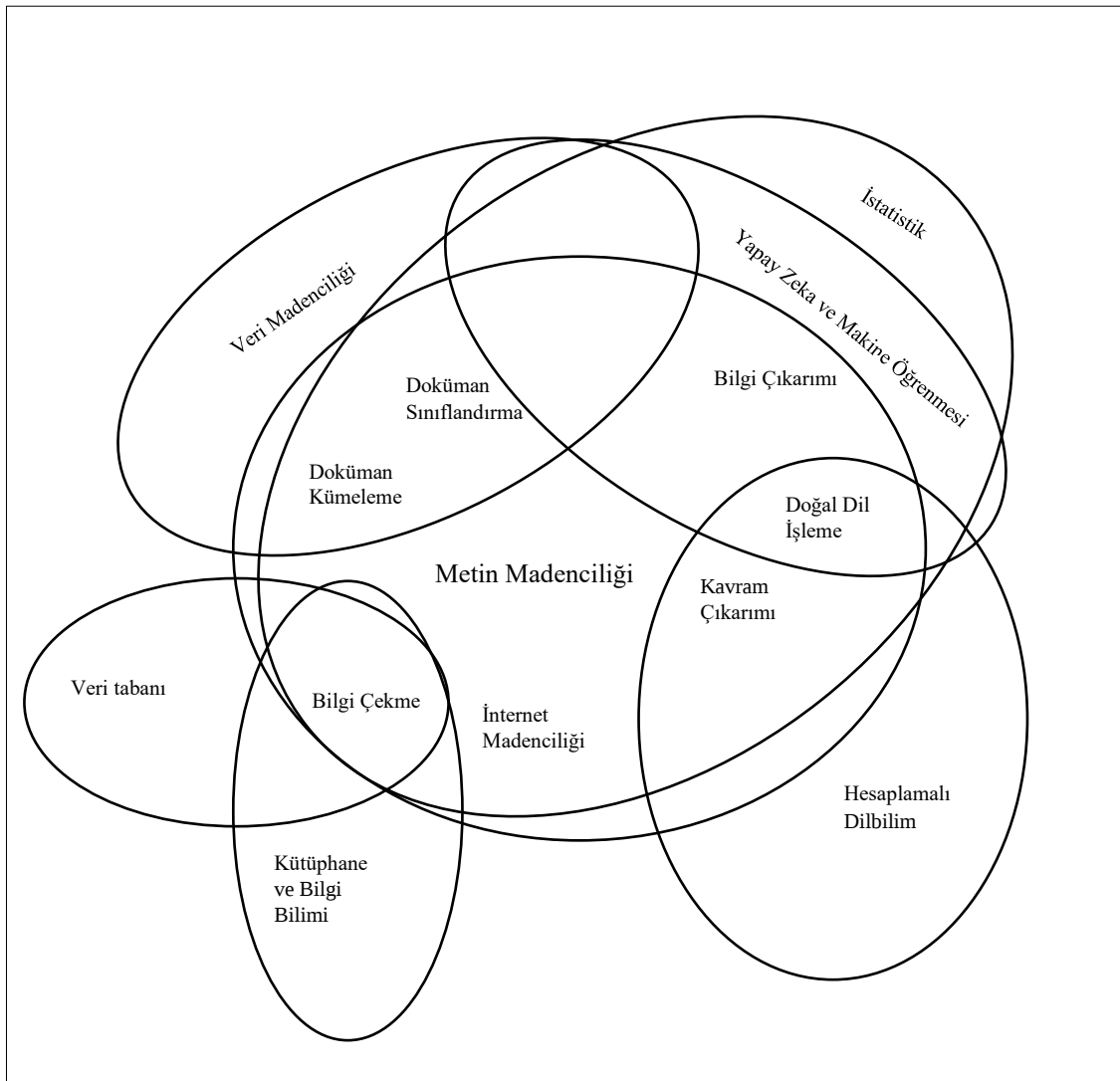


Şekil 1. Metin Madenciliği Adımları

Kaynak: Şeker, 2015: 2

Şekil 1’de genel olarak bir metin madenciliğinin adımları görülmektedir. İlk olarak veriler, bir metin veri tabanından alınır ve özellik çıkarımı işlemleri yapılır. Elde edilen özellikler ile makine öğrenmesi algoritması çalıştırılır ve bu işlemlerin sonucunda yapılandırılmış veri elde edilir (Şeker, 2015: 2).

Metin madenciliği, farklı disiplinler ile ilişkilidir. Şekil 2’de metin madenciliği ile kesişen ana alanlar gösterilmiştir. Metin madenciliğinin doküman sınıflandırma, kavram çıkarma ve doküman kümeleme, internet madenciliği, bilgi çıkarma ve doğal dil işleme, arama ve bilgi alma olmak üzere yedi uygulama alanı bulunmaktadır (Miner vd., 2012: 31).



Şekil 2. Metin Madenciliğinin İlişkili Olduğu Disiplinler

Kaynak: Miner vd., 2012: 31

Şekil 2’de verilen yedi uygulama alanını, şu şekilde özetlemek mümkündür:

- Bilgi çekme (search and information retrieval): Arama motorları ve anahtar kelime araması dahil olmak üzere metin belgelerinin saklanması ve alınmasıdır.
- Doküman kümeleme (document clustering): Veri madenciliği kümeleme yöntemlerini kullanarak terimleri, parçacıkları, paragrafları veya belgeleri gruplandırma ve kategorilere ayırma işlemidir.
- Doküman sınıflandırma (document classification): Etiketli örnekler üzerinde eğitilen modellere dayalı olarak veri madenciliği sınıflandırma yöntemlerini kullanarak parçacıkları, paragrafları veya belgeleri gruplandırma ve kategorilere ayırma işlemidir.
- İnternet madenciliği (web mining): İnternette veri ve metin madenciliği, özellikle internetin ölçeğine ve birbirine bağlılığına odaklanır.
- Bilgi çıkarımı (information extraction): Yapılandırılmamış veya yarı yapılandırılmış metinden ilgili gerçeklerin ve ilişkilerin belirlenmesi ve çıkarılması; yapılandırılmamış ve yarı yapılandırılmış metinden yapılandırılmış veri yapma sürecidir.
- Doğal dil işleme (natural language processing): Düşük seviyeli dil işleme ve anlama görevleridir (örneğin, konuşmanın bir bölümünü etiketleme). Genellikle hesaplama dilbilimi ile eş anlamlı olarak kullanılır.
- Kavram çıkarımı (concept extraction): Kelimelerin ve cümlelerin anlamsal olarak benzer gruplara göre gruplandırılmasıdır (Miner vd., 2012: 32).

Metin madenciliğinde ilk olarak, analizi yapılacak olan metin verisi toplanır. Veri; kitap, gazete, makale, dergi, film ve ürün yorumları, sosyal medya yorumları vb. herhangi bir yazılı kaynaktan toplanabilir. Toplanan metin verilerinin her kaydı, o örneğe atıfta bulunmak için kullanılabilecek benzersiz bir tanımlayıcıya sahip olmalıdır. Metin analizinde bu örnekler, belge (document) olarak bilinir. Bir belge, genellikle birçok karakterden oluşur. Birçok belge, bir belge koleksiyonunu (document collection) veya derlemi (corpus) oluşturur (Anandaraajan vd., 2019: 45). Toplanan metin verileri; metin ön işleme, metin belgesinin modellerle gösterimi ve model performansının değerlendirilmesi aşamalarından geçer. Bu aşamalar, alt başlıklarda detaylı bir şekilde açıklanmıştır.

1.1.1. Metin Ön İşleme Süreci

Toplanan belgelerin analiz edilmeden önce bir dizi işlemten geçirilmesi gerekmektedir. Tüm bu işlemler, metin ön işleme olarak adlandırılmaktadır. Bu anlamda düşünüldüğünde ön işleme, metin madenciliğinde ilk adımdır (Vijayarani vd., 2015: 4). Ön işleme sürecinin amacı; her belgeyi, metni tek tek kelimelere bölen bir özellik vektörü olarak temsil etmektir. Özellik vektöründe bulunan kelimelerin sıklıkları hesaplanarak ilgili dokümanın sınıflandırılması yapılır (Kadhim, 2018: 3). Metin ön işleme teknikleri; tokenizasyon, standardizasyon ve temizleme, durak kelimelerin temizlenmesi ve kök bulma ve lemmalaştırma olarak belirtilebilir.

Tokenizasyon (belirteçleme, dizgeciklere ayırma), incelenen metin belgesinin analiz edilecek en küçük anlamlı birimlere ayrılması işlemidir (Ramachandran ve Parvathi, 2019: 4). Bu analiz birimleri en basit haliyle harf grubu, bir kelime, kelime grubu ya da cümlelerin kendisi olabilir. Örnek olarak bir cümlelerin kelime tokenlerine ayrılmış listesi şu şekilde verilebilir:

Belge 1: Bugün orman ile denizi birleştiren ara yoldan eve geldim.

Token listesi:

Belge 1: [Bugün] [orman] [ile] [denizi] [birleştiren] [ara] [yoldan] [eve] [geldim] [.]

Verilen örnekte, unigram tokenizasyon kullanılmıştır. Tek bir kelime, bir unigram olarak bilinir. Bu yöntem *N-Gram* olarak isimlendirilir. *N*-gramlar tokenlerdir; başka bir deyişle *n* uzunluğunda ardışık kelime dizileridir. Örneğin bigramlar, yan yana iki kelimeden oluşan tokenlerdir. *N*-gramlar, bitişik sözcükleri aynı simgede gruplandırdıkları için sözcüklerin birlikte ortaya çıkışıyla ilgili bilgileri tutar (Anandaraajan vd., 2019: 48).

Standardizasyon, metin belgesinde bulunan terimlerin karşılaştırılabilir hale getirilmesinde kullanılır. Bir belgeye ait tüm kelimelerin küçük harf karakterine sahip olması, standardizasyona bir örnektir. *Temizleme* aşamasında ise tüm kelimeler küçük harfe çevrildikten sonra sayılar, noktalama işaretleri ve özel karakterler belgeden temizlenir (Anandaraajan vd., 2019: 50).

Bir cümlelerin standardizasyon ve temizleme işlemlerinden önceki ve sonraki haline bir örnek, şu şekilde verilebilir:

Standardizasyon ve temizleme işlemlerinden önce:

Belge 1: Bugün orman ile denizi birleştiren ara yoldan eve geldim.

Standardizasyon ve temizleme işlemlerinden sonra:

Belge 1: [bugün] [orman] [ile] [denizi] [birleştiren] [ara] [yoldan] [eve] [geldim]

Durak kelimelerin temizlenmesi aşamasında durak ya da dolgu kelimeleri, metin içerikleriyle ilgisiz olan dil bilgisi sözcükleri olup model verimliliğinin artırılması için kaldırılmalıdır (Jo, 2019: 24). “Ama”, “ancak”, “bu”, “çok”, “ile”, “ne” vb. ifadeler, Türk diline ait durak kelimelere örnek olarak verilebilir.

Standardizasyon ve temizleme aşamalarından geçmiş bir cümle için durak kelimelerin temizlenmesi aşamasına ilişkin bir örnek, şu şekilde verilebilir:

Belge 1: [bugün] [orman] [ile] [denizi] [birleştiren] [ara] [yoldan] [eve] [geldim]

Durak kelimeleri temizlendikten sonra cümle:

Belge 1: [bugün] [orman] [denizi] [birleştiren] [ara] [yoldan] [eve] [geldim]

Kök bulma (stemming), özniteliklerden ekleri (ön ve son ekler) çıkarma işlemidir başka bir deyişle çekimli (veya bazen türetilmiş) kelimeleri, köklerine indirgemek için yapılan işlemdir. Kökün, kelimenin orijinal morfolojik köküyle tanımlanması gerekmez ve genellikle benzer kökle kelime haritası aracılığıyla yeterince ilişkilendirilir. Bu süreç, öznitelik uzayındaki öznitelik sayısını azaltmak ve özniteliklerin farklı biçimleri tek bir öznitelige dönüştürüldüğünde sınıflandırıcının performansını artırmak için kullanılır (Kadhim, 2018: 4).

Örneğin “yaşlı”, “yaştı” ve “yaşlanmak” sözcüklerinin kökü, “yaş” kelimesidir. Örnek olarak incelenen “Belge 1”, kök bulma işlemi sonrası şu şekle dönüşmüştür:

Belge 1: [bugün] [orman] [deniz] [bir] [ara] [yol] [ev] [gel]

Lemmalaştırma (lemmatization), temel olarak sözcüklerin çekimli biçimlerini “otomatik” olarak sözlükte geçen madde başı biçimine dönüştürme/indirgeme işlemi ve bir tür sınıflandırmadır. Lema (lemma) sözcüğü, sözlükte madde başı biçimi olarak açıklanmaktadır (Tahiroğlu, 2021 : 2). Kök bulmada ve genel olarak metin analitiğinde karşılaşılan zorluklardan biri, tek bir kelimenin, kelimenin bağlamına veya konuşmanın bir kısmına bağlı olarak birden fazla anlama sahip olabilmesidir. Lemmalaştırma, kelime köklerini gruplandıran kurallara konuşmanın bir kısmını dahil ederek bu problemle ilgilenir. Bu dahil etme, konuşmanın bölümüne bağlı olarak birden çok anlama sahip kelimeler için ayrı kurallar sağlar (Anandarajan vd., 2019: 56).

Sözcük türü etiketleme (Part-of-Speech Tagging), tokenleri ya da kelimeleri konuşma bölümlerine göre etiketleme işlemidir. Bu aşamada kelimeler, cümledeki görevleri dikkate alınarak gruplandırılır. Başka bir deyişle kelime kökleri bulunduktan sonra devam eden kelimenin türü ve kelimelerin anlamsal içerikleri belirlenmektedir. Sözcük türü etiketi kategori sayısı, dil yapısına göre farklılaşabilmektedir (Değer, 2017: 92).

1.1.2. Metin Belgesinin Gösterimi

Metin madenciliği çalışmalarında belgeler, doğrudan işlenmeyip analiz için uygun biçime getirilir. Ön işleme aşamasında temizlenen ve tokenlerine ayrıştırılan belgeler, bu sayede daha kolay işlenebilir bir hale dönüştürülür. Metin ön işlemenin temel amacı, metin belgelerindeki veri kümelerinden öznitelikler veya anahtar terimleri elde etmek ve kelime ile belge arasındaki uygunluğu ve kelime ile sınıf arasındaki uygunluğu geliştirmektir (Kadhim, 2018: 23). Ön işleme aşamasından sonra belgeler, genellikle özellik vektörleriyle (feature vectors) temsil edilir. Bir öznitelik, özellik uzayında bir boyutu temsil eder. Bir belge ise bir dizi öznitelik ve bunların ağırlıklarıyla vektör uzayında bir vektör olarak temsil edilir (Feldman ve Sanger, 2007: 68). Bu çalışmada vektör uzayı modeli ve kelime çantası modeli olmak üzere iki metin gösterim yöntemi açıklanmıştır.

Vektör uzayı modeli (vector space model), ilk kez Salton vd. (1975) tarafından önerilmiştir. Model, en basit metin gösterim yöntemidir. Model tanımına geçmeden önce kullanılan temel tanımlar şu şekildedir:

Metin: Metin, doğal bir dilde yazılmış cümle veya paragraf grubu olarak tanımlanır (Jo, 2019: 5). Her bir metin belgesi, d ile gösterilir. Terimler, belgeler, kavramlar vb. vektör uzayındaki vektörlerdir (Raghavan ve Wong, 1986: 2).

Terim: Terim, vektör uzay modelindeki en küçük ayrılmaz dil birimidir ve karakterleri, kelimeleri, tümceleri vb. gösterebilir. Vektör uzay modelinde bir metin parçası, bir terimler topluluğu olarak kabul edilir. t_i ; i . terimi göstermek üzere terimler topluluğu $(t_1, t_2, \dots, t_n; i = 1, 2, \dots, n)$ şeklinde ifade edilir (Zong vd., 2021: 32).

Terim ağırlığı: Terimlerin oluşturulmasından sonra vektör uzayı sabitlenir. Ardından metin belgesi, elde edilen özniteliklerin terim ağırlığı hesaplanarak vektör uzayında bir vektör olarak temsil edilir. Terim ağırlıklandırma, her kelimenin ağırlığını önem derecesi olarak hesaplama ve atama sürecini ifade eder. n tane terim içeren metin için, her t terimine belirli ilkelere göre o terimin metindeki önemini ve uygunluğunu gösteren bir w ağırlığı verilir. Bu şekilde metin, karşılık gelen ağırlıkları ile bir terimler koleksiyonu ile temsil edilebilir: $(t_1: w_1, t_2: w_2, \dots, t_n: w_n)$ kısaltılmış haliyle $(w_1, w_2, \dots, w_n; i = 1, 2, \dots, n)$ olarak ifade edilebilir (Jo, 2019: 25-26).

Vektör uzay modelinde terimler kümesinin tasarımı ve terim ağırlıkları hesaplama yöntemleri önemlidir. Bazı yaygın terim ağırlıklandırma yöntemleri şu şekildedir:

- *İkili Ağırlıklandırma*: Bu yöntem, geçerli belgede bir özneliliğin görünüp görünmediğini gösterir. Eğer öznelilik belgede varsa, ağırlık 1; aksi halde 0 değerini alır (Salton ve Buckley, 1988: 6; Lan vd., 2005: 2-3). d belgesindeki ikili ağırlık t_i , Eşitlik 1'deki gibi gösterilir:

$$BOOL_i \begin{cases} 1 & t_i, d \text{ dokümanında bulunuyorsa} \\ 0 & \text{aksi halde} \end{cases} \quad (1)$$

- *Terim Frekansı Ağırlıklandırılması (term frequency, TF)*: Bu ağırlık, geçerli belgedeki bir terimin sıklığını temsil eder. Terim frekansı, Eşitlik 2'de verildiği gibi ifade edilebilir (Manning vd., 2008: 107):

$$tf_i = N(t_i, d) \quad (2)$$

Burada tf_i , i kelimesinin sıklığını; N ise derlemdeki tüm belgelerin sayısını göstermektedir.

- *Ters Doküman Frekansı Ağırlıklandırması (inverse document frequency, IDF)*: *Ters doküman frekansı*, terimlerin derlemdeki önemini yansıtan küresel bir istatistiksel özelliktir. Belge sıklığı (document frequency), derlemde belirli bir terimi içeren belge sayısını belirtir. Bir terimin belge sıklığı ne kadar yüksekse, içerdiği etkili bilgi miktarı o kadar düşük olur. Yani daha az sayıda belgede görülen bir terimin ayırt edici özelliği daha fazladır (Dumais, 1991: 5; Nguyen, 2014: 99). Ters doküman frekansı, Eşitlik 3'teki gibi tanımlanır:

$$idf_i = \log \frac{N}{df_i} \quad (3)$$

Burada df_i , i kelimesinin belge sıklığını; idf_i ise i kelimesinin ters doküman sıklığını göstermektedir.

- *Terim Frekansı - Ters Doküman Frekansı Ağırlıklandırması (term frequency-inverse document frequency, TF-IDF)*: En çok kullanılan terim ağırlıklandırma yöntemidir (Jo, 2019: 25) ve Eşitlik 4'te verildiği gibi gösterilebilir:

$$tf - idf_i = tf_i \cdot idf_i \quad (4)$$

TF-IDF ağırlıklandırma yönteminin ayırt edici özelliği, mevcut belgede sıklıkla ve diğer belgelerde nadiren görünen özellikleri tespit etmesidir (Zong vd., 2021: 35).

Terim ağırlıkları hesaplanmadan önce belgelere metin ön işleme teknikleri uygulanmalı ve belgeden öznitelikler elde edilmelidir. Genellikle belgeyi oluşturan kelimeler, vektör uzay modelinde terim olarak kullanılmaktadır. Bu nedenle terimler, bir kelime çantası olarak görülebilir. Vektör uzay modeli, aynı zamanda kelime çantası (bags of words) modeli olarak da adlandırılır (Zong vd., 2021: 34).

Kelime çantası modeli, metin madenciliğinde kullanılan en temel ve en basit metin belgesi gösterimidir. Bu yöntemde yalnızca belge içinde geçen kelimelerin frekanslarına bakılarak işlem yapılır. Kelimelerin dil bilgisi ve kelime sırası dikkate alınmaz. Bu tür kelime çantası yaklaşımlarının popüleritesi, çok yüksek doğruluk elde etme yeteneğine sahip olmalarına rağmen, esas olarak basitliklerinden ve verimliliklerinden kaynaklanmaktadır. Kelime çantası oluşturulurken belge, bir kelime koleksiyonu olarak ele alınır. Daha sonra bu koleksiyon, belgeyi sınıflandırmak için kullanılır (Barry, 2017: 3).

1.1.3. Modelin Performansının Değerlendirilmesi

Bir makine öğrenmesi modelinin geliştirilmesinin son ve en önemli adımı, performans değerlendirilmesidir. Bu aşamada modelin performansı ve etkinliği farklı ölçütler kullanılarak değerlendirilir. Performans değerlendirme için seçilecek ölçütler, problemin tanımına ve içerdiği veri tipine (metin verisi, sayısal veri, dengeli veya dengesiz veri) göre değişkenlik göstermektedir (Birjali vd., 2021: 18).

Farklı değerlendirme ölçütlerinin altında yatan mekanizmalar farklılık gösterebileceğinden, bu ölçütlerin her birinin tam olarak neyi temsil ettiğini ve ne tür bilgiler aktarmaya çalıştığını anlamak, karşılaştırılabilirlik açısından çok önemlidir. Bu ölçütlere örnek olarak geri çağırma, kesinlik, doğruluk, F-ölçümü verilebilir. Bu ölçütler, karışıklık matrisine dayanmaktadır (Kowsari vd., 2019: 45). Tablo 1`de karışıklık matrisi gösterilmektedir.

Tablo 1. Karışıklık Matrisi

		Gerçek Sınıf	
		<i>Pozitif</i>	<i>Negatif</i>
Sınıf Tahmini	<i>Pozitif</i>	TP (True Positive)	FP (False Positive)
	<i>Negatif</i>	FN (False Negative)	TN (True Negative)

Kaynak: Birjali vd., 2021: 19

Tablo 1’de *TP*, sınıflandırıcı tarafından pozitif olarak doğru tahmin edilen pozitif örneklerin sayısını; *FP*, sınıflandırıcı tarafından pozitif olarak yanlış tahmin edilen negatif örneklerin sayısını; *FN*, sınıflandırıcı tarafından negatif olarak yanlış tahmin edilen pozitif örneklerin sayısını; *TN*, sınıflandırıcı tarafından negatif olarak doğru tahmin edilen negatif örneklerin sayısını temsil etmektedir.

Doğruluk (accuracy), sınıflandırma problemlerinde en sık kullanılan performans ölçütü olup Eşitlik 5’te verilmiştir. Doğruluk, doğru tahmin edilen örneklerin toplam örnek sayısına oranını göstermektedir.

$$\text{Doğruluk} = \frac{TP+TN}{TP+TN+FP+FN} \quad (5)$$

Veri kümesindeki sınıflar neredeyse dengeli olduğunda doğruluk, makine öğrenmesinde duygu sınıflandırması için iyi bir seçimdir (Birjali vd., 2021:19).

Kesinlik (precision), pozitif olarak doğru tahmin edilen örneklerin, tahmin edilen toplam pozitif örnek sayısına oranını temsil etmekte olup Eşitlik 6’da verilmiştir (Dalianis, 2018: 47). Başka bir deyişle kesinlik, tam olma kalitesini ölçmektedir.

$$\text{Kesinlik} = \frac{TP}{TP+FP} \quad (6)$$

Bu metrik, tahminin güvenilir olması gereken problemler için uygundur (Birjali vd., 2021:19).

Hassasiyet (recall), pozitif olarak doğru tahmin edilen örneklerin tüm pozitif örneklere oranı olarak tanımlanmakta olup Eşitlik 7’de verilmiştir (Dalianis, 2018: 47). Hassasiyet, modelin yanlış sınıflandırılmasını ölçer (Birjali vd., 2021: 19).

$$Hassasiyet = \frac{TP}{TP+FN} \quad (7)$$

F1 Skor, bu iki metriğin harmonik ortalamasını hesaplayarak hem hassasiyeti hem de kesinliği ölçmeye yardımcı olmaktadır (Vujovic, 2021: 4). *F1* skoru, Eşitlik 8’de verildiği gibi hesaplanmaktadır.

$$F1\ skor = \frac{2 \times Kesinlik \times Hassasiyet}{Kesinlik + Hassasiyet} \quad (8)$$

1.2. Doğal Dil İşleme

Doğal dil işleme (natural language processing), teoriden ilham alan, insan dillerinin otomatik analizi ve temsili için bir dizi hesaplama tekniğidir (Hirschberg ve Manning, 2015: 1). İnternette artan veri ile doğal dil işlemenin önemi giderek artmakta ve farklı alanlarda kullanımı yaygınlaşmaktadır. Bu kullanım alanları; büyük metinlerin indekslenmesi ve aranması, bilgi alma, metin sınıflandırma, bilgi çıkarma, otomatik dil çevirisi, metinlerin otomatik özetlenmesi, soru cevaplama, bilgi edinimi, metin ve diyalog oluşturmadır (Chowdhary, 2020: 609).

Metin verisi üzerinde yapılan tüm doğal dil işleme çalışmaları ana başlık olarak aynı olsa da metni oluşturan dilin zorluğuna bağlı olarak farklılık göstermektedir. Farklı dillerde oluşturulan metinler, doğal dil işleme aşamasında dilin kendine özgü yapısı gereği farklı uygulamalar ortaya çıkarmaktadır. Bu tez çalışmasında Türkçe metin verileri kullanılarak bir uygulama yapılmıştır.

Türkçe ele alındığında; Moğolca, Tunguzca, Korece ve Japonca ailelerini de içeren Altay dilleri ailesinin Türk ailesine mensup bir dildir. Modern Türkçe, çoğunlukla Türkiye, Orta Doğu ve Batı Avrupa ülkelerinde yaklaşık 60 milyon kişi tarafından konuşulmaktadır. Bazıları nesli tükenmiş yaklaşık 40 dili kapsayan Türk dilleri, çok daha geniş bir coğrafyada 165-200 milyon kişi tarafından ana dil olarak konuşulmaktadır (Ofłazer ve Saraçlar, 2018: 1). Türkçe, sondan eklemeli bir dildir. Tüm ekler, kelimenin sonuna gelerek kelimenin anlamını koruyan ya da farklı bir anlam taşıyan yeni kelimeler oluşturabilir. Türkçe’de ön ek alan ya da cinsiyet bildiren kelimeler bulunmaz. Türkçe’deki kelimeler, cümle içinde kullanıldıklarında çok sayıda çekim ve türetme eki alabilirler. Bu ekler kullanılarak kelimenin kök anlamında farklı yeni kelimeler elde edilebilir. Ayrıca Türkçe’ye; Arapça, Farsça, İtalyanca ve Fransızca gibi birçok farklı

dilden girmiş kelimeler bulunmaktadır. Bu nedenle Türkçe, kelime dağarcığı ve anlam açısından oldukça zengin bir dildir (Tohma ve Kutlu, 2020: 3).

Teknolojinin gelişmesi ile yaygınlaşan çeviri sistemleri, akıllı asistanlar, sohbet robotları günlük hayatta hızlı bir şekilde kullanmaya başlanan doğal dil işleme araçlarıdır. Tüm bu araçların geliştirilmesi; ilgili dilin anlaşılması, işlenmesi ve bir model yardımıyla makinelere öğretilmesi süreçlerini kapsamaktadır. İnsan ile etkileşime girdikçe gelişen bu araçlar, sürekli bir öğrenme süreci içinde performanslarını artırmaktadır (Chen ve Lui, 2018: 3).

Doğal dil işleme ile metin madenciliğinin ortak çalışma alanı olan duygu analizi, son yıllarda popüler olan diğer bir çalışma alanıdır.

1.3. Duygu Analizi

Duygu analizi (sentiment analysis), metinde ifade edilen fikirlerin, duyguların ve anlamların çıkarılması çalışmasıdır (Liu, 2010: 3). Duygu analizi; literatürde fikir madenciliği, duygu durum analizi, duygu sınıflandırma, özellik analizi ve yorum madenciliği gibi farklı isimlerle de yer almaktadır. Özellikle akademik çalışmalarda sıklıkla duygu analizi ya da fikir madenciliği olarak yer almaktadır. Bu tez çalışmasında yaygın olarak duygu analizi adı kullanılacaktır.

Eski çağlardan beri doğru karar verme, insan yaşamının en önemli olgusudur. İnsanlar, bilgilerden ve tecrübelerden yararlanarak bir görüş edinmekte ve bu görüşler, tüm faaliyetlerinin merkezinde yer almaktadır. Fakat insanlar, karar vermek gerektiğinde sadece kendi görüşlerine değil başkalarının görüşlerine de ihtiyaç duymaktadır. Örneğin bir firma ürettiği ürün ile ilgili geliştirme çalışmalarında müşterilerinin görüşlerine ihtiyaç duymaktadır. Ya da bir müşteri için, ilk defa satın alacağı bir ürün ile ilgili daha önceden o ürünü kullanmış kişilerin görüşleri oldukça önemlidir. Benzer şekilde bir topluluğun siyasi görüşünün belirlenmesi, gelecek seçimlerde hangi partinin kazanacağını belirlenmesinde oldukça önemlidir. Belirtilen bu durumlar, duygu analizine bir örnek olup duygu analizi ile ilgili farklı uygulamalara rastlamak mümkündür. D'Andrea vd. (2015), duygu analizi uygulamalarını şu şekilde sınıflandırmıştır:

- İş hayatı: Tüketici sesi, marka itibarı analizi, çevrimiçi reklamcılık, çevrimiçi ticaret.
- Politika: Oy tavsiyesi uygulamaları, politikacıların pozisyonlarının netleştirilmesi.

- Kamu eylemleri: Gerçek dünya olaylarını izleme, hukuki konular, politika veya hükümet düzenleme önerileri, akıllı ulaşım sistemleri.
- Finans: Emtia fiyatları ve hisselerin gelişimi, finansal riskin bireyselleştirilmesi (D`Andrea vd., 2015: 6).

Duygu analizi ile ilgili ilk çalışmalar, 2000`li yılların başlarına rastlamaktadır. İnternet kullanımının artmasıyla kişilerin görüşlerini paylaşabildiği sosyal platformlar ortaya çıkmış ve duygu analizinde kullanılabilecek büyük verilere ancak ulaşılmıştır. Duygu analizinin başlangıcı ve hızlı büyümesi, sosyal medya ile örtüşmektedir. Bu nedenle, duygu analizi araştırmaları sadece doğal dil işleme üzerinde önemli bir etkiye sahip olmakla kalmamakta, aynı zamanda insan görüşlerinin yoğun olduğu yönetim bilimleri, siyaset bilimi, ekonomi ve sosyal bilimler üzerinde de ciddi bir etkiye sahiptir (Liu, 2012: 8).

Duygu analizinde; doküman seviyesinde, cümle seviyesinde, özellik/görüş seviyesinde sınıflandırma olmak üzere üç farklı sınıflandırma yaklaşımı vardır:

Doküman seviyesinde sınıflandırma: Dokümanların kendi içlerinde duygularını etiketleme ile ilgilidir. Doküman düzeyinde sınıflandırmada belgenin tamamı, ya olumludur ya da olumsuzdur (Kharde ve Sonawane, 2016: 10).

Cümle seviyesinde sınıflandırma: Cümle seviyesindeki duygu analizi, cümleleri kendi duyarlılık kutuplarıyla etiketleme ile ilgilidir. Cümle düzeyi duyarlılık sınıflandırması cümleyi; pozitif, negatif veya nötr sınıf içinde sınıflandırabilir (Kharde ve Sonawane, 2016: 10).

Özellik/Görüş seviyesinde sınıflandırma: Bir metin belgesinde tamamıyla tek bir görüş belirtilmeyip hem olumlu hem de olumsuz görüş aynı anda bulunabilmektedir. Görüş seviyesinde sınıflandırma ile metnin hangi görüşleri belirttiği, hangi bakış açısına göre olumlu veya olumsuz olduğu incelenmektedir (Şeker, 2016: 21).

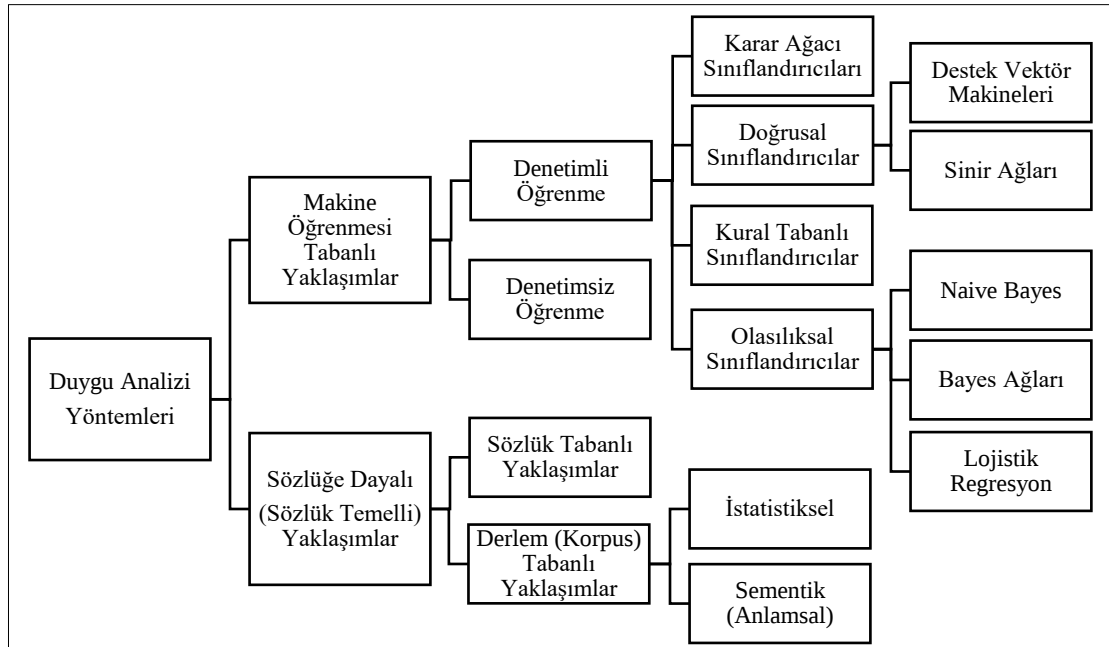
Duygu analizi uygulamalarında modelin performansını etkileyen farklı zorluklar ile karşılaşmaktadır. Bu zorlukları şu şekilde özetlemek mümkündür:

- Büyük miktarlardaki verilerde gürültü, eksik veri ve tutarsızlık gibi etkenler analizlerin doğruluğunu etkilemektedir. Duygu analizi sınıf tahmininin doğruluğu, metin belgesindeki bu hatalar düzeltilerek iyileştirilebilir (Wankhade vd., 2022: 39).
- Belirli bir dil veya kültürden gelen verilerle eğitilen duygu analizi modelleri, diğer dillerden veya kültürlerden gelen metinlerde ifade edilen duyguları doğru bir şekilde yakalayamayabilir. Bunun nedeni dil, deyimler ve ifadelerin kullanımının kültürler

arasında büyük ölçüde değişebilmesi ve bir kültürde eğitilen bir duygu analizi modelinin, başka bir kültürdeki duyguyu analiz etmede etkili olmayabilmesidir (Tan vd., 2023: 16).

- Duygu analizi algoritmaları, genellikle insanların bir metnin duygusunu elle etiketlediği açıklamalı verilere dayanmaktadır. Ancak büyük bir veri kümesini toplamak ve etiketlemek, zaman alıcı olabilmekte ve yoğun kaynak gerektirebilmektedir (Liu, 2010:13; Xu vd., 2022: 11).
- Bir duygu analizi modelini eğitmek için kullanılan eğitim verileri, önyargılı olabilir ve bu önyargı, modelin doğruluğunu ve yeni verilere genellemesini etkileyebilir. Örneğin, ağırlıklı olarak bir cinsiyet veya ırktan metinlerden oluşan bir veri kümesi, o gruba karşı önyargılı bir modele yol açabilir ve bu da diğer gruplardan gelen metinler için yanlış tahminlerle sonuçlanabilir (Tan vd., 2023: 17) .

Duygu analizi uygulamalarında kullanılan yöntemler, sözlük tabanlı yöntemler ve makine öğrenmesi tabanlı yöntemler olarak iki temel başlıkta ele alınabilir. Literatürde bulunan bazı çalışmalarda, bu iki yöntemin birlikte ele alındığı hibrid yöntemlere rastlamak mümkündür. Şekil 3'te duygu analizi yöntemleri, genel bir ayırım çerçevesinde verilmiştir. Bu tez kapsamında bu iki temel başlık ele alınmıştır.



Şekil 3. Duygu Analizi Yöntemleri

Kaynak: Medhat vd., 2014: 1095

1.3.1. Sözlüğe Dayalı (Sözlük Temelli) Yaklaşımlar

Sözlük tabanlı yaklaşımlar, duygu analizi için fikir kelimelerinin bir fonksiyonuna dayanmaktadır. Fikir kelimeleri, olumlu ya da olumsuz duyguları ifade etmek için yaygın olarak kullanılmaktadır. Fikir kelimelerine örnek olarak, "iyi" ve "kötü" kelimeleri verilebilir. Yaklaşım, genel olarak duygu yönelimini tanımlamak ve belirlemek için bir fikir kelimesi sözlüğü kullanır (pozitif, negatif veya nötr). Sözlüğe, fikir sözlüğü adı verilir. Bu yaklaşım metni belge, cümle veya varlık düzeyinde analiz etmek için kullanılabilir. Sözlüğe dayalı yaklaşım, olumlu ya da olumsuz duyguları ifade eden sözcükler olan fikir (ya da duygu) sözcüklerine dayanır. İstenen bir durumu kodlayan kelimeler (örneğin, büyük ve iyi) pozitif bir kutupluluğa sahipken, istenmeyen bir durumu kodlayan kelimeler negatif kutuplara sahiptir (örneğin, kötü ve korkunç) (Zhang vd., 2011: 2). Sözlük tabanlı yaklaşımlar, sözlük tabanlı ve derlem (korpus) tabanlı yaklaşımlar olmak üzere ikiye ayrılabilir:

- *Sözlük tabanlı yaklaşım*, bağlamsal duygu yöneliminin her bir kelimenin veya ifadenin duygu yöneliminin toplamı olduğu varsayımına dayanır (Palanisamy vd., 2013: 1-2). Öncelikle duygu kelime listelerini derleyerek bir duygu sözlüğü oluşturulur, ardından sözlükte tanımlanan olumlu ve olumsuz göstergelere dayanarak belgenin kutupluluk puanı belirlenir.
- *Derlem (korpus) tabanlı yaklaşım*, bir metnin içeriğine özgü duygu kelimelerini bulmayı amaçlar. Bu yaklaşımda büyük bir metin verisindeki özel duygu kelimelerini bulmak için duygu kelimelerinin bir kök listesi ile ortaya çıkan kalıplara bakılır. Bir dildeki tüm kelimeleri kapsayacak kadar büyük bir dizin elde etmek oldukça zor olduğu için derlem tabanlı yaklaşımının tek başına kullanımı, sözlük temelli bir yaklaşım kadar etkili değildir (Shehu, 2019: 22-23). İstatistiksel yaklaşım veya semantik yaklaşım olmak üzere iki başlık altında incelenmektedir.

1.3.2. Makine Öğrenmesi Tabanlı Yöntemler

Makine öğrenmesi tabanlı yöntemlerde, metnin sınıflandırılması için makine öğrenmesine dayalı sınıflandırma tekniği kullanılır. Denetimsiz ve denetimli öğrenme olmak üzere başlıca iki tür makine öğrenmesi tekniği vardır:

- *Denetimsiz öğrenme*: Denetimsiz öğrenmede sistem eğitilirken etiketsiz veriler kullanılmaktadır (Özdeş, 2017: 33). Bu nedenle denetimsiz öğrenme, etiketli veri kümesinin bulunmasının zor olduğu durumlarda kullanılır (Shehu, 2019: 15). Bu tür öğrenme biçiminde amaç, tanıma ve sınıflandırma olmayıp kümeleme, olasılık

yoğunluk tahmini ve boyut indirgemedir. Denetimsiz öğrenme algoritmalarının performansı, denetimli öğrenme algoritmalarına göre daha düşüktür. Ancak denetimsiz öğrenme algoritmaları, denetimli öğrenme algoritmalarında karşılaşılan alan bağımlılık problemini ortadan kaldırmaktadır (Özdeş, 2017: 33).

- *Denetimli öğrenme:* Denetimli yöntemlerde algoritmalar öncelikle etiketlenmiş bir eğitim verisi ile eğitilir, bu eğitimle algoritmaya öğrenme yetisi kazandırılmış olur. Öğrenen algoritma yeni gelen sınıf etiketi belli olmayan verileri, karar mekanizması kullanarak sınıflandırır (Aydın vd., 2018: 4). Duygu analizinde temel olarak denetimli öğrenme teknikleri uygulanmaktadır. Bir makine öğrenmesi tabanlı yöntemde iki veri kümesine ihtiyaç vardır:

- Eğitim verisi
- Test verisi (Kharde ve Sonawane, 2016: 4)

Denetimli öğrenme genellikle naive Bayes, lojistik regresyon, yapay sinir ağları, destek vektör makineleri, karar ağaçları ve k-ortalama sınıflandırma yöntemleri kullanılmaktadır.

Bu tez çalışmasının uygulama bölümünde, denetimli makine öğrenmesi tabanlı yöntemlerden naive Bayes ve lojistik regresyon olasılıksal sınıflandırıcılar kullanılmıştır. Naive Bayes yöntemi, bu tez çalışması kapsamında ele alınan yaşam boyu duygu sınıflandırması için oldukça uygundur. Başka bir deyişle, geçmiş alanlardan elde edilen bilgiler, naive Bayes sınıflandırma yöntemi başlangıç sınıf olasılıkları olarak kabul edildiğinde, yeni görevin tahmini için eğitilen modelin öncül olasılıkları olarak kullanılabilir (Chen ve Liu, 2018: 48). Lojistik regresyon yöntemi, ikili metin sınıflandırma problemlerinde kullanılan oldukça popüler bir yöntemdir (Al-Amrani vd., 2017: 1). Duygu analizi çalışmalarında; sosyal medya verileri, müşteri davranışları ve film yorumları gibi veri kaynakları kullanılarak, son yıllarda hala çalışılan (örneğin; İrfan vd. (2025), Montaser vd. (2025)) bir sınıflandırma yöntemidir. Yaşam boyu duygu sınıflandırması model performanslarının karşılaştırılması amacıyla bu iki istatistiksel sınıflandırma yöntemi tercih edilmiş ve aşağıdaki başlıklarda sadece bu yöntemler açıklanmıştır.

1.3.2.1. Naive Bayes yöntemi

Naive Bayes sınıflandırma yöntemi, özellikler arasında güçlü (naif) bağımsızlık varsayımlarıyla Bayes teoreminin uygulanmasına dayanan olasılıksal bir sınıflandırıcıdır (Vembandasamy vd., 2015: 442).

Naive Bayes yönteminin amacı; eğitilmiş bir veri kümesi ve bu veri kümesine ait değişken vektörü verildiğinde, gelecek etiketsiz verinin sınıflandırılmasını sağlamaktır. Problem, bir denetimli sınıflandırma problemidir. Naive Bayes yöntemini oluşturmak ve uygulamak kolaydır. Büyük veri kümelerine kolayca uygulanabilir. Yorumlanması kolaydır, bu nedenle karmaşık sınıflandırma algoritmalarının aksine sınıflandırmanın nasıl yapıldığının anlaşılması oldukça kolaydır. Model sonuçları genellikle iyidir. Model performansı olarak en iyi sınıflandırıcı olmasa bile güvenilir bir performans göstermektedir (Wu vd., 2008: 24-25).

Eşitlik 9'da Bayes kuralı verilmiştir. Burada X , bir gözlem ve C , bu gözlemin bulunabileceği sınıf kümesini göstermektedir. X gözleminin hangi sınıfta olduğu, Bayes kuralı ile tahmin edilir. En yüksek koşullu olasılık değeri, X gözleminin sınıfını vermektedir (Taheri ve Mammadov, 2013: 788).

$$P(C|X) = \frac{P(C)P(X|C)}{P(X)} \quad (9)$$

Naive Bayes sınıflandırıcısında, verilen sınıfta $X_1, X_2, \dots, X_n$ özelliklerinin birbirinden bağımsız olduğu varsayımı altında formül, Eşitlik 10'daki gibidir (Taheri ve Mammadov, 2013: 788).

$$P(C|X) = \frac{P(C) \prod_{i=1}^n P(X_i|C)}{P(X)} \quad (10)$$

Bu formül ile bir X gözleminin hangi sınıfa ait olduğunun olasılığı hesaplanmaktadır.

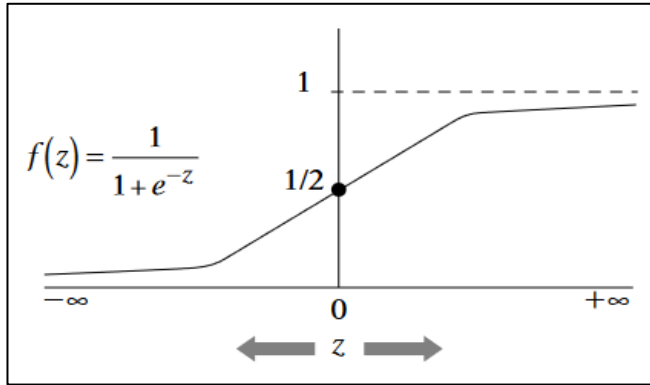
1.3.2.2. Lojistik regresyon yöntemi

Regresyon yöntemleri, bir bağımlı değişken ile bir veya daha fazla bağımsız değişken arasındaki ilişkiyi tanımlamakla ilgilenen herhangi bir veri analizinin ayrılmaz bir parçası haline gelmiştir. Bağımsız değişken niceldir. Çoğu zaman bağımlı değişken, ayrıktır ve iki veya daha fazla olası değer alır. Lojistik regresyon modeli, bu verilerin

analizi için en sık kullanılan regresyon modelidir (Hosmer vd., 2013: 1). Bağımlı değişken ikili değer alıyorsa lojistik regresyon, bu tür verilere uygulamak için iyi bir sınıflandırma algoritmasıdır (Komarek, 2004: 22). Lojistik regresyonun basitliği ve birlikte çalışabilirliği, zaman zaman topluluk öğrenmesi veya destek vektör makineleri gibi diğer karmaşık doğrusal olmayan modellerden daha iyi performans göstermesine yol açabilir (Kirasich vd., 2018: 8).

Lojistik regresyon, matematiksel olarak logit fonksiyonuna dayanmaktadır. Bu fonksiyon, Eşitlik 11’de gösterilmektedir. Şekil 4’te lojistik fonksiyon grafiği görülmektedir. Fonksiyonda z değişkeni, $-\infty$ ve $+\infty$ arasında değer almaktadır. Model, her zaman 0 ile 1 arasında bir sayı olan bir olasılığı tanımlamak için tasarlanmıştır. z değişkeninin alacağı değerler sonucunda fonksiyon s şeklinde bir Sigmoid fonksiyondur (Kleinbaum vd., 2002: 6).

$$f(z) = \frac{1}{1+e^{-z}} \quad (11)$$



Şekil 4. Lojistik Fonksiyon Grafiği

Kaynak: Kleinbaum vd., 2002: 5

Lojistik modelde z , doğrusal regresyon denklemdir. z , y bağımlı değişkenini göstermektedir.

$$y = \beta_0 + \beta_1 \cdot x_1 + \beta_2 x_2 + \dots + \beta_k x_k \quad (12)$$

Eşitlik 12’de β_0 , kesişim noktası veya sapma terimidir. $\beta_1, \dots, \beta_k$; katsayı veya ağırlıkları ifade etmektedir. $X = x_1, x_2, \dots, x_k$ olmak üzere, X , metin verisi vektörünü

göstermektedir. Bu çalışma kapsamında her bir alan için toplanan ürün yorumuna ait kelime vektörüdür.

Sınıflandırma probleminin sonucu iki sınıftan birine ait olabilir; $y = 1$ ya da $y = 0$ 'dır. Her iki durum için lojistik model tahminleri, Eşitlik 13 ve Eşitlik 14'te verilmiştir (Kirasich vd., 2018: 9).

$$p(y_i = 1|x_i) = 1 - \frac{1}{1 + e^{\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k}} \quad i = 1, 2, \dots, k \quad (13)$$

$$p(y_i = 0|x_i) = \frac{1}{1 + e^{\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k}} \quad i = 1, 2, \dots, k \quad (14)$$

Lojistik regresyon sınıflandırma algoritmasında pozitif sınıfa ait olma olasılığı hesaplanır. Problem gerçek değer ile tahmin edilen değerler arasındaki farklara ilişkin Logloss değerini diğer bir adıyla kayıp fonksiyonunu minimum yapabilecek ağırlıkları bulma yöntemiyle çözülür. Eşitlik 15'te kayıp fonksiyonu verilmiştir.

$$Logloss = -\frac{1}{m} \sum_{i=1}^m p_i \log(y_i) + (1 - p_i) \log(1 - y_i) \quad i = 1, 2, \dots, m \quad (15)$$

Burada m , veri kümesindeki toplam örnek sayısını; p_i , i . metin verisi gerçek sınıf etiketini göstermektedir. Amaç, gradyan inişi veya stokastik gradyan inişi gibi optimizasyon teknikleri kullanılarak tanımlanan eğitim kümesi üzerinde negatif logaritmik olasılığın en aza indirildiği bir ağırlık kümesi bulmaktır (Ng ve Ma, 2023: 22). Negatif logaritmik olasılığı en aza indirmek aynı zamanda doğru sınıfı seçme olasılığını en üst düzeye çıkarmak anlamına gelir. Test verisinin gerçek sınıf etiketi ve tahmin edilen sınıf etiketi arasındaki farkı ölçen kayıp fonksiyonuna çapraz entropi denir (Kirasich vd., 2018: 10).

1.4. Türkçe'yi Konu Alan Duygu Analizi Çalışmalarına İlişkin Literatür Taraması

Duygu analizi ile ilgili yapılan çalışmalara bakıldığında genellikle İngilizce veri kümelerinin sıklıkla kullanıldığı görülmektedir. İngilizce'nin uluslararası dil olarak sıklıkla kullanılması buna bağlı olarak İngilizce veri kümelerine kolaylıkla erişim sağlanması, bunun en önemli nedenidir. Bununla birlikte giderek farklı dilleri kapsayan çalışmalar da literatürde yerini almaktadır. Bu dillerden bir tanesi de Türkçe'dir. Doğal dil işleme alanının bir parçası olan duygu analizi, sınıflandırılmak istenen verinin ait

olduğu dil farklılaştıkça dilin sahip olduğu özellikler göz önüne alınarak farklı ön işleme ve analiz süreçleri gerektirmektedir. Bu kapsamda bu bölümde bu çalışmanın konusunu oluşturan Türkçe duygu analizi konusu içerisinde Türk dilini ele alan duygu analizi çalışmaları incelenmiştir. Literatürde bulunan çalışmaları; kullanılan yöntem, veri kümesi, terim doküman matrisi oluşturma, kelime dizgicikleme vb. alanlara ayırmak mümkün olabilmektedir. Bu tez çalışması kapsamında literatürde bulunan çalışmalar ilk olarak yöntem bazında ele alınmış ve her çalışma, ayrıntılı olarak incelenmiştir.

Türkçe duygu analizini konu alan ve bilinen ilk çalışma, Eroğul (2009) tarafından yapılmıştır. Çalışmada, Türkçe film yorumlarından oluşan veri kümesi kullanılarak İngilizce veri kümelerine uygulanan duygu analizi yöntemlerinin denenmesi hedeflenmiştir. Ayrıca Türkçe veri kümesine özgü yöntemler önerilmiştir. Film yorumlarının sınıflandırılmasında, denetimli öğrenme yaklaşımlarından destek vektör makineleri sınıflandırma yöntemi kullanılmıştır.

Çetin ve Amasyalı (2013), telekom sektöründe faaliyet gösteren iki farklı firmaya ait tweetleri ele almıştır. Bunlar; pozitif, negatif ve nötr olmak üzere 3 sınıfa ayrılmıştır. Eğitim kümesi oluşturulurken dengeli bir veri kümesi oluşturulmuştur. Metinlerin temsil edilmesi için terim frekansı ve ters doküman matrisleri kullanılmıştır. Bu yöntemlerdeki parametreler değiştirilerek denetimli öğrenme algoritmaları ile kurulan modellerin performansları karşılaştırılmıştır.

Türkçe ile ilgili yapılan duygu analizi çalışmalarında kitap yorumları, film yorumları ve ürün yorumlarını içeren farklı veri kümeleri kullanılarak yapılan çalışmalara örnek olarak Çelik (2013) tarafından yapılan çalışma verilebilmektedir. Çalışmada lojistik regresyon, naive Bayes, destek vektör makineleri ve karar ağaçları sınıflandırma yöntemleri kullanılarak farklı veri kümelerinde model performansları karşılaştırılmıştır. Nanğır'ın (2013) yaptığı çalışmada ise aynı veri kümesi kullanılarak Türk dili için çoklu sınıflandırıcı makine öğrenmesi algoritmaları uygulanmıştır. Denetimli öğrenme algoritmaları kullanılarak sınıflandırıcı birlikteliği (ensemble of classifiers) kavramından yararlanılmış, çoklu sınıflandırıcılar yardımıyla sınıflandırma performansının artırılması hedeflenmiştir. Bu özgün sınıflandırıcı yaklaşımının yanı sıra, makine öğrenmesi algoritmalarının parametre optimizasyonu gerçekleştirilerek performans artırılmıştır.

Farklı alanların kullanıldığı başka bir çalışmada Meral ve Diri (2014), Twitter üzerinden dokuz farklı alanla ilgili toplanan iletileri olumlu, nötr ve olumsuz olmak üzere üç sınıfa ayırmıştır. Eğitim veri kümesi, gönüllü kişilerin iletileri okuyup etiketlemesi ile oluşturulmuştur. Verilerin işlenmesinde kelime tabanlı ve n-gram tabanlı dil işleme

yöntemleri kullanılmıştır. Öznitelik uzayının büyük olmasını engellemek için korelasyon tabanlı öznitelik seçimi (correlation-based feature selection) kullanılmıştır.

Güran, Uysal ve Doğrusöz'e (2014) ait çalışmada kullanılan veri kümeleri; Twitter verilerinden, blog yazılarından ve film yorumlarından oluşmaktadır. Tweetler; olumlu, olumsuz ve nötr sınıflara, blog yazıları; neşeli, sinirli, üzgün, karışık ruh hali belirten durumlara ve film yorumları ise olumlu, olumsuz ve nötr sınıflara ayrılmıştır. Çalışmada, destek vektör makineleri sınıflandırma yöntemi parametrelerinin optimize edilmesi amaçlanmaktadır. Öznitelikler; 2-gram, 3-gram, 4-gram ve 5-gramlar birlikte kullanılarak oluşturulmuştur. Korelasyon tabanlı öznitelik seçimi kullanılarak öznitelik sayısı azaltılmıştır.

Akba (2014), film yorumlarından oluşturulmuş veri kümesini kullanmıştır. Kullanılan veri kümesi üzerinde farklı öznitelik seçme yöntemlerinin, Türkçe metinlerdeki ayrıştırıcı terimlerin tespitine ve destek vektör makineleri sınıflandırma yöntemi performansına etkisi incelenmiştir. Veri kümesi, yorum puanları ile etiketlenmiştir. 0,5-5 arasında puanlanan yorumlar ele alındığında 0,5-2 arası olumsuz, 2,5-3,5 arası nötr, 4-5 arası olumlu olarak etiketlenmiştir. Çalışmada, sadece olumlu-olumsuz sınıflar kullanılmıştır.

Duygu analizi sınıflandırması denetimli öğrenme algoritmalarının yanı sıra denetimsiz öğrenme algoritmaları kullanılarak da yapılabilmektedir. Uçan'a (2014) ait çalışmada, Türkçe film ve otel yorumlarından oluşturulan veri kümeleri kullanılarak sözlük tabanlı bir sınıflandırma gerçekleştirilmiştir. Sözlüksel benzeşim yöntemiyle Türkçe duygu sözlüğü oluşturulmuş ve ele alınan konu için bağımsız, kolay ve hızlı bir şekilde duygu analizi yapılması hedeflenmiştir.

Farklı alanlara ait veri kümelerinin kullanıldığı sınıflandırma çalışmalarında beklenen başarı sağlanamamaktadır. Örneğin, otomotiv alanındaki yorumların (hedef alan) duygu sınıflandırmasında, etiketli film yorumları (kaynak alan) ile eğitilmiş bir sınıflandırıcının kullanılması sınıflandırma başarısını oldukça düşürmektedir. Bu durumun nedeni, alan adaptasyonu problemi başka bir deyişle kaynak ve hedef alanlar arasındaki öznitelik dağılımlarının farklı olmasıdır. Alan adaptasyonuna ilişkin Pak (2015) tarafından yapılan çalışmada yeni bir yaklaşım geliştirilmiştir. Bu yaklaşımın varsayımı, birden fazla kaynak alana ait etiketli eğitim verisi olduğu ve sınıflandırılacak hedef alan verisinin hangi alana ait olduğunun bilinmediğidir. Bu durumda, iki yaklaşım önerilmiştir: İlk yaklaşım; kaynak alan olarak film, kitap ve bilgisayar alanlarına ait tüm eğitim verilerinin birleştirilmesi ve eğitilmesi, hedef verinin birleştirilen sınıflandırıcı ile

sınıflandırılmasıdır. İkinci yaklaşım ise alanı bilinmeyen hedef verinin öncelikle mevcut kaynak alanlardan hangisine ait olduğunu belirlenmesi ve belirlenen alanın etiketli eğitim verisi ile hedef veriyi sınıflandırılmasıdır. Çalışmada N adet kaynak alan, içerdiği etiketli eğitim verileri kullanılarak eğitilmiş ve N adet duygu sınıflandırıcısı oluşturulmuştur. Bu duygu sınıflandırıcıları ile alanı belirlenen verinin, pozitif ve negatif sınıflarından hangisine ait olduğu belirlenmiştir.

Ekinci, Özcan ve Amasyalı (2016) tarafından yapılan çalışma, Türkçe metinlerde zaman etkisini ele alan ilk çalışmadır. Bu kapsamda TTNET ve Turkcell ile ilgili 2012, 2013, 2014 ve 2015 yıllarında toplanan Türkçe tweetler kullanılmıştır. Bu tweetlerden 2012 yılına ait veriler eğitim verisi, diğer yıllara ait veriler ise test verisi olarak ayrılmıştır. 3-gram ve Zemberek kütüphanesi ile elde edilen kelime kökleri öznitelik olarak kullanılmıştır. Her iki veri kümesi üzerinde de zaman etkisini gözlemleyebilmek için farklı deneyler tasarlanmıştır. İlk olarak iki veri kümesi içinde tweetler yıllara göre kendi içlerinde sınıflandırılmıştır. Bu sınıflandırmalardan en iyi sonuç veren sınıflandırma yöntemleri, destek vektör makineleri ve rastgele ormandır. Çalışmaya bu sınıflandırıcılarla devam edilmiştir. 2012 yılına ait veri kümesi, eğitim verisi olarak kullanılmış ve diğer yıllar, belirlenen bu iki yöntemle ayrı ayrı sınıflandırılmıştır. Ayrıca çalışmada zaman etkisini azaltmak için aktif öğrenme metotları ve bu metotlar ile daha az eğitim verisi ile sınıflandırma sonuçlarının iyileşmesi amaçlanmıştır. Bununla ilgili olarak detaylı bir yapı oluşturulmuştur. Oluşturulan eğitim verisi senaryosuna göre yıllara göre sınıflandırıcı sonuçları hesaplanmıştır.

Oğul ve Ercan (2016), Türkçe otel yorumlarını kullanarak duygu analizi gerçekleştirmiştir. Olumlu-olumsuz iki sınıftan oluşan veri kümesinde, öznitelik ağırlıklandırma yöntemi olarak doküman terim matrisi ve TF-IDF kullanılmıştır. Bu iki yönteme ek olarak, sözlük tabanlı sınıflandırma yöntem olarak önerilen ve bu iki ağırlıklandırma yönteminden farklı olan SentiTFIDF yöntemi kullanılmıştır.

Veri kümesi olarak farklı alanları kullanan bir diğer çalışma ise Topaçan`a (2016) aittir. Çalışmada elektronik, telekom, beyaz eşya, otomotiv ve bankacılık sektörlerine ait Twitter verileri kullanılmış ve verilerin, duygu yönelimlerine göre üç sınıfa ayrılması amaçlanmıştır. Ön işleme aşamasında kullanılan 5 farklı boyut ağırlıklandırma algoritmasının (ikili kodlama, terim frekansı, TF-Log norm, ters doküman frekansı, TF-IDF), sınıflandırmaya etkisi araştırılmıştır. Kullanılan denetimli öğrenme algoritmalarının gösterdikleri performanslar, elde edilen özniteliklerin tümü ve öznitelik azaltım yöntemlerinde bilgi kazancı ve ki-kare yöntemleri kullanılarak araştırılmıştır.

Yapılan tüm karşılaştırmalar sonucunda, incelenen her aşamada en iyi sınıflandırma sonucunu veren yöntemler belirlenmiş ve Türkçe metinlerden anlamlı bilgiler çıkarmak için bir yöntem olarak önerilmiştir.

Ekinci ve Omurca (2017) tarafından yapılan çalışmada otel ile ilgili Türkçe kullanıcı yorumları üzerinden ürün özelliklerinin çıkartılması amaçlanmıştır. Bu kapsamda özellik tabanlı (aspect based) duygu analizi, konu modelleme (topic modelling) yöntemlerinden biri olan gizli dirichlet ayırımı kullanılmıştır. Benzer şekilde Ahmetoğlu ve Daş (2020), Türkçe otel yorumlarını kullanmıştır. Yorumların sınıflandırılması için derin öğrenme ve tekrarlayan sinir ağları kullanılmıştır.

Akın ve Şimşek (2018) tarafından yapılan çalışmada sözlük temelli duygu analizi ele alınmıştır. Çalışmada Twitter'daki bir kanalda 8 ay boyunca yayınlanan programlara ilişkin seyirci iletileri uygulama verisi olarak kullanılmıştır. Bu iletilerin; olumlu, olumsuz ya da nötr olarak sınıflandırılan duygulardan hangisini içerdiği duygu analizi ile belirlenmiştir. Benzer şekilde Atan ve Çınar (2019) da sözlük temelli duygu analizini ele almıştır. Çalışmada Borsa İstanbul'da işlem gören Borsa İstanbul şirketlerine ilişkin 2014 yılında farklı kaynaklarda yayınlanmış 14.108 haber veri kümesi olarak kullanılmış ve metin madenciliği yöntemleri ile yıllık ve çeyreklik periyotta haber sayıları belirlenmiştir. Haber içeriklerindeki ifadeler, Türk diline çevrilmiş bir “Duygu Sözlüğü” ile sayısal değerlere dönüştürülmüş ve aynı dönemde piyasada oluşan şirket değerleri arasındaki ilişkiler bu değerler ile analiz edilmiştir. Ayrıca toplanan haberlerin pozitif veya negatif olarak sınıflandırılması, olumlu ve olumsuz duygu sınıfları için hazırlanmış iki ayrı sözlük ile yapılmış ve sözlükteki kelimelerin, haberlerde geçme sıklığı dikkate alınmıştır.

Duygu analizinde sıklıkla karşılaşılan sorunlardan biri, dengesiz veri kümeleridir. Denetimli öğrenme algoritmalarında eğitim veri kümesi içerisinde sınıfların dengeli bir şekilde dağılması, sınıflandırıcının tahmin performansını doğrudan etkilemektedir. Oğul ve Güran (2019) tarafından yapılan çalışmada hem Türkçe hem de İngilizce metinlerin bulunduğu 3 farklı duygu analizi veri kümesi üzerinde kelime tabanlı *N*-gram yapıları kullanılarak, veri kümelerinin dengeli hale getirilmesini sağlayan rastgele örneklem artırma (random over sampling), sentetik azınlık örneklem artırma tekniği (synthetic minority oversampling technique), rastgele örneklem azaltma (random under sampling) ve yakın aile (nearmiss) algoritmaları kullanılmıştır. Bu algoritmaların, lojistik regresyon sınıflandırma yönteminin performansına etkileri incelenmiştir. Dengesiz veri kümesini ele alan bir diğer çalışma, Santur (2020) tarafından ortaya konmuştur. Çalışmada, müşteri yorumları veri kümesi kullanılmıştır. Müşteri yorumlarında genellikle ürünü beğenme ya

da beğenmeme şeklinde görüşler bulunmaktadır. Veri kümesinin yaklaşık % 98'i pozitif, % 2'si ise negatif yorumları içermektedir. Çalışmada ikili (pozitif-negatif) sınıflandırma performansının iyileştirilmesi için derin öğrenme yaklaşımı kullanılmıştır.

Toçoğlu, Çelikten, Aygün ve Alpkoçak (2019), farklı bölgelerde yaşayan ve farklı yaşlardaki 5.000 kişiden anket yöntemi ile altı duygu sınıfına (mutluluk, korku, öfke, üzüntü, tikslenme, şaşırma) ilişkin anılarını ve deneyimlerini metin olarak paylaşmalarını istemiştir. Veri kümesi olarak 4.709 katılımcı dokümanı ve bu dokümanlardan çıkarılan 27.350 adet belge kullanılmıştır. Toplanan veri kümesinde kök bulma işleminde sabit önek (fixed prefix stemming) ve öznitelik seçiminde ise karşılıklı bilgi (mutual information) yöntemleri kullanılmıştır.

Farklı alanlardan veri kümesi oluşturularak yapılan çalışmalardan biri de Pervan (2019) tarafından yapılmıştır. Bu çalışmada iki farklı online alışveriş sitesinden alınan ürün yorumları kullanılmıştır. İkili sınıflandırma yapıldığı için yorumların aldığı yıldız sayılarına göre sadece 1 ve 5 puan alan yorumlar ele alınmış, diğer yorumlar göz ardı edilmiştir. Ürün yorumları; elektronik, ev eşyaları, kozmetik vb. alanlarda toplanmıştır. Elde edilen veri kümesi kullanılarak rastgele orman, uzun kısa süreli bellek (long short-term memory) ve genişletilmiş evrişimsel sinir ağı (dilated convolutional neural networks) denetimli öğrenme algoritmaları eğitilmiş ve sonuçlar karşılaştırılmıştır.

Sağlam (2019), ana akım haber medyalarındaki haber metinlerini kullanarak Türkçe duygu sözlüğü geliştirmiştir. Bu kapsamda, polaritesi bilinen haberlerden büyük bir derlem oluşturulmuştur. Bu metinlerden elde edilen terimlerin ton ve polarite değerleri belirlenerek, mevcut bir Türkçe duygu sözlüğü ile birleştirilmiş ve yeni bir Türkçe duygu sözlüğü oluşturulmuştur. Bu sözlüğün daha önceki çalışmalarda oluşturulan sözlüklerden farkı, başka dillerden çeviri ile oluşturulmamasıdır.

Aytekin ve Bayram (2021), farklı 3 alan olarak ürün, film ve kitap yorumlarını seçmiş ve bu alanlardaki kullanıcı yorumlarını ve bunlara ilişkin kullanıcı puanlamalarını (yıldız veya 1-5 aralığında puanlar) veri olarak kullanmıştır. Yapılan analizler sonucunda her bir modelin, kendi alanındaki olumlu ve olumsuz yorum verileriyle test edildiğinde daha yüksek, farklı alanlardan alınan verilerle test edildiğinde ise daha düşük bir başarı gösterdiği görülmüştür. Bu nedenle çalışmada alanlara ait duygu analizi modellerinin de farklı olması gerektiği belirtilmiş ancak bunun pratik bir yöntem olmadığı vurgulanmıştır. Bu açığı kapatmak amacıyla çok sayıda farklı alan için uygulanabilecek bir karma veri uygulaması geliştirilmiştir.

Yavuz (2023), Trendyol online alışveriş sitesi üzerinden 15 farklı markaya ait müşteri yorumlarını kullanarak sözlük tabanlı duygu sınıflandırması yapmıştır. Toplanan veri kümeleri, ön işlemeden geçirilmiş ve Python sözlük tabanlı kütüphaneleri olan TextBlob ve Vader yardımıyla verilerin duygu sınıfları bulunmuştur. Bu sonuçlara göre müşterilerin markaya göre duygu sıralamaları, müşterilerin ürünlere verdikleri yıldızlara göre metin özetleme bulguları ve müşterilerin verdikleri yıldız sayıları ile yaptıkları yorum duygu sınıfı arasındaki benzerlikler araştırılmıştır.

Müşteri yorumları kullanılarak yapılan diğer bir çalışma Kılıçer ve Şamlı'ya (2023) aittir. Hepsiburada, Trendyol ve N11 gibi online alışveriş sitelerinden elektronik cihazlar hakkında toplanan yorumlar, ön işlemeden geçirilerek k-en yakın komşu, k-star, naive Bayes, lojistik regresyon, destek vektör makinesi, karar ağacı ve rastgele orman sınıflandırma yöntemleri kullanılarak pozitif, negatif ve nötr olarak sınıflandırılmıştır. Kullanılan yöntemlerin model performansları karşılaştırılmıştır.

Mayda ve Uğurlu (2024) tarafından yapılan çalışmada Trendyol online alışveriş sitesinden 20 farklı ürün için müşteri yorumları toplanmıştır. Toplanan yorumlar, ön işlemeden geçirilmiş ve Zemberek kütüphanesi kullanılarak yorumlardaki dil bilgisi ve yazım hataları düzeltilmiştir. Yorumlar, yazarlar tarafından faydalı ve faydasız olarak etiketlenmiş ve BERT modeli kullanılarak yorumların sınıflandırılması yapılmıştır.

Bayrak, Beğen, Öner, Kaplan, Demirdağ ve Sayan (2024), müşterilerden toplanan ürün ve film yorumlarını kullanarak dengesiz veri kümelerini ele almıştır. Çalışmada pozitif, negatif ve nötr sınıflar kullanılmıştır. Veri kümesindeki dengesizliği ortadan kaldırmak için sentetik azınlık örneklem artırma tekniği ve OpenAI yazılımı olan ChatGPT-2 modeli kullanılmıştır. Veri kümesini dengeli hale getirmek için eksik olan sınıflara yeni yorumlar üretilmiştir. Elde edilen dengeli veri kümeleri ile karar ağaçları, rastgele orman, destek vektör makinesi, yapay sinir ağları ve uzun kısa süreli bellek sınıflandırma yöntemleri eğitilmiş ve model performansları karşılaştırılmıştır. Benzer şekilde Görmez, Arslan ve Atak (2024), müşteri ve film yorumlarından oluşan veri kümelerini kullanmış, veri kümelerine ön işleme yapmış ve derin öğrenme yöntemlerini kullanarak duygu sınıflandırması yapmıştır. Ön işleme öncesi ve sonrası için derin öğrenmede parametre optimizasyonu yapılmış ve model performansları karşılaştırılmıştır.

Şen, (2025), Trendyol online alışveriş sitesinden dört farklı alanda ait müşteri yorumlarını kullanarak metin sınıflandırması gerçekleştirmiştir. Ürün yorumları; fiyat performans, iade edilen, paketleme iyi-hızlı teslimat, kalitesiz-kusurlu-kötü paketleme,

tavsiye edilen-kaliteli ürünler şeklinde sınıflandırılmıştır. Yapılan sınıflandırma yalnızca yorumların içeriklerini değil, müşteri yorumlarının olumlu ya da olumsuz duygu durumlarını da yansıtarak çalışmanın duygu analizi alanına önemli katkılar sunması sağlanmıştır. Sınıflandırma yöntemi olarak lojistik regresyon, rastgele orman sınıflandırma yöntemlerinin yanı sıra, derin öğrenme tabanlı sınıflandırma yöntemleri kullanılmıştır. Özmen (2025), Trendyol online alışveriş sitesinden 50 kozmetik markasına ait müşteri ürün yorumlarını kullanarak duygu analizi gerçekleştirmiştir. Toplanan yorumlara ön işleme ve el ile etiketleme yapılarak veri kümesi, analize hazır hale getirilmiştir. K-en yakın komşu, destek vektör makineleri, karar ağacı, lojistik regresyon, rastgele orman sınıflandırma yöntemleri kullanılarak ikili sınıflandırma (pozitif ve negatif) yapılmıştır. Model sonuçları karşılaştırmalı olarak değerlendirilmiştir.

Martin (2025) tarafından yapılan çalışmada, Kadın Destek uygulamasına yapılan kullanıcı yorumları ele alınmıştır. Yorumlar, üç (pozitif, nötr ve negatif) ve iki (pozitif ve negatif) etiketli olarak sınıflandırılmıştır. Her iki sınıflandırma grubu için dengeli ve dengesiz veri kümeleri kullanılarak değerlendirme yapılmıştır. Kullanıcı yorumlarının sınıflandırılmasında; naive Bayes, k-en yakın komşu ve sıralı minimal algoritması yöntemleri kullanılmıştır.

Tablo 2’de bu bölümde ele alınan duygu analizi çalışmalarının bir özeti sunulmaktadır. Ele alınan çalışmalar; uygulanan yöntem, dikkate aldıkları veri kümelerine, duygu analiz seviyelerine ve kullanılan yazılıma göre değerlendirilmiştir.

Tablo 2. Türkçe Dilini Ele Alan Duygu Analizi Çalışmalarının Bir Özeti

Yazar(lar)	Yöntem(ler)	Veri Kaynağı	Duygu Analizi Seviyesi	Yazılım(lar)
Eroğul (2009)	Destek Vektör Makineleri	Beyazperde	Cümle	Python, Zemberek
Çetin ve Amasyalı (2013)	Naive Bayes, Karar Destek Makinesi, Karar Ağacı, Rastgele Orman, K-En Yakın Komşu	Twitter	Cümle	Weka
Çelik (2013)	Lojistik Regresyon, Naive Bayes (Multinomial), Destek Vektör Makineleri, Karar Ağaçları	Kitap.antoloji Beyazperde, Hepsiburada	Cümle	Weka, Zemberek
Nanğır (2013)	Destek Vektör Makineleri, Karar Ağaçları, Naive Bayes, Çoğunluk Oylaması Kuralı	Literatürden alınan veri kümesi (Çelik, 2013)	Cümle	Weka
Meral ve Diri (2014)	Naive Bayes, Destek Vektör Makineleri, Rastgele Orman	Twitter	Cümle	Weka
Nizam ve Akın (2014)	Naive Bayes, Rastgele Orman, Sıralı Minimal Optimizasyon, Karar Ağacı ve En Yakın Komşu	Twitter	Cümle	Weka
Çatak (2014)	Destek Vektör Makineleri	Twitter	Cümle	Java
Güran vd. (2014)	Destek Vektör Makineleri	Twitter, Blog yazısı, Film Yorumları	Cümle	Weka
Akba (2014)	Destek Vektör Makineleri	Beyazperde	Cümle	C# Zemberek Weka
Uçan (2014)	Duygu sözlüğü oluşturma	Beyazperde, Otelpuan	Cümle	Zemberek Google Translate WordNet
Eliaçık ve Erdoğan (2015)	Destek Vektör Makineleri	Twitter	Cümle	Weka
Nalçakan vd. (2015)	Naive Bayes, Rastgele Orman, Destek Vektör Makineleri, J48, Kstar	Twitter	Cümle	Weka
Pak (2015)	Destek Vektör Makineleri, Naive Bayes	Hepsiburada, Beyazperde	-	Zemberek

Türkmenoğlu (2015)	Karar Destek Makineleri, Naive Bayes, Karar Ağaçları	Twitter Beyazperde	Cümle	Zemberek Weka
Çoban ve Özyer (2016)	K-En Yakın Komşu, Destek Vektör Makineleri, Naive Bayes Maksimum Entropi	Twitter	Cümle	-
Ekinci vd. (2016)	Naive Bayes, Rastgele Orman, J48, Destek Vektör Makineleri, K-En Yakın Komşu	Twitter	Cümle	Zemberek Weka Python
Oğul ve Ercan (2016)	Naive Bayes, Rastgele Orman, Destek Vektör Makineleri	Tripadvisor booking.com	-	-
Gözükara ve Özel (2016)	Destek Vektör Makineleri	İnternet	Cümle	-
Topaçan (2016)	Naive Bayes, Destek Vektör Makinesi, Karar Ağaçları	Twitter	Cümle	Zemberek
Parlar (2016)	Multinomial Naive Bayes, Destek Vektör Makinesi, Karar Ağaçları, Lojistik Regresyon	Beyazperde Hepsiburada İngilizce film ve ürün yorumları	-	Zemberek Python
Yiğit (2017)	Destek Vektör Makineleri, Karar Ağaçları, K-En Yakın Komşu, Lojistik Regresyon, Rastgele Orman, Regresyon Modeli	Türk Telekom Grubu	Cümle	-
Hayran ve Sert (2017)	Destek Vektör Makineleri	Twitter	Cümle	-
Ekinci ve Omurca (2017)	Gizli Dirichlet Ayırımı	otelpuan	Özellik	Zemberek
Yıldırım ve Yüksel (2017)	ZeroR, J48, K-En Yakın Komşu, Naive Bayes	Twitter Borsa İstanbul	Cümle	Python Weka
Onan (2017)	Destek Vektör Makinaları, Naive Bayes, Lojistik Regresyon, K-En Yakın Komşu, Rastgele Orman	Twitter	Cümle	Python Weka
Akın ve Şimşek (2018)	Sözlük Tabanlı Duygu Analizi	Twitter	Cümle	R, Excel, VBA

Yıldırım ve Yıldız (2018)	Destek Karar Makineleri, Lojistik Regresyon, K-En Yakın Komşu, Karar Ağacı, Çok-Katlı Naive Bayes (Multinomial), Yapay Sinir Ağları	https://github.com/denopas/TTC-3600/blob/master/TTC-3600_Orj.rar http://www.kemik.yildiz.edu.tr/veri_kumelerimiz.html	Doküman	Python Weka R
Kızılkaya (2018)	Karar Ağaçları, Rastgele Orman, Naive Bayes	Twitter	Cümle	R
Atan ve Çınar (2019)	Sözlük Tabanlı Duygu Analizi	Google News	Doküman	Zemberek R, Java
Ayan vd. (2019)	Lineer Ridge Regresyonu, Naive Bayes	Twitter	Cümle	MYSQL Python
Oğul ve Güran (2019)	Lojistik Regresyon Algoritması	SemEval YTÜ Kemik Grubu Amerikan Havayolu Şirketi	-	-
Toçoğlu vd. (2019)	Yapay Sinir Ağları, Destek Vektör Makineleri, Rastgele Orman, K-En Yakın Komşu	TREMO veri kümesi	Doküman	-
Pervan (2019)	Rastgele Orman, Uzun Kısa Dönemli Bellek, Genişletilmiş Evrimsel Sinir Ağları	Hepsiburada, N11	Paragraf/Cümle	Python
Sağlam (2019)	Sözlük Tabanlı Duygu Analizi	GDELT veri kümeleri	Paragraf/Cümle	
Tuna (2019)	Lojistik Regresyon, Rastgele Orman, Karar Ağaçları, K-En Yakın Komşu, Destek Vektör Makineleri, Doğrusal Diskriminant Analizi, Naive Bayes	TripAdvisor	Cümle	Python Zemberek
Aydoğan ve Karcı (2019)	Naive Bayes, Destek Vektör Makineleri	İnternet	Doküman	Python Zemberek
Emekli ve Selvi (2020)	Evrimsel Sinir Ağları, Tekrarlayan Sinir Ağları, Uzun Kısa Süreli Bellek	Twitter	Cümle	Python
Ahmetoğlu ve Daş (2020)	Derin öğrenme ve Tekrarlayan Sinir Ağları	İnternet	-	Python BeautifulSoup

Onan (2020)	Evrişimli Sinir Ağları	Twitter	Cümle	Python
Santur (2020)	Evrişimli Sinir Ağları, Tekrarlayan Sinir Ağları, Uzun Kısa Süreli Bellek	Hepsiburada	Paragraf/Cümle	-
Kilimci (2020)	Evrişimli Sinir Ağları, Tekrarlayan Sinir Ağları, Uzun Kısa Süreli Bellek	Twitter	Cümle	Python
Aytekin ve Bayram (2021)	Denetimli ve denetimsiz öğrenme algoritmaları	Yorumbudur Hepsiburada Beyazperde Kitapyurdu	Paragraf/Cümle	Python
Yavuz (2023)	Python Sözlük tabanlı TextBlob ve Vader kütüphaneleri	Trendyol	Cümle	Python
Kılıçer, Şamlı (2023)	K-En yakın komşu, K-Star, Naive Bayes, Lojistik Regresyon, Destek Vektör Makinesi, Karar Ağacı, Rastgele Orman	Hepsiburada, Trendyol ve N11	Cümle	Weka
Mayda, Uğurlu (2024)	BERT	Trenyol	Cümle	Python
Bayrak vd. (2024)	Karar Ağaçları, Rastgele Orman, Destek Vektör Makinesi, Yapay Sinir Ağları, Uzun Kısa Süreli Bellek	Hepsiburada, Trendyol ve N11 Film Yorumları	Cümle	-
Görmez vd. (2024)	Evrişimsel Sinir Ağları, Derin Öğrenme	Vinwoker Beyazperde	Cümle	Python
Şen (2025)	Yapay Sinir Ağları Tabanlı Derin Öğrenme Yöntemleri, Lojistik Regresyon, Rastgele Orman	Trendyol	Cümle	Python
Marttin (2025)	Naive Bayes, K-En Yakın Komşu, Sıralı Minimal Algoritması	Apple Store, Google Play, Youtube	Cümle	Weka
Özmen (2025)	K-En Yakın Komşu, Destek Vektör Makineleri, Karar Ağacı, Lojistik Regresyon, Rastgele Orman	Trendyol	Cümle	-

İKİNCİ BÖLÜM

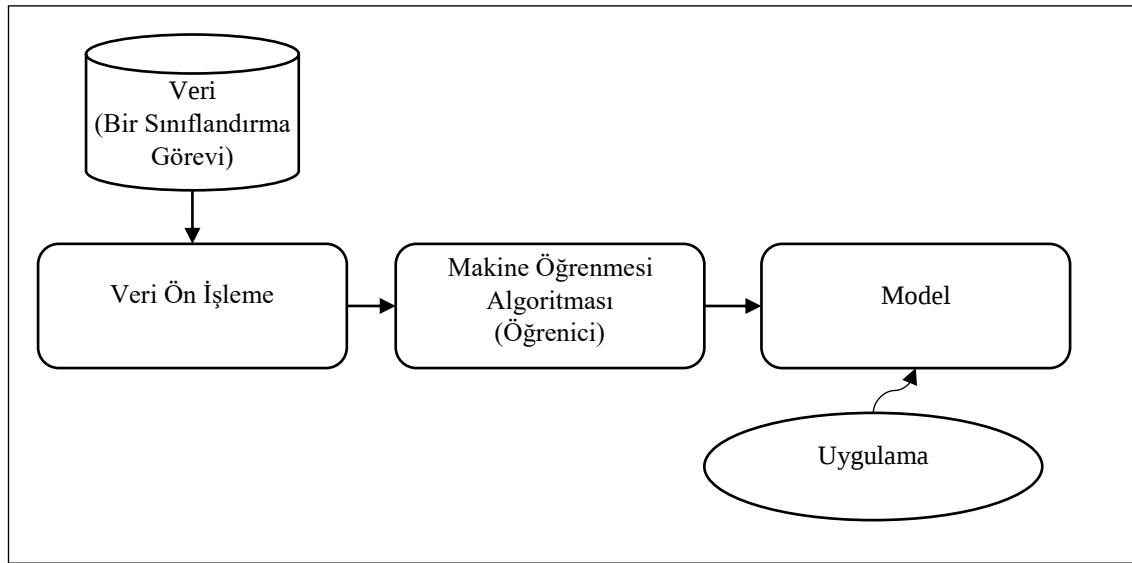
YAŞAM BOYU ÖĞRENME YAKLAŞIMI

Öğrenme, edinilen deneyimlere göre bir davranışın değiştirilmesi olarak tanımlanabilmektedir. Bu nedenle öğrenme, davranışın ayırt edici özelliklerini içermektedir. Makine öğrenmesi; öğrenme yeteneğine sahip algoritmaların, bilgisayar uygulamalarının ve sistemlerin nasıl geliştirileceğinin incelenmesidir. Böylece makine öğrenmesi, deneyim yoluyla bazı görevlerdeki performanslarını iyileştirmektedir (Lampropoulos ve Tsihrintzis, 2015: 31). Makine öğrenmesi, karmaşık sistemlerin güvenilir bir şekilde modellenmesini sağlayabilmekte ve birçok uygulamaya ilişkin etkin sonuçlar veren süreci otomatikleştirebilmektedir (Khan, 2019: 1). Makine öğrenmesi için mevcut süreç, bir model oluşturmak için belirli bir veri kümesi üzerinde bir algoritma çalıştırmaktır. Daha sonra model, gerçek yaşam görevlerinde uygulanır. Bu süreç, izole öğrenme olarak adlandırılmaktadır (Liu, 2017: 1). İzole edilmiş/izole öğrenme tekniğindeki temel problem, eski bilginin tutulmaması ve saklanmamasıdır. Bu durum, insana özgü öğrenme/insan öğrenmesi ile tamamen ters bir durumdur. Başka bir deyişle insanlar, izole öğrenme yoluyla öğrenmede iyi olmayıp genellikle geçmiş bilgilerini (deneyimlerini) kullanmakta ve bu deneyimleri, gelecekteki öğrenme ve karar verme süreçleri için kullanmaktadırlar (Singh, 2019: 2). Diğer yandan yaşam boyu makine öğrenimi (lifelong machine learning) veya yaşam boyu öğrenme (lifelong learning) ise insan öğrenme sürecini ve yeteneğini taklit etmeyi amaçlamakta olup (Liu, 2017: 1), literatürde farklı alanlarda kullanılmaktadır.

2.1. Yaşam Boyu Öğrenme

Mevcut istatistiksel öğrenme algoritmalarının çoğu, naive Bayes, destek vektör makinaları, maksimum entropi, derin sinir ağları gibi yöntemler, izole edilmiş bir veri kümesine ve tek bir göreve dayalı olarak öğrenme gerçekleştirmektedir. Yani öğrenme süreci, tek bir dönemde izole edilmiş bir veri kümesi üzerinde yürütülür. Bu şekildeki izole bir öğrenme modelinin birçok sınırlamaları bulunmaktadır. İlk olarak, geçmiş görevlerde öğrenilen bilgi, gelecekteki görevlere uygulanmaz. İkincisi, ön bilgi eksikliğinden dolayı, normal olarak daha genelleştirilmiş ve uyarlanabilir bir modeli eğitmek için çok sayıda eğitim örneği gerekmektedir (Xia vd., 2017: 3). Başka bir deyişle, bilgiyi tutma yeteneği olmadan, bir makine öğrenmesi modelinin etkili bir şekilde öğrenmek için büyük veri kümelerine ihtiyacı olacaktır. Denetimli öğrenmede, genellikle

eğitim verilerinin etiketlenmesi manuel olarak yapılır. Bu manuel işlem, yoğun emek gerektirir ve zaman alıcıdır. Veri kümesi çok büyük veya çok karmaşıksa, bu tür veri kümelerini etiketlemek neredeyse imkânsız hale gelir. Verilerdeki değişkenlerin, dünya çapında meydana gelen değişiklikler nedeniyle sürekli değişimini takip etmek de önemlidir. Bu, etiketlemenin değişen trendlerle sürekli olarak güncellenmesi gerektiği anlamına gelir. Denetimsiz öğrenme durumunda bile, modeli eğitmek için büyük veri kümelerinin toplanması mümkün olmayabilir (Singh, 2019: 2).



Şekil 5. Klasik Makine Öğrenmesi Paradigması

Kaynak: Chen ve Liu, 2018: 12

Şekil 5’te klasik makine öğrenmesi paradigması verilmiştir. Klasik makine öğrenmesinde bir sınıflandırma görevine ait sadece bir veri kümesi bulunmaktadır. Veri kümesi, ön işleme aşamasını geçtikten sonra kullanılan makine öğrenmesi algoritması tarafından öğrenilir ve modelin doğruluğu sınanır. En iyi performansı gösteren model elde edilinceye kadar iyileştirme süreci devam eder. Eğitilen modele tek bir alandan elde edilen test verisi verilerek modelin sınıflandırma yapması beklenir. Bu süreçte test verisinin, başlangıç veri kümesi ile aynı alana sahip olması sınıflandırma doğruluğunu etkilemektedir. Farklı alana ait test verisinin sınıflandırılması istendiğinde elde edilen sonuçların doğruluğu düşük olacaktır. Ayrıca model performansının yüksek olması istenen durumlarda, çok daha fazla eğitim verisine ihtiyaç vardır. Toplanan eğitim verisinin etiketlenmesi hem zaman kaybına hem de maliyete sebep olmaktadır. Tüm bu

durumlar göz önüne alındığında yaşam boyu öğrenme yaklaşımı, klasik makine öğrenmesi paradigmasının eksik kaldığı konularda daha etkilidir (Chen ve Liu, 2018: 12).

Yaşam boyu makine öğrenimi veya yaşam boyu öğrenme, insanın öğrenme sürecini ve yeteneğini taklit etmeyi amaçlamaktadır (Liu, 2017: 1). İnsanlar, tek başlarına öğrenmeyip sürekli olarak öğrenmekte, geçmişteki problemlerden öğrendikleri bilgileri saklamakta böylelikle gelecekteki karmaşık problemleri çözmeleri kolaylaşmaktadır (Arif vd., 2019: 1). Yaşam boyu bir öğrenme, oldukça doğaldır çünkü çevremizde olan her şey, yakından birbiriyle ilişkilidir ve bağlantılıdır. Geçmişte öğrenilen kavramlar ve bunların ilişkileri, yeni bir konunun daha iyi anlaşılmasına ve öğrenilmesine yardımcı olabilir. Örneğin; bir makine öğrenmesi algoritmasının bir filmle ilgili olumlu ve olumsuz yorumları tanımak için doğru bir sınıflandırıcı oluşturması gerekmektedir, insanlar filmlere ait 1000 olumlu ve 1000 olumsuz yoruma ihtiyaç duymamaktadır (Liu, 2017: 1). Diğer yandan günümüzde fiziksel robotlar, sohbet botları ve akıllı asistanlar gibi insanlarla veya sistemlerle gerçek zamanlı etkileşim kuran yapay zeka uygulamalarının da yaşam boyu öğrenme yeteneklerine ihtiyacı vardır (Arif vd., 2019: 1).

Yaşam boyu öğrenme yaklaşımı terimi, ilk olarak 1995 yılında kullanılmıştır (Thrun ve Mitchell, 1995, Thrun, 1996b). Thrun ve Mitchell (1995) tarafından yapılan tanım, öğrenen robotlar ile ilgilidir. Yapılan çalışmada mevcut robot sistemlerinin ağırlıklı olarak sadece bir öğrenme görevine odaklandıkları belirtilerek bir robot ajanın karşılaştığı yaşam boyu öğrenme problemleri ele alınmıştır. Thrun (1996b)'ya ait diğer bir çalışmada, öğrenici (öğrenen, öğrenci, learner) olarak adlandırılan bir makinenin tüm öğrenme problemleri ile karşılaştığında kullanılabileceği öğrenme yaklaşımları ele alınmıştır. Bu yaklaşımın temeli, bilgi aktarımı ve yaşam boyu öğrenmedir. Başka bir deyişle bir öğrenici, n . öğrenme görevi ile karşılaştığında, doğruluğu artırmak için $n - 1$ öğrenme görevinden topladığı bilgileri kullanmaktadır (Thrun, 1996b: 2). Yaşam boyu öğrenmeye ilişkin ilk çalışmalardan sonra bu yaklaşım; yaşam boyu denetimli öğrenme, derin sinir ağlarında sürekli öğrenme, yaşam boyu denetimsiz öğrenme, yaşam boyu yarı denetimli öğrenme, yaşam boyu pekiştirmeli öğrenme gibi çeşitli alanlarda çalışılmıştır (Silver vd., 2013: 49-50).

Yaşam boyu öğrenme, aynı zamanda “sürekli öğrenme” (continual learning) ile temel olarak aynı anlamdadır. Ancak iki isim altında yapılan geçmiş araştırmalar, aynı sorunun farklı yönleri üzerine eğilmiştir. Sürekli öğrenme terimi, derin öğrenme alanında yaşam boyu öğrenmeden daha yaygın olarak kullanılmaktadır. Sürekli öğrenmenin odak noktası, katastrofik unutmayı (catastrophic forgetting) çözmektir (Liu, 2020: 2). Diğer

yandan yaşam boyu öğrenmenin beş temel özelliği vardır. Bunlar; sürekli öğrenme süreci, bilgi tabanındaki bilginin birikimi ve korunması, gelecekteki öğrenme için birikmiş geçmiş bilgiyi kullanma yeteneği, yeni görevler keşfetme yeteneği ve çalışırken öğrenme veya iş başında öğrenme yeteneğidir. Bu yetenekler olmadan bir makine öğrenmesi sistemi, dinamik bir ortamda öğrenemeyecek ve gerçek anlamda akıllı olamayacaktır (Thrun, 1996b: 2).

Yaşam boyu öğrenme yaklaşımının ilişkili olduğu alanlardan biri de doğal dil işlemedir. Yaşam boyu öğrenme yaklaşımı; büyük veriyi gerektirmekte, büyük veri de metinler ve görsellere dayanmaktadır. Bir dilde ifade edilen anlamsal seviyedeki bilgilerin insanlar tarafından anlaşılması kolaydır. Kelime anlamı ve dil bilgisi, farklı görevlerde ve farklı alanlarda aynıdır. Bu özellikler sayesinde Bayesci yöntemler gibi bazı geleneksel yaklaşımlar etkin çalışmakta ve bilgi sunumu ve aktarımı mümkün olmaktadır (Hong vd., 2018: 5). Başka bir deyişle doğal dil, insanlar tarafından bilgi biriktirmede, bilgiyi işlemede, yeni bilgi türetmede (bilişsel hesaplama) ve bilgiyi zekaya dönüştürmede oldukça etkili bir şekilde kullanılır. Dilin tüm bu avantajları düşünüldüğünde doğal dil temelli insan öğrenme yaklaşımını, yaşam boyu öğrenme yaklaşımı ile birleştirmenin yaşam boyu öğrenmeyi daha etkili ve başarılı hale getirebileceğine inanılmaktadır (Khan, 2019: 2).

Yaşam boyu öğrenme ve doğal dil işleme alanlarının birlikte kullanımının önemini şu şekilde özetlemek mümkündür:

- Kelimeler ve kelime öbekleri, tüm alanlarda ve görevlerde neredeyse aynı anlama sahiptir.
- Her alandaki cümleler, aynı söz dizimini veya dil bilgisini takip eder.
- Neredeyse tüm doğal dil işleme sorunları birbirleriyle yakından ilişkilidir. Bu durum, sorunların birbirleriyle bağlantılı olduğu ve bazı durumlarda birbirlerini etkilediği anlamındadır (Chen ve Liu, 2018: 2).

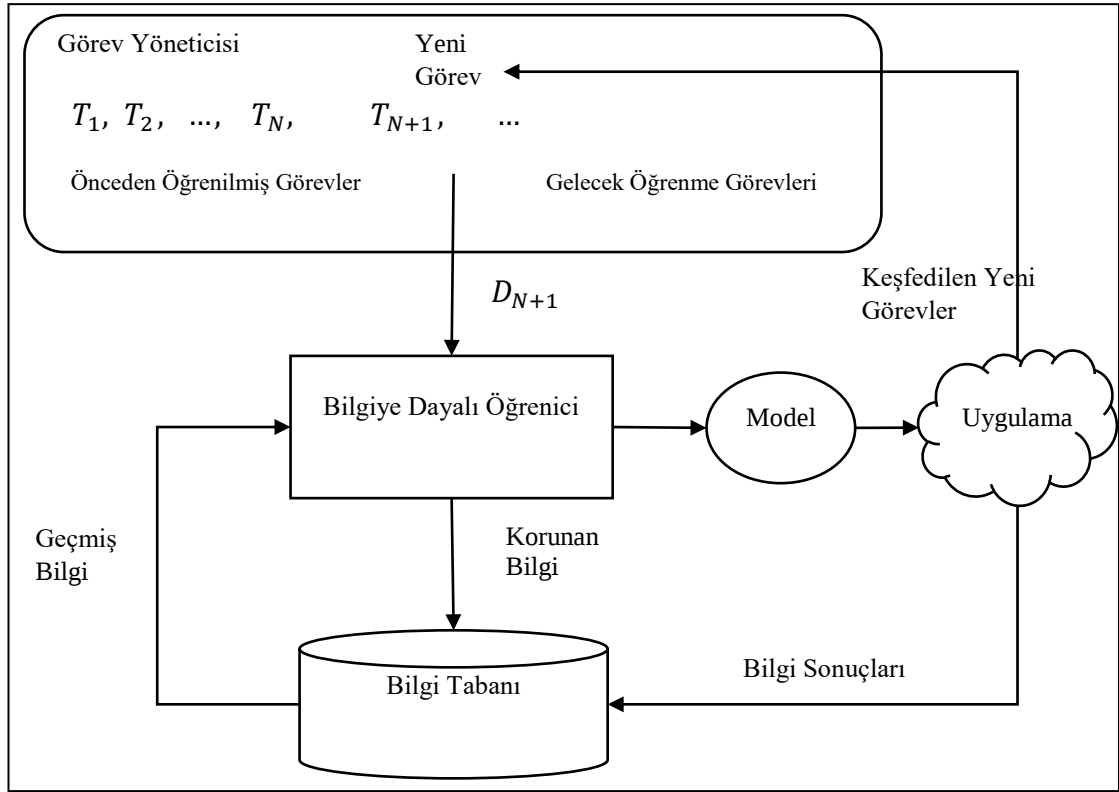
Yukarıda maddeler halinde verilen cümlelere göre aynı ifadelerin ve anlamların aynı söz dizimini paylaşması nedeniyle öğrenilen bilgi, alanlar ve görevler arasında kullanılabilir. Bu nedenle insanların, yeni bir uygulama alanıyla karşılaştığında dili yeniden öğrenmesi (veya yeni bir dil öğrenmesi) gerekmez. Örneğin, bir kişinin hiç psikoloji çalışmadığını ve psikoloji çalışmak istendiği varsayılırsa bu kişinin psikoloji alanına özgü kavramlar dışında psikoloji metninde kullanılan dili öğrenmesine gerek yoktur. Çünkü dilin kendisi ile ilgili tüm bilgiler, herhangi bir alandaki ile aynıdır (Chen ve Liu, 2018: 2). Bu kapsamda bakıldığında yaşam boyu öğrenme yaklaşımı için metin

verisi, geçmiş bilgilerin saklanması ve gelecek yeni alanlara ait verinin tahmininde kullanılmasına oldukça uygundur. Başka bir deyişle, aynı dilin kullanıldığı metin verilerinden öğrenilen bilgiler, benzer kapsama sahip diğer alanlar için de kullanılabilir.

2.2. Yaşam Boyu Öğrenme Süreci

Yaşam boyu öğrenme, bir sürekli öğrenme süreci olup genel hatlarıyla süreçte gerçekleştirilen işlemler şu şekilde özetlenebilir: Herhangi bir zamanda öğrenci, birbirinden farklı ve sıralı N öğrenme görevini ($T_1, T_2, T_3, \dots, T_N$) gerçekleştirmiştir. Önceki görevler (previous task) olarak da adlandırılan bu görevler, $D_1, D_2, \dots, D_N$ olmak üzere farklı görevlere karşılık gelen veri kümelerine sahiptir. Görevler, farklı türlerde (types) ve farklı alanlardan (domain) olabilir. $(N + 1)$. görev (*yeni, mevcut ve şimdiki görev*) ile karşılaşıldığında öğrenci, bilgi tabanındaki eski görevlerden öğrendiği bilgileri, T_{N+1} . görevi öğrenmek için kullanır. Bu görev, sistemin kendisi tarafından verilebilir ya da keşfedilebilir. Yaşam boyu öğrenmenin amacı, genellikle yeni görevin veya önceki görevler arasından herhangi bir görevin performansının optimize edilmesidir. Bu anlamda bilgi tabanı, T_{N+1} görevinin öğrenimi tamamlandıktan sonra, T_{N+1} 'in öğreniminden kazanılan bilgi (örneğin; orta düzey ve nihai sonuçlar) ile güncellenir. Bir yaşam boyu öğrenci, aynı zamanda öğrenilecek yeni görevleri keşfedebilir ve uygulama veya test esnasında modelin performansını iyileştirmeyi öğrenebilir (Chen ve Liu, 2018: 9).

Yaşam boyu öğrenme yaklaşımında daha önce de belirtildiği gibi görevlerin aynı alandan olması gerekmez. Bazı araştırmacılar çalışmalarında, her alandan yalnızca bir görev olduğu için alan ve görev terimlerini birbirinin yerine kullanabilmektedir (Chen ve Liu, 2018: 11).



Şekil 6. Yaşam Boyu Öğrenme Sistem Yapısı

Kaynak: Chen ve Liu, 2018: 13

Yaşam boyu öğrenme süreci, birçok bileşeni içermektedir. Şekil 6’da süreç bileşenleri ve bileşenlerin bağlantıları verilmiştir. Mevcut sistemler, daha basit bir yapıya sahip olabilmekte ve tüm bileşenleri kullanmayabilmektedir. Yaşam boyu öğrenme süreci; görev yöneticisi, bilgi tabanı, bilgiye dayalı öğrenici, model ve uygulama olmak üzere beş bileşenden oluşmaktadır. Bu bileşenler şu şekilde özetlenebilir:

Görev yöneticisi: Gelen görevlerin listesini sisteme depolar. Ana işlevi, görev geçişlerini yapmak ve görevleri, bilgiye dayalı öğreniciye atamaktır. Görev yöneticisi, performansı iyileştirmek için her görevi, birden çok kez atanmış olsa bile, benzersiz olarak görür (Abulaish vd., 2024: 6).

Bilgi tabanı (Knowledge Base): Önceden öğrenilen bilgiyi depolamak için kullanılmaktadır. Birkaç alt bileşenden oluşmaktadır (Chen vd., 2015: 1):

- Geçmiş bilgi deposu (past information store):* Geçmiş öğrenmelerden ortaya çıkan model sonuçları, pattern ve diğer farklı çıktılar saklanabilir. Ayrıca her bir geçmiş görevin orijinal verisi, ara sonuçları ve son model sonuçları alt bileşen olarak da saklanabilir.

- b. *Bilgi madencisi (Knowledge miner)*: Bilgi tabanında saklanan geçmiş bilgiyi kullanır. Bu madencilik, geçmiş görevlerden elde edilen bilgilerden bilgi öğrendiği için bir meta öğrenme süreci olarak kabul edilebilir (Chen vd., 2015: 1). Bu meta öğrenme süreci; varlık ismi tanıma (named entity recognition), duygu sözcükleri bulma ve benzeri gibi görevlerden ileri eğitime katkı sağlayacak meta bilginin çıkarılmasıdır (Hong vd., 2020: 9).

Bilgiye dayalı öğrenici: Geçmiş bilgileri, bilgi tabanından alır ve bir görev için modeli oluşturmak üzere ana algoritmayı yürütür. Ayrıca yeni edinilen bilgileri de içerecek şekilde bilgi tabanını günceller (Abulaish vd., 2024: 6)

Model: Denetimli öğrenmede bir tahmin modeli veya sınıflandırıcı, denetimsiz öğrenmede kümeler veya konular olabilen bir öğrenilmiş modeldir (Chen ve Liu, 2018: 14).

Uygulama: Modelin gerçek dünyadaki uygulamasıdır. Model uygulaması sırasında sistemin hala yeni bilgiler (yani "sonuçlardaki bilgi") öğrenebileceğini ve muhtemelen öğrenilecek yeni görevleri keşfedebileceğini unutmamak önemlidir. Uygulama, aynı zamanda modelin iyileştirilmesi için bilgiye dayalı öğreniciye geri bildirim de verebilir (Chen ve Liu, 2018: 14).

Yaşam boyu makine öğrenmesi, bir veya daha fazla alandan birçok görevi öğrenebilen sistemleri kapsamaktadır. Öğrenilen bilgiler, verimli ve etkili bir şekilde korunmakta ve yeni görevleri daha verimli ve etkili bir şekilde öğrenmek için kullanılmaktadır (Silver vd., 2013: 51). Başka bir deyişle yaşam boyu öğrenmede ilk amaç, geçmişte öğrenilen bilgilerin yeni görev öğrenimine yardımcı olmak için seçici bir şekilde aktarılmasıdır. Burada bilgi seçimi, kritik öneme sahiptir. İkinci amaç ise yeni görevi öğrenirken geçmişte öğrenilen bilgiyi korumaktır. Çünkü bu aşamadaki güncellemeler nedeniyle birçok önceki bilgi parçası bozulursa gelecekteki görevler bunlardan yararlanamayacaktır (Qin vd., 2020: 2). Geçmiş bilgi, eğitilmiş modelden veya önceki görevlerden çıkarılan diğer yararlı bilgileri içerir. Mevcut görev öğrenimi sırasında, meta bilgi olarak da adlandırılan verilerden bazı ek bilgiler çıkarılabilir. Meta bilgi, genellikle eğitime yardımcı olmak için yeni girdi özellikleri olarak kabul edilir. Eğitimden sonra oluşturulan model, aynı zamanda mevcut görevden üretilen bilgiyi de içerir. Sonuçlar veya tahminler de bilgiye sahiptir. Bu nedenle, modelin kendisi veya onun ürettiği tahmin yeni bilgi olarak kabul edilebilir (Hong vd., 2020: 7).

Klasik izole öğrenme yöntemlerinden farklı olarak yaşam boyu öğrenme yaklaşımında bir öğrenme metodolojisi vardır. Bu metodolojide, bir dizi görev

bulunmakta ve öğrenmede gelişmeler olması beklenmektedir. Yaşam boyu öğrenme algoritmasının deneysel değerlendirilmesi genellikle aşağıda belirtilen adımlar kullanılarak gerçekleştirilmektedir (Chen ve Liu, 2018: 16-18) :

Önceki görevlerde bulunan verilerin çalıştırılması: Algoritma, ilk önce bilgi tabanında bulunan ve önceki görevlerden elde edilen veriler üzerinde belirli bir sırayla teker teker çalıştırılır ve elde edilen bilgiler saklanır. Bu adımda algoritmanın farklı versiyonları (farklı bilgi türleri kullanarak ya da daha fazla/az bilgi kullanarak) ile denemeler yapılabilir.

Yeni veriler kullanılarak çalıştırılması: Bilgi tabanındaki bilgiler kullanılarak algoritma yeni görev verileri üzerinde çalıştırılır.

Temel algoritmaların çalıştırılması: Sonuçların karşılaştırılması için bazı temel algoritmalar kullanılır. Mevcut yaşam boyu öğrenme algoritması ve izole öğrenme algoritmasından elde edilen sonuçlar karşılaştırılır.

Sonuçların analiz edilmesi: Bu adım, ikinci ve üçüncü adımların sonuçlarının karşılaştırmalı analizini yapar. Örneğin ikinci adımdaki yaşam boyu öğrenme algoritmasından elde edilen sonuçlar, üçüncü adımdaki temellerden elde edilen sonuçlardan daha üstün olduğunu göstermek için kullanılabilir.

2.3. Yaşam Boyu Denetimli Öğrenme

Denetimli öğrenmede, etiketli bir eğitim verisi kullanılarak bir fonksiyonun öğretilmesi ve tahmin yapabilmesi amaçlanmaktadır. Yaşam boyu denetimli öğrenmede de benzer şekilde bir dizi etiketli sıralı görev sisteme gelmektedir. Bu öğrenme şeklinde de amaç, her yeni gelen görev için tahmin yapmaktır (Sodhani vd., 2022: 12-13). Yaşam boyu denetimli öğrenme, ilk olarak yapay sinir ağları alanında kullanılmıştır (Thrun, 1996a: 13). Çalışmada bir robotun karşılaştığı ve çeşitli kontrol görevlerine karşılık öğrenmesi gereken kontrol politikasına ilişkin yaşam boyu öğrenme problemlerine çözüm aranmıştır. Bu çalışmadan sonra yaşam boyu denetimli öğrenmeye ilişkin farklı konularda ve alanlarda uygulamalar yapılmış ve çözüm algoritmaları önerilmiştir.

Yaşam boyu denetimli öğrenme, bir sürekli öğrenme sürecidir. Öğrenici model, bir sıra halinde kendisine verilen N adet sıralı denetimli öğrenme görevini $(T_1, T_2, \dots, T_N)$, öğrenir ve öğrenilen bilgileri, bilgi tabanında saklar. Yeni bir T_{N+1} görevi geldiğinde öğrenici, yeni bir f_{N+1} modelinin öğrenmesine yardımcı olmak için bilgi tabanındaki geçmiş bilgiyi kullanır. Böylelikle T_{N+1} 'e ait test verisi D_{N+1} sınıflandırılır. T_{N+1} 'in

öğrenilmesinden sonra bilgi tabanı, T_{N+1} 'den öğrenilen bilgiler ile güncellenir (Chen ve Liu, 2018: 35).

Bu tez çalışmasında yaşam boyu denetimli öğrenme yaklaşımı altında ürün yorumlarını ele alan duygu analizi sınıflandırmasına ilişkin bir uygulama yapılmıştır. Bu nedenle aşağıdaki başlıklarda, çalışmanın uygulama bölümünde kullanılacak yöntemler anlatılmıştır.

2.3.1. Yaşam Boyu Naive Bayes Sınıflandırma Yöntemi

Naive Bayes algoritması, istatistiksel bir sınıflandırma yöntemidir. Sınıflandırma problemlerinde sıklıkla kullanılan algoritma, metin madenciliği alanında yapılan çalışmalarda da kullanılmaktadır (Kim vd., 2006: 1).

Naive Bayes metin sınıflandırma yöntemine ilişkin ilk çalışma McCallum ve Nigam (1998) tarafından yapılmıştır. Önerilen yöntem baz alınarak yapılan çalışmalara örnek olarak; Chen vd. (2015), Ha vd. (2016), Wang vd. (2019), Hong vd. (2020) verilebilir. Bu tez çalışmasında kullanılan algoritmalarındaki notasyonlar (indisler ve kümeler), Chen vd. (2015) tarafından verilen notasyonlardan uyarlanmıştır. Temel olarak verilen her bir c_j ($j=1,2$) sınıfında bulunan her bir w kelimesinin $P(w|c_j)$ ile gösterilen olasılığını hesaplamaktadır. Her bir c_j sınıfının başlangıç olasılığı, $P(c_j)$ 'dir. Bu olasılıklar, $P(c_j|d)$ ile gösterilen ve verilen bir d belgesinin c_j sınıfında olma olasılığının hesaplanmasında kullanılmaktadır. c_j , pozitif (+) ya da negatif (−) olabilir. $D = \{d_1, d_2, \dots, d_{|D|}\}$ eğitim verilerini içeren veri kümesi olmak üzere, V sözlük (D veri kümesindeki birbirinden farklı kelimeler/terimler kümesi) ve $C = \{c_1, c_2, \dots, c_{|C|}\}$ D veri kümesi ile ilişkili sınıflar kümesini göstermektedir. Naive Bayes sınıflandırma yöntemi verilen her bir sınıf için, her bir $w \in V$ kelimesinin koşullu olasılığını hesaplayarak sınıflandırıcıyı eğitir.

$P(w|c_j)$, Eşitlik 16 ile şu şekilde hesaplanmaktadır:

$$P(w|c_j) = \frac{\lambda + N_{c_j,w}}{\lambda|V| + \sum_{v=1}^{|V|} N_{c_j,v}} \quad j = 1,2 \quad (16)$$

Burada $N_{c_j,w}$, w kelimesinin c_j sınıfına ait belgelerde geçme sayısıdır. λ ($0 \leq \lambda \leq 1$), düzgünleştirme katsayısıdır. $\lambda = 1$ olduğunda Laplace düzgünleştirmesi

olarak isimlendirilmektedir. Her bir sınıf için $P(c_j)$ başlangıç olasılığı, Eşitlik 17 ile şu şekilde hesaplanmaktadır:

$$P(c_j) = \frac{\sum_{i=1}^{|D|} P(c_j|d_i)}{|D|}, \quad i = 1, 2, \dots, |D|, j = 1, 2 \quad (17)$$

Eğer $P(c_j|d_i) = 1$ ise c_j , eğitim verisi d_i 'nin sınıfını gösterir, aksi halde bu ifadenin değeri, 0'dır.

Bir test verisi d verildiğinde naive Bayes sınıflandırma yöntemi, her bir sınıf için $P(c_j|d)$ olasılığını hesaplar ve olasılığı en yüksek sınıfı seçer. Bu da sınıflandırma algoritmasının sonucudur ve Eşitlik 18 ile şu şekilde hesaplanır:

$$P(c_j|d) = \frac{P(c_j)P(d|c_j)}{P(d)} \quad j = 1, 2 \quad (18)$$

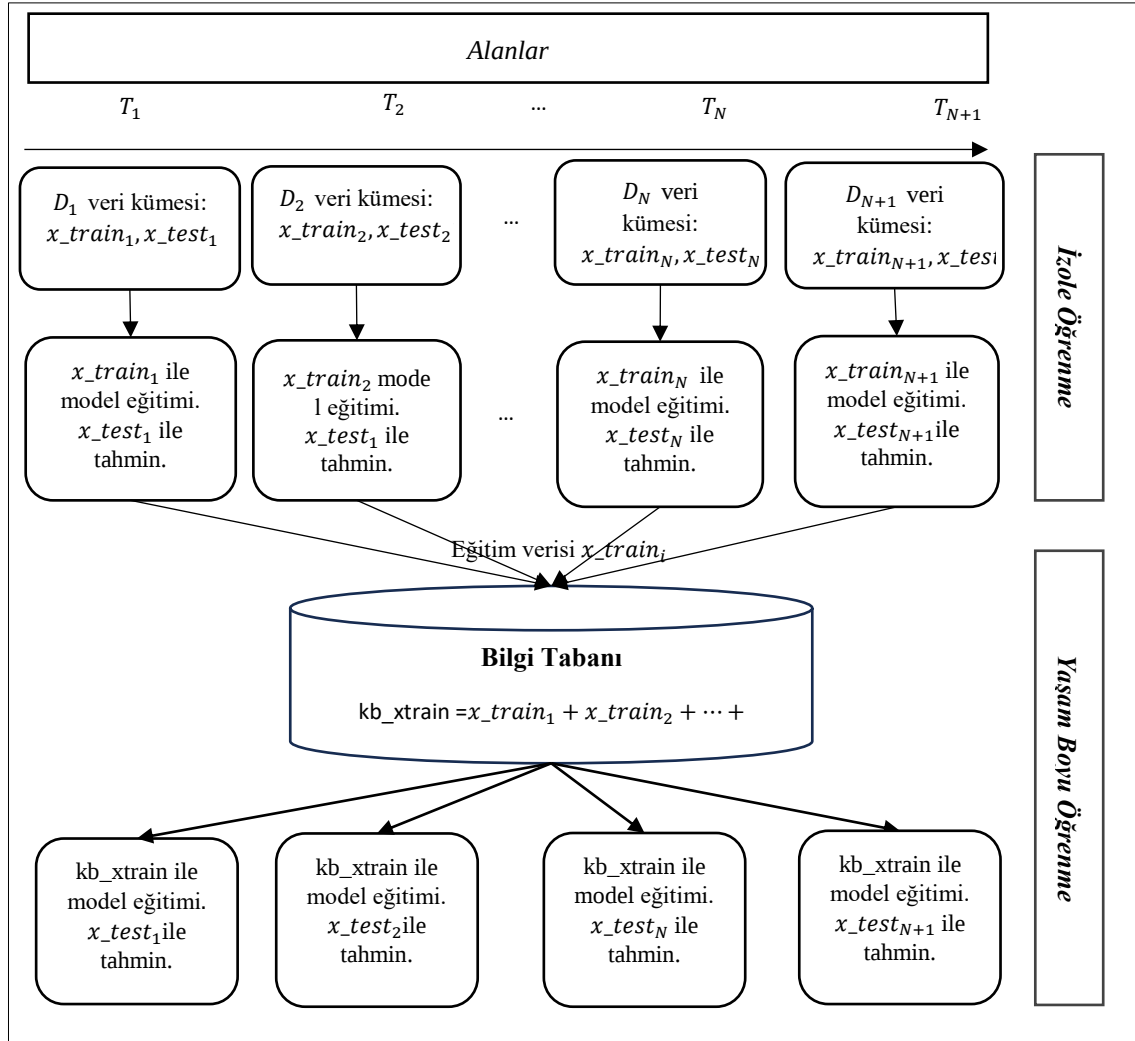
Eşitlik 19'da ise Eşitlik 18'in bir d test verisinde bulunan kelimeler için düzenlenmiş hali verilmiştir.

$$P(c_j|d) = \frac{P(c_j) \prod_{w \in d} P(w|c_j)^{n_{w,d}}}{\sum_{r=1}^{|C|} P(c_r) \prod_{w \in d} P(w|c_r)^{n_{w,d}}}, \quad j = 1, 2 \quad (19)$$

Burada n_w, d ; w kelimesinin d test verisinde görünme sayısıdır.

Bu tez çalışmasında, naive Bayes sınıflandırma yöntemi kullanılarak “Yaşam Boyu Naive Bayes Sınıflandırması” yapısı oluşturulmuş ve bu yapı, Şekil 7’de gösterilmiştir. Yaşam boyu naive Bayes sınıflandırmasına ait Python kodu Ek 1’de verilmiştir. Şekil 7’ye göre alanlar, model kurucu tarafından belirlenen sırayla modele verilir. Burada sisteme giriş yapan her alan, yeni alanı ve sistemde yeni alandan önce bulunan alanlar da geçmiş alanları ifade eder. Örneğin Şekil 7’de, T_1 alanı, ilk alandır. Alana ait D_1 veri kümesi, x_{train_1} eğitim verisi ve x_{test_1} test verisi olmak üzere ikiye ayrılır. Eğitim verisinden kelime vektörü oluşturulur. Buraya kadar yapılan işlemler, bir izole naive Bayes algoritmasında uygulanan işlemlerdir. Yaşam boyu sınıflandırma modelini eğitmek için bilgi tabanında bulunan geçmiş alanlara ait verilerin eğitim kümeleri kullanılmaktadır. T_1 alanı için bilgi tabanında sadece kendi eğitim kümesi bulunmaktadır. Sonraki alanlarda, her gelen eğitim kümesi içindeki kelimeler ve bunlara

ait sıklıklar kullanılır. Bilgi tabanında bulunan ve geçmiş tüm alanları içeren eğitim kümesi, kb_xtrain olarak isimlendirilmiştir. kb_xtrain veri kümesi kullanılarak bir $kb_xtrain_vektör$ oluşturulur. Yaşam boyu naive Bayes sınıflandırmasının eğitilmesinde $kb_xtrain_vektör$ kullanılmaktadır. Tahmin için o anda mevcut alanın test veri kümesi kullanılmaktadır. Bu sistem, alanların hepsi sınıflandırılincaya ya da model kurucu sistemi durduruncaya kadar çalışmaya devam etmektedir.

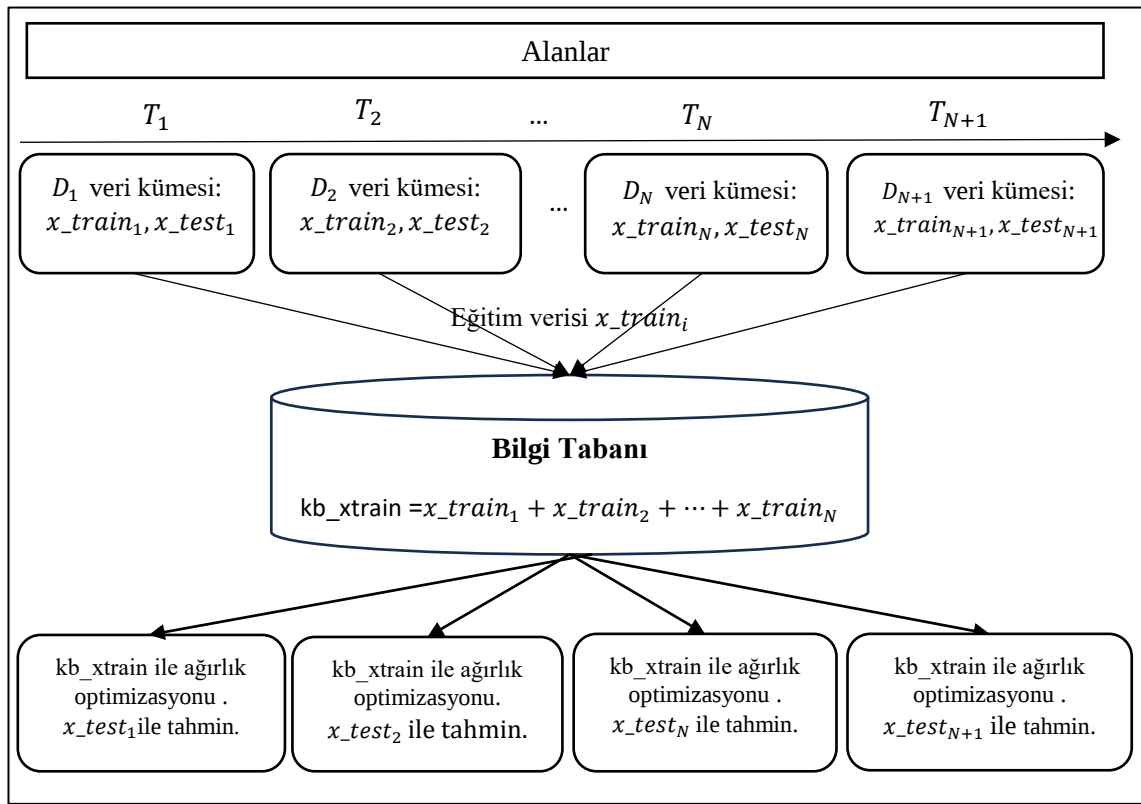


Şekil 7. Yaşam Boyu Naive Bayes Sınıflandırması

2.3.2. Yaşam Boyu Lojistik Regresyon Sınıflandırma Yöntemi

Lojistik regresyon yöntemi, metin sınıflandırmasında etkili sonuçlar veren bir yöntemdir (Onan, 2021: 3). Lojistik regresyon bir sınıflandırma modeli için yalnızca koşullu olasılık sonuçları üretmeyip logit adı verilen bu olasılıkların matematiksel bir dönüşümünü de model kurucuya vermektedir (Garson, 2014: 18).

Şekil 8`de alanlar ($T_1, T_2, \dots, T_N$) model kurucu tarafından belirlenen sırayla modele verilir. Mevcut alana ait veri kümesi (D_N), eğitim (x_{train_N}) ve test verisi (x_{test_N}) olarak ikiye bölünür. Ardından eğitim veri kümesi kullanılarak kelime vektörü oluşturulur. Burada metin vektörünü oluşturan kelimeler, regresyon denklemindeki x_i ($i = 1, 2, \dots, k$) katsayılarıdır. Regresyon denklemi katsayıları veya ağırlık olarak isimlendirilen β_i ($i = 1, 2, \dots, k$) değerleri, gerçek sınıf ve tahmin arasındaki farkı en aza indirecek parametre değerleridir. Model eğitiminde optimum ağırlıklar hesaplandıktan sonra test verisi sınıflandırılır. Buraya kadar yapılan işlemler, izole bir lojistik regresyon sınıflandırması ile aynıdır.



Şekil 8. Yaşam Boyu Lojistik Regresyon Sınıflandırması

Lojistik regresyon ile yaşam boyu duygu sınıflandırılmasında alanlar arasında aktarılabilecek bilgi, kelime sıklıklarıdır. Her bir mevcut alan için, geçmiş alanda ve kendi eğitim verisinde bulunan kelimelerin sıklıkları kullanılarak her bir kelime için ağırlık parametreleri hesaplanır. Ağırlıkların bulunmasında, literatürde farklı algoritmalar kullanılmaktadır. Bu tez çalışmasında ise stokastik gradyan inişi optimizasyon algoritması kullanılmıştır. Gradyan inişi algoritması, lojistik regresyon sınıflandırma yönteminde en iyi regresyon ağırlıklarını bulmak için sıklıkla kullanılan bir yöntemdir.

(Manogaran ve Lopez, 2018: 11). Gradyan inişi, bir modelin kayıp fonksiyonunu minimum yapan en uygun parametreleri bulmaya yarayan iteratif bir yöntemdir (Ruder, 2017: 2). Bu ağırlıklar kullanılarak mevcut alanın test verisi sınıflandırılır. Bu sistem, alanların hepsi sınıflandırılincaya ya da model kurucu sistemi durduruncaya kadar çalışmaya devam eder. Yaşam boyu lojistik regresyon sınıflandırmasına ait Python kodu, Ek 2`de verilmiştir.

2.3.3. Toplama Dayalı Yaşam Boyu Duygu Sınıflandırması

Bu tez çalışmasının uygulama bölümünde kullanılan model; Chen ve Liu (2015) tarafından yapılan çalışma temel alınarak ve Xia vd. (2017) ile Wang vd. (2019) tarafından yapılan çalışmalardan esinlenilerek hazırlanmıştır. Chen ve Liu (2015) tarafından yapılan çalışmada, ürün yorumları için pozitif ve negatif sınıflandırma yapılmıştır. Bu sınıflandırma yöntemine, yaşam boyu duygu analizi (lifelong sentiment classsification) adı verilmiştir. Yaşam boyu duygu analizinde yer alan temel öğeler, bilgi tabanı ve bilgi tabanında tutulan model parametrelerinden oluşmaktadır. Çalışmada bu veriler, yaşam boyu duygu analizi sistemine entegre edilerek bir naive Bayes yaklaşımı önerilmiştir. Xia, Jiang ve He (2017) tarafından yapılan çalışmada, İngilizce tweetler ve Çince ürün yorumları kullanılarak duygu sınıflandırması yapılmıştır. Duygu sınıflandırması, yaşam boyu öğrenme yaklaşımı kapsamında ele alınmıştır. Geçmiş alanlardan öğrenilen bilgiler saklanarak mevcut alan model performansının iyileştirilmesinde kullanılmıştır. Geçmiş alan bilgisi, geçmiş alanlara ait kelime sıklıklarını ifade etmektedir. Mevcut (yeni, şimdiki) alanın kelime sıklıkları ve geçmiş alanlardan elde edilen kelimelere ait sıklıklar toplanarak güncel kelime sıklıkları oluşturulmuştur. Sınıflandırma modelinin eğitiminde bu sıklıklar, girdi olarak kullanılmıştır. Bu mantık, toplama dayalı modelin temelini oluşturmaktadır. Wang, Liu, Wang, Ma ve Yang (2019) tarafından yapılan çalışmada ise yaşam boyu duygu sınıflandırması kapsamında geçmiş ve mevcut alanlar arasında aktarılan bilginin yönüne dikkat çekilmiştir. Yaşam boyu öğrenmede bilgi aktarımı, ileriye doğru (geçmiş alanlardan elde edilen bilgilerin, mevcut alanın performansının iyileştirmesinde kullanılması) olabileceği gibi geriye doğru da (herhangi bir geçmiş alanın performansının iyileştirilmesinde kullanılması) olabilmektedir. Çalışmada, metin verisini oluşturan her bir kelime ele alınmıştır. Geçmiş alanlarda bulunan kelimeler ve mevcut alanın kelimeleri karşılaştırılarak aynı kelimeler için sıklıklar toplanmış ve tüm kelimeler, modelin eğitimi için kullanılmıştır.

Naive Bayes sınıflandırma yöntemi, yaşam boyu öğrenme için oldukça uygundur. Geçmiş bilgiler, yeni görevin olasılıkları için öncül görevi görebilir. Yaşam boyu duygu analizi, bu fikirden ortaya çıkmıştır. Ancak yaşam boyu duygu analizi kapsamında cevaplanması gereken iki soru bulunmaktadır. Bunlardan ilki, şu anki görevin (test görevinin ya da $N + 1$. görevin) etiketli eğitim verisine sahip olduğu göz önüne alındığında, geçmiş alanlardan öğrenilen bilgilerin yeni görev sınıflandırmasına neden katkıda bulunabileceğidir. Bu sorunun cevabı, modeli eğitmek için etiketli eğitim verisinden alınan örneklemin örnek seçim yanlılığı ya da az miktarda eğitim veri kümesi nedeniyle test verisini tam olarak temsil etmeyebileceğidir. Örneğin bir duygu analizi uygulamasında test verisi, mevcut eğitim verilerinde bulunmayan ancak geçmiş bazı görevlerdeki verilerde görünen bazı duygu sözcükleri içerebilir. Böylece geçmiş bilgi, mevcut yeni görev için önceki duygu polarite bilgisini sağlayabilir. İkinci soru, geçmiş alanlar çok çeşitli olsa ve mevcut alana pek benzemese bile geçmiş bilginin neden mevcut veri kümesinin sınıflandırılmasına yardımcı olabileceğidir. Bunun temel nedeni, duygu sözcüklerinin ve ifadelerinin büyük ölçüde alanlardan bağımsız olmasıdır. Ürün yorumlarını ele alan bir duygu sınıflandırması için yapılan yorumlarda bulunan olumlu ve olumsuz sözcük ve ifadeler, alandan bağımsız olup alanlar arasında paylaşılmaktadır. Bu durumda çok sayıda önceki alan ile çalışmak, sistemin daha fazla olumlu ve olumsuz kelime öğrenmesine katkı sağlar. Bu nedenle yaşam boyu duygu analizi çalışmalarında çok sayıda birbiri ile ilişkili olmayan alanlardan çok sayıda veri kullanmak önemlidir (Chen ve Liu, 2018: 48- 49).

Chen ve Liu (2018) tarafından yapılan tanımlar ve bu çalışmada kullanılan naive Bayes yaklaşımı şu şekildedir:

Bölüm 2.4'te naive Bayes algoritması ve algoritmaya ait parametreler verilmiştir. Naive Bayes algoritmasının sonuçlarını etkileyen temel parametre $P(w, c_j)$, w 'nin c_j sınıfında bulunma sıklığına ($N_{c_j, w}$) göre hesaplanmaktadır. $N_{c_j, w}$ 'nin deneysel ifadesi, ikili sınıflandırmada $N_{+, w}$ ve $N_{-, w}$ olarak gösterilir. Bu, sınıflandırmayı iyileştirmek için bu sayıların uygun şekilde yenileyebileceğini göstermektedir. Yaşam boyu öğrenme duygu analizi sistemi kapsamında verilen temel öğeler kullanılarak algoritma güncellenmiştir. Bilgi tabanındaki geçmiş görevlerden toplanan bilgiler kullanılarak yeni bir toplama dayalı naive sınıflandırma önerilmektedir. Bu sınıflandırma yaklaşımında sanal sayılar (virtual counts) kullanılmaktadır. İkili bir sınıflandırma probleminde sanal sayılar, $X_{+, w}$ ve $X_{-, w}$ şeklinde gösterilebilir ve Eşitlik 20 ve Eşitlik 21 ile hesaplanır.

Sanal sayılar, bilgi tabanında bulunan geçmiş görevlere ait bilgiyi mevcut (şimdiki) görevden elde edilen bilgi ile birleştirilerek hesaplanmaktadır.

$$X_{+,w} = N_{+,w}^t + N_{+,w}^{KB} \quad (20)$$

$$X_{-,w} = N_{-,w}^t + N_{-,w}^{KB} \quad (21)$$

Sınıflandırıcıyı oluştururken Eşitlik 16'daki $N_{+,w}$ ve $N_{-,w}$, $X_{+,w}$ ve $X_{-,w}$ ile değiştirilir.

Bu çalışmada kullanılan model, her bir alana ait metin verisini oluşturan kelime sıklıklarının toplanmasına dayalıdır. Farklı alanlara ait metin verisi, kendi alanını temsil eden kelimelerin yanı sıra diğer tüm alanlarla ortak paylaşılan kelimeleri de içermektedir (Chen ve Liu, 2018: 2). Ayrıca olumlu veya olumsuz bir kelime, bir alanın eğitim verilerinde görünmeyebilir ancak bu alanın test verilerinde görünebilir. Bu durumlarda, diğer alanlardan gelen geçmiş bilginin kullanılması fayda sağlamaktadır. Çünkü çoğu kelime, alanlar arasına benzer anlama sahiptir (Wang vd., 2019: 3).

Şekil 9'da "Toplama Dayalı Yaşam Boyu Duygu Sınıflandırması" algoritması verilmiştir. Bu algoritmaya göre veri kümeleri sisteme, model kurucunun belirlediği sırayla verilir. Başka bir deyişle, T_i alanına ait D_i metin verisi sisteme i . sırada gelmektedir. D_i metin veri kümesi, eğitim ve test verisi olarak ikiye ayrılır. Eğitim verisi, geçmiş alanların ve mevcut T_i alanı eğitim verisinin içinde bulunduğu *eğitimveriseti* içerisinde saklanır. Ardından *eğitimveriseti* içinde bulunan pozitif ve negatif yorumlar belirlenir. Bu yorumlar, kelimelerine ayrılarak *pozitifkelimeler* ve *negatifkelimeler* listeleri ile bu kelimelere ait sıklıkların bulunduğu *pozitifkelimesıklıkları* ve *negatifkelimesıklıkları* listeleri oluşturulur. Bu listeler oluşturulurken geçmiş alanlar ve mevcut T_i alanı eğitim verisinde bulunan ortak kelimeler için sıklıklar birleştirilerek toplanır. Bu toplama işlemi yapılırken naive Bayes sınıflandırma yönteminin temelinde yatan koşullu olasılıkları hesaplayabilmek için kelimelerin pozitif ve negatif sınıflarda bulunma sıklıkları ayrı ayrı hesaplanmıştır. Bu yapı, bir kelimenin pozitif yorumda ve negatif yorumda geçme sıklığına bağlı olarak iki farklı sıklık değerine sahip olmasını gerektirir. Bir kelimenin sıklığı, pozitif yorumlarda ve negatif yorumlarda bulunma durumuna göre ayrı ayrı hesaplanmakta ve *bilgideposunda* saklanmaktadır. Ardından *bilgideposunda* bulunan her bir kelime için *pozitifkoşulluolasılık* ve *negatifkoşulluolasılık* parametreleri hesaplanır. Bu parametreler, test verisini

oluşturan her bir ürün yorumu için *pozitif_x_test_cümleolasılığı* ve *negatif_x_test_cümleolasılığı*nın hesaplanmasında kullanılır. Olasılık değerlerinin karşılaştırılması sonucu, cümlelerin duygu sınıfı belirlenmektedir. Toplama dayalı yaşam boyu duygu sınıflandırmasın ait Python kodu, Ek 3`te verilmiştir.

Girdi: Yeni bir alana ait metin verisi

Çıktı: Yeni alan test verisi duygu sınıfı

for yeni alan T_i için;

$veri_i \leftarrow T_i$ alanı metin verisi D_i

$etiket_i \leftarrow D_i$ metin verisi duygu sınıfı

$x_{train_i}, x_{test_i} \leftarrow veri_i$ eğitim ve test verisi bölme

$y_{train_i}, y_{test_i} \leftarrow etiket_i$ eğitim ve test verisi bölme

$eğitimveriseti[x_{train}] = x_{train_i}$

$eğitimveriseti[y_{train}] = y_{train_i}$

$pozitifyorumlarindeks = [eğitimveriseti[y_{train}] == 1]$

$negatifyorumlarindeks = [eğitimveriseti[y_{train}] == 0]$

$pozitifyorumlar = eğitimveriseti[pozitifyorumlarindeks]$

$negatifyorumlar = eğitimveriseti[negatifyorumlarindeks]$

$pozitifkelimesıklıkları \leftarrow pozitifyorumlar$ kelime sıklık hesaplama

$pozitifkelimeler \leftarrow pozitifyorumlar$

$negatifkelimesıklıkları \leftarrow negatifyorumlar$ kelime sıklık hesaplama

$negatifkelimeler \leftarrow negatifyorumlar$

$kelime \leftarrow pozitifkelimeler \cup negatifkelimeler$

$bilgideposu \leftarrow kelime, pozitifkelimesıklıkları, negatifkelimesıklıkları$

$pozitifkoşulluolasılık \leftarrow bilgideposu(kelime, pozitifkelimesıklıkları)$

$negatifkoşulluolasılık \leftarrow bilgideposu(kelime, negatifkelimesıklıkları)$

$bilgideposu \leftarrow pozitifkoşulluolasılık$

$bilgideposu \leftarrow negatifkoşulluolasılık$

$x_{test_kelimeler} \leftarrow x_{test_i}$

$pozitif_x_test_cümleolasılığı \leftarrow x_{test_kelimeler}, bilgideposu$

$negatif_x_test_cümleolasılığı \leftarrow x_{test_kelimeler}, bilgideposu$

if $pozitif_x_test_cümleolasılığı > negatif_x_test_cümleolasılığı$:

$x_{test} \leftarrow pozitif$

else:

$x_{test} \leftarrow negatif$

end

Şekil 9. Toplama Dayalı Modele Ait Algoritma

2.4. Yaşam Boyu Öğrenme Yaklaşımına ve Duygu Analizine İlişkin Literatür Taraması

Yaşam boyu öğrenme yaklaşımı, uzun yıllar önce önerilmiş olsa da konuyu ele alan uygulamalar henüz yaygınlaşmaya başlamıştır (Chen ve Liu, 2018: 8). Literatürde bulunan çalışmalara bakıldığında çoğunlukla İngilizce metin verisi kullanılmıştır. Bunun nedeni, İngilizce'nin dünya genelinde yaygın konuşulan bir dil olması (Zeng ve Yang, 2024: 1) ve internet üzerinde yer alan metin verilerinin de çoğunlukla İngilizce kullanılarak oluşturulmasıdır (Dunn, 2020: 5-6). Bununla birlikte az da olsa Çince (Xia vd., 2017) ve Vietnamca (Ha vd., 2016) gibi farklı dillerde yapılan çalışmalar bulunmaktadır. Ulusal ve uluslararası literatür tarandığında, Türkçe metin verisi kullanılarak yaşam boyu öğrenme yaklaşımı kapsamında yapılan bir çalışmanın henüz bulunmadığı görülmüştür.

Yapılan çalışmalarda genellikle Amazon online alışveriş sitesinden toplanan ürün yorumları kullanılmıştır. Çince veri kümesi ile yapılan çalışmada Weibo.com sitesinden, Vietnamca veri kümesi ile yapılan çalışmada Tiki.vn sitesinden alınan ürün yorumları kullanılmıştır. Ayrıca film yorumları kullanılarak yapılan çalışmalar da mevcuttur. Bazı çalışmalarda veri kaynağı olarak Twitter sosyal medya platformu kullanılmıştır. Veri kaynakları ya da metin verisinin toplandığı alanlar farklı olsa da yapılan çalışmalar, problemin çözümü için yeni modellerin önerilmesi ve model performanslarının iyileştirilmesini amaçlamaktadır. Bu nedenle elde edilen sonuçlar, teoride kalmış ve gerçek dünya uygulamalarına yönelik çalışmaya rastlanmamıştır.

Yaşam boyu öğrenme yaklaşımı kapsamında literatürdeki ilk çalışma, Chen ve Liu (2015) tarafından yapılmıştır. Çalışmanın amacı, daha iyi ve geniş kapsamlı eğitim verileri oluşturmaktır. Geçmiş alanlardan elde edilen eğitim verileri kullanılarak hedef alan (target domain) model performansının artırılmasını sağlamaktır. Çalışmada, naive Bayes ve destek vektör makineleri algoritmaları kullanılarak dengelenmiş ve dengelenmemiş veri kümeleri için uygulama yapılmıştır. Veri kümesi olarak Amazon online alışveriş sitesinden 20 farklı ürüne ilişkin İngilizce müşteri yorumları alınmıştır. Yöntemlerdeki amaç fonksiyonunun optimizasyonunda stokastik gradyan inişi yöntemi kullanılmıştır.

Ha, Pham, Nguyen, Vuong ve Tran Nguyen (2016), Chen vd. (2015) tarafından önerilen yaşam boyu öğrenme yaklaşımını kullanmıştır. Bu yaklaşıma ek olarak öznitelik seçiminde bigram, unigram ve bag of bigrams yöntemleri uygulanmış ve aralarındaki farklar incelenmiştir. Çalışmada iki farklı veri kümesi kullanılmıştır. Bu çalışma, veri

kümesinin Vietnamca dilinde olması bakımından ilktir. Vietnamca veri kümesi, Tiki.vn sitesinden 10 alanda toplanan 15.394 Vietnamca yorumlardan oluşmaktadır. İngilizce yorumlardan oluşan veri kümesi, dengelenmiş veri kümesi iken Vietnamca yorumlardan oluşan veri kümesi, dengelenmemiş veri kümesidir.

Yaşam boyu öğrenme modellerinde, geçmiş görevlerden öğrenilen bilgiler saklanmakta ve bu bilgiler, yeni görevin sınıflandırılmasında kullanılmaktadır. Bu bilgilerin öğrenilmesi ve saklanması, belirli kurallar ile sağlanmaktadır. Khan, Yar ve Khalid (2016) tarafından yapılan çalışmada, oluşturulan kuralların doğruluğunu ve kalitesini artırmak için histogram tabanlı kural doğrulama algoritması kullanılmıştır. Çalışmada veri kümesi olarak Amazon online alışveriş sitesindeki İngilizce ürün yorumları kullanılmıştır.

Uzaktan denetim (distant supervision), metin içinde bulunan duygu ifadelerini (emoji) doğal duygu etiketi olarak kabul eden bir tekniktir. İnternette giderek artan miktarda veri olmasına karşın etiketli eğitim verisi bulmak ya da oluşturmak oldukça zordur. Xia, Jiang ve He (2017), bu zorluğun aşılması için uzaktan denetlenen yaşam boyu duygu analizi modeli oluşturmuştur. Çalışmada, iki farklı veri kümesi kullanılmıştır. İlki, literatürdeki bir çalışmadan alınmış Twitter İngilizce veri kümesi olup 9.000.000 tweet içermektedir. Diğer veri kümesi ise Weibo.com sitesinden alınan 1.000.000 Çince yorumlardır. İngilizce veri kümesi 6 farklı alanı, Çince veri kümesi ise 3 farklı alanı içermektedir.

Veri kaynağı olarak Twitter sosyal medya platformunun kullanıldığı bir diğer çalışma, Shende ve Waghmare (2019)'a aittir. Twitter'dan toplanan filmlere ait tweetler ve Kaggle sitesinde bulunan hazır film yorumları kullanılarak bir yaşam boyu öğrenme duygu sınıflandırması önerilmiştir. Naive Bayes ve destek vektör makineleri algoritmaları kullanılarak denetimli öğrenme modeli kurulmuştur. Çalışmada iki farklı veri kümesinin birleştirilerek kullanılması ile yöntemin etkinliği kanıtlanmıştır.

Wang, Liu, Wang, Ma ve Yang (2019), yaşam boyu öğrenme yaklaşımında geçmiş bilginin tersine akışı ile ilgilenmiştir. Yaşam boyu öğrenmede geçmiş görevlerden elde edilen bilgi, gelecek/mevcut görevin sınıflandırılmasında kullanılmaktadır. Bu, ileri yönlü bir bilgi akışını göstermektedir ve tek yönlüdür. Fakat yaşam boyu öğrenme, ileri ve geri yönlü bilgi akışı sağlayabilmelidir. Bu nedenle geçmiş bir görevin, kendi eğitim verisi kullanılmadan bilgi tabanı kullanılarak sınıflandırma performansının iyileştirilmesi hedeflenmiştir. Bunun için naive Bayes algoritması kullanılmıştır. İzole naive Bayes, izole destek vektör makinaları, basit yaşam boyu naive Bayes ve destek vektör makinaları

ve önerilen yeni naive Bayes sınıflandırma yöntemi kullanılarak model performansları karşılaştırılmıştır.

Xu, Pan ve Xia (2020) tarafından yapılan çalışmada üç farklı veri kümesi kullanılmıştır. Bunlar; film, iletişim ve sağlık alanlarında verilerin bulunduğu 3 sınıflı (pozitif, nötr, negatif) Amazon veri kümesi, Cornell film yorumlarından oluşan 2 sınıflı (pozitif, negatif) veri kümesi ve kitap, dvd, mutfak, elektronik alanlarında verilerin bulunduğu 2 sınıflı Amazon veri kümesidir. Bu veri kümeleri kullanılarak büyük boyutlu ve çeşitli alanlarda e-ticaret sitelerinde bulunan ürün yorumları için sürekli öğrenme modeli önerilmiştir. Geleneksel naive Bayes algoritmasının yüksek hesaplama etkinliğini korumak için algoritmanın parametre tahmin mekanizması genişletilmiştir. Mevcut ve önerilen haliyle naive Bayes algoritmasına ait sonuçlar, karşılaştırmalı olarak verilmiştir.

Hong, Guan, Man ve Wong (2020), Chen ve Liu'nun (2015) çalışmalarında kullandığı Amazon ürün müşteri yorumları veri kümesini kullanmıştır. Ayrıca Amazon veri kümesinden oldukça farklı olan bir film yorumları veri kümesi de kullanılmıştır. Film yorumları veri kümesinin kullanılmasının nedeni, araştırmacıların görev benzerliğinin yaşam boyu öğrenmeyi nasıl etkilediğini anlamalarına yardımcı olmaktır. Veri kümesi, ön işlemeden geçirilmiş ve tokenlarına ayrılmıştır. Çalışmada yaşam boyu öğrenme kapsamında bir BERT modeli önerilmiştir. Naive Bayes algoritmasının, Chen ve Liu (2015) tarafından önerilen yaşam boyu öğrenme algoritmasının ve önerilen BERT modelinin performansları karşılaştırılmıştır.

Wang, Mazumder, Wang, Liu, Yang ve Li (2020) tarafından yapılan çalışmada Chen ve Liu'nun (2015) çalışmalarında kullandığı Amazon ürün müşteri yorumları veri kümesi ve internet ortamında bulunan büyük boyutlu Amazon ürün yorumları veri kümesi kullanılmıştır. Çalışmada, yapay sinir ağları ve yeni önerilen Bayes destekli yaşam boyu dikkat ağı algoritması kullanılmıştır. Önerilen yeni modelin temel fikri, dikkat bilgisini öğrenmek için naive Bayes'in üretken parametrelerinden yararlanmaktır. Modelde parametrelerin elde edilmesi ve güncellenmesi, yapay sinir ağları ile gerçekleştirilmiştir. Naive Bayes, yapay sinir ağları ve önerilen yaşam boyu duygu sınıflandırması yöntemlerinin performansları karşılaştırılmıştır.

Geng, Yang, Yuan, Wang, Ao ve Xu (2021), yaşam boyu öğrenme için yeni bir yöntem önermiştir. Bu yöntem, derin öğrenme temelli ağ budama yöntemidir. Yöntemin amacı, büyük derin ağlardaki gereksiz parametreleri kaldırmaktır. Önerilen yöntemin temelinde BERT algoritması yer almaktadır. Veri kümesi olarak Amazon ürün yorumları

ve film yorumları kullanılmıştır. Önerilen yeni algoritma, literatürde bulunan farklı algoritmalarla karşılaştırılmıştır.

Zhang, Wang, Yuan, Geng ve Yang (2023), Amazon ürün yorumlarını kullanarak yaşam boyu duygu analizi sınıflandırması için Bayes tabanlı bir online öğrenme modeli önermiştir. Önerilen model, BERT modeline entegre edilerek sınıflandırma performansının iyileştirilmesi amaçlanmıştır. Önerilen model, aynı zamanda farklı metinlerden oluşan veri kümeleri için de kullanılarak performans sonuçları değerlendirilmiştir.

Tablo 3'te bu bölümde ele alınan yaşam boyu duygu sınıflandırması çalışmalarının bir özeti sunulmaktadır. Ele alınan çalışmalar; uygulanan yöntem, veri kaynağına ve metin verisinin diline göre değerlendirilmiştir.

Tablo 3. Yaşam Boyu Duygu Sınıflandırmasını Ele Alan Çalışmaların Bir Özeti

Yazar(lar)	Yöntem(ler)	Veri Kaynağı	Veri Kümesinin Dili
Chen, Ma, Liu (2015)	Naive Bayes, Destek Vektör Makineleri	Amazon	İngilizce
Ha, Hoang, Nghiem (2016)	Naive Bayes	Amazon, Tiki.vn	İngilizce Vietnamca
Khan, Yar, Khalid (2016)	Histogram Tabanlı Algoritma	Amazon	İngilizce
Xia, Jiang, He (2017)	Naive Bayes, Lojistik Regresyon	Twitter, Weibo.com	İngilizce Çince
Shende, Waghmare (2019)	Naive Bayes, Destek Vektör Makineleri	Twitter, Kaggle	İngilizce
Wang, Liu, Wang, Ma ve Yang (2019)	Naive Bayes, Destek Vektör Makineleri	Amazon	İngilizce
Xu, Pan, Xia (2020)	Naive Bayes	Amazon, Film Yorumları	İngilizce
Hong, Guan, Man, Wong (2020)	Naive Bayes, BERT	Amazon, Film Yorumları	İngilizce
Wang, Mazumder, Wang, Liu, Yang Li (2020)	Naive Bayes, Yapay Sinir Ağları	Amazon	İngilizce
Geng, Yang, Yuan, Wang, Ao, Xu (2021)	Derin Öğrenme, BERT	Amazon, Film Yorumları	İngilizce
Zang, Wang, Yuan, Geng, Yang (2023)	BERT, Adaptive Uncertainty Regularization	Amazon	İngilizce

Literatürde yaşam boyu öğrenme yaklaşımını cümle ve belge seviyesinde duygu analizi olarak ele alan çalışmaların yanı sıra yaşam boyu öğrenme yaklaşımında hedef tabanlı duygu analizi ve konu modelle yaklaşımlarını dikkate alan çalışmalar da bulunmaktadır. Yaşam boyu öğrenme yaklaşımının konu alan hedef tabanlı duygu analizi çalışmaları için Khan vd. (2016a), Shu vd. (2017), Wang vd. (2018), Amplayo vd. (2018) ve Ke vd. (2021), Ramos ve Fuentes (2023) incelenebilir.

ÜÇÜNCÜ BÖLÜM

MÜŞTERİ ÜRÜN YORUMLARININ

YAŞAM BOYU DUYGU ANALİZİ İLE SINIFLANDIRILMASINA

İLİŞKİN BİR UYGULAMA

Metin madenciliği, yapılandırılmamış ya da yarı yapılandırılmış metin verisinden anlamlı ilişkiler ortaya çıkarma ve analiz etme sürecidir (Vijayarani vd., 2015: 7). Duygu analizi ise metin verilerinin ifade ettiği duygunun analiz edilmesini amaçlayan bir sınıflandırma yöntemidir. Makine öğrenmesi yöntemleri, bir sınıflandırma görevi olarak duygu analizi gerçekleştirmektedir. Etiketli metin verisi kullanılarak metin vektörleri oluşturulur. Bu vektörler kullanılarak makine öğrenmesi yöntemleri eğitilir ve sınıflandırma yapılır (Jo, 2019: 96). Bu tez çalışması kapsamında duygu sınıflandırması problemi için, metin sınıflandırmasında etkili olduğu bilinen denetimli makine öğrenmesi yöntemlerinden naive Bayes ve lojistik regresyon sınıflandırma yöntemleri kullanılmıştır (Hasanli ve Rustamov, 2019: 3). En basit haliyle makine öğrenmesi yöntemleri izole şekilde öğrenmektedir. Başka bir deyişle model, sadece bir alandan (konudan, içerikten) elde edilmiş veri kümesi ile o alana özgü olarak eğitilir ve model performansında düşüş olmaması için genellikle farklı bir alandan elde edilmiş veri kümesinin sınıflandırılmasında kullanılmaz (Aytekin ve Bayram, 2021: 11). Diğer yandan günümüzdeki teknolojik gelişmeler, artan veri miktarı, zaman ve para kaybı gibi sebeplerle her alan için farklı bir model oluşturulması mümkün değildir. Bu nedenle izole öğrenmeyi başka bir deyişle her alan için farklı model oluşturmayı ortadan kaldıracak bir sürekli öğrenme yaklaşımı uygun olacaktır. Son yıllarda gittikçe artan bu görüş, yaşam boyu öğrenme yaklaşımını ortaya koymuştur. Bu tez kapsamında da sürekli öğrenmeyi temel alan yaşam boyu öğrenme yaklaşımı ele alınmış ve yaşam boyu öğrenmeye ilişkin bir uygulama yapılmıştır.

3.1. Veri Kümesi

Bu çalışmanın temel konusu, duygu analizi ve duygu analizinin bir uzantısı olan yaşam boyu öğrenme yaklaşımıdır. Bu çalışmada yaşam boyu öğrenmeye ilişkin bir uygulama yapılmıştır. Uygulamanın temel amacı, bir online alışveriş sitesinde müşteriler tarafından ürünlere yapılan yorumların metin madenciliği yöntemlerinden biri olan duygu analizi ile sınıflandırılmasıdır. Genel olarak yapılan uygulamada farklı alanlardan

toplanan müşteri yorumları, doğal dil işleme yöntemleri kullanılarak işlenmiştir. Elde edilen veriler; naive Bayes, lojistik regresyon ve toplama dayalı model kullanılarak sınıflandırılmış ve tüm modellerin performansları karşılaştırılmıştır.

Çalışmada bir online alışveriş sitesinden elde edilen müşteri yorumları kullanılmıştır. Bu kapsamda alışveriş sitesinden rastgele seçilen on farklı alandan, ürünü satın alan müşterilerin ürün hakkında yaptıkları yorumlar toplanmıştır. Seçilen 10 farklı alan; “Airfryer”, “Bebek bezi”, “Çanta”, “Çikolata”, “Gömlek”, “Gözlük”, “Parfüm”, “Peynir”, “Saat” ve “Toz deterjan”dır. Veriler, Şubat 2024 ve Haziran 2024 zaman aralığında toplanmıştır. Bu zaman aralığında toplanan yorum sayısı, 30.000 adettir. Müşteri yorum sayıları, her alanda farklılık göstermektedir. Başka bir deyişle, alanlardan toplanan yorum sayısı eşit değildir. Yorumlar, Python yazılım programı BeautifulSoup¹ ve Selenium² kütüphaneleri aracılığıyla toplanmıştır. Bu kütüphaneler yardımıyla ürüne ait yorum, yorum bilgisi ve ürün yıldız sayıları toplanmıştır. Toplanan yorumlar, Microsoft Excel belgesi olarak saklanmış olup Tablo 4’te, oluşturulan veri yapısı örneği görülmektedir.

Tablo 4. Yorum Veri Yapısı Örneği

Sıra	Yorumlar	Yorum Bilgisi	Yıldız Sayısı
0	Ürünlerini defalarca aldığım bir satıcı.Ürünleri çok güzel lezzetli, paketlemesi güzel, aynı görüldüğü şekilde geldi...Fiyatı uygun, teslimatı hızlı ve sorunsuzdu Teşekkür ederim. 🍏	Y** B**Elite Üye 21 Haziran 2023	5
1	Çocuklar için aldım gayet güzel	k** u** 29 Ocak 2023	5
2	Çocuklar bayıldı🍏 teşekkürler 🍏	G** K** 2 Şubat 2023	5
3	çocukların severek tükettiği aşırı sevimli kahve yani e atıştırılabilirlik olarak indirimli oldukça stoklanacak bı ürün	A** Ö** 1 Ekim 2023	5
4	aşırı lezzetsiz. bunu fenomenler nasıl övüyor anlamıyorum. yemediklerine eminim...	A** k** 16 Mayıs 2023	1

3.2. Veri Ön İşleme

Toplanan yorumlarda kelime hataları, eksik ve fazla harf kullanımı, sayı ve özel karakter kullanımı, emojiler, linkler veya yabancı sözcükler gibi model eğitimi sırasında

¹ <https://pypi.org/project/beautifulsoup4/>

² <https://pypi.org/project/selenium/#description>

fayda sağlamayan, model performansını olumsuz yönde etkileyen durumlar ile karşılaşılmaktadır. Bu durumların ortadan kaldırılması için metin ön işleme yapılmıştır.

Veri ön işleme aşamasında öncelikle dengeli bir veri kümesi oluşturmak için pozitif ve negatif olmak üzere, her iki sınıftan eşit sayıda veri alınmıştır. Yorumların seçilmesinde yorumun, anlam bakımında farklı kelimeler içermesine ve yorumda bulunan kelime sayısının fazla olmasına dikkat edilmiştir

Denetimli öğrenme modellerinde, model doğruluğunu etkileyen önemli etkenlerden biri, eğitim verisinin etiketlenmesidir. Verilerin etiketlenmesinde ürün yıldız sayıları kullanılmıştır. Literatürdeki bazı çalışmalarda (örneğin; Chen vd., 2015; Ha vd., 2018; Wang vd., 2019; Pervan, 2019; Aytekin ve Bayram, 2021; Yavuz, 2023) ürün yıldız sayılarının, doğrudan duygu etiketi olarak kullanıldığı örneklere rastlanmaktadır. Ancak bazı durumlarda olumsuz yorum yapan kişilerin ürüne yüksek yıldız verdikleri ya da olumlu yorum yapan kişilerin ürüne düşük yıldız verdikleri durumlar ile karşılaşılmaktadır. Diğer çalışmalardan farklı olarak bu çalışmada yıldız sayılarına ek olarak yorumlar, tek tek okunmuş ve el ile pozitif ve negatif olarak etiketlenmiştir. Yorumların el ile etiketlenmesi sonucu her bir alandan 500 pozitif, 500 negatif olmak üzere 1.000 adet yorum elde edilmiştir.

Model doğruluğunu etkileyen diğer önemli etken ise metin belgesindeki yazım yanlışlarıdır (Devi ve Saleema, 2024: 1). Bu nedenle veri kümelerinde bulunan yazım yanlışları ve harf hataları, el ile düzeltilmiş ve yabancı dillerde yapılan yorumlar, Türkçe'ye çevrilmiştir. El ile yapılan ön işleminin ardından Python NLTK³ kütüphanesi kullanılarak tüm sayı ve rakamlar, emojiler, noktalama işaretleri silinmiştir. Tüm yorumlar küçük harfe dönüştürülmüştür. Cümle içerisinde anlam ifade etmeyen durak kelimeler silinmiştir. Ancak durak kelimelerin silinmesinde NLTK kütüphanesi yetersiz kalmaktadır. Bu nedenle Türkçe durak kelime listesi oluşturulmuş ve modele eklenerek ön işleme yapılmıştır. Oluşturulan durak kelime listesi, Ek 4`te verilmiştir.

Veri kümesinde nadir de olsa tüm yorumun alınamadığı durumlar söz konusudur. Bu durum, uzun yorumlarda (genellikle yorumda bulunan kelime sayısı elli ve üzerinde ise) ortaya çıkmaktadır. El ile ön işleme aşamasında yorum düzenlenirken, yorumun bildirdiği duygu durumu değişmeyecek şekilde en anlamlı biçimde kısaltılmıştır.

³ <https://pypi.org/project/nltk/>

Ön işleme ve el ile etiketleme sonunda her bir alanda 500 pozitif ve 500 negatif yorum olmak üzere toplam 1000 adet yorum bulunmaktadır. 10 alanın toplam yorum sayısı ise 10.000 adettir.

3.3. Model Eğitimi İçin Verilerin Hazırlanması

Veri ön işleme aşamasından sonra modelin eğitilmesi için verilerin hazırlanması aşamasına geçilmiştir. Bunun için öncelikle yorumların tokenlarına ayrılması gerekmektedir. Bu işlem için Python NLTK kütüphanesi kullanılmıştır. Elde edilen tokenların köklerine ayrılmasında, Zeyrek⁴ kütüphanesi kullanılmıştır. Zeyrek, Türkçe kelimelerin kelime dağarcığını oluşturmak ve analiz etmek için Zemberek kütüphanesinin Python'a kısmi bir aktarımıdır.

Köklerine ayrılmış tokenlara örnekler şu şekildedir:

Kelime: aralıkta

- (aralık_Noun)(-)(aralık:noun_S + a3sg_S + pnon_S + ta:loc_ST)
- (ara_Adj)(-)(ara:adjectiveRoot_ST + lık:ness_S + noun_S + a3sg_S + pnon_S + ta:loc_ST)

Kelime: hızlı

- (hızlı_Adj)(-)(hızlı:adjectiveRoot_ST)
- (hızlı_Adv)(-)(hızlı:advRoot_ST)
- (hız_Noun)(-)(hız:noun_S + a3sg_S + pnon_S + nom_ST + lı:with_S + adjectiveRoot_ST)

Kelime: beğendi/beğendik

- (beğenmek_Verb)(-)(beğen:verbRoot_S + di:vPast_S + k:vA1pl_ST)
- (beğenmek_Verb)(-)(beğen:verbRoot_S + dik:vPastPart_S + adjAfterVerb_S + aPnon_ST)

3.4. Metin Vektörlerinin Oluşturulması

Model eğitimi için kullanılacak metin vektörleri, Python Sklearn⁵ kütüphanesinde bulunan CountVectorizer ve TfidfVectorizer yöntemleri kullanılarak oluşturulmuştur. CountVectorizer vektör oluşturma yönteminde, bir metin belgesi token sayısı matrisine dönüştürülür. Kelime çantası yöntemi kullanılarak kelimelerin metin içinde geçme

⁴ <https://github.com/obulat/zeyrek>

⁵ <https://pypi.org/project/scikit-learn/>

sıklıkları kullanılır (Barry, 2017: 3). TfidfVectorizer vektör oluşturma yönteminde ise kelimelerin metinde geçme sıklıklarına ek olarak bir kelimenin ne kadar önemli olduğunu dikkate alan ters belge sıklığı kullanılarak hesaplanan TF-IDF ağırlıkları kullanılmaktadır (Zong vd., 2021: 35)

Metin vektörleri oluşturulurken herhangi bir öznitelik seçme yöntemi kullanılmamıştır. Çünkü yaşam boyu öğrenme yaklaşımında saklanması gereken en önemli bilgi, metin verisini oluşturan kelimelerdir. Bir kelime, metin içerisinde az görülebilir ancak bu kelime, diğer alanlarda da kullanılabileceğinden öğrenme süreci boyunca bulunma sıklığı artacaktır. Bu anlamda en sık geçen kelimeler, en çok öğrenilen kelimelerdir (Chen ve Liu, 2018: 3). Eğitim vektörünü oluşturan kelime listesi, Ek 5`te verilmiştir.

3.5. Araştırma Modeli

Bu çalışmada bir online alışveriş sitesinde müşteriler tarafından 10 farklı alanda ürünlere yapılan yorumların metin madenciliği yöntemlerinden biri olan yaşam boyu duygu analizi ile sınıflandırılmasında üç farklı analiz gerçekleştirilmiştir. İlk iki analizde, naive Bayes ve lojistik regresyon yöntemlerine yaşam boyu öğrenme yaklaşımı entegre edilmiştir. Bu uygulama, en basit haliyle yapılmıştır. Son olarak toplama dayalı yaşam boyu duygu sınıflandırması modeli kullanılmıştır. Bu model, öğrenme süreci boyunca tüm veri kümelerinde bulunan kelime sıklıklarının toplama yoluyla birleştirilmesini temel almaktadır.

3.6. Uygulama Çıktıları

Çalışmanın bu bölümünde yaşam boyu öğrenme yaklaşımı altında üç farklı sınıflandırma yöntemi kullanılmış ve duygu sınıflandırması yapılmıştır. İlk olarak naive Bayes ve lojistik regresyon yöntemleri kullanılarak sınıflandırma yapılmıştır. İkinci olarak ise toplama dayalı yaşam boyu duygu sınıflandırması modeli ile duygu sınıflandırması yapılmıştır. Tüm yöntemlerden elde edilen sonuçlar verilmiş ve yöntemlerin performansları karşılaştırılarak analiz edilmiştir.

Yaşam boyu duygu sınıflandırması kapsamında on farklı alana ait metin verileri kullanılarak uygulama yapılmıştır. Yapılan uygulamalarda kelime vektörü oluşturma yöntemleri olarak CountVectorizer ve TF-IDF kullanılmıştır. Oluşturulan metin vektörlerinin büyüklüğü, ilk olarak maksimum öznitelik sayısı kadar alınmıştır. Ayrıca öznitelik sayının 1000 ve 500 olarak ele alındığı metin vektörleri ile denemeler de

yapılmıştır. Bunun yanı sıra kelime kökleri bulunarak kelime köklerinin, model performansına etkisi incelenmiştir.

Yaşam boyu öğrenmede model performansını etkileyen bir diğer etken ise alanların sisteme hangi sıra ile verildiğidir (Chen ve Liu, 2018: 17). Tablo 5`te örnek olarak Alan Sıralaması 1 verilmiştir. Bu sıralama; airfryer, gömlek, çanta, toz deterjan, bebek bezi, saat, gözlük, parfüm, peynir, çikolatadır. Tablo 6`da ise Alan Sıralaması 1`deki sıra ile sisteme verilen alanlar için yaşam boyu naïve Bayes ve yaşam boyu lojistik regresyon sınıflandırma yöntemlerinin model performansları sunulmuştur.

Tablo 5. Alan Sıralaması 1

Alan Sırası	Alanlar
1	Airfryer
2	Gömlek
3	Çanta
4	Toz Deterjan
5	Bebek Bezi
6	Saat
7	Gözlük
8	Parfüm
9	Peynir
10	Çikolata

Tablo 6. Naive Bayes ve Lojistik Regresyon Sınıflandırma Yöntemlerinin Sonuçları (Alan Sıralaması 1)

Model	Naive Bayes		Lojistik Regresyon	
Kelime Vektörü	CountVectorizer	TF-IDF	CountVectorizer	TF-IDF
Maksimum Öznitelik Sayısı	-	-	-	-
1000 Öznitelik	0,510	0,500	0,480	0,500
500 Öznitelik	0,580	0,580	0,525	0,495
Maksimum Öznitelik Sayısı + Ngram(1,3)	-	-	-	-
1000 Öznitelik + Ngram (1,3)	0,475	0,495	0,435	0,460
500 Öznitelik + Ngram (1,3)	0,485	0,590	0,490	0,530
500 Öznitelik + Kök Ayırma	0,485	0,590	0,490	0,530

Tablo 6`da Alan Sıralaması 1 gözetilerek sisteme verilen veri kümelerinin, naive Bayes ve lojistik regresyon sınıflandırma yöntemleri ile elde edilen model performans sonuçları incelendiğinde en yüksek doğruluk değerinin % 59 olduğu görülmektedir. Bu doğruluk değerine; yaşam boyu naive Bayes sınıflandırma yöntemi ile ulaşılmıştır. Bu analizde metin, köklerine ayrılmış ve TF-IDF kelime vektöründe 500 öznitelik kullanılmıştır.

Tablo 6`da görüldüğü gibi “Maksimum Öznitelik Sayısı” ve “Maksimum Öznitelik Sayısı + Ngram (1,3)” sonuç değeri, sınıflandırma modelinin hata vermesi nedeniyle elde edilememiştir. Çünkü sınıflandırma modeli eğitilirken öznitelik sayısı sürekli artmakta ancak test verisi vektöründe bulunan öznitelik sayısı sınırlı kalmaktadır. Başka bir deyişle, eğitim ve test vektörlerinin boyutları uyuşmamakta ve model, bir süre sonra hata vermektedir.

Tablo 7. Toplama Dayalı Yaşam Boyu Duygu Sınıflandırması Model Doğruluğu (Alan Sıralaması 1)

Toplama Dayalı Yaşam Boyu Duygu Sınıflandırması		
Kelime Vektörü Boyutu	CountVectorizer	TF-IDF
Maksimum Öznitelik Sayısı	0,945	0,945
1000 Öznitelik	0,935	0,925
500 Öznitelik	0,935	0,930
Maksimum Öznitelik Sayısı + Kök Ayırma	0,930	0,925
1000 Öznitelik + Kök Ayırma	0,935	0,930
500 Öznitelik + Kök Ayırma	0,945	0,930

Tablo 7`de toplama dayalı model doğrulukları verilmiştir. Alan sıralaması 1 için en yüksek model doğruluğu, % 94,5`tir. Bu doğruluk değerine CountVectorizer ve TF-IDF kelime vektörü maksimum öznitelik sayısında ulaşılmıştır. Ayrıca metin köklerine ayrılmış durumda, CountVectorizer kelime vektörü, 500 öznitelik sayısında model doğruluğu, % 94,5`tir.

Tablo 7`ye bakıldığında modeli eğitiminde metin vektörleri kullanılmadığı için tüm denemelerde sonuçların elde edildiği görülmektedir. Toplama dayalı yaşam boyu duygu sınıflandırması modelinde, modeli eğitmek için metin vektörleri yerine geçmiş alan ve mevcut alan eğitim kümesinde bulunan tüm kelimelere ait sıklık değerleri kullanılmaktadır. Bu sıklık değerleri, bilgi tabanında saklanmaktadır.

Tablo 8. Toplama Dayalı Model Bilgi Tabanı Örneği

Kelime	Pozitif Sıklık	Negatif Sıklık	Toplam Sıklık	Pozitif Koşullu Olasılık	Negatif Koşullu Olasılık
güzel	8735	1708	10443	0,06159	0,01024
aldım	6864	2009	8873	0,04840	0,01204
ürün	3875	6251	10126	0,02732	0,03746
ederim	3288	117	3405	0,02319	0,00070
iyi	2897	755	3652	0,02043	0,00453
beğendim	2694	79	2773	0,01900	0,00047
gayet	2556	172	2728	0,01802	0,00103
geldi	2554	4500	7054	0,01801	0,02697
tam	2339	362	2701	0,01649	0,00217
teşekkür	2165	65	2230	0,01527	0,00039
tavsiye	1934	1295	3229	0,01364	0,00776
teşekkürler	1799	121	1920	0,01269	0,00073
büyük	1783	993	2776	0,01257	0,00595
almıştım	1742	649	2391	0,01228	0,00389
alın	1692	292	1984	0,01193	0,00175
harika	1573	45	1618	0,01109	0,00027
kaliteli	1510	435	1945	0,01065	0,00261
hızlı	1220	186	1406	0,00860	0,00112
indirimdeyken	1202	38	1240	0,00848	0,00023
memnun	1182	763	1945	0,00834	0,00457
hediye	1153	351	1504	0,00813	0,00210
kullanışlı	1123	44	1167	0,00792	0,00026
paketlenme	1063	452	1515	0,00750	0,00271
kullanıyorum	1053	156	1209	0,00743	0,00094
indirimde	1047	79	1126	0,00738	0,00047
stok	1023	26	1049	0,00721	0,00016
indirimden	968	102	1070	0,00683	0,00061
küçük	951	1243	2194	0,00671	0,00745

Tablo 8’de bilgi tabanında bulunan örnek bir kelime kümesi ve bu kelimelere ait saklanan bilgiler verilmiştir. Bilgi tabanında her bir kelime için pozitif ve negatif kelime sıklıkları, pozitif ve negatif koşullu olasılıklar saklanır. Bu sayede modele daha fazla girdi ile daha fazla bilgi öğretilir. Bilgi tabanında bulunan bu bilgiler kullanılarak test verisini oluşturan her bir yorum için pozitif cümle olasılığı ve negatif cümle olasılığı hesaplanmaktadır. Yorumda bulunan her bir kelime için önce koşullu olasılık değerleri

hesaplanır. Eğer test verisinde bulunan bir kelime bilgi tabanında yoksa, bu kelimenin koşullu olasılığına çok küçük bir değer atanır. Bu sayede kelimenin geçtiği cümlelerin olasılığının sıfır olması engellenir.

Yaşam boyu öğrenme kapsamında ele alınan üç analizden en iyi model sınıflandırma doğruluğu, toplama dayalı yaşam boyu duygu sınıflandırması modeli kullanıldığında elde edilmiştir. Bunun nedeni; diğer iki modelin aksine bu modelde, bilgi kaybı yaşanmamasıdır. Yaşam boyu naive Bayes ve yaşam boyu lojistik regresyon sınıflandırma yöntemleri ile yapılan analizlerde, model eğitiminde kullanılan metin vektörleri en sık geçen kelimelerle sınırlandırılmaktadır. Örneğin; kelime vektörü 1000 ya da 500 öznitelikten oluşmaktadır. Kelime vektörü öznitelik sayısı artırılabilir fakat maksimum öznitelik sayısına ulaşıldığında model hata vermektedir. Bu durumda öznitelik sayısının artırılması bir yere kadar fayda sağlayacaktır. Toplama dayalı yaşam boyu duygu sınıflandırması modelinde ise metin verisinde bulunan her kelime saklanmaktadır. Bu sayede bilgi kaybı söz konusu değildir. Bir alanda düşük sıklıkla geçen bir kelimenin yaşam boyu öğrenme süreci boyunca farklı alanlarda görülme sıklığı artacaktır. Bu da model performansının iyileşmesine katkı sağlamaktadır.

3.7. Model Performansının Değerlendirilmesi

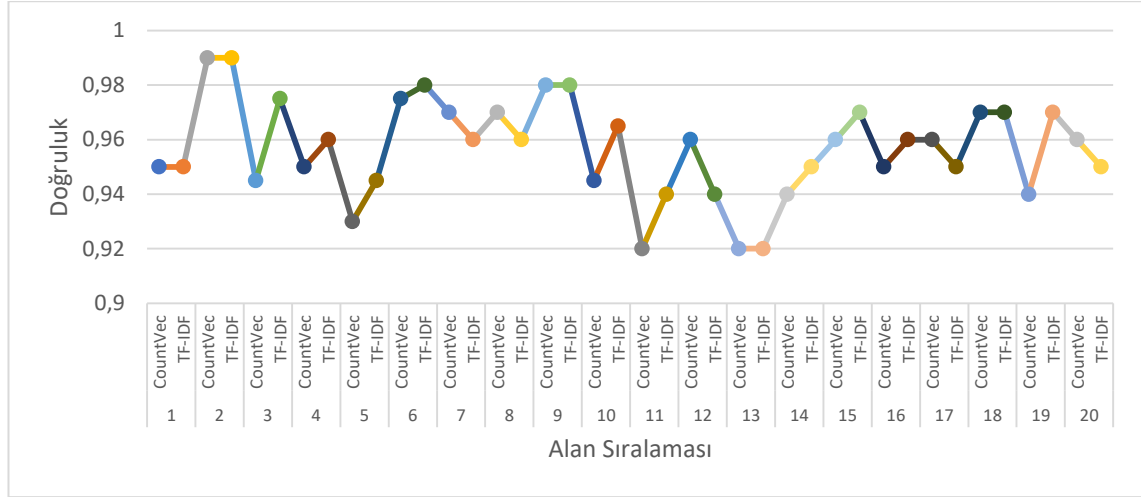
Yaşam boyu duygu sınıflandırması model performansını etkileyen etkenler aşağıda detaylı bir biçimde incelenmiştir.

3.7.1. Alan Sıralamasının Model Performansına Etkisi

Yaşam boyu öğrenme yaklaşımında model performansına etkileyen önemli etkenlerden biri, alanların modele öğretilme sırasıdır. Farklı alan sıralamaları için yaşam boyu öğrenme duygu sınıflandırması modelinin performansında değişimler görülebilmektedir. Alan sıralamasının etkisini analiz etmek için birkaç rastgele alan sıralaması denenebilir ve sonuçlar karşılaştırılabilir. Literatürdeki çalışmalar (örneğin; Chen vd., 2015; Ha vd., 2018; Xia vd., 2017; Xu vd., 2020) çoğunlukla yalnızca bir rastgele sıralama kullanmaktadır (Chen ve Liu, 2018: 17). Alan sıralamasının model performansına etkisini dikkate alan bir çalışma Wang vd. (2019) tarafından yapılmıştır. Çalışmada 20 farklı alana ait farklı 10 alan sıralaması kullanılarak model performansları incelenmiştir.

Bu tez çalışması kapsamında farklı alan sıralamaları kullanılarak alan sıralamalarının model performansına etkisi incelenmiştir. Bu çalışmada toplam 10 farklı

alan kullanılmıştır. Bu alanlar için rastgele farklı 20 alan sıralaması oluşturulmuş ve bu sıralamalar, Ek 6'da verilmiştir. Bu sıralamalar altında toplama dayalı duygu sınıflandırması model doğruluğu incelenmiş ve elde edilen model doğruluk değerleri, Ek 7'de verilmiştir.

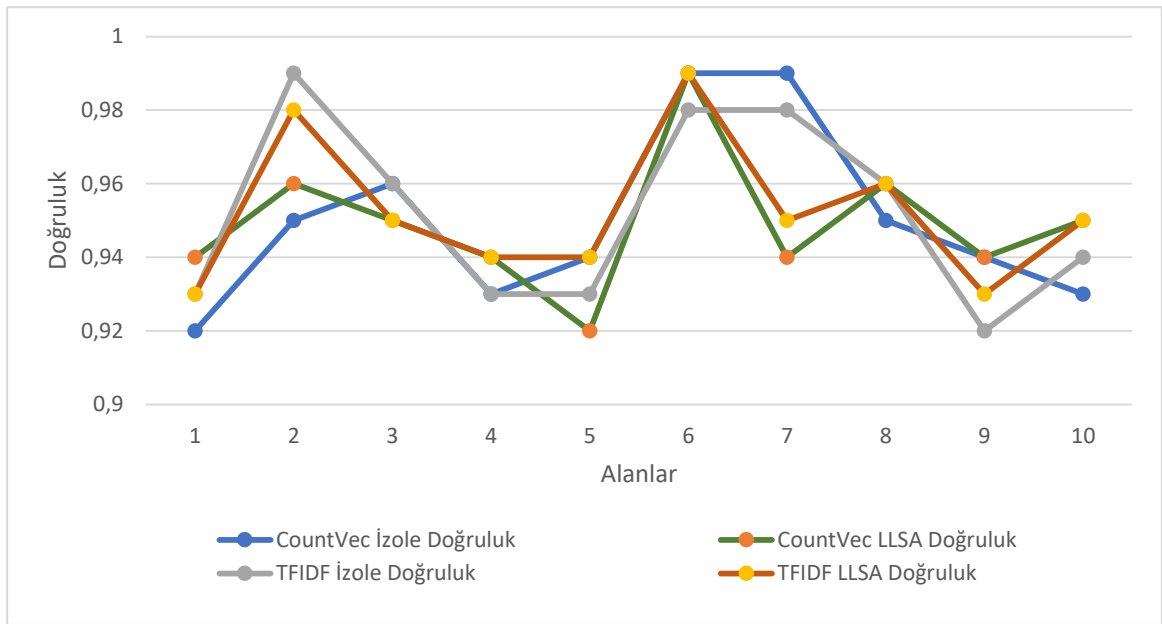


Şekil 10. Alan Sıralamasının Model Performansına Etkisi

Şekil 10'da sisteme veri kümelerinin farklı alan sıralaması ile verildiğinde model performansındaki değişim gösterilmiştir. Model performansı, maksimum öznitelik sayısında, CountVectorizer ve TF-IDF kelime vektörü yöntemleri kullanılarak hesaplanmıştır. Şekil 10 incelendiğinde maksimum öznitelik sayısında farklı kelime vektörü yöntemlerinin ve alan sıralamasının model performansında değişime neden olduğu görülmektedir. Başka bir deyişle, alanların sisteme veriliş sırası, toplama dayalı modelin öğrenme davranışını etkilemektedir. Örneğin; Alan Sıralaması 1 (Tablo 5) ele alındığında; sisteme son olarak verilen alan *çikolata*dır. *Çikolatadan* önce gelen alan ise *peynir*dir. *Peynir* alanında geçen koku, tat, lezzet gibi sözcükler *çikolata* alanında da bulunmaktadır. Geçmiş alandan öğrenilen bu kelimeler, bilgi tabanında saklandığı için *çikolata* alanı sisteme geldiğinde model, bu geçmiş alan bilgilerini kullanarak test aşamasında bu sözcükleri daha kolay tahmin edecektir. *Çikolata* alanında bulunan antepfıstığı, frambuaz, bitter gibi bu alana özgü kelimeler geçmiş alanlarda bulunmayan kelimelerdir. Bu kelimeler, yeni öğrenilecek ve test aşamasında bu sözcükleri tahmin etmek daha zor olacaktır. Benzer alanların art arda sisteme verilmesi, ortak kelime bulundurma olasılığını arttırmakta ve model performansını olumlu yönde etkilemektedir.

Diğer yandan iyi, güzel, kötü, ürün, kargo gibi her alana ait yorumlarda sıklıkla karşılaşılan kelimeler modelin tahmin performansını arttırmaktadır.

Ek 8`de farklı alan sıralamalarının model doğruluğuna etkisinin yanı sıra her bir alan için izole model doğrulukları da verilmiştir. Alanların her biri için birbirinden bağımsız bir biçimde hesaplanan izole model performanslarına bakıldığında model doğruluklarının % 90`ın üzerinde olduğu görülmektedir. Şekil 11`de Alan Sıralaması 1`e ait izole model ve toplama dayalı model doğrulukları karşılaştırmalı olarak verilmiştir.



Şekil 11. İzole Modelin ve Toplama Dayalı Modelin Doğruluğu (Alan Sıralaması 1)

Öğrenme süreci boyunca izole modelin performansı, kimi zaman toplama dayalı modelin performansının önüne geçmiştir. Bunun nedeni, yaşam boyu öğrenme modelindeki bilgi eksikliğidir. Daha açık bir ifade ile model, daha önce hiç görmediği yeni bir alan ile karşılaşmıştır ve bilgi tabanındaki bilgi, bu yeni alanı yeterince iyi (izole model kadar) tahmin edecek düzeyde değildir. Fakat model, yaşam boyu öğrenme süreci boyunca yeni bilgiler öğrenmeye devam eder. Öğrenme süreci tamamlandığında toplama dayalı model performansının, sisteme son gelen alanın izole model performansından yüksek olduğu görülmektedir. Bu durum, yaşam boyu öğrenme sürecinin başarılı bir şekilde tamamlandığını göstermekte olup yaşam boyu öğrenmenin temel amacı olan daha iyi öğrenen modellerin geliştirildiğini göstermektedir.

Diğer yandan Alan Sıralaması 3, Alan Sıralaması 11 ve Alan Sıralaması 14 ile elde edilen toplama dayalı modelin performansı, ilgili alan sıralamasında en sonda

bulunan alanın izole modellerin performanslarının gerisinde kalmıştır. Bu durumun nedeninin, bu alan sıralamalarında sonuncu sırada yer alan alanların kendilerine özgü kelimelerinin bulunması olduğu düşünülmektedir. Örneğin; Alan Sıralaması 11 için sisteme verilen son alan *peynir*dir. Bu alana ait veri kümesinde küflü, ekşi, tuzlu gibi alana özgü kelimeler bulunmaktadır. Bu ve benzeri kelimeler, geçmiş alanlarda bulunmadığı için yaşam boyu öğrenme model doğruluğu izole model doğruluğunun gerisinde kalmıştır.

3.7.2. Doğrulama Yöntemlerinin Model Performansına Etkisi

Makine öğrenmesi modellerinin performanslarının değerlendirilmesinde doğrulama (validasyon) yöntemleri kullanılmaktadır (Dwivedi, 2018: 4; Straub, 2021: 5). Doğrulama yöntemleri, model eğitimi boyunca aşırı öğrenmeyi ortadan kaldırarak modelin gerçek dünya verilerine uyum sağlamasını kolaylaştırmaktadır.

Doğrulama yöntemleri; yalnızca bir öğrenme algoritmasının performans özelliklerini anlama açısından değil, aynı zamanda model seçimi yoluyla algoritma performansını optimize etme açısından da önemlidir. Etiketli bir veri kümesinin tamamının kullanımı, modelin doğruluğunu büyük ölçüde arttırdığı için model eğitiminde kullanılmamaktadır. Çünkü sınıflandırmanın temel amaçlarından biri, etiketli veriler ile eğitilen bir modelin daha önce görülmemiş test verisi örnekleri ile sınanmasıdır. Bu nedenle bir veri kümesi, her zaman eğitim verisi, doğrulama verisi ve test verisi olmak üzere üç parçaya bölünmelidir (Aggarwal, 2018: 221-222).

Veri kümesi öncelikle eğitim verisi ve test verisi olarak iki parçaya ayrılır. Test verisi, model eğitimi boyunca saklanır. Eğitim veri kümesi, model eğitiminde kullanılmak amacıyla tekrar eğitim ve doğrulama veri kümesi olarak ikiye bölünür.

Bu tez çalışmasında model doğrulama yöntemlerinden çapraz doğrulama (cross validation) ve bootstrap (yeniden örnekleme) yeniden örnekleme yöntemleri kullanılmıştır. Yeniden örnekleme yöntemleri seçilirken karşılaştırılabilir olması amacıyla, ortak noktaları olan bu iki yöntem seçilmiştir. Her iki yeniden örnekleme yöntemi de model uyumu ve model değerlendirmesi için mevcut tüm veri kümesini kullanır. İki yöntemde hesaplama maliyeti yüksektir. İkisi de veri kümesinin çok büyük olmadığı veri kümelerine uygulanabilmektedir (Hawkins ve Kraker, 2010: 6). Ayrıca bu iki yöntem model doğrulama sonuçlarının karşılaştırmasında (örneğin; Borra ve Di Ciaccio, 2010; Lasfar ve Tóth, 2024) sıklıkla kullanılmaktadır.

Çapraz doğrulama yöntemi, makine öğrenmesi modellerinin değerlendirilmesinde kullanılan bir yeniden örnekleme yöntemidir. Çapraz doğrulama yönteminin amacı, model performansının tarafsız bir tahminini sağlamaktır. Çapraz doğrulama yöntemlerinden biri *k katlı çapraz doğrulama* (k-fold cross validation). Bu yöntemde veri noktaları, k tane eşit büyüklükteki parçaya ayrılır. Yinelemeli bir süreçte, k . parça doğrulama kümesi olarak seçildiğinde, kalan $k-1$ grup birleştirilir ve eğitim kümesi olarak kullanılır. Bu süreç, k kez yinelenir (Hastie vd., 2009: 241-242). Böylece her grup, doğrulama kümesi olarak bir kez seçilir. k yinelemesi boyunca performans ölçümlerinin ortalaması, doğrulama hatasının tahmini olarak kullanılır. Model değerlendirmesi k kez gerçekleştirildiği için performans ölçümünün varyansı azalmakta ve ortaya çıkan tahminin daha güvenilir olduğu düşünülmektedir (Maleki vd., 2020: 437).

Bootstrap yöntemi, bir veri kümesinin örnekleme dağılımını tahmin etmek için kullanılan bir yeniden örnekleme yöntemidir (Raschka, 2020: 16). Bootstrap yöntemi, standart hatayı tahmin etmek için bilgisayar tabanlı bir yöntem olarak 1979'da tanıtılmış olup önyüklemeli örneklemeyle bağlıdır (Efron ve Tibshirani, 1993: 45). Bootstrap yöntemi, olağan yaklaşımların geçersiz olduğu durumlarda doğru çözümler üretebilmektedir (Wehrens vd., 2000: 1). Bootstrap yönteminde bir veri kümesinden rastgele örneklemler seçilerek, ilgili veri kümesinin bir versiyonu ya da yeniden örnekleme oluşturulur. Yeni oluşturulan örneklem kullanılarak örneklem ortalaması ve standart hata gibi örneklem istatistikleri hesaplanarak veri kümesine ilişkin çıkarımlar yapılır (Efron ve Tibshirani, 1993: 46, Efron ve Tibshirani, 1995: 10-11).

Bu çalışmada sisteme belirli bir sırayla verilen her bir alana ait veri kümesi, % 80 eğitim verisi ve % 20 test verisi olmak üzere ikiye ayrılmıştır. Yaşam boyu naive Bayes ve yaşam boyu lojistik regresyon sınıflandırma yöntemleri kullanılarak yapılan uygulamada eğitim verisi, k katlı çapraz doğrulama ve bootstrap yöntemleri kullanılarak model eğitilmiştir. K katlı çapraz doğrulama yönteminde, katman sayısı $k=10$ olarak alınmıştır (Hastie vd., 2009: 260). Eğitim kümesi 10 parçaya bölünerek her bir 10 katmanda bir doğrulama veri kümesi, eğitilen modelin performansının ölçülmesi için ayrılmıştır. Bootstrap yönteminde ise, 50 iterasyon boyunca eğitim veri kümesinden rastgele örneklemler alınmıştır (Efron ve Tibshirani, 1995: 6). Örneklem seçiminde, yerine koyma yöntemi kullanılmıştır. Model performansına ait sonuçlar, Tablo 9'da verilmiştir.

Tablo 9. Naive Bayes ve Lojistik Regresyon Modellerinin Doğrulama Sonuçları

Yöntem	ÇD	Bootsrap	Vektör	Validasyon Doğruluğu		Doğruluk
				ÇD (k=10)	Bootstrap (50)	
NB	-	-	TF-IDF:500	-	-	0,5800
NB	+	-	TF-IDF:500	0,9411	-	0,5850
NB	-	+	TF-IDF:500	-	0,9386 ± 0,0098	0,5750
NB	-	-	TF-IDF:1000	-	-	0,5000
NB	+	-	TF-IDF:1000	0,9527	-	0,4850
NB	-	+	TF-IDF:1000	-	0,9452 ± 0,0091	0,4950
NB	-	-	CountVec.:500	-	-	0,5800
NB	+	-	CountVec.:500	0,9395	-	0,5700
NB	-	+	CountVec.:500	-	0,9382 ± 0,0086	0,5700
NB	-	-	CountVec.:1000	-	-	0,5100
NB	+	-	CountVec.:1000	0,9509	-	0,4950
NB	-	+	CountVec.:1000	-	0,9440 ± 0,0116	0,5100
LR	-	-	TF-IDF:500	-	-	0,495
LR	+	-	TF-IDF:500	0,5400	-	0,4950
LR	-	+	TF-IDF:500	-	0,9044 ± 0,0149	0,5450
LR	-	-	TF-IDF:1000	-	-	0,5000
LR	+	-	TF-IDF:1000	0,5100	-	0,5000
LR	-	+	TF-IDF:1000	-	0,9086 ± 0,0114	0,5050
LR	-	-	CountVec.:500	-	-	0,5250
LR	+	-	CountVec.:500	0,5700	-	0,5250
LR	-	+	CountVec.:500	-	0,9065 ± 0,0141	0,5350
LR	-	-	CountVec.:1000	-	-	0,4800
LR	+	-	CountVec.:1000	0,5000	-	0,4800
LR	-	+	CountVec.:1000	-	0,9077 ± 0,0133	0,4850

NB: Naive Bayes, LR: Lojistik Regresyon, ÇD: Çapraz Doğrulama, CountVec.: CountVectorizer, TF-IDF:Ters Doküman Frekansı Ağırlıklandırması

Tablo 9`da yaşam boyu naive Bayes ve yaşam boyu lojistik regresyon yöntemlerine ait model doğrulama sonuçları incelendiğinde TF-IDF kelime vektörü 500 öznitelik ve 10 katlı çapraz doğrulama yöntemi kullanılarak naive Bayes sınıflandırma yönteminde en yüksek model doğruluğunun % 58,5 olduğu görülmektedir. Lojistik regresyon sınıflandırma yönteminde ise en yüksek model doğruluğu, TF-IDF kelime vektörü 500 öznitelik ve bootstrap yöntemi kullanılarak % 54,5 olarak elde edilmiştir.

Toplama dayalı yaşam boyu duygu sınıflandırmasında 10 katlı çapraz doğrulama yöntemi kullanılmıştır. Model performansına ait sonuçlar, Tablo 10`da verilmiştir.

Tablo 10. Toplama Dayalı Duygu Sınıflandırması Model Doğrulama Sonuçları

Yöntem	ÇD	Vektör	Validasyon Doğruluğu	Doğruluk
			ÇD (k=10)	
Toplama Dayalı Duygu Sınıflandırması	-	TF-IDF Maksimum Öznitelik Sayısı	-	0,93
	+	TF-IDF Maksimum Öznitelik Sayısı	0,98	0,95
	-	TF-IDF: 500	-	0,94
	+	TF-IDF: 500	0,96	0,945
	-	TF-IDF: 1000	-	0,94
	+	TF-IDF: 1000	0,97	0,94
	-	CountVec. Maksimum Öznitelik Sayısı	-	0,96
	+	CountVec. Maksimum Öznitelik Sayısı	0,98	0,95
	-	CountVec.: 500	-	0,94
	+	CountVec.: 500	0,96	0,94
	-	CountVec.: 1000	-	0,93
	+	CountVec.: 1000	0,98	0,94

ÇD: Çapraz Doğrulama, CountVec.: CountVectorizer, TF-IDF: Ters Doküman Frekansı Ağırlıklandırması

Toplama dayalı duygu sınıflandırması modelinin doğrulama sonuçları incelendiğinde (Tablo 10), CountVectorizer kelime vektörü ve maksimum öznitelik sayısı kullanıldığında en yüksek model doğruluğunun % 96 olduğu görülmektedir.

Tablo 9 ve Tablo 10`da verilen sonuçlar incelendiğinde yaşam boyu naive Bayes ve yaşam boyu lojistik regresyon sınıflandırma yöntemlerinde doğrulama yöntemleri kullanılarak iyileşme sağlandığı görülmektedir. Toplama dayalı modelde çapraz doğrulama yönteminin performansa sağladığı katkı yüksek değildir. Bu durum, toplama dayalı duygu sınıflandırması modelinin güçlü bir model olduğunun göstergesidir.

SONUÇ

İnternet kullanımının ve online alışveriş sitelerinden yapılan alışveriş hacminin artmasıyla birlikte ürünlere ait fikir alışverişi ve yorumlar da online olarak yapılmaya başlanmıştır. Sosyal medya platformlarında paylaşılan ürün tavsiyeleri, internet sitelerinde yapılan ürün karşılaştırmaları ve online alışveriş sitelerinde ürünü alan kişiler tarafından yapılan müşteri yorumları, kişilerin ürünü alma süreçlerinde belirleyici olmaktadır. Müşteri yorumları, ürüne ait en tarafsız ve doğru bilgileri barındırır. Bu açıdan bakıldığında yapılan yorumlar, ürünü alacak yeni bir müşteri için olumlu ya da olumsuz bir algı oluşturacak ve satın alma kararını etkileyecektir.

Duygu analizi; makine öğrenmesi, metin madenciliği ve doğal dil işleme alanlarının kesişiminde bulunan bir yöntemdir. Son yıllarda duygu analizi ve yaşam boyu öğrenme yaklaşımının sıklıkla bir arada kullanıldığı çalışmalara rastlamak mümkündür. Yaşam boyu öğrenme; insana özgü öğrenme biçimini konu almaktadır. İnsana özgü öğrenme; bir insanın hayatı boyunca yaptığı davranışlar sonucu elde ettiği, karşılaştığı bilgileri birikimli olarak öğrenmesi ve tüm bilgileri, yeni olaylar ve durumlar karşısında kullandığı öğrenme biçimidir. Bu özelliğin makinelere aktarılmasıyla, sürekli bir öğrenme ile performansı yüksek ve daha akıllı modeller üretilmesi amaçlanmaktadır. Yaşam boyu öğrenme yaklaşımı duygu analizi ile ele alındığında; tek bir alana ait veri kümesi ile eğitilen duygu analizi modelleri yerine birden farklı alan ile eğitilmiş modeller kullanılarak model performansının artırılması hedeflenmektedir.

Bu çalışmada müşteri ürün yorumlarının duygu sınıflarının tespiti için duygu analizi yöntemi kullanılmıştır. Bu amaçla bir online alışveriş sitesinde herkese açık olarak bulunan müşteri ürün yorumları kullanılmıştır. Ürün yorumlarının toplanıp saklanmasında Python yazılım programı kullanılmıştır. Yorumlar, Microsoft Excel belgesi olarak saklanmıştır. Ürün yorumları, ilk olarak ön işleme sürecinden geçirilmiştir. Sayı ve özel karakter kullanımı, emojiler, linkler veya yabancı sözcükler ve durak kelimeler, Python programlama dili kullanılarak temizlenmiştir. Kelime hataları el ile manuel olarak düzeltilmiştir. Ayrıca her bir ürün yorumu, tek tek el ile cümlelerin belirttiği duygu durumuna göre pozitif veya negatif olarak etiketlenmiştir. Ön işleme aşamasının ardından her bir alan, belirlenen sırayla sisteme verilmiştir. Sisteme gelen mevcut alan, makine öğrenmesi modelinin eğitilmesi için eğitim ve test verisi olarak ikiye ayrılmıştır. Bu aşamada yaşam boyu bilgi aktarımı için sisteme gelen mevcut alanın eğitim veri kümesi, kendinden önceki alanların eğitim kümeleriyle birleştirilmiştir. Birleştirilmiş

eğitim veri kümesi, model eğitiminde kullanılmıştır. Modelin test edilmesinde mevcut alan test kümesi kullanılarak tahmin doğruluğu hesaplanmıştır.

Bu çalışmada kullanılan sınıflandırma yöntemleri; naive Bayes, lojistik regresyon ve toplama dayalı yaşam boyu duygu sınıflandırması modelidir. Karşılaştırmalı sonuçlarda; sisteme verilen alan sıralamasının, kullanılan kelime vektörü yönteminin, kelime vektörü öznitelik sayısının ve metinlerin köklerine ayrılmasının yöntemlerin performanslarına etkileri analiz edilmiştir. Sınıflandırma sonuçlarına bakıldığında model, farklı alan sıralamaları ile eğitildiğinde performansta değişiklik gözlenmektedir. Bu durum, model performansının doğrudan alan sıralamasından etkilendiğini göstermektedir. Bunun nedeni, alanlar arasındaki benzer kelimelerin varlığıdır. Birbirine benzer kelimeler bulunduran alanlar ardı ardına geldiklerinde toplam kelime sıklıkları artmaktadır. Bu durum, model tahmin doğruluğunu artırmaktadır. Benzer kelime bulundurmeyen alanlar ardı ardına geldiğinde ise sistemde farklı kelimeler ve bu kelimelere ait düşük sıklıklar olduğu için model tahmin performansında düşüş yaşanmaktadır. Bu durum, toplama dayalı yaşam boyu duygu sınıflandırmasında gözlemlenebilmektedir. Her alanda bulunan tüm kelimelere ait sıklıklar kullanılarak eğitilen model, Alan Sıralaması 2’de % 99 doğruluğa ulaşmıştır.

Bu tez çalışmasının literatüre katkıları şu şekilde özetlenebilir:

- Literatürde yaşam boyu öğrenmeyi ve duygu sınıflandırmasını konu alan çalışmalar genellikle teorik çalışmalardır. Bu tez çalışmasında ise hem duygu analizini hem de yaşam boyu öğrenmeyi dikkate alarak teorik alt yapı oluşturulmuş ve detaylı bir literatür taraması sunulmuştur. Ayrıca teorik bilgilere ek olarak uygulamalı bir analiz yapılmıştır.
- Günümüzde e-ticaret, ülke ekonomileri için oldukça önemli hale gelmiştir. E-ticaret kapsamında online alışveriş sitelerindeki ürünlere ilişkin kullanıcı yorumlarının sayısı hem kullanıcılar hem de üreticiler ve site sahibi tarafından kullanılabilir olması bakımından oldukça fazladır. Bu çalışmada bir online alışveriş sitesinde kullanıcıların ürün yorumlarının duygu sınıflandırılmasının yapılması uygulama bölümünün temelini oluşturmaktadır. Bu anlamda çalışmanın uygulama bölümünde ele alınan problem, güncel ve e-ticaret için önemli bir problemdir.
- Yaşam boyu duygu sınıflandırması kapsamında yapılan uygulamada görece olarak büyük sayıda farklı alan dikkate alınmıştır. Alan sayısının artması, sınıflandırma modelinin karmaşıklığını arttırmakta ve sınıflandırma modelinin performansını etkilemektedir. Farklı ve çok sayıda alandan oluşan ürün yorumu verisi; yaşam boyu

naive Bayes, yaşam boyu lojistik regresyon ve yaşam boyu toplama dayalı sınıflandırma yöntemleri ile incelenmiş ve yöntemlerin performansları karşılaştırılmıştır.

- Bu çalışma; yaşam boyu öğrenme yaklaşımında, Türkçe veri kümesi ile yapılan ilk çalışmadır. Türkçe ürün yorumları kullanılarak farklı alanlara ait veri kümeleri arasında bilgi aktarımı sağlanmıştır.
- Çalışmada sisteme verilen alan sırasının değişiminin ve doğrulama yöntemlerinin yaşam boyu duygu sınıflandırması modelinin performansına etkisi detaylı bir şekilde analiz edilmiştir.
- Uygulamadan elde edilen sonuçlar, gerçek hayatta kullanılabilir niteliktedir. Bu çıktılar aracılığıyla online alışveriş sitesinde kullanıcıların almış olduğu ürünler geçmiş alan olarak belirlenebilmekte ve her bir kullanıcı için bu geçmiş alanlara ait yorumlar kullanılarak toplama dayalı model eğitilebilmektedir. Bu sayede alışveriş sitesi tarafından kullanıcıya özel ürün değerlendirmesi yapılabilmekte ve yeni alınacak ürünün yorumlarının kullanıcının okumasına gerek kalmadan analiz edilmesi sağlanabilmekte ve ürün hakkındaki genel duygu yönelimi belirlenebilmektedir.
- Bu çalışma, yaşam boyu duygu sınıflandırması analizi için etkin modeller belirlemeye yönelik önemli bilgiler sunmakta ve e-ticaret işletmelerinin müşteri geri bildirimlerine dayanan veri odaklı kararlar almasına olanak sağlamaktadır.

Diğer yandan yapılan bu tez çalışmasının bazı kısıtları bulunmaktadır. Çalışmanın en önemli kısıtı, Türkçe için ön işlenmiş veri kümesi azlığıdır. Sonraki çalışmalarda, farklı alanlardan ön işlenmiş veri kümeleri kullanılarak alan sayısının model performansına etkisi incelenebilir. Ayrıca farklı bir alan olarak film yorumları, kitap yorumları, blog yorumları vb. kullanılabilir. Farklı online alışveriş sitelerinden ürün yorumları toplanarak denemeler yapılabilir. Kullanılan toplama dayalı modelin performansı N-gramlar kullanılarak değerlendirilebilir. Ayrıca Türkçe veri kümesi kullanılarak yaşam boyu öğrenme yaklaşımı kapsamında; yapay sinir ağları ve derin öğrenme gibi son yıllarda popüler olan yöntemler ile çalışmalar yapılabilir ve sonuçlar karşılaştırılabilir.

KAYNAKLAR

- Abulaish, M., Wasi, N. A., Sharma, S. (2024). "The Role of Lifelong Machine Learning in Bridging the Gap Between Human and Machine Learning: A Scientometric Analysis", *WIREs Data Mining and Knowledge Discovery*, 14/2, e1526.
- Aggarwal, C. C. (2018). *Machine Learning for Text*, Springer International Publishing, Cham.
- Ahmetoğlu, H., Daş, R. (2020). "Türkçe Otel Yorumlarıyla Eğitilen Kelime Vektörü Modellerinin Duygu Analizi ile İncelenmesi", *Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, 24/2, 455-463.
- Akba, F. (2014). *Duygu Analizinde Öznitelik Seçme Metriklerinin Değerlendirilmesi: Türkçe Film Eleştirileri*, (Basılmamış Yüksek Lisans Tezi), Hacettepe Üniversitesi Fen Bilimleri Enstitüsü, Ankara.
- Akın, B. K., Şimşek, U. T. G. (2018). "Adaptif Öğrenme Sözlüğü Temelli Duygu Analiz Algoritması Önerisi", *Bilişim Teknolojileri Dergisi*, 11/3, 245-253.
- Al-Amrani, Y., Lazaar, M., Elkadiri, K. E. (2017). "Sentiment Analysis Using Supervised Classification Algorithms", *Proceedings of the 2nd International Conference on Big Data, Cloud and Applications*, 1-8.
- Amplayo, R. K., Lee, S., Song, M. (2018). "Incorporating Product Description to Sentiment Topic Models for Improved Aspect-Based Sentiment Analysis", *Information Sciences*, 454, 200-215.
- Anandarajan, M., Hill, C., Nolan, T. (2019). *Text Preprocessing, Practical Text Analytics*, Springer International Publishing, Cham.
- Arif, M. H., Iqbal, M., Li, J. (2019). "Extracting and Reusing Blocks of Knowledge in Learning Classifier Systems for Text Classification: A Lifelong Machine Learning Approach", *Soft Computing*, 23/23, 12673-12682.
- Atan, S., Çınar, Y. (2019). "Borsa İstanbul'da Finansal Haberler ile Piyasa Değeri İlişkisinin Metin Madenciliği ve Duygu (Sentiment) Analizi ile İncelenmesi", *Ankara Üniversitesi SBF Dergisi*, 74/1, 1-34.
- Artsın, M. (2020). "Bir Metin Madenciliği Uygulaması: VOSVIEWER", *Eskişehir Teknik Üniversitesi Bilim ve Teknoloji Dergisi B-Teorik Bilimler*, 8/2, 344-354.
- Ayan, B., Kuyumcu, B., Ceylan, B. (2019). "Twitter Üzerindeki İslamofobik Twitlerin Duygu Analizi ile Tespiti", *Gazi Üniversitesi Fen Bilimleri Dergisi Part C: Tasarım ve Teknoloji*, 7/2, 495-502.
- Aydın, I., Salur, U. M., Başkaya, F. (2018). "Duygu Analizi için Çoklu Popülasyon Parçacık Sürü Optimizasyonu", *Türkiye Bilişim Vakfı Bilgisayar Bilimleri ve Mühendisliği Dergisi*, 11/1, 52-64.
- Aydoğan, M., Karcı, A. (2019). "Kelime Temsil Yöntemleri ile Kelime Benzerliklerinin İncelenmesi", *Çukurova Üniversitesi Mühendislik-Mimarlık Fakültesi Dergisi*, 34/2, 181-196.
- Aytekin, Ç., Bayram, M. A. (2021). "Türkçe Metinler için Duygu Analizi Yaklaşımı İle İletişimde Bağlamdan Bağımsız Modellerin Geliştirilmesi Üzerine Bir Araştırma: Karma Veri Uygulaması Önerisi", *Electronic Journal of New Media*, 5/1, 12-25.
- Barry, J. (2017). "Sentiment Analysis of Online Reviews Using Bag-of-Words and LSTM Approaches", *Artificial Intelligence and Cognitive Science*, 272-274.
- Bayrak, A. T., Beğen, E., Öner, S. C., Kaplan, A., Demirdağ, M., Sayan, İ. U. (2024). "ChatGpt ile Veri Üretimi: Duygu Analizi Sınıflandırması Üzerine Bir İnceleme Data Generation with ChatGpt: A Review on Sentiment Analysis Classification", *2024 International Congress on Human-Computer Interaction, Optimization and Robotic Applications (HORA)*, 1-5.

- Birjali, M., Kasri, M., Beni-Hssane, A. (2021). "A Comprehensive Survey on Sentiment Analysis: Approaches, Challenges and Trends", *Knowledge-Based Systems*, 226. 107134.
- Borra, S., Di Ciaccio, A. (2010). "Measuring the Prediction Error. A Comparison of Cross-Validation, Bootstrap and Covariance Penalty Methods", *Computational Statistics and Data Analysis*, 54/12, 2976-2989.
- Chen, Z., ve Liu, B. (2018). *Lifelong Machine Learning*, Springer International Publishing, Cham.
- Chen, Z., Ma, N., Liu, B. (2015). "Lifelong Learning for Sentiment Classification", *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 2: Short Papers)*, 750-756.
- Chowdhary, K. R. (2020). *Fundamentals of Artificial Intelligence*, Springer India, New Delhi.
- Çatak, F. Ö. (2014). "Eşle/İndirge Yöntemi Kullanılarak Destek Vektör Makinesi Algoritması ile Yüksek Boyutlu Sosyal Medya Mesajlarının Kutupsal Değerinin Ölçülmesi", *Yönetim Bilişim Sistemleri Kongresi*, 2014.
- Çelik, H. (2013). *Sentiment Analysis for Turkish Language*, (Basılmamış Yüksek Lisans Tezi), İstanbul Kültür Üniversitesi Fen Bilimleri Enstitüsü, İstanbul.
- Çetin, M., Amasyalı, M. F. (2013). "Supervised and Traditional Term Weighting Methods for Sentiment Analysis", In 2013 21st Signal Processing and Communications Applications Conference (SIU). *IEEE*, 1- 4.
- Çoban, Ö., Özzyer, G. T. (2016). "Türkçe Twitter Mesajları için LDA ile Duygu Sınıflandırması", *24th Signal Processing and Communication Application Conference (SIU)*, 129-132.
- D'Andrea, A., Ferri, F., Grifoni, P., Guzzo, T. (2015). "Approaches, Tools and Applications for Sentiment Analysis Implementation", *International Journal of Computer Applications*, 125/3, 26-33.
- Dalianis, H. (2018). *Clinical Text Mining: Secondary Use of Electronic Patient Records*. Springer.
- Değer, N. S. (2017). *Sosyal Medya Mesajlarında Veri Madenciliği ile Duygu Analizi*, (Basılmamış Doktora Tezi), İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, İstanbul.
- Devi, P. I. A., Saleema, A. A. (2024). *Spell Checking and Correcting Techniques in Nlp*, In book: *Futuristic Trends in Computing Technologies and Data Sciences*, IIP Series, 154-161.
- Dumais, S. T. (1991). "Improving the Retrieval of Information from External Sources", *Behavior Research Methods, Instruments and Computers*, 23/2, 229-236.
- Dunn, J. (2020). "Mapping Languages: The Corpus of Global Language Use", *Language Resources and Evaluation*, 54/4, 999-1018.
- Dwivedi, A. K. (2018). "Performance Evaluation of Different Machine Learning Techniques for Prediction of Heart Disease", *Neural Computing and Applications*, 29/10, 685-693.
- Efron, B., Tibshirani, R. (1993). *An Introduction to the Bootstrap*, Chapman and Hall/CRC, Florida.
- Efron, B., Tibshirani, R. (1995). "Cross-Validation and the Bootstrap: Estimating the Error Rate of a Prediction Rule", *USA: Division of Biostatistics*, 92, 548-560.
- Ekinci, E., Omurca, S. İ. (2017). "Ürün Özelliklerinin Konu Modelleme Yöntemi ile Çıkartılması", *Türkiye Bilişim Vakfı Bilgisayar Bilimleri ve Mühendisliği Dergisi*, 9/1, 51-58.

- Ekinci, M., Özcan, M. S., Amasyalı, M. F. (2016). "Duygu Analizinde Zaman Etkisi", In 2016 24th Signal Processing and Communication Application Conference, 209-212.
- Eliaçık, A. B., Erdoğan, N. (2015). "Mikro Bloglardaki Finans Toplulukları için Kullanıcı Ağırlıklandırılmış Duygu Analizi Yöntemi", Proceedings of the 9th Turkish National Software Engineering Symposium, İzmir.
- Emekli, B., Selvi, İ. H. (2020). "Gsm Operatörlerine Yönelik Atılan Türkçe Tweetlerin Derin Öğrenme Yöntemleriyle Duygu Analizi", 4. Uluslararası Marmara Fen Bilimleri Kongresi, İstanbul.
- Eroğul, U. (2009). *Sentiment Analysis in Turkish*, (Basılmamış Yüksek Lisans Tezi), Orta Doğu Teknik Üniversitesi Fen Bilimleri Enstitüsü, Ankara.
- Feldman, R., Sanger, J. (2007). *The Text Mining Handbook Advanced Approaches in Analyzing Unstructured Data*, Cambridge University Press, New York.
- Garson, G. D. (2014). *Logistic Regression: Binary and Multinomial*, Statistical Publishing Associates, New York.
- Geng, B., Yang, M., Yuan, F., Wang, S., Ao, X., Xu, R. (2021). "Iterative Network Pruning with Uncertainty Regularization for Lifelong Sentiment Classification", Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval, 1229-1238.
- Görmez, Y., Arslan, H., Atak, B. (2024). "Türkçe Metinlerde Duygu Analizi: Derin Öğrenme Yaklaşımlarının ve Ön İşlem Süreçlerinin Model Performansına Etkisi", *Fırat Üniversitesi Mühendislik Bilimleri Dergisi*, 36/1, 509-520.
- Gözükara, F., Özel, S. A. (2016). "An Experimental Investigation of Document Vector Computation Methods for Sentiment Analysis of Turkish and English Reviews", *Çukurova Üniversitesi Mühendislik Mimarlık Fakültesi Dergisi*, 31/2, 467-481.
- Gupta, S., Gupta, S. K. (2019). "Natural Language Processing in Mining Unstructured Data from Software Repositories: A Review", *Indian Academy of Sciences*, 44/12.
- Güran, A., Uysal, M., Doğrusöz, Ö. (2014). "Destek Vektör Makineleri Parametre Optimizasyonunun Duygu Analizi Üzerindeki Etkisi", *DEÜ Mühendislik Fakültesi Mühendislik Bilimleri Dergisi*, 16/48, 86-93.
- Ha, Q.-T., Pham, T.-N., Nguyen, V.-Q., Nguyen, T.-C., Vuong, T.-H., Tran, M.-T., Nguyen, T.-T. (2018). "A New Lifelong Topic Modeling Method and Its Application to Vietnamese Text Multi-label Classification". In Asian Conference on Intelligent Information and Database Systems, 200-210, Springer.
- Ha, Q.-V., Nguyen-Hoang, B.-D., Nghiem, M.-Q. (2016). "Lifelong Learning for Cross-Domain Vietnamese Sentiment Classification", *Computational Social Networks*, 9795, 298-308.
- Han, J., Pei, J., Tong, H. (2022). *Data Mining: Concepts and Techniques*, Morgan Kaufmann, San Francisco.
- Hasanli, H., Rustamov, S. (2019). "Sentiment Analysis of Azerbaijani Twits Using Logistic Regression, Naive Bayes and SVM", 2019 IEEE 13th International Conference on Application of Information and Communication Technologies (AICT), 1-7.
- Hastie, T., Tibshirani, R., Friedman, J. (2009). *The Elements of Statistical Learning*, Springer New York, New York.
- Hawkins, D. M., Kraker, J. (2010). "Deterministic Fallacies and Model Validation", *Journal of Chemometrics*, 24, 188-193.
- Hayran, A., Sert, M. (2017). "Sentiment Analysis on Microblog Data Based on Word Embedding and Fusion Techniques", In 2017 25th Signal Processing and Communications Applications Conference, 1-4.

- Hirschberg, J., Manning, C. D. (2015). "Advances in Natural Language Processing", *Science* 349/6245, 261-266.
- Hong, X., Guan, S.-U., Man, K. L., Wong, P. W. H. (2020). "Lifelong Machine Learning Architecture for Classification", *Symmetry*, 12/5, 852.
- Hong, X., Wong, P., Liu, D., Guan, S.-U., Man, K. L., Huang, X. (2018). "Lifelong Machine Learning: Outlook and Direction", *Proceedings of the 2nd International Conference on Big Data Research*, 76-79.
- Hosmer, D. W., Lemeshow, S., Sturdivant, R. X. (2013). *Applied Logistic Regression*, Wiley, New Jersey ve Canada.
- Irfan, M., Venkatesam, A. S., Shanmugasundaram, H., Gourisaria, M. K., Chowdhury, S., Patra, S. S. (2025). "Sentiment Analysis of Movie Reviews: Optimizing Predictions Through TF-IDF Vectorization", *2025 3rd International Conference on Intelligent Data Communication Technologies and Internet of Things (IDCIoT)*, 102-107.
- Jo, T. (2019). *Text Mining: Concepts, Implementation, and Big Data Challenge*, Springer International Publishing, Cham.
- Kadhim, A. (2018). "An Evaluation of Preprocessing Techniques for Text Classification", *International Journal of Computer Science and Information Security*, 16, 22-32.
- Ke, Z., Liu, B., Xu, H., Shu, L. (2021). "CLASSIC: Continual and Contrastive Learning of Aspect Sentiment Classification Tasks", *arXiv preprint arXiv:2112.02714*.
- Khan, D. E. (2019). "Lifelong Machine Learning with Logic, Semantics and Natural Language Processing", *Proceedings on the International Conference on Artificial Intelligence*, 41-47, Atina.
- Khan, M. T., Durrani, M., Khalid, S., Aziz, F. (2016a). "Lifelong Aspect Extraction from Big Data: Knowledge Engineering", *Complex Adaptive Systems Modeling*, 4/1, 1-15.
- Khan, M. T., Yar, S., Khalid, S. (2016). "Histogram Based Rule Verification in Lifelong Learning Models", *2016 19th International Multi-Topic Conference (INMIC)*, 1-5.
- Kharde, V. A., Sonawane, S. S. (2016). "Sentiment Analysis of Twitter Data: A Survey of Techniques", *International Journal of Computer Applications*, 139/11, 5-15.
- Kılıçer, S., Şamlı, R. (2023). "E-Ticaret Sitelerindeki Türkçe Ürün Yorumları Üzerine Makine Öğrenmesi Algoritmaları ile Duygu Analizi", *Veri Bilimi Dergisi*, 6/2, 15-23.
- Kilimci, Z. H. (2020). "Borsa Tahmini için Derin Topluluk Modelleri (DTM) ile Finansal Duygu Analizi", *Gazi Üniversitesi Mühendislik-Mimarlık Fakültesi Dergisi*, 35/2, 635-650.
- Kim, S.-B., Han, K.-S., Rim, H.-C., Myaeng, S. H. (2006). "Some Effective Techniques for Naive Bayes Text Classification", *IEEE Transactions on Knowledge and Data Engineering*, 18/11, 1457-1466.
- Kirasich, K., Smith, T., Sadler, B. (2018). "Random Forest vs Logistic Regression: Binary Classification for Heterogeneous Datasets", *SMU Data Science Review*, 1/3, 9.
- Kızılkaya, Y. M. (2018). *Duygu Analizi ve Sosyal Medya Alanında Uygulama*, (Basılmamış Doktora Tezi), Uludağ Üniversitesi Sosyal Bilimler Enstitüsü, Bursa.
- Kleinbaum, D. G., Klein, M., Pryor, E. R. (2002). *Logistic Regression: A Self-Learning Text*, Springer, New York.
- Komarek, P. (2004). *Logistic Regression for Data Mining and High-Dimensional Classification*, (Basılmamış Doktora Tezi), Carnegie Mellon Üniversitesi Matematik Bilimi Bölümü, Pensilvanya.
- Kowsari, K., Jafari Meimandi, K., Heidarysafa, M., Mendu, S., Barnes, L., Brown, D. (2019). "Text Classification Algorithms: A Survey", *Information*, 10/4, 150.

- Lampropoulos, A. S., Tsihrintzis, G. A. (2015). *Machine Learning Paradigms: Applications in Recommender Systems*, Springer International Publishing, Switzerland
- Lan, M., Sung, S.-Y., Low, H.-B., Tan, C.-L. (2005). "A Comparative Study on Term Weighting Schemes for Text Categorization", *Proceedings. 2005 IEEE International Joint Conference on Neural Networks*, 1, 546-551.
- Lasfar, R., Tóth, G. (2024). "The Difference of Model Robustness Assessment Using Cross-Validation and Bootstrap Methods", *Journal of Chemometrics*, 38/6.
- Liu, B. (2010). "Sentiment Analysis and Subjectivity", *Handbook of Natural Language Processing*, 2, 627-666.
- Liu, B. (2012). *Sentiment Analysis and Opinion Mining*, Springer Cham, Cham.
- Liu, B. (2017). "Lifelong Machine Learning: A Paradigm for Continuous Learning", *Frontiers of Computer Science*, 11/3, 359-361.
- Liu, B. (2020). "Learning on the Job: Online Lifelong and Continual Learning", *Proceedings of the AAAI Conference on Artificial Intelligence*, 34/9, 13544-13549.
- Maleki, F., Muthukrishnan, N., Ovens, K., Reinhold, C., Forghani, R. (2020). "Machine Learning Algorithm Validation", *Neuroimaging Clinics of North America*, 30/4, 433-445.
- Manning, C. D., Raghavan, P., Schütze, H. (2008). *Introduction to Information Retrieval*, Cambridge University Press, İngiltere.
- Manogaran, G., Lopez, D. (2018). "Health Data Analytics Using Scalable Logistic Regression with Stochastic Gradient Descent", *International Journal of Advanced Intelligence Paradigms*, 10/1-2, 118-132.
- Martin, P. M. (2025). *KADES Uygulaması Hakkında Kullanıcı Yorumları Üzerine Web Madenciliği ve Duygu Analizi*, (Basılmamış Yüksek Lisans Tezi), Bilecik Şeyh Edebali Üniversitesi Lisansüstü Eğitim Enstitüsü, Bilecik.
- Mayda, İ., Uğurlu, Y. (2024). "E-ticaret Sitelerindeki Türkçe Müşteri Yorumlarının Faydalılık Tahmini Predicting the Usefulness of Turkish Consumer Reviews on E-commerce Websites", In *2024 Innovations in Intelligent Systems and Applications Conference*, 1-5.
- Medhat, W., Hassan, A., Korashy, H. (2014). "Sentiment Analysis Algorithms and Applications: A survey", *Ain Shams Engineering Journal*, 5/4, 1093-1113.
- Meral, M., Diri, B. (2014). "Twitter Üzerinde Duygu Analizi", *2014 22nd Signal Processing and Communications Applications Conference*, 690-693.
- McCallum, A., Nigam, K. (1998). "A Comparison of Event Models for Naive Bayes Text Classification", In *AAAI-98 Workshop on Learning for Text Categorization*, 752/1, 41-48.
- Miner, G., Delen, D., Elder, J., Fast, A., Hill, T., Nisbet, B. (2012). *Practical Text Mining and Statistical Analysis for Non-structured Text Data Applications*, Elsevier, Amerika Birleşik Devletleri.
- Montaser, M. A. A., Ghosh, B. P., Barua, A., Karim, F., Das, B. C., Shawon, R. E. R., Chowdhury, M. S. R. (2025). "Sentiment Analysis of Social Media Data: Business Insights and Consumer Behavior Trends in the USA", *Edelweiss Applied Science and Technology*, 9/1, 515-535.
- Nalçakan, Y., Bayramoğlu, Ş. S., Tuna, S. (2015). "Sosyal Medya Verileri Üzerinde Yapay Öğrenme ile Duygu Analizi Çalışması", *Trakya Üniversitesi Bilgisayar Mühendisliği Bölümü*, Edirne, 1-10.
- Nanğır, M. (2013). *Türk Dili için Çoklu Sınıflandırıcı Yöntemler ile Duygu Sınıflandırma*, (Basılmamış Yüksek Lisans Tezi), İstanbul Kültür Üniversitesi Fen Bilimleri Enstitüsü, İstanbul.

- Ng, A., Ma, T. (2023) "CS229 Lecture Notes." *CS229 Lecture Notes Stanford University*.
- Nguyen, E. (2014). "Text Mining and Network Analysis of Digital Libraries in R", in *Data Mining Applications with R*, Elsevier, Boston.
- Nizam, H., Akın, S. S. (2014). "Sosyal Medyada Makine Öğrenmesi ile Duygu Analizinde Dengeli ve Dengesiz Veri Setlerinin Performanslarının Karşılaştırılması", XIX. Türkiye'de İnternet Konferansı, 1/6, 873-883.
- Oflazer, K., Saraçlar, M. (2018). *Turkish Natural Language Processing*, Springer Nature, Cham.
- Oğul, B. B., Ercan, G. (2016). "Sentiment Classification on Turkish Hotel Reviews", 2016 24th Signal Processing and Communication Application Conference, 497-500.
- Oğul, H. A., Güran, A. (2019). "Imbalanced Dataset Problem in Sentiment Analysis", 2019 4th International Conference on Computer Science and Engineering, 313-317.
- Onan, A. (2017). "Türkçe Twitter Mesajlarında Gizli Dirichlet Tahsisine Dayalı Duygu Analizi", Akademik Bilişim Konferansı, 8-10, Aksaray.
- Onan, A. (2020). "Evrişimli Sinir Ağı Mimarilerine Dayalı Türkçe Duygu Analizi", *European Journal of Science and Technology*, 374-380.
- Onan, A. (2021). "COVID-19 ile İlgili Sosyal Medya Gönderilerinin Metin Madenciliği Yöntemlerine Dayalı Olarak Zaman-Mekansal Analizi", *European Journal of Science and Technology*, 26, 138-143.
- Özdeş, M. (2017). *Büyük Veri Araçlarını Kullanarak Duygu Analizi Gerçekleştirimi*, (Basılmamış Yüksek Lisans Tezi), Pamukkale Üniversitesi Fen Bilimleri Enstitüsü, Denizli.
- Özmen, C. G. (2025). *Veri Madenciliği ile İşletme Analitiği: Web Temelli Verilerin Duygu Analizi ile İncelenmesi*, (Basılmamış Doktora Tezi), Adana Alparslan Türkeş Bilim ve Teknoloji Üniversitesi Sosyal Bilimler Enstitüsü, Adana.
- Pak, M. Y. (2015). *Metinlerde Duygu Analizi ve Sınıflandırma için Yeni Yöntemler*, (Basılmamış Yüksek Lisans Tezi), Anadolu Üniversitesi Fen Bilimleri Enstitüsü, Eskişehir.
- Palanisamy, P., Yadav, V., Elchuri, H. (2013). "Serendio: Simple and Practical Lexicon Based Approach to Sentiment Analysis". In Second Joint Conference on Lexical and Computational Semantics, 2: Proceedings of the Seventh International Workshop on Semantic Evaluation, 543-548.
- Parlar, T. (2016). *Feature Selection for Sentiment Analysis in Turkish Texts*, (Basılmamış Doktora Tezi), Çukurova Üniversitesi Fen Bilimleri ve Uygulamalı Bilimler Enstitüsü, Adana.
- Pervan, N. (2019). *Derin Öğrenme Yaklaşımları Kullanarak Türkçe Metinlerden Anlamsal Çıkarım Yapma*, (Basılmamış Yüksek Lisans Tezi), Ankara Üniversitesi Fen Bilimleri Enstitüsü, Ankara.
- Qin, Q., Hu, W., Liu, B. (2020). "Using the Past Knowledge to Improve Sentiment Classification", Findings of the Association for Computational Linguistics: EMNLP 2020, 1124-1133.
- Raghavan, V. V., Wong, S. K. M. (1986). "A Critical Analysis of Vector Space Model for Information Retrieval", *Journal of the American Society for Information Science*, 37/5, 279-287.
- Ramachandran, D., Parvathi, R. (2019). "Analysis of Twitter Specific Preprocessing Technique for Tweets", *Procedia Computer Science*, 165, 245-251.
- Ramos, D. L., Fuentes, F. J. A. (2023). "A Model of Continual and Deep Learning for Aspect Based in Sentiment Analysis", *Journal of Automation, Mobile Robotics and Intelligent Systems*, 17/1, 3-12.

- Raschka, S. (2020). "Model Evaluation, Model Selection, and Algorithm Selection in Machine Learning", *arXiv preprint arXiv:1811.12808*.
- Rodríguez-Ibáñez, M., Casáñez-Ventura, A., Castejón-Mateos, F., Cuenca-Jiménez, P. M. (2023). "A Review on Sentiment Analysis from Social Media Platforms", *Expert Systems with Applications*, 223, 119862.
- Ruder, S. (2017). "An Overview of Gradient Descent Optimization Algorithms", *arXiv preprint arXiv:1609.04747*.
- Sağlam, F. (2019). *Otomatik Duygu Sözlüğü Geliştirilmesi ve Haberlerin Duygu Analizi*, (Basılmamış Doktora Tezi), Hacettepe Üniversitesi Fen Bilimleri Enstitüsü, Ankara.
- Salton, G., Buckley, C. (1988). "Term-weighting Approaches in Automatic Text Retrieval", *Information Processing and Management*, 24/5, 513-523.
- Salton, G., Wong, A., Yang, C. S. (1975). "A Vector Space Model for Automatic Indexing", *Communications of the ACM*, 18/11, 613-620.
- Santur, Y. (2020). "Derin Öğrenme ve Aşağı Örnekleme Yaklaşımları Kullanılarak Duygu Sınıflandırma Performansının İyileştirilmesi", *Fırat Üniversitesi Mühendislik Bilimleri Dergisi*, 32/2, 561-570.
- Shehu, H. A. (2019). *Kutupsallık Sözlüğü ve Yapay Zeka Yardımı ile Türkçe Twitter Verileri Üzerinde Duygu Analizi*, (Basılmamış Yüksek Lisans Tezi), Pamukkale Üniversitesi Fen Bilimleri Enstitüsü, Denizli.
- Shende, K., Waghmare, K. (2019). "Lifelong Learning Sentiment Classification of Continuously Increasing Data", *Journal of The Gujarat Society*, 21/1, 119-121.
- Shu, L., Xu, H., Liu, B. (2017). "Lifelong Learning CRF for Supervised Aspect Extraction", *arXiv preprint arXiv:1705.00251*.
- Shu, X., Ye, Y. (2023). "Knowledge Discovery: Methods from Data Mining and Machine Learning", *Social Science Research*, 110, 102817.
- Silver, D. L., Yang, Q., Li, L. (2013). "Lifelong Machine Learning Systems: Beyond Learning Algorithms", *AAAI Spring Symposium: Lifelong Machine Learning*, 13/5.
- Singh, A. (2019). Lifelong Machine Learning Technique, https://www.researchgate.net/publication/346473629_Lifelong_Machine_Learning_Technique (29.04.2025).
- Sodhani, S., Faramarzi, M., Mehta, S. V., Malviya, P., Abdelsalam, M., Janarthanan, J., Chandar, S. (2022). "An Introduction to Lifelong Supervised Learning", *arXiv preprint arXiv:2207.04354*.
- Straub, J. (2021). "Machine Learning Performance Validation and Training Using a 'Perfect' Expert System", *MethodsX*, 8, 101477.
- Şeker, S. E. (2015). "Metin Madenciliği", *Yönetim Bilişim Sistemleri Ansiklopedisi*, 3/2, 30-32.
- Şeker, S. E. (2016). "Duygu Analizi", *Yönetim Bilişim Sistemleri Ansiklopedisi*, 3/3, 21-36.
- Şen, A. (2025). *Türkçe E-Ticaret Yorumlarının Çok Etiketli Analizi için Derin Öğrenme Modellerinin Uygulanması*, (Basılmamış Yüksek Lisans Tezi), Konya Teknik Üniversitesi Lisansüstü Eğitim Enstitüsü, Konya.
- Taheri, S., Mammadov, M. (2013). "Learning the Naive Bayes Classifier with Optimization Models", *International Journal of Applied Mathematics and Computer Science*, 23/4, 787-795.
- Tahiroğlu, B. T. (2021). "Lematizasyon ve Türkçe için Bir Lematizasyon Uygulaması: ElemanTR", *Rumelide Dil ve Edebiyat Araştırmaları Dergisi*, 24, 475-486

- Talib, R., Hanif, M. K., Ayesha, S., Fatima, F. (2016). "Text Mining: Techniques, Applications and Issues", *International Journal of Advanced Computer Science and Applications*, 7/11, 414-418.
- Tan, K. L., Lee, C. P., Lim, K. M. (2023). "A Survey of Sentiment Analysis: Approaches, Datasets, and Future Research", *Applied Sciences*, 13/7, 4550.
- Terzi, Ö., Küçüksille, E. U., Ergin, G., İlker, A. (2011). "Veri Madenciliği Süreci Kullanılarak Güneş Işınımı Tahmini", *SDU International Technologic Science*, 3/2, 29-37.
- Thrun, S. (1996a). *Explanation-Based Neural Network Learning: A Lifelong Learning Approach*, Kluwer Academic Publishers, Massachusetts.
- Thrun, S. (1996b). "Is Learning the n-th Thing Any Easier than Learning the First?", *Advances in Neural Information Processing Systems*, 8, 640-646.
- Thrun, S., Mitchell, T. M. (1995). "Lifelong Robot Learning", *Robotics and Autonomous Systems*, 15/1-2, 25-26.
- Toçoğlu, M. A., Çelikten, A., Aygün, İ., Alpkoçak, A. (2019). "Türkçe Metinlerde Duygu Analizi için Farklı Makine Öğrenmesi Yöntemlerinin Karşılaştırılması", *Dokuz Eylül Üniversitesi Mühendislik Fakültesi Fen ve Mühendislik Dergisi*, 21/63, 719-725.
- Tohma, K., Kutlu, Y. (2020). "Challenges Encountered in Turkish Natural Language Processing Studies", *Natural and Engineering Sciences*, 5(3), 204-211.
- Topaçan, Ü. (2016). *Sosyal Medya Paylaşımlarında Duygu Analizi: Makine Öğrenimi Yaklaşımı Üzerine Bir Araştırma*, (Basılmamış Doktora Tezi), Marmara Üniversitesi Sosyal Bilimler Enstitüsü, İstanbul.
- Tuna, M. F. (2019). *Çevrimiçi Yorum ve Şikayetlerin Otel İşletmeleri Üzerinden Duygu Analizi ile İncelenmesi*, (Basılmamış Doktora Tezi), Erciyes Üniversitesi Sosyal Bilimler Enstitüsü, Erciyes.
- Türkmenoğlu, C. (2015). *Türkçe Metinlerde Duygu Analizi*, (Basılmamış Yüksek Lisans Tezi), İstanbul Teknik Üniversitesi Fen Bilimleri Enstitüsü, İstanbul.
- Uçan, A. (2014). *Otomatik Duygu Sözlüğü Çevirimi ve Duygu Analizinde Kullanımı*, (Basılmamış Yüksek Lisans Tezi), Hacettepe Üniversitesi Fen Bilimleri Enstitüsü, Ankara.
- Vembandasamy, K., Sasipriya, R., Deepa, E. (2015). "Heart Diseases Detection Using Naive Bayes Algorithm", *International Journal of Innovative Science, Engineering and Technology*, 2/9, 441-444.
- Vijayarani, S., Ilamathi, J., Ms., N. (2015). "Preprocessing Techniques for Text Mining-An Overview", *International Journal of Computer Science and Communication Networks*, 5(1), 7-16.
- Vujovic, Ž. (2021). "Classification Model Evaluation Metrics", *International Journal of Advanced Computer Science and Applications*, 12/6, 599-606.
- Wang, H., Liu, B., Wang, S., Ma, N., Yang, Y. (2019). "Forward and Backward Knowledge Transfer for Sentiment Classification", In *Asian Conference on Machine Learning Research*, 101, 457-472.
- Wang, H., Wang, S., Mazumder, S., Liu, B., Yang, Y., Li, T. (2020). "Bayes-enhanced Lifelong Attention Networks for Sentiment Classification", *Proceedings of the 28th International Conference on Computational Linguistics*, 580-591.
- Wang, S., Zhou, M., Mazumder, S., Liu, B., Chang, Y. (2018). "Disentangling Aspect and Opinion Words in Target-based Sentiment Analysis Using Lifelong Learning", *arXiv preprint arXiv:1802.05818*.

- Wankhade, M., Rao, A. C. S., Kulkarni, C. (2022). "A Survey on Sentiment Analysis Methods, Applications, and Challenges", *Artificial Intelligence Review*, 55/7, 5731-5780.
- Wehrens, R., Putter, H., Buydens, L. M. C. (2000). "The Bootstrap: A Tutorial", *Chemometrics and Intelligent Laboratory Systems*, 54/1, 35-52.
- Wu, X., Kumar, V., Ross Quinlan, J., Ghosh, J., Yang, Q., Motoda, H., McLachlan, G. J., Ng, A., Liu, B., Yu, P. S., Zhou, Z.-H., Steinbach, M., Hand, D. J., Steinberg, D. (2008). "Top 10 Algorithms in Data Mining", *Knowledge and Information Systems*, 14/1, 1-37.
- Xia, R., Jiang, J., He, H. (2017). "Distantly Supervised Lifelong Learning for Large-Scale Social Media Sentiment Analysis", *IEEE Transactions on Affective Computing*, 8/4, 480-491.
- Xu, F., Pan, Z., Xia, R. (2020). "E-commerce Product Review Sentiment Classification Based on a Naive Bayes Continuous Learning Framework", *Information Processing and Management*, 57/5, 102221.
- Xu, Q. A., Chang, V., Jayne, C. (2022), "A Systematic Review of Social Media-Based Sentiment Analysis: Emerging Trends and Challenges", *Decision Analytics Journal*, 3, 100073.
- Yavuz, G. (2023). *Web Kazıma ve Duygu Analizi Temelli Ürün Analiz Sistemi*, (Basılmamış Yüksek Lisans Tezi), Aksaray Üniversitesi Sosyal Bilimler Enstitüsü, Aksaray.
- Yıldırım, M., Yüksel, C. A. (2017). "Sosyal Medya ile Hisse Senedi Fiyatının Günlük Hareket Yönü Arasındaki İlişkinin İncelenmesi: Duygu Analizi Uygulaması", *Uluslararası İktisadi ve İdari İncelemeler Dergisi*, 33-44.
- Yıldırım, P., Birant, D. (2018). "Application of Data Mining Techniques in Cloud Computing: A Literature Review", *Pamukkale University Journal of Engineering Sciences*, 24/2, 336-343.
- Yıldırım, S., Yıldız, T. (2018). "Türkçe için Karşılaştırmalı Metin Sınıflandırma Analizi", *Pamukkale University Journal of Engineering Sciences*, 24/5, 879-886.
- Yiğit, İ. O. (2017). "Çağrı Merkezi Metin Madenciliği Yazılım Çerçevesi", *Turkish National Software Engineering Symposium*, Antalya.
- Zeng, J., Yang, J. (2024). "English Language Hegemony: Retrospect and Prospect", *Humanities and Social Sciences Communications*, 11/1, 1-9.
- Zhang, L., Ghosh, R., Dekhil, M., Hsu, M., Liu, B. (2011). "Combining Lexicon-based and Learning-based Methods for Twitter Sentiment Analysis", *HP Laboratories, Technical Report HPL-2011*, 89, 1-8.
- Zhang, L., Wang, S., Yuan, F., Geng, B., Yang, M. (2023). "Lifelong Language Learning with Adaptive Uncertainty Regularization", *Information Sciences*, 622, 794-807.
- Zong, C., Xia, R., Zhang, J. (2021). *Text Data Mining*, Springer Nature Singapore, Singapore.

EKLER

Ek 1: Yaşam Boyu Naive Bayes Sınıflandırması

```

import nltk
import numpy as np
from sklearn.preprocessing import LabelEncoder
import pandas as pd
import glob
from sklearn.feature_extraction.text import CountVectorizer, TfidfVectorizer
from sklearn.model_selection import train_test_split
from sklearn.naive_bayes import MultinomialNB
from sklearn.metrics import classification_report, confusion_matrix, accuracy_score
LE = LabelEncoder()
vectorizer = CountVectorizer()
train_df = pd.DataFrame(columns=['Xtrain', 'ytrain'])
for i, dosya_adı in enumerate(glob.glob("*****.xlsx")):
    veri = pd.read_excel(dosya_adı)
    documents = veri['Yorumlar']
    labels = veri['Polarite']
    y = LE.fit_transform(labels)
    X_train, X_test, train_y, test_y = train_test_split()
    vectorizer.fit(X_train)
    x_train_count = vectorizer.transform(X_train)
    x_test_count = vectorizer.transform(X_test)
    new_train_df = pd.DataFrame({ 'Xtrain': X_train, 'ytrain': train_y })
    train_df = pd.concat([train_df, new_train_df], ignore_index=True)
#-----
# Lifelong Sentiment Classification Model
kb_xtrain_vec = vectorizer.fit_transform(train_df['Xtrain']).toarray()
words_freq = kb_xtrain_vec.sum(axis=0)
names = vectorizer.get_feature_names_out()
kb_y_train = train_df['ytrain']
nb_classifier = MultinomialNB()
nb_classifier.fit(kb_xtrain_vec, kb_y_train)
y_pred = nb_classifier.predict(x_test_count)
accuracy = accuracy_score(test_y, y_pred)

```

Ek 2: Yaşam Boyu Lojistik Regresyon Sınıflandırması

```

import nltk
import numpy as np
from sklearn.preprocessing import LabelEncoder
import pandas as pd
import glob
from sklearn.feature_extraction.text import CountVectorizer, TfidfVectorizer
from sklearn.model_selection import train_test_split
from sklearn.metrics import classification_report, confusion_matrix, accuracy_score
LE = LabelEncoder()
vectorizer = CountVectorizer(max_features=1000)
train_df = pd.DataFrame(columns=['Xtrain', 'ytrain'])
words_freq = []
names = []
for i, dosya_adı in enumerate(glob.glob("*****.xlsx")):
    veri = pd.read_excel(dosya_adı)
    documents = veri['Yorumlar']
    labels = veri['Polarite']
    y = LE.fit_transform(labels)
    X_train, X_test, y_train, y_test = train_test_split(
        vectorizer.fit(X_train)
        X_train_vec = vectorizer.transform(X_train).toarray()
        X_test_vec = vectorizer.transform(X_test).toarray()
        new_train_df = pd.DataFrame({ 'Xtrain': X_train, 'ytrain': y_train })
        train_df = pd.concat([train_df, new_train_df], ignore_index=True)
    # -----
    # Lifelong Sentiment Classification Model
    kb_xtrain_vec = vectorizer.fit_transform(train_df['Xtrain']).toarray()
    words_freq = kb_xtrain_vec.sum(axis=0)
    names = vectorizer.get_feature_names_out()
    kb_y_train = train_df['ytrain']
    #Stokastik gradyan inişi ile ağırlıkların bulunması
    weights_LLSA, bias_LLSA = train_logistic_regression(kb_xtrain_vec, kb_y_train,
learning_rate=0.1, epochs=1000)
    knowledge_base = pd.DataFrame({'Kelime': names, 'Freq': words_freq, 'Weights':
weights_LLSA})
    def predict(X, weights, bias):
        z = np.dot(X, weights) + bias
        probabilities = sigmoid(z)
        return (probabilities >= 0.5).astype(int)
    y_pred_LLSA = predict(X_test_vec, weights_LLSA, bias_LLSA)
    accuracy = accuracy_score(y_test, y_pred_LLSA)

```

Ek 3: Toplama Dayalı Duygu Sınıflandırması

```

import nltk
from sklearn.preprocessing import LabelEncoder
import pandas as pd
import glob
from sklearn.feature_extraction.text import CountVectorizer, TfidfVectorizer
from sklearn.model_selection import train_test_split
from collections import defaultdict
LE = LabelEncoder()
vectorizer = CountVectorizer()
for i, dosya_adı in enumerate(glob.glob("*****.xlsx")):
    veri = pd.read_excel(dosya_adı)
    documents = veri['Yorumlar']
    labels = veri['Polarite']
    y = LE.fit_transform(labels)
    X_train, X_test, train_y, test_y = train_test_split()
    vectorizer.fit(X_train)
    x_train_count = vectorizer.transform(X_train)
    x_test_count = vectorizer.transform(X_test)
    new_train_df = pd.DataFrame({'Xtrain': X_train, 'ytrain': train_y})
    train_df = pd.concat([train_df, new_train_df], ignore_index=True)
    pozitif_kelime_sıklık_df['kelime'] = pozitif_names
    pozitif_kelime_sıklık_df['pos_freq'] = pozitif_freq
    negatif_kelime_sıklık_df['kelime'] = negatif_names
    negatif_kelime_sıklık_df['neg_freq'] = negatif_freq
    knowledge_base = knowledge_base[['kelime', 'pos_freq', 'neg_freq', 'toplam_freq']]
    #-----
    # Lifelong Sentiment Classification Model
    pozitif_olasılık = (len(pozitif_yorumlar) / len(X_train))
    negatif_olasılık = (len(negatif_yorumlar) / len(X_train))
    for row in knowledge_base.itertuples():
        word_freq_pos.append(row[2])
        word_freq_neg.append(row[3])
    pozitif_cond_prob = []
    negatif_cond_prob = []
    for word in word_freq_pos:
        pozitif_cond_prob.append((((0.1 + word) / (0.1 * len(knowledge_base) +
            sum(i for i in kb_df_pozitif['pos_freq'])))))
    for word in word_freq_neg:
        negatif_cond_prob.append(((0.1 + word) / (0.1 * len(knowledge_base) +
            sum(i for i in kb_df_negatif['neg_freq']))))
    knowledge_base['pos_cond_prob'] = pozitif_cond_prob
    knowledge_base['neg_cond_prob'] = negatif_cond_prob
    for row in knowledge_base.itertuples():
        word_pos_cond_prob.append(row[5])
        word_neg_cond_prob.append(row[6])

```

```

word_pos_cond_prob = dict(zip(knowledge_base['kelime'], word_pos_cond_prob))
word_neg_cond_prob = dict(zip(knowledge_base['kelime'], word_neg_cond_prob))
X_test = [item.lower() for item in X_test]
X_test_names = [x.split() for x in X_test]
epsilon = 1e-12

pos_cumle = []
neg_cumle = []
for cümle in X_test_names:
    cumle_olasiligi_pos = pozitif_olasılık
    cumle_olasiligi_neg = negatif_olasılık
    for kelime in cümle:
        cumle_olasiligi_pos *= word_pos_cond_prob.get(kelime,epsilon)
        cumle_olasiligi_neg *= word_neg_cond_prob.get(kelime,epsilon)
    pos_cumle.append(cumle_olasiligi_pos)
    neg_cumle.append(cumle_olasiligi_neg)
test_kb = pd.DataFrame()
test_kb["test_cümle"] = X_test
test_kb["pozitif_olasılık"] = pos_cumle
test_kb["negatif_olasılık"] = neg_cumle
sinif_tahmini = []
for i, row in test_kb.iterrows():
    if pos_cumle[i] > neg_cumle[i]:
        sınıf_tahmini.append(1)
    else:
        sınıf_tahmini.append(0)
test_kb["sınıf tahmini"] = sınıf_tahmini
test_kb["test_y"] = test_y
dogru_tahmin_sayisi = sum(1 for y_pred_LLSA, y_true in zip(sınıf_tahmini, test_y) if
(y_pred_LLSA == y_true))
dogruluk = dogru_tahmin_sayisi/len(test_y)

```

Ek 4. Türkçe Durak Kelime Listesi

'değil', 'değilim', 'değilmiş', 'altmış', 'altı', 'ancak', 'arada', 'ayrıca', 'bana', 'ben', 'benden', 'beni', 'benim', 'beri', 'beş', 'bile', 'bin', 'bir', 'birçok', 'bir kez', 'bir şey', 'bir şeyi', 'bize', 'bizden', 'bizi', 'bizim', 'böyle', 'böylece', 'buna', 'bunda', 'bundan', 'bunlar', 'bunları', 'bunların', 'bunu', 'bunun', 'burada', 'dahi', 'diğer', 'doksan', 'dokuz', 'dolayı', 'dolayısıyla', 'dört', 'edecek', 'eden', 'ederek', 'edilecek', 'ediliyor', 'edilmesi', 'ediyor', 'elli', 'etmesi', 'etti', 'ettiği', 'ettiğini', 'ettiğim', 'göre', 'halen', 'hani', 'hatta', 'henüz', 'herhangi', 'herkesin', 'hiç bir', 'iki', 'ilgili', 'işte', 'itibaren', 'itibariyle', 'kadar', 'karşın', 'katrilyon', 'kendî', 'kendilerine', 'kendini', 'kendisi', 'kendisine', 'kendisini', 'kimden', 'kime', 'kimi', 'kimse', 'kırk', 'milyar', 'milyon', 'mı', 'mi', 'mu', 'mü', 'mıydı', 'miydi', 'müydü', 'muydu', 'nedenle', 'olan', 'olarak', 'oldu', 'olduğu', 'olduğunu', 'olduklarını', 'olmadı', 'olmadığı', 'olmak', 'olması', 'olmayan', 'olmayan', 'olmaz', 'olmaz', 'olsa', 'olsun', 'olup', 'olur', 'olursa', 'oluyor', 'on', 'ona', 'ondan', 'onlar', 'onlardan', 'onları', 'onların', 'onu', 'onun', 'otuz', 'oysa', 'öyle', 'pek', 'rağmen', 'sadece', 'sekiz', 'seksen', 'sen', 'senden', 'seni', 'senin', 'siz', 'sizden', 'sizi', 'sizin', 'şeyden', 'şeyi', 'şeyler', 'şöyle', 'şuna', 'şunda', 'şundan', 'şunları', 'şunu', 'tarafından', 'trilyon', 'üç', 'üzere', 'üzerinde', 'var', 'vardı', 'yapacak', 'yapılan', 'yapılması', 'yapıyor', 'yapmak', 'yaptı', 'yaptığı', 'yaptığını', 'yaptıkları', 'yedi', 'yerine', 'yetmiş', 'yine', 'yirmi', 'yoksa', 'yüz', 'zaten', 'madem', 'altmış', 'altı', 'bazı', 'beş', 'kırk', 'nasıl', 'yetmiş'.

Ek 5. Eğitim Vektörü Kelime Listesi (Kısaltılmış)

['abartıldığı' 'ablama' 'adet' 'adetli' 'adi' 'aile' 'ailecek' 'airfryer' 'akmış' 'aksesuar' 'alabilirsiniz' 'alacak' 'alacağım' 'alakası' 'alakasız' 'alalı' 'alayım' 'aldanmayın' 'aldı' 'aldık' 'aldıktan' 'aldım' 'aldırın' 'aldığım' 'aldığıma' 'aldığımda' 'aldığımdan' 'aldığımız' 'allah' 'almak' 'almam' 'almama' 'almanızı' 'almasaydım' 'almaya' 'almayacağım' 'almayı' 'almayın' 'almış' 'almıştık' 'almıştım' 'almışım' 'alsam' 'alsaydım' 'alsın' 'alt' 'alın' 'alınabilir' 'alınca' 'alınmalı' 'alınır' 'alıp' 'alırken' 'alırım' 'alıyor' 'alıyorum' 'alıyoruz' 'alışveriş' 'alışım' 'ambalaj' 'ambalajı' 'an' 'anca' 'anda' 'anlamadım' 'anlamıyorum' 'anneler' 'annem' 'anneme' 'anında' 'ariel' 'arkadaşlar' 'arkadaşım' 'arkadaşıma' 'artık' 'asla' 'astar' 'astarı' 'attım' 'ay' 'ayatlanmıyor' 'aydır' 'aynı' 'aynısı' 'ayrı' 'açtım' 'açtığımda' 'açık' 'açıkçası' 'açıp' 'ağır' 'aşırı' 'baby' 'bakarak' 'bargello' 'basit' 'baya' 'bayağı' 'bayıldı' 'bayıldım' 'bayılıyorum' 'başarılı' 'başka' 'başladı' 'başladım' 'bebek' 'bebeğim' 'bebeğime' 'bebeğimin' 'bebeğimiz' 'beden' 'bedeni' 'bedenim' 'bedenimi' 'bedeninizi' 'bekledim' 'beklediğim' 'beklediğimden' 'beklemeyin' 'beklemiyordum' 'bekletince' 'bekliyorum' 'belli' 'bence' 'bende' 'benzemiyor' 'berbat' 'bey' 'beyaz' 'bez' 'bezi' 'bezler' 'beğendi' 'beğendik' 'beğendim' 'beğendiğim' 'beğenemedim' 'beğenerek' 'beğenildi' 'beğenmedik' 'beğenmedim' 'bidon' 'bildiğin' 'bildiğiniz' 'bilgi' 'bilmiyorum' 'bilseydim' 'bingo' 'biraz' 'birde' 'birlikte' 'bitmeden' 'bitti' 'bittikçe' 'bol' 'boy' 'boyu' 'boyutu' 'bozuk' 'bozuldu' 'boş' 'boşa' 'boşuna' 'bugün' 'bundan' 'buradan' 'buz' 'bütün' 'büyük' 'büyüklüğü' 'bırakıyor' 'cam' 'camları' 'camı' 'cevap' 'cidden' 'cırtlari' 'dahi' 'dakika' 'dalga' 'damak' 'dan' 'dandik' 'dar' 'dedi' 'dediler' 'dedim' 'defolu' 'değil' 'denedim' 'denemek' 'denk' 'derece' 'derecede' 'derim' 'deterjan' 'devam' 'devamlı' 'değildi' 'değişik' 'değmez' 'dikiş' 'dikişleri' 'dikkat' 'dikkatli' 'direk' 'diyebilirim' 'diyor' 'diğeri' 'dk' 'dokusu' 'doğacak' 'doğan' 'doğru' 'doğum' 'durdu'..... 'olabilirdi' 'olamaz' 'olanlar' 'olduk' 'oldukça' 'oldum' 'olduğum' 'olduğundan' 'olmadan' 'olmasın' 'olmasına' 'olmasını' 'olmuyor' 'olmuş' 'olun' 'olunca' 'olurdu' 'olurmuş' 'oranı' 'orijinal' 'orta' 'oyuncak' 'pahalı' 'paket' 'paketi' 'paketin' 'paketleme' 'paketlemesi' 'paketlenmesi' 'paketlenmiş' 'paketlenmişti' 'para' 'param' 'paranıza' 'paranızı' 'parasını' 'paraya' 'parayı' 'parfüm' 'parfümü' 'parlak' 'parça' 'patates' 'patlak' 'patlamış' 'pazar' 'pembe' 'performans' 'performansı' 'pes' 'peynir' 'peyniri' 'peynirin' 'peynirler' 'philips' 'pili' 'pişik' 'pişiriyor' 'pişirme' 'pişirmiyor' 'pişman' 'pişmanım' 'plastik' 'poşet' 'poşete' 'poşette' 'pratik' 'ptt' 'puan' 'pudra' 'puf' 'rahat' 'rahatlıkla' 'rahatlığıyla' 'rahatsız' 'rengi' 'rengini' 'renk' 'resimde' 'resmen' 'rezalet' 'rezil' 'rezillik' 'saat' 'saati' 'saatim' 'saatin' 'saatçiy'e' 'sabah' 'sahte' 'sakın' 'salamura' 'salaş' 'sanırım' 'zarar' 'zarif' 'zincir' 'zor' 'zorunda' 'çabuk' 'çakma' 'çalışmıyor' 'çalışıyor' 'çanta' 'çantanın' 'çantası' 'çantayı' 'çikolata' 'çizik' 'çizikler' 'çocuk' 'çokta' 'çöp' 'çöpe' 'çıkarmıyor' 'çıkmadı' 'çıktı' 'çıkıyor' 'önce' 'önceden' 'önceki' 'öncelikle' 'öneririm' 'önermiyorum' 'örgü' 'ötesi' 'özellikle' 'özen' 'özenle' 'özenli' 'özenliydi' 'özensiz' 'üniversite' 'ürün' 'üründe' 'üründen' 'ürüne' 'ürünle' 'ürünler' 'ürünleri' 'ürünlerin' 'ürünü' 'ürünüm' 'ürünümü' 'ürünün' 'üst' 'üstelik' 'üstü' 'üstünde' 'üstüne' 'üzdü' 'üzerine' 'üzgünüm' 'üzüldüm' 'ıslak' 'şahane' 'şaka' 'şekerli' 'şekilde' 'şeklinde' 'şiddetle' 'şikayet' 'şimdi' 'şirketi' 'şuan' 'şık'

Ek 6. Farklı 20 Alan Sıralaması

Alan Sırası	Alan Sıralamaları																			
	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20
1	A1	A1	A2	A3	A10	A8	A7	A9	A4	A9	A4	A6	A5	A5	A6	A8	A6	A4	A8	A3
2	A2	A4	A1	A10	A7	A9	A2	A4	A7	A10	A6	A9	A3	A8	A10	A1	A9	A5	A9	A7
3	A3	A5	A3	A5	A4	A6	A1	A6	A2	A1	A2	A3	A10	A6	A5	A6	A3	A6	A6	A2
4	A4	A3	A5	A4	A2	A10	A4	A2	A8	A5	A8	A5	A1	A9	A1	A5	A5	A7	A7	A4
5	A5	A10	A8	A2	A3	A5	A3	A8	A1	A3	A7	A10	A7	A10	A4	A9	A10	A8	A1	A5
6	A6	A2	A4	A1	A8	A1	A10	A7	A10	A7	A3	A1	A8	A3	A8	A10	A8	A9	A2	A10
7	A7	A7	A9	A7	A1	A4	A9	A1	A5	A8	A5	A8	A2	A4	A9	A7	A2	A10	A4	A6
8	A8	A8	A10	A6	A6	A2	A6	A10	A3	A2	A1	A2	A6	A1	A2	A4	A7	A1	A5	A9
9	A9	A9	A6	A9	A5	A7	A8	A3	A9	A6	A10	A7	A4	A2	A7	A2	A4	A2	A10	A8
10	A10	A6	A7	A8	A9	A3	A5	A5	A6	A4	A9	A4	A9	A7	A3	A3	A1	A3	A3	A1

A1: Airfryer, A2: G mlek, A3:  anta, A4: Toz Deterjan, A5: Bebek Bezi, A6: Saat, A7: G zl k, A8: Parf m, A9: Peynir, A10:  ikolata

Ek 7. Farklı Alan Sıralamalarına Ait Model Doğrulukları

Alanlar	CountVectorizer	TF_IDF
Alan Sıralaması 1	0,95	0,95
Alan Sıralaması 2	0,99	0,99
Alan Sıralaması 3	0,945	0,975
Alan Sıralaması 4	0,95	0,96
Alan Sıralaması 5	0,93	0,945
Alan Sıralaması 6	0,975	0,98
Alan Sıralaması 7	0,97	0,96
Alan Sıralaması 8	0,97	0,96
Alan Sıralaması 9	0,98	0,98
Alan Sıralaması 10	0,945	0,965
Alan Sıralaması 11	0,92	0,94
Alan Sıralaması 12	0,96	0,94
Alan Sıralaması 13	0,92	0,92
Alan Sıralaması 14	0,94	0,95
Alan Sıralaması 15	0,96	0,97
Alan Sıralaması 16	0,95	0,96
Alan Sıralaması 17	0,96	0,95
Alan Sıralaması 18	0,97	0,97
Alan Sıralaması 19	0,94	0,97
Alan Sıralaması 20	0,96	0,95

Ek 8. İzole Modelin ve Toplama Dayalı Duygu Sınıflandırması Modelinin Performansları

Alan Sıralaması 1					Alan Sıralaması 2				
	CountVec		TFIDF			CountVec		TFIDF	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Airfryer	0,92	0,94	0,93	0,93	Airfryer	0,92	0,94	0,93	0,93
Gömlek	0,95	0,96	0,99	0,98	Toz Deterjan	0,93	0,94	0,93	0,94
Çanta	0,96	0,95	0,96	0,95	Bebek Bezi	0,94	0,93	0,93	0,94
Toz Deterjan	0,93	0,94	0,93	0,94	Çanta	0,96	0,95	0,96	0,98
Bebek Bezi	0,94	0,92	0,93	0,94	Çikolata	0,93	0,93	0,94	0,94
Saat	0,99	0,99	0,98	0,99	Gömlek	0,95	0,96	0,99	0,97
Gözlük	0,99	0,94	0,98	0,95	Gözlük	0,99	0,94	0,98	0,95
Parfüm	0,95	0,96	0,96	0,96	Parfüm	0,95	0,95	0,96	0,97
Peynir	0,94	0,94	0,92	0,93	Peynir	0,94	0,93	0,92	0,94
Çikolata	0,93	0,95	0,94	0,95	Saat	0,99	0,99	0,98	0,99
Alan Sıralaması 3					Alan Sıralaması 4				
	CountVec		TFIDF			CountVec		TFIDF	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Gömlek	0,94	0,96	0,97	0,965	Çanta	0,955	0,96	0,955	0,965
Airfryer	0,935	0,945	0,935	0,96	Çikolata	0,92	0,935	0,93	0,925
Çanta	0,955	0,96	0,955	0,975	Bebek Bezi	0,955	0,955	0,945	0,96
Bebek Bezi	0,955	0,95	0,945	0,96	Toz Deterjan	0,945	0,95	0,94	0,955
Parfüm	0,955	0,955	0,96	0,96	Gömlek	0,94	0,98	0,97	0,995
Toz Deterjan	0,945	0,955	0,94	0,945	Airfryer	0,935	0,96	0,935	0,955
Peynir	0,93	0,93	0,93	0,94	Gözlük	0,985	0,945	0,98	0,96
Çikolata	0,92	0,94	0,93	0,94	Saat	0,96	0,98	0,955	0,98
Saat	0,96	0,98	0,955	0,985	Peynir	0,93	0,925	0,93	0,945
Gözlük	0,985	0,945	0,98	0,975	Parfüm	0,955	0,95	0,96	0,96
Alan Sıralaması 5					Alan Sıralaması 6				
	CountVec		TFIDF			CountVec		TFIDF	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Çikolata	0,92	0,92	0,93	0,92	Parfüm	0,955	0,96	0,96	0,95
Gözlük	0,985	0,965	0,98	0,97	Peynir	0,93	0,9	0,93	0,925
Toz Deterjan	0,945	0,945	0,94	0,945	Saat	0,96	0,965	0,955	0,965
Gömlek	0,94	0,98	0,97	0,98	Çikolata	0,92	0,945	0,93	0,94
Çanta	0,955	0,97	0,955	0,975	Bebek Bezi	0,955	0,975	0,945	0,975
Parfüm	0,955	0,95	0,96	0,96	Airfryer	0,935	0,965	0,935	0,96

Airfryer	0,935	0,97	0,935	0,96	Toz Deterjan	0,945	0,95	0,94	0,96
Saat	0,96	0,975	0,955	0,975	Gömlek	0,94	0,97	0,97	0,99
Bebek Bezi	0,955	0,97	0,945	0,965	Gözlük	0,985	0,95	0,98	0,975
Peynir	0,93	0,93	0,93	0,945	Çanta	0,955	0,975	0,955	0,98
Alan Sıralaması 7					Alan Sıralaması 8				
	CountVec		TFIDF			CountVec		TFIDF	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Gözlük	0,985	0,98	0,98	0,985	Peynir	0,93	0,91	0,93	0,935
Gömlek	0,94	0,965	0,97	0,98	Toz Deterjan	0,945	0,965	0,94	0,965
Airfryer	0,935	0,94	0,935	0,94	Saat	0,96	0,965	0,955	0,965
Toz Deterjan	0,945	0,945	0,94	0,94	Gömlek	0,94	0,975	0,97	0,985
Çanta	0,955	0,975	0,955	0,975	Parfüm	0,955	0,96	0,96	0,96
Çikolata	0,92	0,945	0,93	0,935	Gözlük	0,985	0,97	0,98	0,975
Peynir	0,93	0,91	0,93	0,94	Airfryer	0,935	0,965	0,935	0,96
Saat	0,96	0,98	0,955	0,98	Çikolata	0,92	0,95	0,93	0,95
Parfüm	0,955	0,95	0,96	0,96	Çanta	0,955	0,98	0,955	0,98
Bebek Bezi	0,955	0,97	0,945	0,96	Bebek Bezi	0,955	0,97	0,945	0,96
Alan Sıralaması 9					Alan Sıralaması 10				
	CountVec		TFIDF			CountVec		TFIDF	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Toz Deterjan	0,945	0,96	0,94	0,945	Peynir	0,93	0,91	0,93	0,935
Gözlük	0,985	0,97	0,98	0,975	Çikolata	0,92	0,945	0,93	0,95
Gömlek	0,94	0,965	0,97	0,965	Airfryer	0,935	0,97	0,935	0,96
Parfüm	0,955	0,955	0,96	0,955	Bebek Bezi	0,955	0,95	0,945	0,95
Airfryer	0,935	0,975	0,935	0,97	Çanta	0,955	0,955	0,955	0,97
Çikolata	0,92	0,95	0,93	0,953	Gözlük	0,985	0,955	0,98	0,965
Bebek Bezi	0,955	0,97	0,945	0,96	Parfüm	0,955	0,97	0,96	0,965
Çanta	0,955	0,98	0,955	0,98	Gömlek	0,94	0,975	0,97	0,995
Peynir	0,93	0,915	0,93	0,935	Saat	0,96	0,985	0,955	0,975
Saat	0,96	0,98	0,955	0,98	Toz Deterjan	0,945	0,945	0,94	0,965
Alan Sıralaması 11					Alan Sıralaması 12				
	CountVec		TFIDF			CountVec		TFIDF	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Toz Deterjan	0,93	0,94	0,93	0,93	Saat	0,99	0,99	0,98	0,98
Saat	0,99	1	0,98	1	Peynir	0,94	0,95	0,92	0,95
Gömlek	0,95	0,97	0,99	0,98	Çanta	0,96	0,95	0,96	0,96

Parfüm	0,95	0,96	0,96	0,96	Bebek Bezi	0,94	0,94	0,95	0,94
Gözlük	0,99	0,93	0,98	0,95	Çikolata	0,93	0,94	0,94	0,95
Çanta	0,96	0,95	0,96	0,97	Airfryer	0,92	0,94	0,93	0,93
Bebek Bezi	0,94	0,92	0,93	0,93	Parfüm	0,95	0,97	0,96	0,97
Airfryer	0,92	0,93	0,93	0,94	Gömlek	0,95	0,98	0,99	0,99
Çikolata	0,93	0,94	0,94	0,94	Gözlük	0,99	0,93	0,98	0,95
Peynir	0,94	0,92	0,92	0,94	Toz Deterjan	0,93	0,96	0,93	0,94
Alan Sıralaması 13					Alan Sıralaması 14				
	CountVec		TFIDE			CountVec		TFIDE	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Bebek Bezi	0,94	0,95	0,93	0,95	Bebek Bezi	0,94	0,95	0,93	0,95
Çanta	0,96	0,95	0,96	0,95	Parfüm	0,95	0,97	0,96	0,97
Çikolata	0,93	0,94	0,94	0,94	Saat	0,99	0,98	0,98	1
Airfryer	0,92	0,96	0,93	0,96	Peynir	0,94	0,94	0,92	0,93
Gözlük	0,99	0,95	0,98	0,96	Çikolata	0,93	0,95	0,94	0,95
Parfüm	0,95	0,97	0,96	0,97	Çanta	0,96	0,94	0,96	0,96
Gömlek	0,95	0,98	0,99	0,99	Toz Deterjan	0,93	0,96	0,93	0,94
Saat	0,99	0,99	0,98	0,99	Airfryer	0,92	0,95	0,93	0,95
Toz Deterjan	0,93	0,94	0,93	0,92	Gömlek	0,95	0,98	0,99	0,99
Peynir	0,94	0,92	0,92	0,92	Gözlük	0,99	0,94	0,98	0,95
Alan Sıralaması 15					Alan Sıralaması 16				
	CountVec		TFIDE			CountVec		TFIDE	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Saat	0,99	0,99	0,98	0,98	Parfüm	0,95	0,96	0,96	0,96
Çikolata	0,93	0,93	0,94	0,93	Airfryer	0,92	0,95	0,93	0,95
Bebek Bezi	0,94	0,93	0,93	0,95	Saat	0,99	0,99	0,98	1
Airfryer	0,92	0,93	0,93	0,92	Bebek Bezi	0,94	0,93	0,93	0,94
Toz Deterjan	0,93	0,92	0,93	0,92	Peynir	0,94	0,93	0,92	0,93
Parfüm	0,95	0,95	0,96	0,96	Çikolata	0,93	0,94	0,94	0,94
Peynir	0,94	0,94	0,92	0,93	Gözlük	0,99	0,93	0,98	0,95
Gömlek	0,95	0,99	0,99	0,99	Toz Deterjan	0,93	0,95	0,93	0,92
Gözlük	0,99	0,94	0,98	0,95	Gömlek	0,95	0,98	0,99	0,99
Çanta	0,96	0,96	0,96	0,97	Çanta	0,96	0,95	0,96	0,96
Alan Sıralaması 17					Alan Sıralaması 18				
	CountVec		TFIDE			CountVec		TFIDE	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk

Saat	0,99	0,99	0,98	0,98	Toz Deterjan	0,93	0,94	0,93	0,93
Peynir	0,94	0,95	0,92	0,95	Bebek Bezi	0,94	0,94	0,93	0,95
Çanta	0,96	0,95	0,96	0,96	Saat	0,99	0,99	0,98	1
Bebek Bezi	0,94	0,94	0,93	0,94	Gözlük	0,99	0,93	0,98	0,96
Çikolata	0,93	0,94	0,94	0,95	Parfüm	0,95	0,97	0,96	0,97
Parfüm	0,95	0,97	0,96	0,97	Peynir	0,94	0,94	0,92	0,93
Gömlek	0,95	0,98	0,99	0,99	Çikolata	0,93	0,93	0,94	0,93
Gözlük	0,99	0,92	0,98	0,94	Airfryer	0,92	0,95	0,93	0,95
Toz Deterjan	0,93	0,96	0,93	0,95	Gömlek	0,95	0,97	0,99	0,99
Airfryer	0,92	0,96	0,93	0,95	Çanta	0,96	0,97	0,96	0,97
Alan Sıralaması 19					Alan Sıralaması 20				
	CountVec		TFIDF			CountVec		TFIDF	
Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk	Alanlar	İzole Doğruluk	LLSA Doğruluk	İzole Doğruluk	LLSA Doğruluk
Parfüm	0,95	0,96	0,96	0,96	Çanta	0,96	0,94	0,96	0,95
Peynir	0,94	0,93	0,92	0,91	Gözlük	0,99	0,94	0,98	0,96
Saat	0,99	0,99	0,98	1	Gömlek	0,95	0,97	0,99	0,99
Gözlük	0,99	0,94	0,98	0,97	Toz Deterjan	0,93	0,97	0,93	0,96
Airfryer	0,92	0,96	0,93	0,96	Bebek Bezi	0,94	0,91	0,93	0,94
Gömlek	0,95	0,96	0,99	0,98	Çikolata	0,93	0,93	0,94	0,93
Toz Deterjan	0,93	0,95	0,93	0,95	Saat	0,99	0,98	0,98	0,98
Bebek Bezi	0,94	0,93	0,93	0,93	Peynir	0,94	0,93	0,92	0,94
Çikolata	0,93	0,95	0,94	0,95	Parfüm	0,95	0,97	0,96	0,97
Çanta	0,96	0,94	0,96	0,97	Airfryer	0,92	0,96	0,93	0,95