



T.C.
NECMETTİN ERBAKAN
ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ



TÜRKÇE FİLM YORUMLARINDA DUYGU
ANALİZİ İÇİN BAĞLAMSAK, KURALLI VE
İSTATİSTİKSEL YÖNTEMLERİN HİBRİT
MODEL YAKLAŞIMI İLE
KARŞILAŞTIRMALI ANALİZİ

Beyza DÜNDAR
YÜKSEK LİSANS TEZİ

Nisan-2026
KONYA
Her Hakkı Saklıdır

TEZ KABUL VE ONAYI

Beyza DÜNDAR tarafından hazırlanan “TÜRKÇE FİLM YORUMLARINDA DUYGU ANALİZİ İÇİN BAĞLAMSAK, KURALLI VE İSTATİSTİKSEL YÖNTEMLERİN HİBRİT MODEL YAKLAŞIMI İLE KARŞILAŞTIRMALI ANALİZİ” adlı tez çalışması 27/04/2026 tarihinde aşağıdaki jüri tarafından oy birliği / oy çokluğu ile Necmettin Erbakan Üniversitesi Fen Bilimleri Enstitüsü Bilgisayar Mühendisliği Anabilim Dalı’nda YÜKSEK LİSANS TEZİ olarak kabul edilmiştir.

Jüri Üyeleri

Başkan

Doç.Dr Mehmet Akif ŞAHMAN

Danışman

Prof.Dr Sabri KOÇER

Üye

Dr. Öğr. Üyesi Alperen EROĞLU

İmza

.....

.....

.....

Fen Bilimleri Enstitüsü Yönetim Kurulu’nun/.../20.. gün ve sayılı kararıyla onaylanmıştır.

Prof. Dr. Havvanur UÇBEYİAY
FBE Müdürü

TEZ BİLDİRİMİ

Bu tezdeki bütün bilgilerin etik davranış ve akademik kurallar çerçevesinde elde edildiğini ve tez yazım kurallarına uygun olarak hazırlanan bu çalışmada bana ait olmayan her türlü ifade ve bilginin kaynağına eksiksiz atıf yapıldığını bildiririm.

DECLARATION PAGE

I hereby declare that all information in this document has been obtained and presented in accordance with academic rules and ethical conduct. I also declare that, as required by these rules and conduct, I have fully cited and referenced all material and results that are not original to this work.

Öğrencinin Adı SOYADI
Beyza DÜNDAR

Tarih: 27/04/2026

ÖZET
YÜKSEK LİSANS TEZİ
TÜRKÇE FİLM YORUMLARINDA DUYGU ANALİZİ İÇİN BAĞLAMSAK,
KURALLI VE İSTATİSTİKSEL YÖNTEMLERİN HİBRİT MODEL YAKLAŞIMI
İLE KARŞILAŞTIRMALI ANALİZİ

Beyza DÜNDAR

Necmettin Erbakan Üniversitesi Fen Bilimleri Enstitüsü
Bilgisayar Mühendisliği Anabilim Dalı

Danışman: Prof.Dr. Sabri KOÇER

Bu tez çalışmasında, doğal dil işleme (NLP) ve duygu analizi yöntemleri kullanılarak Türkçe film yorumlarının otomatik olarak sınıflandırılması amaçlanmıştır. Günümüzde dijital platformlarda kullanıcılar tarafından üretilen metinsel verilerin hızla artması, bu verilerden anlamlı bilgi elde edilmesini önemli hale getirmiştir. Özellikle Türkçe gibi doğal dil işleme kaynakları sınırlı dillerde duygu analizi çalışmaları, akademik ve uygulamalı araştırmalar açısından önemli bir problem alanı oluşturmaktadır. Bu kapsamda çalışmada, Türkçe film yorumları üzerinde farklı makine öğrenmesi, derin öğrenme ve bulanık mantık tabanlı yöntemlerin performansları karşılaştırmalı olarak incelenmiştir.

Çalışmada kullanılan veri seti, Kaggle platformunda yer alan “Turkish Movie Sentiment Analysis Dataset” veri setinden elde edilmiştir. Toplam 80.000 yorumdan oluşan veri seti, işlem maliyetini azaltmak ve çalışma analizlerini daha verimli gerçekleştirebilmek amacıyla split yöntemi kullanılarak 40.000 yorumluk alt veri setine dönüştürülmüştür. Yorumlar olumlu ve olumsuz olarak etiketlenmiş, ardından metin verileri üzerinde gerekli ön işleme adımları uygulanmıştır.

Birinci çalışmada, Naive Bayes, Lojistik Regresyon, Destek Vektör Makineleri (DVM), BERT tabanlı derin öğrenme modeli ve Bulanık Mantık yaklaşımı kullanılarak duygu sınıflandırması gerçekleştirilmiştir. İkinci çalışmada, Mamdani Max–Min sınıflandırma yöntemine dayalı bulanık çıkarım mekanizması kullanılarak bulanık mantık modeli geliştirilmiştir. Üçüncü çalışmada ise, BERT ve Bulanık Mantık yöntemlerinin birlikte kullanıldığı hibrit model ile BERT, Bulanık Mantık ve TF-IDF/Lojistik Regresyon yöntemlerini bir araya getiren üçlü hibrit

yapı önerilmiştir. Tüm modellerin performansları doğruluk, kesinlik, geri çağırma ve F1-skor ölçütleri kullanılarak değerlendirilmiştir.

Elde edilen sonuçlara göre, birinci çalışmada BERT modeli %81 doğruluk ve %82 F1-skor değeri ile en başarılı tekil model olmuştur. Naive Bayes modeli %72 doğruluk ve %70 F1-skor, Lojistik Regresyon modeli %76 doğruluk ve %75 F1-skor, Destek Vektör Makineleri modeli ise %75 doğruluk ve %74 F1-skor değeri elde etmiştir. İkinci çalışmada geliştirilen bulanık mantık tabanlı model %68 doğruluk ve %66 F1-skor başarısı göstermiştir. Üçüncü çalışmada önerilen Bulanık Mantık–BERT hibrit modeli %83 doğruluk ve %84 F1-skor elde ederken, Bulanık Mantık–TF-IDF/Lojistik Regresyon–BERT üçlü hibrit modeli %85 doğruluk ve %86 F1-skor ile tüm çalışmalar içerisindeki en yüksek performansı sağlamıştır.

Sonuç olarak, Türkçe film yorumlarında duygu sınıflandırma başarısının kullanılan model mimarisi ve karar mekanizmasına bağlı olarak değiştiği görülmüştür. Özellikle hibrit yapıların, tekil modellere kıyasla daha yüksek başarı ve daha dengeli sonuçlar sunduğu belirlenmiştir. Çalışma, Türkçe duygu analizi literatürüne karşılaştırmalı çalışma sonuçlarını sunmasının yanı sıra, bulanık mantık tabanlı hibrit yöntemlerin duygu analizi problemlerindeki etkinliğini ortaya koyarak akademik çalışmalara ve karar destek sistemlerine katkı sağlamayı amaçlamaktadır.

Anahtar Kelimeler: Doğal Dil İşleme, Türkçe Duygu Analizi, Makine Öğrenmesi, Film Yorumları

ABSTRACT

MASTER'S THESIS

A COMPARATIVE ANALYSIS OF CONTEXTUAL, RULE-BASED, AND STATISTICAL METHODS FOR EMOTION ANALYSIS IN TURKISH FILM REVIEWS USING A HYBRID MODEL APPROACH

Beyza DÜNDAR

NECMETTİN ERBAKAN UNIVERSITY INSTITUTE OF NATURAL SCIENCES

DEPARTMENT OF COMPUTER ENGINEERING

Advisor: Prof. Dr. Sabri KOÇER

This thesis aims to automatically classify Turkish movie reviews using natural language processing (NLP) and sentiment analysis methods. The rapid increase in text data generated by users on digital platforms today has made extracting meaningful information from this data a critical task. In particular, sentiment analysis studies in languages with limited NLP resources, such as Turkish, represent a significant area of research for both academic and applied studies. In this context, the study comparatively examines the performance of various machine learning, deep learning, and fuzzy logic-based methods on Turkish movie reviews.

The dataset used in the study was obtained from the “Turkish Movie Sentiment Analysis Dataset” available on the Kaggle platform. The dataset, consisting of a total of 80,000 reviews, was converted into a 40,000-review subset using the split method to reduce processing costs and conduct experimental analyses more efficiently. The reviews were labeled as positive or negative, and then the necessary preprocessing steps were applied to the text data.

In the first study, sentiment classification was performed using Naive Bayes, Logistic Regression, Support Vector Machines (SVM), a BERT-based deep learning model, and a Fuzzy Logic approach. In the second study, a fuzzy logic model was developed using a fuzzy inference mechanism based on the Mamdani Max–Min classification method. In the third study, a hybrid model combining BERT and Fuzzy Logic methods was proposed, along with a triple hybrid structure integrating BERT, Fuzzy Logic, and TF-IDF/Logistic Regression methods. The performance of all models was evaluated using accuracy, precision, recall, and F1-score metrics.

According to the results, in the first study, the BERT model was the most successful standalone model with 81% accuracy and an F1-score of 82%. The Naive Bayes model achieved 72% accuracy and a 70% F1-score, the Logistic Regression model achieved 76% accuracy and a 75% F1-score, and the Support Vector Machines model achieved 75% accuracy and a 74% F1-score. The fuzzy logic-based model developed in the second study achieved 68% accuracy and a 66% F1-score. In the third study, the proposed Fuzzy Logic–BERT hybrid model achieved 83% accuracy and an 84% F1-score, while the Fuzzy Logic–TF-IDF/Logistic Regression–BERT triple hybrid model delivered 85% accuracy and an 86% F1-score, providing the highest performance among all studies.

In conclusion, it was observed that the success of sentiment classification in Turkish movie reviews varied depending on the model architecture and decision mechanism used. In particular, it was determined that hybrid architectures offered higher success rates and more balanced results compared to single models. In addition to providing comparative experimental results for the Turkish sentiment analysis literature, this study aims to contribute to academic research and decision support systems by demonstrating the effectiveness of fuzzy logic-based hybrid methods in sentiment analysis problems.

Keywords: Natural Language Processing, Turkish Sentiment Analysis, Machine Learning, Movie Reviews

ÖNSÖZ

Yüksek lisans tez çalışmamın planlanması, yürütülmesi ve tamamlanması sürecinde, değerli bilgi ve tecrübelerini benimle paylaşarak bana her aşamada yol gösteren, vizyonu ile çalışmamı şekillendiren kıymetli danışman hocam Prof. Dr. Sabri KOÇER'e en içten teşekkürlerimi ve saygılarımı sunarım.

Hayatımı inşa ettiğim ilk gençlik yıllarımdan bugüne kadar, eğitim hayatımın her kademesinde maddi ve manevi desteklerini esirgemeyen, her koşulda bana inanan, güvenen ve sevgileriyle beni her zaman güçlü kılan sevgili aileme şükran borçluyum.

Son olarak, bu uzun, yoğun ve akademik açıdan zorlu süreci büyük bir kararlılık, sabır ve azimle sürdürerek çalışmanın başarıyla tamamlanmasında en büyük paya sahip olan kendime, gösterdiğim bu içsel güce ve emeğe yönelik derin bir memnuniyet taşımaktayım. Bu süreçte bana eşlik eden ve destek olan herkese teşekkür ederim.

Beyza DÜNDAR

Konya/2026

İÇİNDEKİLER

ÖZET.....	4
ÖNSÖZ.....	8
1.GİRİŞ.....	12
2. LİTERATÜR TARAMASI	15
3.MATERYAL VE YÖNTEMLER.....	26
3.1. Yapay Zeka	26
3.2 Makine Öğrenmesi	26
3.3 Derin Öğrenme	27
3.4 Doğal Dil İşleme.....	28
3.5 Veriseti.....	29
3.6. Veri Ön işleme	30
3.7 Kullanılan Algoritmalar	31
3.7.1 Hiperparametre ve Konfigürasyon Ayarları	31
3.7.2 TF - IDF (Term Frequency – Inverse Document Frequency / Terim Sıklığı – Ters Belge Sıklığı)	32
3.7.3 Lojistik Regresyon.....	35
3.7.4 Naive Bayes	38
3.7.5 Destek Vektör Makineleri (DVM)	40
3.7.6 Bulanık Mantık (Fuzzy Logic)	42
3.7.6 BERT(Bidirectional Encoder Representations from Transformers /Transformatörlerden İki Yönlü Kodlayıcı Temsilleri)	45
3.8 Çalışmada Kullanılan Başarım Ölçüm Terimleri	52
3.8.1 Duyarlılık Değeri (Precision).....	52
3.8.2. Anma Değeri (Recall).....	53
3.8.3. Doğruluk (Accuracy).....	53
3.8.4. F1 Ölçümü (F1 Measure).....	53
4.ARAŞTIRMA BULGULARI VE TARTIŞMA	54
4.1. Çalışma 1: İstatistiksel Modeller ve BERT Bulguları	54
4.2. Çalışma 2: Bulanık Mantık Sistemi Bulguları	55
4.3. Çalışma 3: Hibrit Modeller ve Entegre Analiz	56
4.4 Genel Değerlendirme ve Tartışma	59
5.SONUÇ VE ÖNERİLER	59
5.1. Sonuçlar.....	60
5.2. Öneriler	60
6.KAYNAKLAR.....	62



KISALTMALAR VE SİMGELER

Kısaltmalar

BERT	: Transformatörlerden Çift Yönlü Kodlayıcı Temsilleri
DA/SA	: Duygu Analizi / Sentiment Analysis
DDİ/NLP	: Doğal Dil İşleme / Natural Language Processing
DVM/SVM	: Destek Vektör Makineleri / Support Vector Machine
BM/ FL	: Bulanık Mantık / Fuzzy Logic
LR	: Lojistik Regresyon / Logistic Regression
MÖ/ML	: Makine Öğrenmesi / Machine Learning
NB	: Naive Bayes
TF-IDF	: Terim Sıklığı – Ters Belge Sıklığı
YZ	: Yapay Zeka

1.GİRİŞ

Geçmişten günümüze teknolojinin gelişmesi ve bilgisayarların günlük yaşama entegre olmasıyla birlikte bireysel ve kurumsal faaliyetler daha hızlı ve sistematik bir biçimde yürütülmeye başlanmıştır. İş süreçlerinin hızlanması, daha kısa sürede daha fazla iş yapılabilmesini mümkün kılmış; bu durum bilgisayar sistemlerinin yalnızca işlem gücü açısından değil, aynı zamanda veri depolama kapasiteleri bakımından da sürekli geliştirilmesini gerekli hâle getirmiştir. Günümüzde temel problem yalnızca verilerin depolanması değil, bu verilerin ileride kullanılmak üzere anlamlı ve işlenebilir biçimde analiz edilmesidir.

Kurumsal yapılarda toplanan veriler; fatura kayıtları, satış bilgileri ve personel verileri gibi sayısal nitelikte olabileceği gibi, kullanıcı şikâyetleri, öneriler ve geri bildirimler gibi duygu içeren metinsel verilerden de oluşabilmektedir. Sayısal verilerin analizine yönelik geliştirilen yazılımlar sayesinde, analiz sonuçlarına dayalı olarak karar verme süreçlerini destekleyen sistemler ortaya çıkmış; yapay zekâ ve makine öğrenmesi tabanlı yaklaşımların bu sistemlere entegre edilmesiyle birlikte daha karmaşık problemlere yönelik çözümler geliştirilebilir hâle gelmiştir.

Teknolojik gelişmelere paralel olarak internetin yaygınlaşması, veri üretiminde önemli bir artışa neden olmuştur. Günümüzde internet, büyük hacimli metinsel verilerin üretildiği ve paylaşıldığı en önemli ortamların başında gelmektedir. İnternet tabanlı platformlarda kullanıcılar, deneyimlerini yorum ve değerlendirmeler aracılığıyla paylaşmaktadır. Bu içerikler hem hizmet sağlayıcılar hem de diğer kullanıcılar için önemli bir bilgi kaynağı oluşturmaktadır. Çevrim içi ortamlarda kullanıcı geri bildirimlerinin, bireylerin karar verme süreçlerinde belirleyici bir rol oynadığı ve güven algısının oluşumuna katkı sağladığı çeşitli çalışmalarda vurgulanmıştır (Islam & Sultana, 2018). Bu durum, kullanıcı yorumlarının sistematik biçimde analiz edilmesinin önemini ortaya koymaktadır.

İnternet tabanlı platformlarda üretilen metinsel verilerin miktarının hızla artması, bu verilerin manuel olarak analiz edilmesini pratik olmaktan çıkarmıştır. Özellikle çevrim içi film inceleme platformlarında kullanıcılar, izledikleri filmler hakkındaki duygu ve düşüncelerini özgür biçimde ifade etmektedir. Bu yorumlar yalnızca izleyici tercihlerini yansıtmakla kalmamakta, aynı zamanda içerik sağlayıcılar için değerli geri bildirimler sunmaktadır. Ancak yorum sayısındaki artış, otomatik analiz yöntemlerine olan ihtiyacı kaçınılmaz hâle getirmiştir.

Bu noktada, metin verilerinin otomatik olarak analiz edilmesini amaçlayan doğal dil işleme (NLP/DDİ) çalışmaları önem kazanmıştır. NLP, yapay zekâ ve makine öğrenmesi alanlarının

önemli bir alt dalı olup; metin sınıflandırma, özetleme, konu modelleme ve duygu analizi gibi birçok uygulama alanında kullanılmaktadır (Jurafsky & Martin, 2023). NLP alanının önemli uygulamalarından biri olan duygu analizi, metinlerde ifade edilen görüşlerin olumlu veya olumsuz olarak sınıflandırılmasını amaçlamaktadır. Literatürde opinion mining olarak da adlandırılan duygu analizi; kullanıcıların ürünler, hizmetler veya içerikler hakkındaki düşüncelerinin sistematik biçimde ortaya konulmasını sağlayarak karar destek sistemleri açısından önemli avantajlar sunmaktadır (Pang & Lee, 2008; Liu, 2022).

Doğal diller kullanılarak oluşturulan metinlerde yer alan duygu, görüş ve düşüncelerin olumlu veya olumsuz olarak belirlenmesi süreci duygu analizi olarak tanımlanmaktadır. Duygu analizinin gerçekleştirilebilmesi için incelenecek metinlerin belirli dilbilgisi kuralları çerçevesinde değerlendirilmesi gerekmektedir. Aksi takdirde metin içerisindeki anlam ve duygu vurgusunun doğru biçimde tespit edilmesi zorlaşabilmektedir. Duygu analizi işlemleri klasik olarak insan tarafından gerçekleştirilebildiği gibi, günümüzde çoğunlukla bilgisayar tabanlı algoritmalar ve makine öğrenmesi yöntemleri aracılığıyla uygulanmaktadır. Metinlerde yer alan duygu ifadelerinin doğru biçimde belirlenebilmesi için öncelikle duygu içeren kelime ve ifadelerin başarılı şekilde tespit edilmesi gerekmektedir. Bu amaçla önceden oluşturulmuş sözlük yapıları kullanılabilen veya öznelik seçme (feature selection) yöntemleriyle duygu belirleyici kelimeler tespit edilebilmektedir. Metin tabanlı duygu analizinde insan gücüne dayalı yöntemlerin yüksek maliyetli olması ve büyük veri kümeleri üzerinde uygulanmasının zorlaşması nedeniyle, bilgisayar tabanlı duygu analizi yaklaşımları üzerine çok sayıda araştırma gerçekleştirilmektedir. Bu bölümde sunulan duygu analizi tanımı ve kullanılan kavramsal çerçeve, literatürde yaygın olarak kabul edilen temel çalışmalara dayanmaktadır (Pang & Lee, 2008; Liu, 2022; Jurafsky & Martin, 2023).

Özellikle film yorumları üzerinden gerçekleştirilen duygu analizi çalışmaları, izleyici memnuniyetinin ölçülmesi ve içerik öneri sistemlerinin geliştirilmesi açısından önemli bir araştırma alanı hâline gelmiştir. Türkçe gibi eklemeli dil yapısına sahip dillerde duygu analizi çalışmaları ise dilin morfolojik yapısı, kelime türetme biçimleri ve anlamsal belirsizlikler nedeniyle çeşitli zorluklar içermektedir. Ayrıca Türkçe doğal dil işleme alanında sınırlı sayıda etiketli veri setinin bulunması, bu alandaki çalışmaların gelişimini kısıtlayan önemli faktörlerden biridir (Ofłazer & Saraçlar, 2018). Bu nedenle Türkçe metinler üzerinde gerçekleştirilen duygu analizi çalışmalarının farklı yöntemlerle ele alınması ve karşılaştırmalı analizlerle desteklenmesi gerekmektedir.

Literatürde Türkçe duygu analizi çalışmalarında Naive Bayes, Lojistik Regresyon ve Destek Vektör Makineleri gibi klasik makine öğrenmesi algoritmalarının yaygın olarak kullanıldığı görülmektedir. Islam ve Sultana (2018), farklı makine öğrenmesi algoritmalarını karşılaştırmalı olarak inceledikleri çalışmalarında, uygun ön işleme ve öznitelik çıkarımı yöntemleriyle birlikte kullanıldığında bu algoritmaların duygu sınıflandırma problemlerinde etkili sonuçlar üretebildiğini göstermiştir. Bununla birlikte, klasik yöntemlerin bağlam bilgisini sınırlı düzeyde temsil edebilmesi, özellikle karmaşık ve örtük duygu içeren metinlerde performans düşüşüne neden olabilmektedir (Tang & Zhang, 2018).

Son yıllarda derin öğrenme tabanlı yaklaşımlar, duygu analizi çalışmalarında daha yaygın olarak kullanılmaya başlanmıştır. Bu yaklaşımlar, metinlerin anlamsal yapısını daha etkin biçimde modelleyerek klasik makine öğrenmesi yöntemlerine kıyasla daha başarılı sonuçlar elde edilmesini mümkün kılmaktadır (Tang & Zhang, 2018). Ayrıca, kesin sınırlarla ifade edilmesi zor olan duygu durumlarının modellenmesinde etkili olan bulanık mantık tabanlı yaklaşımlar, duygu analizi alanında alternatif ve tamamlayıcı bir yöntem olarak değerlendirilmektedir.

Bu tez çalışmasında, Türkçe film yorumları kullanılarak duygu analizi gerçekleştirilmiş ve farklı makine öğrenmesi, derin öğrenme ve bulanık mantık tabanlı yöntemlerin performansları karşılaştırmalı olarak incelenmiştir. Çalışmanın temel amacı, Türkçe duygu analizi alanında kullanılan yöntemlerin güçlü ve zayıf yönlerini ortaya koymak ve film yorumları üzerinden geliştirilebilecek karar destek sistemlerine yönelik çıkarımlar sunmaktır.

2. LİTERATÜR TARAMASI

Literatürde duygu analizi alanında gerçekleştirilen çalışmalar incelendiğinde, farklı veri setleri ve çeşitli makine öğrenmesi, derin öğrenme ve hibrit yöntemlerin kullanıldığı görülmektedir. Özellikle Destek Vektör Makineleri (DVM), Naive Bayes, LSTM, BERT ve bulanık mantık tabanlı yaklaşımlar yaygın olarak tercih edilmektedir.

Pang ve arkadaşları (2002) duygu analizi yapmak için yaptıkları bu çalışmada IMDB veri setini kullanmışlardır. Araştırmacılar yöntem olarak farklı makine öğrenmesi sınıflandırıcıları ve n-gram yöntemini tercih etmişlerdir. Unigram DVM ile %82.9, Biagram ME ile %77.4 ve Unigram + Bigram DVM ile %82.7 oranında doğruluk değerleri elde etmişlerdir.

Matsumoto ve arkadaşları (2005) yaptıkları çalışmada IMDB ve Polarity veri setlerini kullanarak duygu analizi yapmışlardır. Araştırmacılar yaptıkları bu çalışmada DVM ve n-gram yöntemini tercih etmişlerdir. Unigram yönteminde %83.7, Bigram yönteminde %80.4 ve Unigram+Bigram yönteminde %84.6 oranında bir doğruluk değeri elde etmişlerdir.

Türkçe'deki duygu analizi konusunda öncü çalışmalardan biri Eroğul U. (2009) tarafından yapılmıştır. Tezinde, beyazperde.com sitesinden toplanan Türkçe metinler üzerine duygu analizi metotları uygulanmıştır. Bu çalışma sonucunda destek vektör makineleri (DVM) Algoritması ile %85 oranında doğruluk elde etmiştir.

Sevindi (2013) tezinde Türkçe film yorumlarını kullanmıştır. Hem sözlük tabanlı yaklaşım hem de makine öğrenmesi algoritmalarını uygulamıştır. Uygulamış olduğu makine öğrenimi algoritmaları arasında destek vektör makinesi algoritması 0.8258 F puan değeri ile en iyi sonuç veren algoritma olduğunu belirtmiştir. Ayrıca, yapılan çalışmada makinesi öğrenmesi yaklaşımının, sözlük tabanlı yaklaşıma kıyasla daha doğru sonuçlar verdiğini gözlemlemiştir.

Narayanan ve arkadaşları (2013) duygu analizi yapmak için yapmış oldukları çalışmada Naive Bayes algoritmasını kullandıklarını belirtmişlerdir. IMDB veri seti kullanılarak yapılan bu çalışmada Naive Bayes algoritmasının farklı yönlerini araştırdıklarını ve yüksek bir doğruluk değeri elde ettiklerini belirtmişlerdir. Yazarlar, önerilen yöntemin hız ve doğruluğu geliştirmek için bir dizi metin sınıflandırma problemine genelleştirileceğini ifade etmişlerdir. Araştırmacılar Naive Bayes sınıflandırıcısında %88.80 oranında bir doğruluk değeri elde etmişlerdir.

Akba (2014) tez çalışmasında, Türkçe film yorumları üzerinde duygu analizi gerçekleştirilmiş ve özellikle öznitelik seçme metriklerinin sınıflandırma başarımına etkisi incelenmiştir. Çalışmada metinler üzerinde TF-IDF, Information Gain ve Chi-Square gibi öznitelik seçme yöntemleri kullanılarak en ayırt edici terimler belirlenmiştir. Sınıflandırma aşamasında ise başlıca Destek Vektör Makineleri ve karşılaştırma amacıyla Naive Bayes algoritmaları kullanılmıştır. Deneysel sonuçlara göre, yalnızca olumlu ve olumsuz sınıflandırmada sistem %83,9 F1 skoru elde ederken, olumlu, olumsuz ve nötr olmak üzere üç sınıflı problemde başarı %63,3 F1 değerine düşmüştür. Elde edilen bulgular, öznitelik seçme yöntemlerinin duygu analizi performansını önemli ölçüde artırdığını göstermektedir.

Liu ve Chen (2015), mikroblog metinlerinde duygu analizini çok etiketli sınıflandırma yaklaşımıyla ele almış ve BR, CC, RAKEL, ECC, MLkNN ve RF-PCTs yöntemlerini karşılaştırmıştır. Çalışmada üç farklı duygu sözlüğü (DUTSD, NTUSD ve HowNet) kullanılmış ve sözlük seçiminin sınıflandırma performansı üzerinde belirleyici olduğu gösterilmiştir. Bulgular, özellikle RAKEL ve ECC yöntemlerinin daha başarılı sonuçlar verdiğini ve üç sözlük arasında DUTSD'nin en yüksek performansı sağladığını ortaya koymaktadır. Çalışma, çok etiketli sınıflandırma yaklaşımının mikroblog duygu analizinde etkili olduğunu göstermektedir.

Tang ve arkadaşları (2015), belge düzeyinde duygu sınıflandırması için kullanıcı ve ürün bilgilerini modele entegre eden User-Product Neural Network (UPNN) adlı bir derin öğrenme yaklaşımı önermiştir. Çalışmada IMDB ve Yelp veri setleri kullanılmış; karşılaştırma amacıyla Naive Bayes, SVM ve Logistic Regression algoritmaları da uygulanmıştır. Sonuçlar, geleneksel makine öğrenmesi yöntemlerinin belirli bir başarı sağladığını ancak derin öğrenme tabanlı modellerin, özellikle kullanıcı ve ürün temsillerini içeren UPNN'nin daha yüksek performans elde ettiğini göstermiştir.

Türkmenoğlu (2016) tezinde, hem sözlük hem de makine öğrenme algoritmalarını Türkçe metinlere uygulamıştır. Çalışmada kullanılan verileri Twitter ve film incelemeleri sitesi verileri olarak 2 bölüme ayrılmıştır. Twitter verilerinin daha kısa ve yapılandırılmamış veri kümesi olduğunu belirtmiştir. Makine öğrenmesi tarafında Naive Bayes, Support Vector Machines (SVM) ve k-Nearest Neighbor (kNN) algoritmaları kullanılmıştır. Sonuçlar, makine öğrenmesi yöntemlerinin, özellikle SVM'nin, sözlük tabanlı yaklaşıma göre daha yüksek sınıflandırma başarısı sağladığını göstermiştir. Bununla birlikte sözlük tabanlı yöntemlerin eğitim verisi gerektirmemesi ve uygulanabilirliğinin kolay olması önemli bir avantaj olarak değerlendirilmiştir.

Çetin ve Eryiğit (2018) çalışmalarında, Türkçe hedef tabanlı duygu analizi üzerine gerçekleştirmiştir. ABSA 2016 yarışmasında tanımlanan görevler esas alınarak Türkçe restoran yorumları veri kümesi üzerinde kapsamlı çalışmalar yürütülmüştür. Çalışmada hedef kategori belirleme, hedef terim belirleme ve her iki görevin birlikte çözümü için, kelime vektörleri ile sözcük ve cümle düzeyindeki doğal dil işleme çıktılarının kullanıldığı Koşullu Rastgele Alanlar (CRF) tabanlı bir dizilim etiketleme yöntemi önerilmiştir. Bu yaklaşım ile hedef kategori belirlemede %66,7, hedef terim belirlemede %53,2 ve hedef kategori ile hedef terimin aynı anda belirlenmesinde %46,7 F1-skoru elde edilerek literatürdeki en yüksek başarımlar raporlanmıştır. Duygu sınıfı belirleme aşamasında ise hedef terime yakın kelimelerden çıkarılan özelliklere dayalı lineer sınıflandırma yöntemi kullanılmış ve bu görev için %76,1 F1-skoru elde edilmiştir. Elde edilen sonuçlar, kelime vektörleri ve dizilim etiketleme yaklaşımlarının Türkçe hedef tabanlı duygu analizinde etkili olduğunu ortaya koymaktadır.

Huang ve arkadaşları (2018) yapmış oldukları bu çalışmada LSTM ve BiLSTM ağlarını kullanmışlar. Ayrıca LSTM ağlarının metin sınıflandırma sürecinde başarılı bir yöntem olduğunu ifade etmişlerdir. Araştırmacılar önermiş oldukları BiLSTM tabanlı model sayesinde farklı anlamlara gelen kelimelerin de tespit edilebilmesinin mümkün olacağını ifade etmiştir.

Islam ve Sultana (2018) tarafından gerçekleştirilen çalışmada, metin tabanlı duygu sınıflandırması problemi kapsamında iki farklı veri seti kullanılmıştır: Amazon ürün yorumları ve IMDb film yorumları. Çalışmada toplam altı farklı makine öğrenmesi algoritması karşılaştırılmıştır: Multinomial Naive Bayes, Bernoulli Naive Bayes, Logistic Regression, Stochastic Gradient Descent (SGD), Linear Support Vector Machine (Linear SVM) ve Random Forest (RF). Özellik çıkarım sürecinde klasik metin temsil yöntemleri kullanılmış ve modellerin performansı doğruluk oranı üzerinden değerlendirilmiştir. Elde edilen sonuçlara göre, her iki veri setinde de Linear SVM algoritması diğer yöntemlere kıyasla daha yüksek doğruluk performansı göstermiştir. Özellikle yüksek boyutlu ve seyrek (sparse) metin verilerinde marjin maksimize eden yapısı sayesinde SVM'nin daha kararlı ve genellenebilir sonuçlar ürettiği belirtilmiştir.

Rao ve arkadaşları (2018) belge düzeyinde duygu sınıflandırması üzerine gerçekleştirdikleri çalışmada LSTM tabanlı bir model önermiş, karşılaştırma amacıyla SVM ve Naive Bayes algoritmalarını da kullanmıştır. Çalışmalar IMDB Large Movie Review Dataset ve Yelp Review Dataset veri setleri üzerinde yürütülmüştür. Sonuçlar, geleneksel yöntemler arasında

SVM'nin her iki veri setinde de en başarılı algoritma olduğunu, Naive Bayes'in ise daha düşük performans gösterdiğini ortaya koymuştur.

Erşahin ve arkadaşları (2019) çalışmalarında, Türkçe duygu analizi için sözlük tabanlı ve makine öğrenmesi tabanlı yaklaşımları bütünleşik biçimde kullanan bir hibrit duygu analizi modeli önermiştir. Önerilen hibrit modelde, SentiTurkNet (STN) duygu sözlüğü Otomatik Eşanlamlı Sözlük (ASDICT) ile genişletilerek kapsama oranı artırılmış ve elde edilen genişletilmiş sözlük (eSTN) üzerinden metinlerdeki olumlu ve olumsuz kelimelere dayalı kutupsallık tabanlı yeni bir özellik üretilmiştir. Bu sözlük tabanlı özellik, geleneksel kelime tabanlı özniteliklere eklenerek Naive Bayes (NB), Destek Vektör Makineleri (DVM) ve J48 karar ağacı sınıflandırıcıları ile eğitilmiştir. Hibrit model; film yorumları (49.476 örnek), otel yorumları (11.164 örnek) ve Twitter verileri (1.756 tweet) olmak üzere üç farklı Türkçe veri seti üzerinde değerlendirilmiştir. Deneysel sonuçlar, hibrit yaklaşımın tüm veri setlerinde hem yalnızca makine öğrenmesi tabanlı hem de yalnızca sözlük tabanlı yöntemlere kıyasla anlamlı performans artışı sağladığını göstermiştir. En yüksek doğruluk değerleri film verisi için %86,31, otel verisi için %91,96 ve Twitter verisi için %83,37 olarak raporlanmış; genel olarak hibrit modelin ortalama %7 oranında doğruluk artışı sağladığı ortaya konmuştur. Bu sonuçlar, sözlük tabanlı anlamsal bilginin makine öğrenmesi modellerine entegre edilmesinin Türkçe duygu analizinde etkili bir strateji olduğunu göstermektedir.

Vashishtha ve Susan (2019) çalışmalarında, duygu analizinde insan bilişini taklit eden bulanık entropi tabanlı bir yaklaşım önerilmiştir. Modelde, SentiWordNet'ten elde edilen bulanık duygu skorları kullanılarak metin içerisindeki yüksek duygu taşıyan anahtar kelimeler otomatik olarak belirlenmiş, düşük bulanık entropi değerine sahip anlamlı kelimeler seçilerek LSTM tabanlı bir derin öğrenme modeli ile sınıflandırma gerçekleştirilmiştir. Deneysel değerlendirmeler sonucunda, önerilen yöntem IMDB film yorumları veri setinde %89,8, Pang-Lee duygu polarite veri setinde ise %78,5 doğruluk elde ederek literatürdeki birçok güncel yöntemi geride bırakmıştır. Elde edilen bulgular, bulanık mantık temelli kelime seçiminin derin öğrenme modelleriyle birlikte kullanıldığında duygu analizi performansını belirgin biçimde artırdığını göstermektedir.

Haque ve arkadaşları (2019) yaptıkları çalışmada IMDB veri setini kullanmışlardır. IMDB veri seti kullanıcıların filmler hakkında yapmış oldukları yorumlardan oluşmaktadır. Araştırmacılar bu veri setinden duygu çıkarımı yapmak için ESA ve Long Short-Term Memory (LSTM) ağlarını kullanmışlardır. Önerilen modelde %91 oranında bir F1 değeri elde edilmiştir. Yapılan

bu çalışmada araştırmacılar ESA ağlarının LSTM ağlarından daha iyi sonuç ürettiğini belirtmişlerdir.

Bedi ve Khurana (2019) tarafından gerçekleştirilen çalışmada, bulanık mantık ile derin öğrenmenin birlikte kullanıldığı (fuzzy–deep learning) bir duygu analizi yaklaşımı önerilmiştir. Deneysel sonuçlara göre, önerilen hibrit model %88.91 doğruluk değerlerine ulaşarak, yalnızca derin öğrenme tabanlı modellerin ve klasik makine öğrenmesi yaklaşımlarının üzerinde bir performans sergilemiştir. Çalışmada, bulanık mantık bileşeninin özellikle belirsiz ve yoruma açık ifadelerde sınıflandırma başarımını artırdığı, bu sayede hem doğruluk oranlarında hem de sınıflar arası dengeyi yansıtan ölçütlerde iyileşme sağlandığı rapor edilmiştir. Elde edilen sonuçlar, duygu analizinde fuzzy–derin öğrenme tabanlı hibrit yaklaşımların daha kararlı ve yüksek başarımlar sunduğunu göstermektedir.

Ahmetoğlu ve Daş (2020) çalışmalarında Türkçe otel yorumları üzerinde gerçekleştirilen bir çalışmada, Word2Vec yöntemiyle elde edilen farklı kelime vektörlerinin duygu analizi performansına etkisi incelenmiştir. Elde edilen kelime temsilleri, Geçitli Tekrarlayan Birimler (GRU) modeli kullanılarak sınıflandırılmış ve özel, genel ve rasgele oluşturulmuş kelime vektörleri karşılaştırılmıştır. Deneysel sonuçlar, geniş kapsamlı ve önceden eğitilmiş kelime vektörlerinin daha yüksek doğruluk ve daha düşük kayıp değerleri sağladığını göstermiştir. Rasgele kelime vektörleriyle oluşturulan modellerin ise anlamsal ilişki kuramaması nedeniyle daha düşük sınıflandırma başarımını sergilediği belirtilmiştir.

Açıkalın ve arkadaşları(2020) tarafından gerçekleştirilen çalışmada, Türkçe duygu analizinde BERT tabanlı modellerin performansı incelenmiştir. Film ve otel yorumlarından oluşan veri setleri üzerinde yapılan deneylerde Çok Dilli BERT, Temel BERT ve Büyük BERT modelleri karşılaştırılmıştır. Film verisi ile eğitilip film verisi üzerinde test edilen senaryoda Büyük BERT modeli %93,32 doğruluk oranı ile en başarılı sonucu elde etmiştir. Otel yorumları üzerinde gerçekleştirilen deneylerde fastText-TR modeli %70,55 doğruluk sağlarken, Çok Dilli BERT modeli %69,80 doğruluk elde etmiştir. Film verisi ile eğitilip otel verisi üzerinde test edilen çapraz alan deneylerinde ise fastText-EN modeli %80,80 doğruluk oranına ulaşmıştır. Otel verisi ile eğitilip film verisi üzerinde test edilen deneylerde ise Büyük BERT modeli %91,96 doğruluk ile en başarılı yöntem olmuştur. Elde edilen sonuçlar, transformer tabanlı BERT modellerinin Türkçe duygu analizi problemlerinde yüksek başarı sağladığını göstermektedir.

Vashishtha ve Susan (2020) alıřmalarında, denetimsiz duygu analizi iin szck kutuplařma skorlarının bulanık mantık ile yorumlanmasına dayalı bir yaklařım nerilmiřtir. Yntem, IMDB film yorumları, Pang-Lee Movie Reviews ve Hotel Reviews veri seti olmak zere  farklı veri seti zerinde deęerlendirilmiřtir. Deneysel sonular, veri setlerinin yapısal zelliklerine baęlı olarak sınıflandırma bařarımının deęiřtięini gstermiřtir. zellikle Hotel Reviews veri seti zerinde %76.2 doęruluk elde edilirken, Pang-Lee %65.45 ve IMBD veriseti %70.06 doęruluk oranları raporlanmış ve bulanık toplama yaklařımının tm veri setlerinde tekil denetimsiz yntemlere kıyasla daha kararlı sonular rettięi belirtilmiřtir.

elik ve arkadaşları (2021) alıřmalarında, duygu analizi kapsamında Support Vector Machine (SVM), K-Nearest Neighbor (kNN), Naive Bayes (NB), Decision Tree (DT) ve Random Forest (RF) olmak zere beř farklı sınıflandırma algoritmasını karřılařtırmıřtır. Veri boyutunu ve eřitlilięini artırmak amacıyla, her biri 500 olumlu ve 500 olumsuz yorumdan oluřan toplam 1000 rnek ieren  farklı veri seti birleřtirilerek daha kapsamlı bir veri yapısı oluřturulmuřtur. Elde edilen bulgular, birleřtirilmiş ve daha byk veri setleriyle eęitilen modellerin daha yksek performans sergiledięini gstermektedir. zellikle Support Vector Machine (SVM), yaklařık %85 doęruluk oranı ile dięer yntemlere kıyasla daha etkili bir sınıflandırma yaklařımı olarak rapor edilmiřtir.

Erkartal ve Yılmaz (2022) alıřmalarında sosyal medya verileri zerinde gerekleřtirmiřtir. Twitter zerinden toplanan ekonomik ierikli paylařımlar kullanılarak duygu analizi yapılmıř ve bulanık mantık destekli Destek Vektr Makineleri (ANFIS-SVM) ile LSTM modellerinin performansları karřılařtırılmıřtır. Aynı n iřleme adımlarının uygulandıęı deneysel alıřmada, ANFIS-SVM yaklařımının LSTM modeline kıyasla daha yksek doęruluk saęladıęı, her iki modelin de olumlu duygu sınıfında daha bařarılı sonular rettięi belirtilmiřtir. Olumsuz duygu sınıflarındaki performans dřřnn ise ironi ve alay gibi dilsel yapıların sınıflandırmayı zorlařtırmasından kaynaklanabileceęi ifade edilmiřtir.

Sarıgil (2023) tez alıřmasında, Twitter verileri zerinde duygu analizi ve kategori tahminleme amacıyla doęal dil iřleme ve makine ęrenmesi yntemleri kullanılmıřtır. Metinlerin sayısal temsili iin TF-IDF yaklařımından yararlanılmış; duygu sınıflandırması ve kategori analizi ařamasında Naive Bayes, Destek Vektr Makineleri (DVM) ve Lojistik Regresyon algoritmaları uygulanmıřtır. Elde edilen sınıflandırma ıktıları, iř zekası araları ile btnleřtirilerek duygu daęılımları, etkileřim ltleri ve konumsal bilgiler zerinden analitik

ve görsel çıkarımlar yapılmıştır. Çalışma, NLP tabanlı sınıflandırma algoritmalarının sosyal medya verilerinin iş zekası süreçlerinde etkin biçimde kullanılabileceğini göstermektedir.

Sharma ve arkadaşları (2023) çalışmalarında, alaycılık tespitini daha yüksek doğrulukla gerçekleştirmek amacıyla kelime ve kelime öbeği temsillerini bir araya getiren hibrit bir topluluk modeli önerilmiştir. Modelde Word2Vec, GloVe ve BERT tabanlı gömme yöntemleri kullanılmış, her bir yaklaşım yoğun katmanlar aracılığıyla öğrenilerek sınıflandırma olasılıkları üretilmiştir; elde edilen bu olasılıklar ise bulanık mantık temelli evrimsel bir karar mekanizmasıyla birleştirilmiştir. Önerilen yaklaşım, kamuya açık SARC, Headlines ve Twitter veri kümeleri üzerinde değerlendirilmiş ve sırasıyla %85,38, %90,81 ve %86,80 doğruluk oranları elde edilmiştir.

Çelik (2023) tez çalışmasında, hakaret içerikli cümleler, saldırgan tavrı cümleler ve herhangi bir olumsuz tavrı olmayan cümleleri veri setine göre analiz edilip, hakaret içerikli cümlelerin ortaya çıkarılarak duygu analizi uygulanması sağlanmıştır. Çalışma kapsamında beş farklı sınıflandırma algoritmaları kullanılmıştır. RF, LR, DT, NB, KNN algoritmaları kullanıldı. Etiketlenmiş veri seti üzerinden çıkan doğruluk (accuracy) oranlarına göre Logistic Regression modeli uygulanmıştır. Kullanılan algoritmalar doğrultusunda veri setindeki belirtilen cümleleri tahmin ederek (predict) doğru sonuca ulaşması kapsamında LR modeli çalışmada uygun görüldüğü analiz sonucu ortaya çıkmıştır. Bu kapsamda Logistic Regression Accuracy değeri yaklaşık olarak 0.89 çıkarak, başarılı sonuç elde edilmiştir.

Gündüz (2023) çalışmasında, Türkçe film yorumları üzerinde duygu analizi amacıyla beyazperde.com'dan derlenen veriler kullanılmış ve ön eğitilmiş BERTurk modeli temel alınmıştır. Çalışmada üç farklı deneysel yaklaşım değerlendirilmiştir. İlk olarak, BERTurk modelinin sondan bir önceki dönüştürücü katmanından elde edilen bağlamsal temsiller Destek Vektör Makineleri (DVM/SVM) ile sınıflandırılmıştır. İkinci aşamada, BERTurk modeli doğrudan ince ayar (fine-tuning) yöntemiyle görev özelinde eğitilmiştir. Son aşamada ise ince ayarlı BERTurk modelinden elde edilen temsiller tekrar DVM ile sınıflandırılmıştır. Deneysel sonuçlara göre en yüksek doğruluk oranı %98,4 ile ince ayarlı BERTurk modelinde elde edilmiştir. İnce ayar işleminin sınıflandırma performansını yaklaşık %10 oranında artırdığı, BERTurk temsillerinin DVM ile birlikte kullanımının ise yaklaşık %5 ek iyileşme sağladığı rapor edilmiştir. Bu bulgular, Türkçe duygu analizinde BERT tabanlı derin bağlamsal temsillerin ve hibrit yaklaşımların yüksek performans sunduğunu göstermektedir.

Pekel (2024) tez çalışmasında, belirli ürün gruplarına yönelik e-ticaret sitelerinde yapılan yorumlar üzerinde derin öğrenme algoritmaları olan, Uzun kısa vadeli bellek modeli (LSTM), Evrişimsel sinir ağları algoritması (CNN), Geçitli tekrarlayan birim (GRU), Transformatörlerin çift yönlü kodlayıcı gösterimleri (BERT) algoritmasıyla oluşturulan tek algoritmali ve hibrit modeller ile duygu analizi gerçekleştirilmiştir. Modelin eğitim sonuçlarına göre, her bir modelin doğruluk değeri başta olmak üzere, bütün modellerin sonuçları birbirlerine çok yakın bulunmuştur. En düşük doğruluk sonucu %90 ile LSTM ve CNN algoritmaları ile oluşturulmuş iki ayrı modelin verdiği görülmüştür. En yüksek doğruluk sonucu %94 ile GRUCNN hibrit modelinin sonucu olarak bulunmuştur.

Dündar ve Koçer (2025) çalışmalarında, Türkçe film yorumları üzerinde duygu analizi amacıyla Beyazperde kaynaklı hazır bir veri seti kullanılarak doğal dil işleme kapsamında analiz gerçekleştirilmiştir. Çalışmada Logistic Regression algoritması uygulanmış ve yaklaşık %89 doğruluk oranı elde edilmiştir. Ayrıca çalışmada, veri toplama sürecine yönelik örnek uygulamalar kapsamında web crawler yöntemine de yer verilmiştir. Bu sonuçlar, Lojistik regresyonun Türkçe metinler üzerinde etkili bir yöntem olduğunu göstermektedir.

Yapılan çalışmaların yöntemleri, kullanılan veri setleri ve elde edilen başarı oranları *Tablo 2.1*'de özetlenmiştir.

Yazar / Yıl	Veri Seti	Kullanılan Yöntem / Algoritma	Başarı Sonucu
Pang vd. (2002)	IMDB	DVM, ME, n-gram	Unigram DVM %82.9, Bigram ME %77.4, Unigram+Bigram DVM %82.7
Matsumoto vd. (2005)	IMDB, Polarity	DVM, n-gram	Unigram %83.7, Bigram %80.4, Unigram+Bigram %84.6
Uğur Eroğul (2009)	Beyazperde Türkçe yorumlar	DVM	%85 doğruluk

Sevindi (2013)	Türkçe film yorumları	Sözlük tabanlı, DVM	En başarılı DVM F-skoru: 0.8258
Narayanan vd. (2013)	IMDB	Naive Bayes	%88.80 doğruluk
Akba (2014)	Türkçe film yorumları	TF-IDF, IG, Chi-Square, DVM, NB	İkili sınıf F1: %83.9, üçlü sınıf F1: %63.3
Liu ve Chen (2015)	Mikroblog verileri	RAkEL, ECC, MLkNN	En başarılı yöntemler: RAkEL ve ECC
Tang vd. (2015)	IMDB, Yelp	UPNN, NB, SVM, LR	UPNN en başarılı model
Türkmenoğlu (2016)	Twitter, film yorumları	NB, SVM, kNN	SVM en başarılı yöntem
Çetin ve Eryiğit (2018)	Türkçe restoran yorumları	CRF tabanlı ABSA	Hedef kategori %66.7 F1, duygu sınıfı %76.1
Huang vd. (2018)	Metin verileri	LSTM, BiLSTM	BiLSTM başarılı sonuçlar vermiştir
Islam ve Sultana (2018)	Amazon, IMDB	NB, LR, SGD, Linear SVM, RF	Linear SVM en yüksek doğruluk
Rao vd. (2018)	IMDB, Yelp	LSTM, SVM, NB	SVM en başarılı klasik yöntem
Erşahin vd. (2019)	Film, otel, Twitter	Hibrit model (eSTN + DVM/NB/J48)	Film %86.31, Otel %91.96, Twitter %83.37
Vashishtha ve Susan (2019)	IMDB, Pang-Lee	Bulanık mantık + LSTM	IMDB %89.8, Pang-Lee %78.5
Haque vd. (2019)	IMDB	ESA, LSTM	F1: %91
Bedi ve Khurana (2019)	Duygu veri setleri	Bulanık mantık + derin öğrenme	%88.91 doğruluk
Ahmetoğlu ve Daş (2020)	Türkçe otel yorumları	Word2Vec + GRU	Ön eğitilmiş vektörler daha başarılı

Açıkalın vd. (2020)	Film, otel yorumları	Çok Dilli BERT, Temel BERT, Büyük BERT	En yüksek başarı Büyük BERT %93.32
Vashishtha ve Susan (2020)	IMDB, Hotel Reviews	Bulanık mantık tabanlı denetimsiz analiz	Hotel %76.2, IMDB %70.06
Çelik vd. (2021)	Birleştirilmiş yorum veri setleri	SVM, kNN, NB, DT, RF	En yüksek başarı SVM %85
Erkartal ve Yılmaz (2022)	Twitter ekonomi verileri	ANFIS-SVM, LSTM	ANFIS-SVM daha başarılı
Sarıgil (2023)	Twitter	TF-IDF, NB, DVM, LR	NLP tabanlı modeller başarılı sonuç vermiştir
Sharma vd. (2023)	SARC, Headlines, Twitter	Word2Vec, GloVe, BERT + bulanık mantık	En yüksek başarı GloVe %90.81
Çelik (2023)	Hakaret içerikli veri seti	RF, LR, DT, NB, KNN	En yüksek başarı LR %89
Gündüz (2023)	Türkçe film yorumları	BERTurk + SVM	%98.4 doğruluk
Pekel (2024)	E-ticaret yorumları	LSTM, CNN, GRU, BERT, hibrit modeller	En yüksek başarı GRUCNN %94
Dündar ve Koçer (2025)	Türkçe film yorumları	Logistic Regression	%89 doğruluk

Tablo 2.1. Literatürde duygu analizi üzerine yapılan çalışmaların karşılaştırılması

Bu çalışmanın literatürdeki diğer araştırmalardan ayrılan temel farkları ve özgün değerleri şu şekildedir:

- **Metodolojik Çeşitlilik ve Üçlü Hibrit Mimari:** Literatürdeki çalışmalar genellikle ikili karşılaştırmalar (örn: SVM vs. BERT) veya basit topluluk (ensemble) oylamaları ile sınırlıdır. Bu tez çalışmasında ise ilk kez; istatistiksel yaklaşım (TF-IDF + Lojistik Regresyon), bağlamsal derin öğrenme (BERTurk) ve kural tabanlı sistem (Mamdani

Tipi Bulanık Mantık) üçlü bir hibrit kararlaştırma mekanizması altında bir araya getirilmiştir.

- **Bulanık Mantığın Karar Düzeltici Olarak Konumlandırılması:** Mevcut literatürde bulanık mantık genellikle tek başına bir sınıflandırıcı olarak test edilmiştir. Bu çalışmada ise bulanık mantık, tek başına kullanılmak yerine, BERT ve Lojistik Regresyon modellerinden gelen olasılık skorlarını girdi kabul eden bir "üst karar verici / dengeleyici" katmanı olarak tasarlanmıştır. Bu sayede, derin öğrenme modellerinin ironik veya çelişkili ifadelerde düştüğü kararsızlıklar, bulanık kurallar yardımıyla esnek bir karar sınırına oturtulmuştur.
- **Grid Search ile Optimize Edilmiş Ağırlıklandırma:** Literatürde hibrit modeller kurgulanırken ağırlıklar genellikle sezgisel veya eşit oranlı (örn: %50-%50) dağıtılmaktadır. Bu çalışmada ise jüri heyetinin yönlendirmeleri doğrultusunda, modellerin hata payını en aza indiren optimum ağırlıklar (%50 BERT, %30 LR, %20 Bulanık Mantık) Grid Search optimizasyonu ile bilimsel olarak belirlenmiş ve doğrulanmıştır.
- **Veri Temsiliyet Kalitesi (Tabakalı Örnekleme):** Büyük veri kümeleriyle çalışılan birçok araştırmada rastgele örneklem seçimi yapılarak sınıf dengesi göz ardı edilmektedir. Bu çalışmada ise 80.000 satırlık ana kütleden 40.000 satıra düşülürken **Tabakalı Örnekleme (Stratified Sampling)** kullanılmış, böylece literatürdeki "veri azaltırken dağılımın bozulması ve modelin yanlılık (bias) geliştirmesi" sorununun önüne geçilmiştir.

Sonuç olarak bu tez; sadece modelleri birbiriyle yarıştırmakla kalmamış, istatistiksel, bağlamsal ve kurallı yaklaşımların güçlü yönlerini birleştirerek Türkçe duygu analizinde kararlılığı ve doğruluğu artıran yeni ve bütüncül bir metodolojik çerçeve sunmuştur.

3.MATERYAL VE YÖNTEMLER

3.1. Yapay Zeka

Yapay zeka (YZ); öğrenme, muhakeme, problem çözme, algılama ve doğal dil işleme gibi insan bilişinin temel unsurlarını hesaplama sistemleri aracılığıyla simüle edilmesini hedefleyen çok disiplinli bir bilim dalıdır. Alanın teorik temelleri, 1956 yılında düzenlenen ve disiplinin akademik bir meşruiyet kazanmasında milat kabul edilen Dartmouth Konferansı'nda John McCarthy tarafından atılmıştır (McCarthy vd., 2006). Ortaya çıkışından itibaren bilgisayar bilimi, bilişsel psikoloji, dilbilim ve mühendislik disiplinlerinin kesişim kümesinde gelişim gösteren yapay zeka, günümüz teknolojik paradigmasının merkezinde yer alan temel itici güçlerden biri haline gelmiştir.

Geleneksel deterministik (kural tabanlı) sistemlerin aksine, modern yapay zeka mimarileri, verilerden otonom çıkarım ve dinamik karar verme mekanizmaları etrafında tasarlanmıştır. Sabit algoritmalara dayanmak yerine, büyük veri kümeleri içindeki kalıpları ve korelasyonları belirleyerek çevresel değişikliklere uyum sağlayan bir yapı sergilerler. Bu uyum sağlama kapasitesi, temel olarak aşağıdaki alt alanların sinerjisiyle elde edilir:

- **Makine Öğrenimi (Machine Learning):** İstatistiksel yöntemlerle veriden öğrenen algoritmalar.
- **Derin Öğrenme (Deep Learning):** Çok katmanlı yapay sinir ağları aracılığıyla karmaşık özellik çıkarımı.
- **Doğal Dil İşleme (NLP):** İnsan dilinin bilgisayarlar tarafından çözümlenmesi ve üretilmesi.

Günümüzde yapay zeka teknolojileri; sağlık, finans, eğitim, savunma ve endüstriyel üretim gibi stratejik sektörlerde operasyonel verimliliği artırmak amacıyla yaygın olarak entegre edilmektedir. Özellikle dilsel görevlerde sergilediği performans; metin sınıflandırma, tavsiye sistemleri, nöral makine çevirisi ve konuşma tanıma gibi alanlarda devrim niteliğinde ilerlemeler kaydedilmesini sağlamıştır. Bu durum, yapay zekanın yapılandırılmamış veri kümelerini işleme ve anlamlandırma konusundaki sofistike yetkinliğini açıkça ortaya koymaktadır.

3.2 Makine Öğrenmesi

Makine öğrenimi (ML), yapay zekâ alanının bir parçasıdır ve hesaplama sistemlerinin, her eylem için doğrudan programlama gerektirmeden, geçmiş verilerden bilgi edinerek belirlenen

görevleri yerine getirmesini sağlar (Mitchell, 1997). Sabit talimatlara dayanan geleneksel kural tabanlı çerçevelerin aksine, ML algoritmaları deneyimsel öğrenme yoluyla yeteneklerini geliştirir, verilerdeki temel kalıpları ve korelasyonları tanır ve bu bilgileri yeni, görülmemiş durumlara uygular (Bishop, 2006). Makine öğrenimi algoritmalarının temel prensibi, girdi ve çıktı değişkenleri arasındaki ilişkilerin matematiksel temsillerini oluşturmaktır (Hastie, Tibshirani & Friedman, 2009). Öğrenme süreci boyunca, bir model, verilerde belirlediği düzenlilikleri yansıtacak şekilde iç parametrelerini ince ayarlayarak, eğitim kümesi olarak adlandırılan bir veri kümesi ile etkileşime girer. Modelin genelleme yeteneği, daha önce karşılaşmadığı verilerden oluşan ayrı bir test kümesi kullanılarak değerlendirilir (Bishop, 2006). Bu metodoloji, katı programlama paradigmalarından, veri maruziyeti yoluyla dinamik olarak gelişen sistemlere doğru bir geçişi işaret eder (Russell & Norvig, 2021).

Makine öğrenimi yaklaşımları genellikle üç ana kategoriye ayrılır: denetimli öğrenme, denetimsiz öğrenme ve pekiştirmeli öğrenme (Russell & Norvig, 2021). Denetimli öğrenmede, modeller, her girdinin önceden tanımlanmış bir çıktıya karşılık geldiği etiketli veri kümeleri kullanılarak geliştirilir ve bu süreçte Lojistik Regresyon, Rastgele Orman ve Destek Vektör Makineleri (SVM) gibi algoritmalar sıklıkla kullanılır (Cortes & Vapnik, 1995). Aksine, denetimsiz öğrenme, Principal Component Analysis (PCA) ve t-distributed Stochastic Neighbor Embedding (t-SNE) gibi teknikleri kullanarak, kümeleme ve boyut indirgeme gibi yöntemlerle gizli yapıları veya kümeleri keşfetmek için etiketlenmemiş verilerle çalışır (van der Maaten & Hinton, 2008). Buna karşılık, pekiştirmeli öğrenme, çevreleriyle etkileşime girerek ve genellikle ödül veya ceza şeklinde geri bildirimlere yanıt olarak davranışlarını değiştirerek yinelemeli olarak uyum sağlayan ajanları içerir. Bu strateji, otonom robotik ve stratejik oyun geliştirme gibi alanlarda etkili olduğu kanıtlanmıştır (Mnih et al., 2015).

3.3 Derin Öğrenme

Derin öğrenme, çok katmanlı yapay sinir ağları aracılığıyla veriden hiyerarşik özellik temsilleri öğrenmeyi amaçlayan bir makine öğrenmesi yaklaşımıdır (LeCun, Bengio & Hinton, 2015). Geleneksel makine öğrenmesi yöntemlerinden farklı olarak, özellik çıkarımı sürecini büyük ölçüde otomatikleştirerek ham veriden doğrudan öğrenme gerçekleştirmektedir.

Bu yaklaşımın temelinde, doğrusal olmayan dönüşümler içeren çok katmanlı ağ yapıları yer almaktadır. Derin sinir ağları, her bir katmanda daha soyut ve anlamlı temsiller oluşturarak karmaşık veri örüntülerini modelleyebilmektedir. Bu sayede özellikle görüntü işleme, konuşma

tanıma ve doğal dil işleme gibi yüksek boyutlu ve karmaşık veri içeren alanlarda yüksek performans elde edilmektedir (Goodfellow, Bengio & Courville, 2016).

Derin öğrenme modellerinin başarısı; büyük veri setlerinin mevcut olması, grafik işlem birimlerinin (GPU) sağladığı hesaplama gücü ve gelişmiş optimizasyon algoritmaları ile doğrudan ilişkilidir. Özellikle evrimsel sinir ağları (CNN), tekrarlayan sinir ağları (RNN) ve Transformer tabanlı mimariler, farklı veri türlerine uygun çözümler sunmaktadır.

Son yıllarda geliştirilen Transformer tabanlı modeller, dikkat mekanizması (attention) sayesinde bağlamsal ilişkileri daha etkin şekilde modelleyebilmekte ve özellikle doğal dil işleme görevlerinde önemli başarılar sağlamaktadır (Vaswani et al., 2017). Bu yönüyle derin öğrenme, günümüzde yapay zekâ uygulamalarının temel bileşenlerinden biri haline gelmiştir.

3.4 Doğal Dil İşleme

Doğal Dil İşleme (NLP/DDİ), insanların günlük yaşamda kullandıkları doğal dillerin bilgisayarlar tarafından işlenmesini, anlaşılmasını ve yorumlanmasını amaçlayan disiplinlerarası bir araştırma alanıdır. Bu alan; dilbilim, bilgisayar bilimi ve yapay zekâ yöntemlerini bir araya getirerek metin veya konuşma biçimindeki verilerden anlamlı bilgi elde etmeyi hedeflemektedir. Doğal dil işleme çalışmaları kapsamında; yazılı veya sözlü metinlerin analiz edilmesi, metinlerin sınıflandırılması, özetlenmesi, duygu analizinin gerçekleştirilmesi ve soru-cevap sistemlerinin geliştirilmesi gibi çeşitli uygulamalar yer almaktadır.

Doğal dillerin bilgisayar ortamında işlenebilmesi için dilin temel bileşenlerinin modellenmesi gerekmektedir. Bu bileşenler; sesbilim (phonology), biçimbilim (morphology), sözdizimi (syntax) ve anlambilim (semantics) olarak dört ana başlık altında incelenmektedir. Özellikle Türkçe gibi eklemeli (bitişken) dillerde, biçimbilimsel yapı oldukça karmaşık olduğundan sözcük kökü–ek ayrımı, lemmatizasyon ve olumsuzluk yapılarının doğru biçimde ele alınması büyük önem taşımaktadır (Adalı,2012).

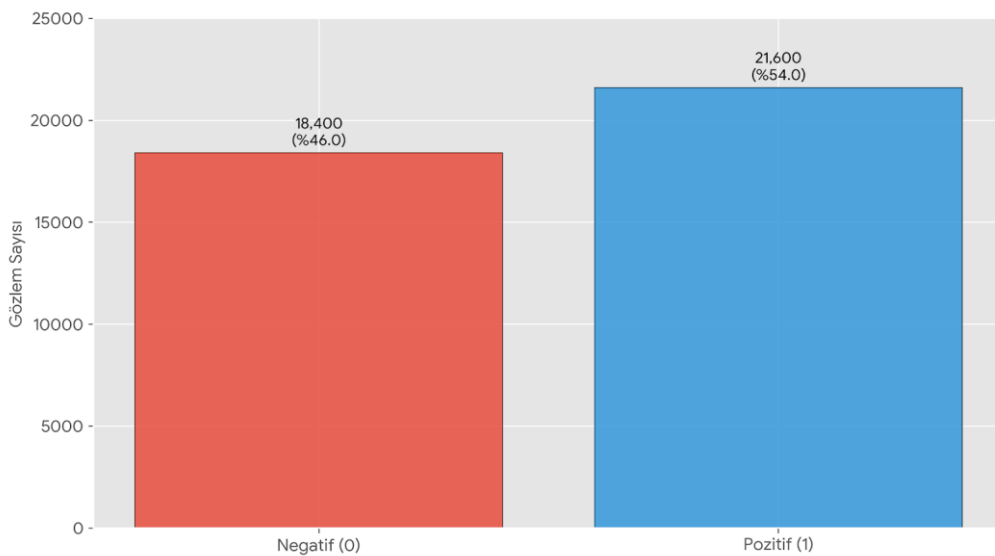
Modern doğal dil işleme yaklaşımları, klasik kural tabanlı yöntemlerin yanı sıra istatistiksel ve makine öğrenmesi temelli yöntemlerden de yararlanmaktadır. Bu kapsamda metinler, öncelikle çeşitli ön işleme adımlarından geçirilmekte; tokenizasyon, normalizasyon, durak kelime (stop-word) temizleme ve kök bulma (stemming/lemmatization) işlemleri uygulanmaktadır. Elde edilen temsil biçimleri daha sonra makine öğrenmesi veya derin öğrenme tabanlı modellerde kullanılmaktadır. Güncel çalışmalarda, bağlamsal kelime temsillerine dayalı derin öğrenme modellerinin, geleneksel yöntemlere kıyasla daha başarılı sonuçlar ürettiği raporlanmaktadır (Zadeh vd., 2017).

3.5 Veriseti

Bu çalışmada kullanılan veri seti, Mustafa Keskin tarafından Kaggle platformunda paylaşılan *Turkish Movie Sentiment Analysis Dataset* veri kümesidir (Keskin, 2019) . Veri seti, çevrim içi film inceleme platformlarından elde edilen toplam 83.238 adet kullanıcı yorumundan oluşmaktadır. Söz konusu yorumlar, 7.722 farklı filme ait olup, farklı tür ve dönemlere ait film değerlendirmelerini kapsayacak şekilde geniş bir içerik çeşitliliği sunmaktadır. Bu durum, veri setinin temsiliyet gücünü artırmakta ve duygu analizi çalışmalarında daha genellenebilir sonuçlar elde edilmesine olanak sağlamaktadır.

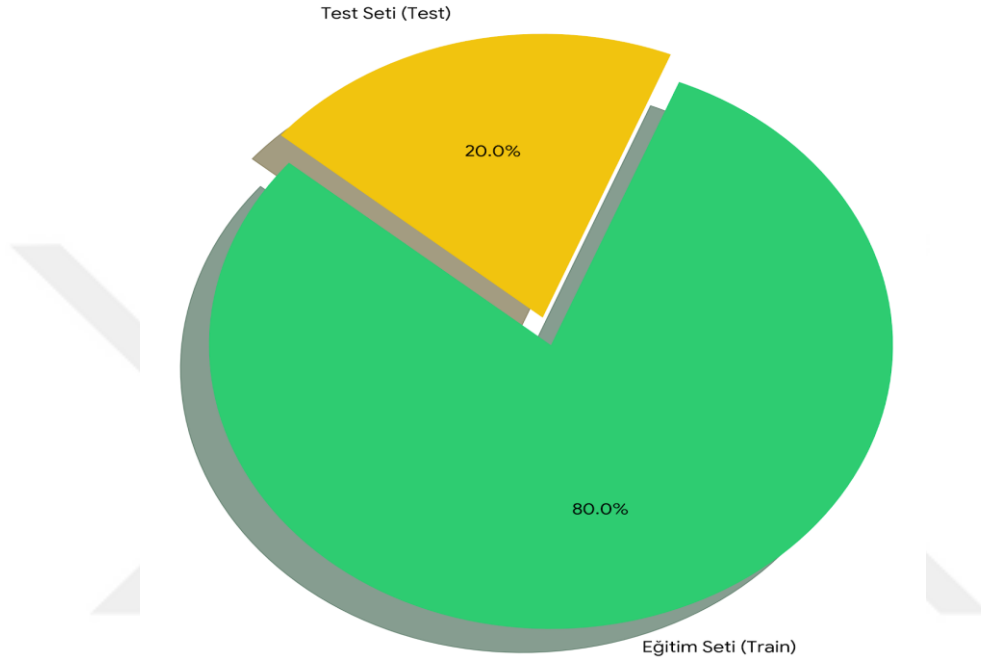
Bu kapsamda, orijinal veri setinden Tabakalı Örneklem (Stratified Sampling) yöntemi kullanılarak 40.000 satırlık bir alt küme oluşturulmuştur. Bu yöntemin tercih edilmesindeki temel amaç, veri miktarı azaltılırken pozitif ve negatif sınıfların ana veri setindeki dağılım oranlarının korunmasıdır. Böylece veri setinin istatistiksel temsil gücü korunmuş ve modelin belirli bir sınıfa yönelik yanlılık geliştirmesi önlenmiştir.

Veri setinde yer alan kullanıcı yorumları, 1 ile 5 arasında verilen puanlamalara göre duygu sınıflarına dönüştürülmüştür. Bu doğrultuda 1 ve 2 puanlı yorumlar “Negatif” (0), 4 ve 5 puanlı yorumlar ise “Pozitif” (1) olarak etiketlenmiştir. Duygu yönelimi açısından belirsizlik oluşturabilecek 3 puanlı yorumlar ise veri setinden çıkarılmıştır. Böylece ikili sınıflandırma problemine daha uygun ve tutarlı bir veri yapısı elde edilmiştir. Örneklem işlemi sonrasında oluşturulan 40.000 gözlemden oluşan veri setinin %46’sı negatif, %54’ü ise pozitif yorumlardan oluşmaktadır. Veri setine ait sınıf dağılımı *Şekil 3.1*’de gösterilmektedir.



Şekil 3.1 Örneklem Üzerinden duygu sınıf dağılımı

Model performansının objektif biçimde değerlendirilebilmesi amacıyla veri seti eğitim (train) ve test olmak üzere iki bölüme ayrılmıştır. Veri setinin %80'i eğitim verisi, %20'si ise test verisi olarak kullanılmıştır. Bu doğrultuda eğitim seti 32.000 örnekten, test seti ise 8.000 örnekten oluşmaktadır. Bölme işlemi sırasında sınıf dağılımlarının korunmasına dikkat edilmiş ve her iki veri kümesinde de benzer oranların devam etmesi sağlanmıştır. Eğitim ve test veri kümelerinin oranları Şekil 3.2'de gösterilmektedir.



Şekil 3.2 Veri seti bölümlenme oranları

3.6. Veri Ön işleme

Bu çalışmada veri ön işleme süreci, kullanılan model türlerinin yapısal gereksinimlerine göre özelleştirilmiş; ancak tüm aşamalar metodolojik tutarlılığı sağlamak amacıyla bütüncül bir yaklaşımla ele alınmıştır.

Çalışmaların hesaplama maliyetini optimize etmek ve özellikle derin öğrenme tabanlı modellerin eğitim sürelerini jüri heyetinin önerileri doğrultusunda yönetilebilir kılmak amacıyla, 80.000 satırlık ana veri kümesinden Tabakalı Örneklem (Stratified Sampling) yöntemi kullanılarak 40.000 örneklik bir alt küme oluşturulmuştur. Bu yöntem sayesinde veri boyutu yarıya indirilirken, orijinal veri setindeki olumlu ve olumsuz duygu sınıflarının dağılım oranları titizlikle korunmuştur.

Genel ön işleme aşamasında; yinelenen kayıtlar kaldırılmış, eksik veriler temizlenmiş ve tüm metinler küçük harfe dönüştürülmüştür. URL ifadeleri, HTML etiketleri ve özel karakterler

regex tabanlı işlemlerle ayıklanmıştır. Metin temizleme sürecinde "çokkk" gibi tekrarlanan karakterler "çokk" formuna normalize edilerek gürültü azaltılmıştır.

Klasik Makine Öğrenmesi Modelleri İçin Hazırlık: Lojistik Regresyon, Naive Bayes ve SVM modelleri için metinler TF-IDF yöntemiyle sayısal forma dönüştürülmüştür. Vektörleştirme aşamasında unigram ve bigram özellikleri birlikte kullanılmış, maksimum özellik sayısı 5.000 ile sınırlandırılmıştır. Veri kümesi %80 eğitim ve %20 test oranında ayrılarak modeller bu temsiller üzerinden eğitilmiştir.

BERT Modeli İçin Hazırlık: BERT modeli için bağlamsal bilgi kaybını önlemek amacıyla agresif temizleme işlemlerinden kaçınılmış, yalnızca temel karakter temizliği yapılmıştır. Metinler HuggingFace kütüphanesi üzerinden dbmdz/bert-base-turkish-cased tokenizer'ı ile alt kelime (subword) birimlerine ayrılmıştır. Maksimum token uzunluğu 128 olarak belirlenmiş, AdamW optimizasyon algoritması ve çapraz entropi kaybı fonksiyonu ile ince ayar (fine-tuning) süreci tamamlanmıştır.

Bulanık Mantık ve Hibrit Yapılandırma: Bulanık mantık sistemi için modellerden elde edilen olasılık skorları (0–1 aralığında) ve metin tabanlı duygu yoğunlukları giriş değişkenleri olarak tanımlanmıştır. Üyelik fonksiyonları üçgensel yapıda tanımlanmış; düşük, orta ve yüksek üyelik dereceleri oluşturulmuştur. Çıkarım mekanizması Mamdani tipi max–min kompozisyonu ile gerçekleştirilmiş ve durulaştırma (defuzzification) aşamasında merkez (centroid) yöntemi kullanılmıştır.

Performans değerlendirmesinde doğruluk (accuracy), F1-skor ve karmaşıklık matrisi (confusion matrix) metrikleri kullanılarak metodolojik tutarlılık sağlanmıştır.

3.7 Kullanılan Algoritmalar

3.7.1 Hiperparametre ve Konfigürasyon Ayarları

Modellerin tekrarlanabilirliğini sağlamak amacıyla kullanılan temel parametreler *Tablo 3.1*'de gösterilmiştir.

Model	Parametre	Ayar/Değer
Lojistik Regresyon	C (Düzenleştirme), Solver	1.0, 'lbfgs'
Naive Bayes	Alpha (Smoothing)	1.0 (Laplace)
DVM	Kernel, C	Linear, 1.0
BERT (Turkish)	Model Versiyonu, LR, Batch	dbmdz/bert-base-turkish-cased, 2e-5, 32

Bulanık Mantık	Çıkarım Tipi, Durulaştırma	Mamdani, Centroid (Merkez)
----------------	----------------------------	----------------------------

Tablo 3.1. Modellere ait hiperparametre ve konfigürasyon bilgileri

3.7.2 TF - IDF (Term Frequency – Inverse Document Frequency / Terim Sıklığı – Ters Belge Sıklığı)

TF-IDF (Term Frequency – Inverse Document Frequency), metin madenciliği ve doğal dil işleme çalışmalarında kelimelerin bir belge içerisindeki önem derecesini ve ayırt ediciliğini sayısal olarak ifade etmek amacıyla kullanılan istatistiksel bir ağırlıklandırma yöntemidir. Bu yöntem, bir terimin belirli bir belge içerisindeki görülme sıklığını (yerel önem) ve aynı terimin tüm belge koleksiyonu içerisindeki yaygınlığını (küresel nadirlik) birlikte dikkate almaktadır. TF-IDF yaklaşımı, bilgi erişim sistemleri bağlamında ilk kez Salton ve Buckley tarafından sistematik biçimde tanımlanmıştır (Salton & Buckley, 1988).

TF-IDF'in temel varsayımı, bir kelimenin bir belgede sık geçmesinin o belge için ayırt edici bir özellik olduğu; ancak kelimenin tüm belgelerde yaygın biçimde bulunmasının (örneğin bağlaçlar ve durak kelimeler) ayırt edicilik değerini azalttığı yönündedir (Manning, Raghavan & Schütze, 2008). Bu nedenle TF-IDF, sık geçen ancak bilgi değeri düşük terimlerin etkisini azaltırken, belgeye özgü terimlerin ağırlığını artırmayı amaçlamaktadır.

3.7.2.1 TF-IDF Çalışma Yöntemi ve Matematiksel Altyapı

TF-IDF ağırlığı, iki temel bileşenin çarpımı ile hesaplanmaktadır:

- **Terim Frekansı (TF):**

Terim frekansı, t teriminin d belgesi içerisindeki göreceli baskınlığını ölçmektedir. Ham geçme sayısı yerine genellikle belge uzunluğuna göre normalize edilmiş değerler kullanılmaktadır:

$$TF(t, d) = \frac{f_{t,d}}{\sum_k f_{k,d}}$$

Burada $f_{t,d}$, t teriminin d belgesindeki ham geçme sayısını; payda ise belgedeki toplam terim sayısını ifade etmektedir.

- **Ters Belge Frekansı (IDF):**

IDF, bir terimin tüm belge koleksiyonu (D) içerisindeki ayırt ediciliğini ölçmektedir.

Bir terim ne kadar az belgede yer alıyorsa, IDF değeri o kadar yüksek olmaktadır. Hesaplama sırasında logaritmik ölçek kullanılarak aşırı frekansların etkisi dengelenmektedir:

$$IDF(t, D) = \log\left(\frac{N}{df_t}\right)$$

Burada N toplam belge sayısını, df_t ise t terimini içeren belge sayısını göstermektedir. Uygulamada sıfıra bölme problemini önlemek amacıyla düzeltme (smoothing) işlemleri uygulanabilmektedir.

- **TF-IDF Skoru:**

Bir terimin nihai ağırlığı, TF ve IDF değerlerinin çarpımı ile elde edilmektedir:

$$TF - IDF(t, d, D) = TF(t, d) * IDF(t, D)$$

3.7.2.2 TF-IDF Mimarisinin Çalışma Prensibi

TF-IDF tabanlı metin temsil süreci, doğrusal bir işlem hattı (pipeline) mimarisi üzerinden ilerlemektedir. *Şekil 3*'te gösterildiği üzere bu süreç aşağıdaki adımlardan oluşmaktadır:

1. **Metin Ön İşleme:**

Ham metin verileri küçük harfe dönüştürme, noktalama işaretlerinin kaldırılması ve durak kelimelerin ayıklanması gibi işlemlerden geçirilir.

2. **Tokenizasyon:**

Metinler kelime birimlerine veya n-gram yapılarına ayrıştırılır.

3. **Vektörizasyon:**

Her belge için yerel TF değerleri hesaplanırken, tüm belge koleksiyonu üzerinden küresel IDF değerleri belirlenir.

4. **Normalizasyon:**

Belge uzunluklarından kaynaklanan ölçek farklılıklarını azaltmak amacıyla TF-IDF vektörleri genellikle L2 normalizasyonuna tabi tutulur.

5. **Seyrek Matris Oluşumu:**

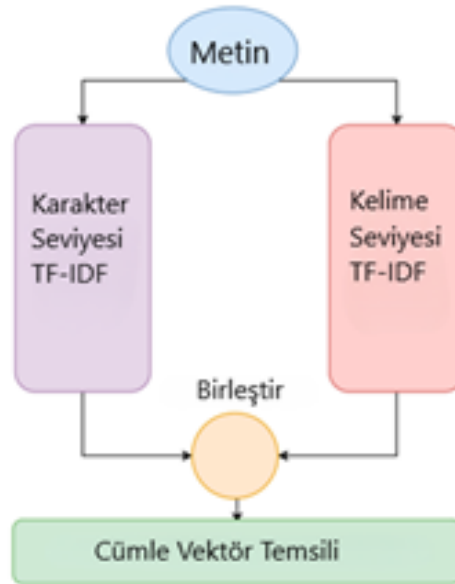
Süreç sonunda her belge, yüksek boyutlu ve seyrek (sparse) bir özellik vektörü ile temsil edilerek makine öğrenmesi algoritmalarına girdi olarak sunulur (Salton, Wong & Yang, 1975).

- **Karakter ve Kelime Düzeyinde TF-IDF Tabanlı Birleşik Metin Temsili**

Klasik TF-IDF yaklaşımında metinler çoğunlukla yalnızca kelime düzeyinde temsil edilmektedir. Ancak özellikle eklemeli bir dil yapısına sahip olan Türkçede, yalnızca kelime temelli temsiller bazı anlamsal ve biçimbilimsel bilgilerin göz ardı edilmesine neden olabilmektedir. Bu nedenle, bu tez çalışmasında metinlerin hem karakter düzeyinde hem de kelime düzeyinde TF-IDF temsilleri ayrı ayrı oluşturulmuş ve elde edilen özellik vektörleri birleştirilerek daha zengin bir cümle temsili elde edilmiştir.

Şekil 3.3'te gösterildiği üzere, ham metin girdisi iki paralel kanaldan işlenmektedir. İlk kanalda karakter düzeyinde TF-IDF temsili oluşturularak yazım farklılıkları, ek yapıları ve morfolojik varyasyonların modele yansıtılması hedeflenmiştir. İkinci kanalda ise kelime düzeyinde TF-IDF temsili kullanılarak anlamsal içerik korunmuştur. Her iki temsilden elde edilen özellik vektörleri birleştirme (concatenation) işlemi ile tek bir vektör hâline getirilmiş ve sonuçta her metin için birleşik bir cümle vektör temsili elde edilmiştir.

Bu birleşik temsil yaklaşımı, metnin hem yüzeysel (karakter tabanlı) hem de anlamsal (kelime tabanlı) özelliklerini aynı anda modele sunarak, özellikle duygu analizi gibi bağlama duyarlı görevlerde sınıflandırma başarımını artırmayı amaçlamaktadır.



Şekil 3.3. Karakter düzeyi ve kelime düzeyi TF-IDF temsillerinin birleştirilerek cümle vektörüne dönüştürülmesi süreci (Tussupov et al.,2025)

3.7.2.3. TF-IDF'in Bu Çalışmadaki Kullanımı

Bu tez çalışmasında TF-IDF temsili, yalnızca tek başına bir metin temsil yöntemi olarak değil; aynı zamanda bulanık mantık ve derin öğrenme tabanlı hibrit modellerin bir bileşeni olarak

kullanılmıştır. TF-IDF yöntemi ile elde edilen kelime ağırlıkları, özellikle bulanık mantık tabanlı sınıflandırma sistemlerinde giriş değişkenleri olarak değerlendirilmiştir.

Çalışmanın aşamalarında, TF-IDF temsiline dayalı özellikler öncelikle Mamdani tipi Max–Min bulanık çıkarım sistemi ile kullanılarak bulanık sınıflandırma gerçekleştirilmiştir. Bu yapı, metinlerdeki duygu yoğunluğunun kesin sınırlardan ziyade dereceli üyelikler ile modellenmesine olanak sağlamıştır.

Ayrıca, TF-IDF tabanlı özellikler, BERT modelinden elde edilen bağlamsal temsil çıktıları ile birleştirilerek TF-IDF–BERT–Fuzzy şeklinde üçlü (triple hybrid) bir mimari oluşturulmuştur. Bu hibrit yaklaşımda, TF-IDF yöntemi metnin istatistiksel ve kelime frekansına dayalı bilgisini temsil ederken; BERT modeli bağlamsal anlamsal bilgiyi sağlamış, Mamdani tipi bulanık mantık sistemi ise bu farklı bilgi kaynaklarını birleştirerek nihai duygu sınıflandırma kararını üretmiştir.

Bu bütünleşik yapı sayesinde, hem istatistiksel hem anlamsal hem de belirsizlik temelli bilgi birlikte değerlendirilmiş ve özellikle bulanık mantık tabanlı çıkarım mekanizmasının Max–Min yaklaşımı ile daha kararlı sonuçlar elde edilmiştir.

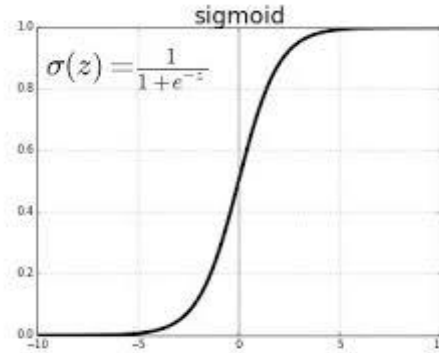
TF-IDF yöntemi, metin verilerinden istatistiksel özellikler elde etmek amacıyla hem geleneksel makine öğrenmesi modellerinde hem de bulanık mantık ve BERT ile oluşturulan hibrit yapılarda girdi temsili olarak kullanılmıştır.

3.7.3 Lojistik Regresyon

Lojistik Regresyon, en az sayıda değişken kullanarak bağımlı değişken ile bağımsız değişkenler arasındaki ilişkiyi en iyi uyuma sahip olacak şekilde tanımlamayı amaçlayan, olasılıksal temelli bir sınıflandırma modelidir. Modelin temel amacı, elde edilen sonuçların yalnızca istatistiksel olarak anlamlı değil, aynı zamanda biyolojik ve mantıksal olarak kabul edilebilir olmasını sağlamaktır. Lojistik modelin biyolojik deneylerin analizinde kullanılmasına ilişkin ilk öneri Berkson (1944) tarafından yapılmış olup, bu çalışma lojistik regresyonun modern istatistiksel modelleme yaklaşımlarının temel taşlarından biri olarak kabul edilmektedir (Berkson, 1944).

Lojistik regresyon, denetimli öğrenme kapsamında değerlendirilen ve özellikle ikili sınıflandırma problemleri için yaygın olarak kullanılan bir yöntemdir. Lineer regresyondan farklı olarak, çıktı değişkeninin sürekli değil, kategorik olduğu durumlarda tercih edilmektedir. Model, bir gözlemin belirli bir sınıfa ait olma olasılığını tahmin eder ve bu olasılığı 0 ile 1

aralığında sınırlamak amacıyla lojistik (sigmoid) fonksiyonu kullanır (Hosmer, Lemeshow & Sturdivant, 2013). Sigmoid Fonksiyonu Şekil 3.4'te gösterilmiştir.



Şekil 3.4. Sigmoid fonksiyonu

Lojistik regresyonun temel varsayımı, bağımsız değişkenler ile sınıf olasılığının log-olasılık oranı (log-odds) arasında doğrusal bir ilişkinin bulunduğuudur. Bu ilişki matematiksel olarak aşağıdaki şekilde ifade edilmektedir:

$$\log \left(\frac{P(y=1 | x)}{1-P(y=1 | x)} \right) = \beta_0 + \beta_1 x_1 + \dots + \beta_n x_n$$

Burada x_i bağımsız değişkenleri, β_i ise model tarafından öğrenilen katsayıları temsil etmektedir. Bu yapı, lojistik regresyonun yorumlanabilirliğini artırmakta ve değişkenlerin sınıf tahmini üzerindeki etkisinin açık biçimde analiz edilmesine olanak tanımaktadır (Hastie, Tibshirani & Friedman, 2009).

Modelin eğitimi sırasında parametreler, Maksimum Olabilirlik Tahmini (Maximum Likelihood Estimation – MLE) yaklaşımı kullanılarak belirlenir. Optimizasyon sürecinde yaygın olarak kullanılan maliyet fonksiyonu, negatif log-olabilirlik olarak da bilinen lojistik kayıp (log-loss / cross-entropy) fonksiyonudur (Bishop, 2006). Bu fonksiyon, tahmin edilen olasılıkların gerçek sınıf etiketleriyle uyumunu ölçmektedir.

Lojistik regresyon, özellikle yüksek boyutlu ve seyrek veri yapılarında etkin çalışması, hesaplama maliyetinin düşük olması ve elde edilen sonuçların kolay yorumlanabilmesi nedeniyle doğal dil işleme, biyoinformatik ve tıbbi veri analizi gibi alanlarda yaygın olarak kullanılmaktadır. Bu özellikleri sayesinde, birçok çalışmada daha karmaşık makine öğrenimi modelleri için temel karşılaştırma (baseline) algoritması olarak tercih edilmektedir (Alpaydın, 2020).

3.7.3.1 Lojistik Regresyonun Çalışma Mimarisi

Lojistik regresyon, doğrusal bir karar fonksiyonu ile olasılık tabanlı sınıflandırma gerçekleştiren, ileri beslemeli ve tek katmanlı bir makine öğrenmesi algoritmasıdır. Bu modelde herhangi bir gizli katman bulunmamakta olup, girdi özellikleri doğrudan çıktıya bağlanmaktadır. Lojistik regresyonun çalışma mimarisi Şekil 5'te görsel olarak sunulmuştur.

Modelin girdi katmanında, metin verilerinden elde edilen TF-IDF temsilleri, kelime frekansları veya gömme (embedding) vektörleri gibi sayısal özellikler yer almaktadır. Bu özellikler, model tarafından öğrenilen ağırlık katsayıları ile çarpılarak doğrusal bir kombinasyon oluşturulmaktadır. Bu işlem sonucunda elde edilen doğrusal skor aşağıdaki şekilde ifade edilmektedir:

$$z = \beta_0 + \sum_{i=1}^n \beta_i x_i$$

Burada x_i girdi özelliklerini, β_i öğrenilen ağırlık katsayılarını ve β_0 bias (sabit terim) değerini temsil etmektedir.

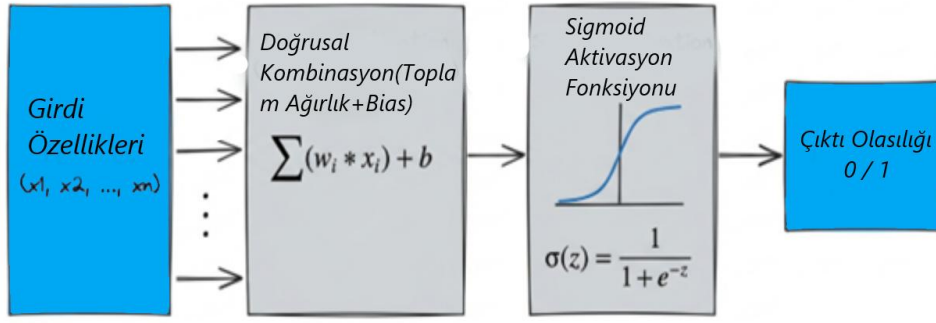
Elde edilen doğrusal skor, bir sınıfa ait olma olasılığını hesaplayabilmek amacıyla sigmoid (lojistik) aktivasyon fonksiyonuna aktarılmaktadır. Sigmoid fonksiyon, doğrusal çıktıyı 0 ile 1 aralığına dönüştürerek olasılıksal bir yorum yapılmasına olanak sağlamaktadır:

$$\sigma(z) = \frac{1}{1 + e^{-z}}$$

Modelin çıktı katmanında her bir örnek için pozitif sınıfa ait olma olasılığı hesaplanır. Bu değer şu şekilde ifade edilmektedir:

$$P(y = 1 | x) = \sigma(z)$$

Şekil 3.5'te de görüldüğü üzere, lojistik regresyon mimarisi; girdi özelliklerinin doğrusal bir birleşimini, ardından sigmoid aktivasyon fonksiyonu aracılığıyla olasılık çıktısına dönüştüren sade ve hesaplama açısından verimli bir yapı sunmaktadır. Bu özellikleri sayesinde lojistik regresyon, özellikle yüksek boyutlu ve seyrek veri uzaylarında (örneğin TF-IDF temsilleri) duygu analizi gibi metin sınıflandırma problemlerinde yaygın olarak tercih edilmektedir (Mitchell, 1997).



Şekil 3.5. Lojistik regresyon çalışma mimarisi

Lojistik Regresyon modeli, TF-IDF tabanlı özellik vektörleri üzerinde doğrusal bir sınıflandırıcı olarak uygulanmış ve sonuçları karşılaştırmalı analizlerde referans alınmıştır.

3.7.4 Naive Bayes

Naive Bayes (NB) sınıflandırıcısı, Bayes teoremine dayanan ve denetimli öğrenme kapsamında değerlendirilen olasılıksal bir sınıflandırma algoritmasıdır. Algoritma, bir örneğin belirli bir sınıfa ait olma olasılığını, gözlemlenen özelliklerin koşullu olasılıklarını kullanarak hesaplamaktadır. Naive Bayes yaklaşımının temel varsayımı, özelliklerin sınıf etiketi verildiğinde birbirinden bağımsız olduğu kabulüdür. Bu varsayım çoğu gerçek dünya probleminde tam olarak sağlanmasa da, algoritmanın pratikte başarılı sonuçlar ürettiği literatürde yaygın biçimde belirtilmektedir.

Naive Bayes sınıflandırıcısının performansı üzerine yapılan en kapsamlı ampirik çalışmalardan biri Rish tarafından 2001 yılında gerçekleştirilmiştir. Bu çalışmada, algoritmanın farklı veri setleri üzerindeki davranışı incelenmiş ve özellikler arasındaki bağımsızlık varsayımının ihlal edilmesine rağmen Naive Bayes'in birçok durumda rekabetçi performans sergilediği gösterilmiştir. Özellikle yüksek boyutlu ve sınırlı örnek içeren veri setlerinde Naive Bayes'in basit yapısına rağmen etkili bir sınıflandırıcı olduğunu vurgulamaktadır (Rish, 2001).

Doğal dil işleme alanında Naive Bayes algoritması, özellikle metin sınıflandırma problemlerinde yaygın olarak kullanılmaktadır. Bunun temel nedenleri arasında kelime tabanlı özelliklerle doğal uyumu, düşük hesaplama maliyeti ve hızlı eğitim süreci yer almaktadır. McCallum ve Nigam 1998 yılında, Naive Bayes algoritmasının metin sınıflandırma

görevlerinde özellikle Multinomial Naive Bayes modeli ile başarılı sonuçlar verdiğini ve yüksek boyutlu kelime uzaylarında kararlı performans sergilediğini ortaya koymuştur (McCallum & Nigam, 1998).

Benzer şekilde, Manning, Raghavan ve Schütze (2008), Naive Bayes sınıflandırıcısının metin madenciliği ve bilgi erişim problemlerinde güçlü bir temel yöntem olduğunu ve çoğu durumda daha karmaşık sınıflandırma algoritmalarına kıyasla karşılaştırılabilir başarımlar sunduğunu belirtmektedir (Manning et al., 2008). Duygu analizi çalışmalarında ise Naive Bayes, metinlerdeki duygu yüklü kelimelerin sınıf tahminine doğrudan katkı sağlaması nedeniyle sıklıkla tercih edilmektedir. Pang, Lee ve Vaithyanathan (2002), Naive Bayes algoritmasının duygu sınıflandırma problemlerinde etkili bir başlangıç yöntemi olduğunu ve temel bir karşılaştırma modeli olarak kullanılabileceğini göstermiştir (Pang et al., 2002).

Genel Bayes formülü, olasılık kuramında koşullu olasılıkların hesaplanmasına olanak tanıyan temel bir yaklaşımdır ve Naive Bayes algoritmasının matematiksel temelini oluşturmaktadır. Bayes teoremi aşağıdaki şekilde ifade edilmektedir:

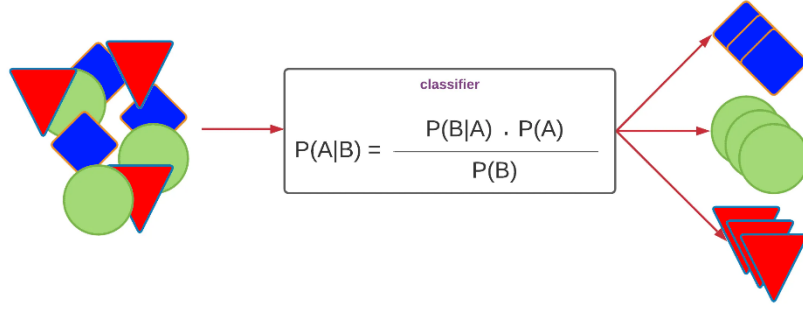
$$P(A|B) = \frac{P(B|A)P(A)}{P(B)}$$

Denklem üzerinde yer alan ifadeler şu anlamlara gelmektedir:

- $P(A | B)$, B olayı gerçekleştiğinde A olayının gerçekleşme olasılığını;
- $P(A)$, A olayının gerçekleşme olasılığını (önsel olasılık – *prior*);
- $P(B | A)$, A olayı gerçekleştiğinde B olayının gerçekleşme olasılığını (olabilirlik – *likelihood*);
- $P(B)$ ise B olayının gerçekleşme olasılığını ve normalizasyon terimini ifade etmektedir.

Bu formül, bir olayın olasılığının yeni bir kanıt elde edildiğinde nasıl güncelleneceğini açıklamakta olup, Bayesçi istatistiğin temelini oluşturmaktadır (Bayes, 1763; Bishop, 2006). Naive Bayes sınıflandırıcısı, Bayes teoremini özellikler arasında koşullu bağımsızlık varsayımı altında kullanarak sınıflandırma işlemini gerçekleştirmektedir.

Naive Bayes sınıflandırma yöntemi Şekil 3.6’da gösterilmiştir.



Şekil 3.6. Naive Bayes sınıflandırma

Bu çalışmada Naive Bayes algoritması, klasik olasılıksal yaklaşımı temsil eden bir temel (baseline) yöntem olarak, diğer modellerin performanslarının karşılaştırılması amacıyla kullanılmıştır.

3.7.5 Destek Vektör Makineleri (DVM)

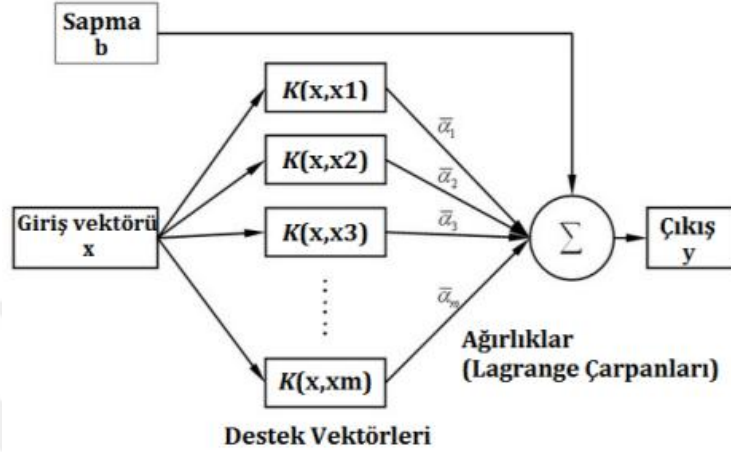
1998 yılında John C. Platt tarafından çok terimli çekirdek kullanarak DVM'yi eğitmek için geliştirilmiştir. Sıralı Minimal Optimizasyon Algoritmasını (Platt 1998) uygular. Bu uygulama, tüm kayıp değerlerini yeni değerlerle küresel olarak değiştirir ve nominal özellikleri ikili özelliklere dönüştürür. Ayrıca, önceden tanımlanmış değerlerle tüm özellikleri normalleştirir. Uygun olasılık tahminleri elde etmek için, lojistik regresyon modellerinin DVM çıktılarına uymasını sağlayan seçenekler kullanılmıştır. Birden fazla sınıf olması durumunda, tahmin edilen olasılıklar Hastie ve Tibshirani tarafından 1998 yılında geliştirilen çift eşleştirme yöntemi kullanılarak eşleştirilmiştir (Hastie & Tibshirani 1998)

Destek Vektör Makineleri (DVM), denetimli öğrenme kapsamında yer alan ve sınıflandırma problemlerinde yaygın olarak kullanılan güçlü bir makine öğrenmesi yöntemidir. DVM'nin temel amacı, farklı sınıflara ait veri noktalarını ayıran ve sınıflar arasındaki mesafeyi (marjin) maksimum yapan optimum ayırıcı hiperdüzlemi belirlemektir (Cortes & Vapnik, 1995). Bu yaklaşım, modelin genelleme yeteneğini artırarak daha önce görülmemiş veriler üzerinde de başarılı sonuçlar elde edilmesini sağlamaktadır (Vapnik, 1998).

DVM algoritmasında karar sınırı, yalnızca ayırıcı hiperdüzleme en yakın veri noktaları olan destek vektörleri tarafından belirlenmektedir. Bu durum, modelin gereksiz veri noktalarından etkilenmesini engelleyerek daha kararlı bir sınıflandırma yapısı sunmaktadır.

3.7.5.1 Destek Vektör Makinelerinin (DVM/SVM) Çalışma Mimarisi

Destek Vektör Makineleri (Support Vector Machines – SVM), sınıflar arasındaki ayırımı maksimize eden optimal bir hiper-düzlem öğrenmeyi amaçlayan, güçlü ve teorik temellere dayanan bir makine öğrenmesi yöntemidir. DVM'nin kernel tabanlı çalışma yapısı Şekil 3.7'de gösterilmektedir.



Şekil 3.7. Destek vektör makinesi (DVM) ağ mimarisi

Modelin girişinde her örnek, $x=(x_1,x_2,...,x_n)$ biçiminde sayısal bir özellik vektörü ile temsil edilmektedir. Doğal dil işleme uygulamalarında bu vektörler genellikle TF-IDF veya Bag of Words gibi yöntemlerle elde edilmektedir (Joachims, 1998). DVM'nin mesafe temelli yapısı nedeniyle, bu özellikler çoğunlukla normalizasyon veya standartlaştırma işlemlerinden geçirilerek modele sunulmaktadır (Hsu, Chang & Lin, 2003).

Şekil 3.7'de görüldüğü üzere, giriş vektörü ile her bir destek vektörü x_i arasında $K(x, x_i)$ biçiminde çekirdek fonksiyonu hesaplanmaktadır. Bu işlem, verilerin açıkça dönüştürülmesine gerek kalmadan daha yüksek boyutlu bir özellik uzayında benzerlik hesaplanmasını sağlayan kernel hilesi (kernel trick) ile gerçekleştirilmektedir.

Elde edilen bu çekirdek çıktıları, her bir destek vektörüne karşılık gelen ağırlık katsayıları (α_i , Lagrange çarpanları) ile çarpılmakta ve tüm katkılar toplanarak bir karar fonksiyonu oluşturulmaktadır. Ayrıca modele bir bias (b) terimi de eklenmektedir. Bu süreç matematiksel olarak şu şekilde ifade edilmektedir:

$$f(x) = \sum_{i=1}^n \alpha_i K(x, x_i) + b$$

Burada α_i katsayıları yalnızca destek vektörleri için sıfırdan farklı değerler almakta olup, modelin karar sınırını belirleyen temel bileşenlerdir.

DVM, bu karar fonksiyonunu öğrenirken sınıflar arasındaki marjini maksimize etmeyi ve aynı zamanda sınıflandırma hatalarını minimize etmeyi amaçlayan bir optimizasyon problemi çözmektedir. Bu denge, düzenlileştirme parametresi C aracılığıyla kontrol edilmektedir (Cortes & Vapnik, 1995)

Son aşamada, elde edilen karar fonksiyonunun işaretine göre örneklerin sınıf etiketleri belirlenmektedir. *Şekil 3.7*, giriş vektörü ile destek vektörleri arasındaki kernel tabanlı benzerlik hesaplamalarının ağırlıklı toplamı üzerinden karar çıktısının nasıl elde edildiğini açık biçimde göstermektedir.

3.7.6 Bulanık Mantık (Fuzzy Logic)

Bulanık mantık, klasik mantıkta kullanılan kesin doğru–yanlış (0–1) yaklaşımının aksine, gerçek dünyadaki belirsizlikleri ve kesin olmayan bilgileri modellemek amacıyla geliştirilmiş bir yapay zekâ yaklaşımıdır. Bu yöntemde bir elemanın bir kümeye aitliği, kesin sınırlar yerine 0 ile 1 arasında değişen üyelik dereceleri ile ifade edilmektedir. Bulanık mantık kavramı ilk kez Lotfi A. Zadeh tarafından 1965 yılında ortaya atılmış ve belirsizlik içeren problemlerin matematiksel olarak modellenmesine olanak sağlamıştır (Zadeh, 1965).

Bulanık mantık tabanlı sistemler; kontrol sistemleri, karar destek sistemleri ve doğal dil işleme uygulamaları gibi alanlarda, insan benzeri karar verme süreçlerini taklit edebilme yetenekleri nedeniyle yaygın olarak kullanılmaktadır (Klir & Yuan, 1995).

Bu çalışmada bulanık mantık tabanlı sınıflandırma işlemi, Max–Min bulanık çıkarım yöntemine dayalı olarak gerçekleştirilmiştir. Max–Min yöntemi, özellikle Mamdani tipi bulanık çıkarım sistemlerinin temelini oluşturan ve kuralların değerlendirilmesinde minimum ve maksimum operatörlerini kullanan bir yaklaşımdır (Mamdani, 1977).

Bu çalışmada bulanık çıkarım mekanizması olarak Mamdani tipi bulanık sistem tercih edilmiştir. Mamdani yaklaşımı;

1.Girdi Verisinin Alınması:

Sistemin girdisi, sınıflandırma problemine ait sayısal veya kategorik özelliklerden oluşmaktadır. Doğal dil işleme uygulamalarında bu girdiler genellikle kelime ağırlıkları, duygu

skorları veya makine öğrenmesi modellerinden elde edilen ara çıktılar şeklinde temsil edilmektedir (Herrera, 2008).

2. Bulanıklaştırma (Fuzzification)

Keskin (crisp) giriş değerleri, önceden tanımlanmış üyelik fonksiyonları aracılığıyla bulanık kümelere dönüştürülmektedir. Bu aşamada üçgensel, yamuksal veya Gauss üyelik fonksiyonları yaygın olarak kullanılmaktadır (Ross, 2010).

3. Kural Tabanı (IF–THEN Kuralları)

Bulanık kural tabanı, uzman bilgisine veya deneysel gözlemlere dayalı olarak oluşturulan IF–THEN kurallarından oluşmaktadır. Bu kurallar, giriş değişkenleri ile sınıf etiketleri arasındaki ilişkileri bulanık biçimde tanımlamaktadır (Klir & Yuan, 1995).

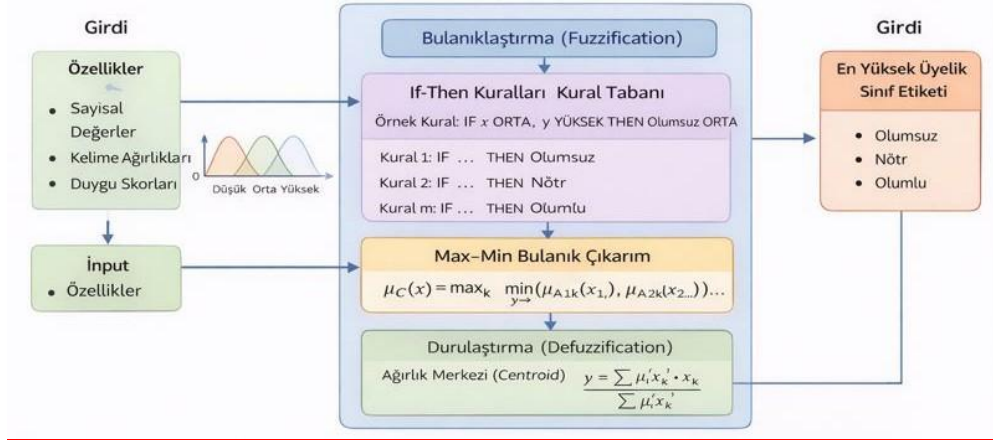
4. Max–Min Bulanık Çıkarım

Her bir kuralın etkinliği, öncül kısımdaki üyelik derecelerinin minimum (Min) operatörü ile birleştirilmesi sonucu elde edilmektedir. Aynı sınıfa ait birden fazla kuraldan üretilen bulanık çıktılar ise maksimum (Max) operatörü kullanılarak birleştirilmektedir. Bu yaklaşım Max–Min bulanık çıkarım yöntemi olarak adlandırılmaktadır (Mamdani, 1977; Ross, 2010).

Bu süreç matematiksel olarak aşağıdaki şekilde ifade edilebilir:

$$\mu_C(x) = \max_k (\min(\mu_{A1k}(x_1), \mu_{A2k}(x_2) \dots))$$

Bulanık mantık tabanlı sınıflandırma sürecinin genel işleyişi, girdi değişkenlerinin bulanıklaştırılması, kural tabanı üzerinden Max–Min çıkarımın gerçekleştirilmesi ve durulaştırma adımlarını içeren mimari yapı ile *Şekil 3.8*'de sunulmaktadır.



Şekil 3.8. Mamdani tipi Max-Min bulanık çıkarım sisteminin çalışma mimarisi

5. Durulaştırma (Defuzzification)

Son aşamada elde edilen bulanık çıktı, tek bir keskin değere dönüştürülmektedir. Bu çalışmada yaygın olarak kullanılan ağırlık merkezi (centroid) yöntemi tercih edilmiştir (Ross, 2010).

Girdi değişkenleri için tanımlanan üyelik fonksiyonları, düşük, orta ve yüksek gibi dilsel terimlerle ifade edilmiştir. Çıkarım sürecinde max-min yöntemi kullanılarak kurallar aktive edilmiş ve sonuçlar birleştirilmiştir. Nihai sınıf etiketi, durulaştırma aşamasında elde edilen sayısal çıktıya göre belirlenmiştir.

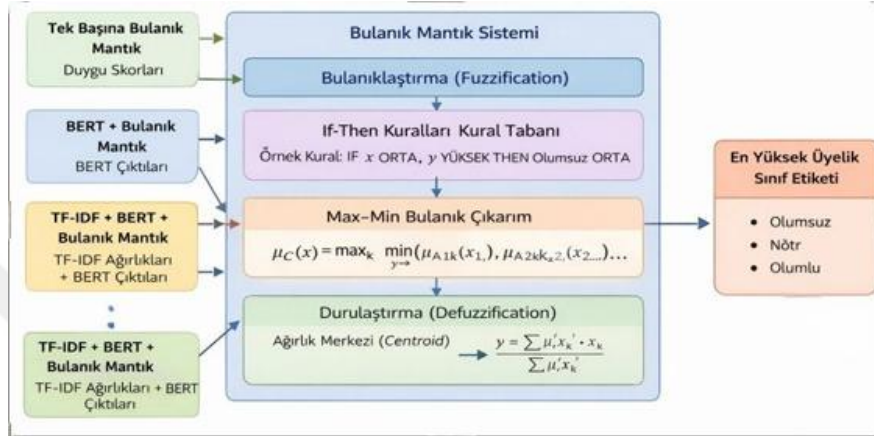
3.7.6.1 Çalışmadaki Kullanım Şekilleri

Bu çalışmada bulanık mantık yaklaşımı üç farklı yapı altında değerlendirilmiştir:

- **Tek Başına Bulanık Mantık:** Metinlerden elde edilen duygu skorları veya istatistiksel özellikler doğrudan bulanık sisteme girdi olarak verilmiş ve sınıflandırma yalnızca Mamdani bulanık çıkarım sistemi üzerinden gerçekleştirilmiştir.
- **BERT-Fuzzy Hibrit Model:** BERT modeli tarafından üretilen bağlamsal temsil veya olasılık skorları, bulanık mantık sistemine girdi olarak kullanılmış; böylece derin öğrenmenin anlamsal gücü ile bulanık mantığın belirsizlik modelleme yeteneği birleştirilmiştir.
- **TF-IDF-BERT-Fuzzy Üçlü Hibrit Model:** Bu yapıda TF-IDF ve BERT tabanlı özellikler birlikte değerlendirilmiş, elde edilen çıktılar Mamdani bulanık çıkarım sistemine aktarılmış ve nihai karar bulanık mantık aracılığıyla verilmiştir. Bu yaklaşım,

hem istatistiksel hem bağlamsal bilgiyi aynı karar mekanizmasında bütünleştirmeyi amaçlamaktadır.

Bu çalışmada, metinlerden elde edilen TF-IDF ve BERT tabanlı özellik temsilleri, Mamdani tipi Max–Min bulanık çıkarım mekanizmasına girdi olarak verilerek hibrit bir sınıflandırma yapısı oluşturulmuştur. Söz konusu hibrit yaklaşımın genel mimarisi Şekil 3.9’da sunulmaktadır.



Şekil 3.9. Çalışmada kullanılan hibrit bulanık mantık tabanlı sınıflandırma mimarisi

3.7.6 BERT(Bidirectional Encoder Representations from Transformers /Transformatörlerden İki Yönlü Kodlayıcı Temsilleri)

Bu çalışmada, Türkçe metinlerin anlamsal derinliğini yakalamak amacıyla HuggingFace kütüphanesi üzerinden erişilen "dbmdz/bert-base-turkish-cased" (BERTTurk) modeli kullanılmıştır. Bu modelin tercih edilme nedeni, 12 katmanlı ve 768 gizli birimli mimarisi sayesinde kelimelerin cümle içindeki sağ ve sol bağlamını aynı anda değerlendirebilmesidir. Eğitim sürecinde maksimum dizi uzunluğu 128 token olarak belirlenmiş ve model 3 epoch boyunca ince ayar (fine-tuning) işlemine tabi tutulmuştur.

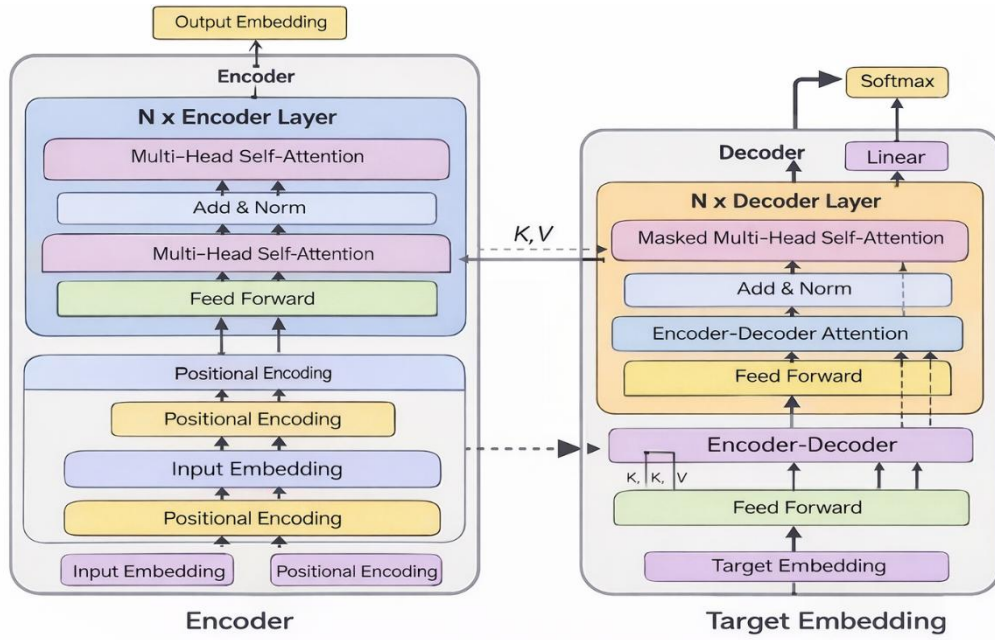
BERT (Bidirectional Encoder Representations from Transformers), doğal dil işleme problemlerinde bağlam bilgisini çift yönlü olarak modelleyebilmek amacıyla geliştirilen, transformer tabanlı bir derin öğrenme modelidir. Model, Google AI Research tarafından Jacob Devlin ve arkadaşları tarafından 2018 yılında önerilmiş ve kısa sürede doğal dil işleme alanında önemli bir dönüm noktası hâline gelmiştir (Devlin et al., 2019). BERT’in temel yeniliği, bir kelimenin anlamını yalnızca kendisinden önce gelen sözcüklerle değil, hem sol hem de sağ bağlamı aynı anda dikkate alarak öğrenmesidir.

Geleneksel dil modelleri genellikle tek yönlü bağlam kullanırken, BERT'in çift yönlü yapısı sayesinde cümle içerisindeki anlamsal ilişkiler daha doğru ve kapsamlı biçimde temsil edilebilmektedir. Bu özellik, BERT'i metin sınıflandırma, duygu analizi, adlandırılmış varlık tanıma ve soru-cevap sistemleri gibi pek çok doğal dil işleme görevinde etkili bir model hâline getirmektedir (Vaswani et al., 2017, Devlin et al., 2019). Ayrıca BERT, Word2Vec ve GloVe gibi statik kelime gömme yöntemlerinden farklı olarak bağlama duyarlı (contextualized) kelime temsilleri üretmekte, böylece çok anlamlı kelimelerin farklı bağlamlarda doğru biçimde temsil edilmesini sağlamaktadır (Aftan & Shah, 2023).

3.7.6.1 Transformer Tabanlı Mimari Yapı

BERT modeli, Transformer mimarisinin yalnızca encoder bloğunu temel almaktadır. Transformer mimarisi, tekrarlayan sinir ağları (RNN) veya evrimsel sinir ağları (CNN) yerine self-attention (öz-dikkat) mekanizmasını kullanarak uzun menzilli bağımlılıkları paralel biçimde öğrenebilmektedir. Bu yaklaşım, özellikle uzun metinlerde bağlam bilgisinin korunması açısından önemli avantajlar sunmaktadır.

Transformer mimarisinin genel yapısı, encoder ve decoder bloklarından oluşan çok katmanlı bir mimari olarak tasarlanmıştır. *Şekil 3.10'*da gösterildiği üzere, her encoder katmanı çok başlı öz-dikkat (multi-head self-attention) mekanizması ve ileri beslemeli sinir ağlarından oluşmakta; artık bağlantılar (residual connections) ve katman normalizasyonu (layer normalization) ile desteklenmektedir. Bu yapı sayesinde model, diziler arasındaki uzun menzilli bağımlılıkları paralel ve verimli bir biçimde öğrenebilmektedir. Decoder bloğu ise hedef diziyi üretmek amacıyla maskeleye mekanizması içeren öz-dikkat ve encoder–decoder dikkat katmanlarını kullanmaktadır (Vaswani et al., 2017).



Şekil 3.10. Transformer tabanlı mimari yapı

Bu yapı sayesinde model, cümle içerisindeki her bir kelimenin diğer kelimelerle olan ilişkisini aynı anda değerlendirebilmekte ve bağlama duyarlı derin temsiller üretebilmektedir (Vaswani et al., 2017).

3.7.6.2 BERT'in Çalışma Yöntemi

BERT tabanlı bir modelin çalışma süreci aşağıdaki temel aşamalardan oluşmaktadır:

- **Girdi Temsili:**

BERT'e verilen giriş, üç farklı gömme vektörünün toplamı şeklinde temsil edilmektedir:

- Token embedding,
- Segment embedding,
- Position embedding.

Bu yapı, modelin hem kelime sırasını hem de cümle segmentlerini ayırt edebilmesini sağlamaktadır. Ayrıca sınıflandırma görevlerinde cümle başına eklenen [CLS] belirteci, tüm diziye temsil eden özet bir vektör olarak kullanılmaktadır. Cümle veya cümle çiftlerini ayırmak amacıyla ise [SEP] belirteci kullanılmaktadır (Devlin et al., 2019; Aftan & Shah, 2023).

- **Çift Yönlü Öz-Dikkat Mekanizması**

Self-attention mekanizması sayesinde her bir kelime, cümledeki diğer tüm kelimelerle olan ilişkisini aynı anda değerlendirebilmektedir. Bu sayede BERT, yalnızca yerel değil küresel bağlamı da dikkate alarak bağlama duyarlı kelime temsilleri üretmektedir (Vaswani et al., 2017).

- **Ön Eğitim (Pre-training)**

BERT modeli, büyük ölçekli metin koleksiyonları üzerinde iki temel görev kullanılarak ön eğitime tabi tutulmaktadır:

Masked Language Model (MLM): Cümledeki kelimelerin belirli bir kısmı maskelenir ve modelden bu kelimeleri tahmin etmesi beklenir.

Next Sentence Prediction (NSP): İki cümlemin birbirini mantıksal olarak takip edip etmediği tahmin edilir.

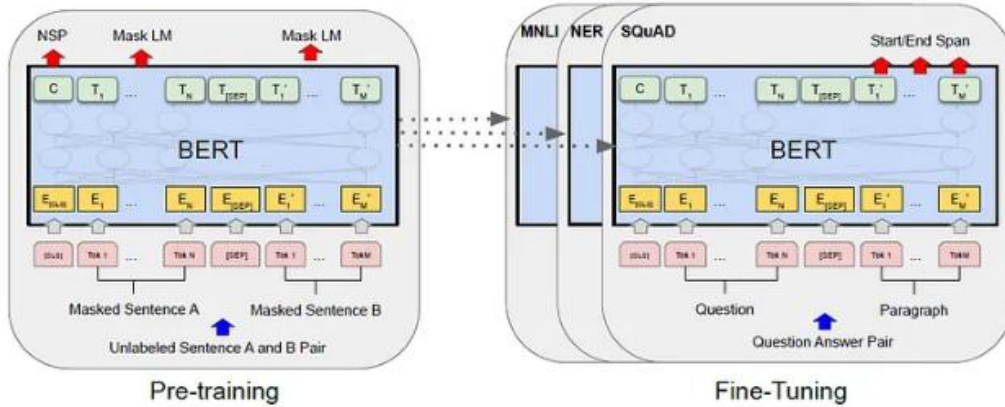
Bu ön eğitim süreci, modelin dilin genel sözdizimsel ve anlamsal yapısını öğrenmesini sağlamaktadır (Devlin et al., 2019).

- **İnce Ayar (Fine-Tuning) ve Kullanım Yaklaşımları**

Ön eğitilmiş BERT modeli, duygu analizi veya metin sınıflandırma gibi belirli görevler için etiketli veri kullanılarak ince ayar (fine-tuning) aşamasından geçirilmektedir. Bu aşamada genellikle BERT mimarisinin üstüne basit bir sınıflandırma katmanı eklenmektedir (Devlin et al., 2019).

Survey çalışmalarda belirtildiği üzere BERT, yalnızca uçtan uca ince ayar yöntemiyle değil, aynı zamanda bir özellik çıkarıcı (feature extractor) olarak da kullanılabilir. Bu yaklaşımda BERT'ten elde edilen bağlamsal temsiller, geleneksel makine öğrenmesi veya farklı karar mekanizmalarına girdi olarak sunulmaktadır (Aftan & Shah, 2023). BERT modelinin ön eğitim ve görev-özel ince ayar süreçleri *Şekil 3.11*'de görsel olarak sunulmaktadır. Ön eğitim aşamasında model, etiketlenmemiş büyük ölçekli metinler üzerinde Masked Language Model (MLM) ve Next Sentence Prediction (NSP) görevleri aracılığıyla dilin genel yapısını öğrenmektedir. Bu aşamada, cümle çiftleri ve maskelenmiş kelimeler kullanılarak bağlama duyarlı kelime temsilleri elde edilmektedir. İnce ayar sürecinde ise önceden eğitilmiş BERT modeli, duygu analizi, metin sınıflandırma veya soru-cevap gibi belirli görevler için

etiketli veri kullanılarak yeniden eğitilmekte ve mimarının üstüne genellikle basit bir sınıflandırma katmanı eklenmektedir (Devlin et al., 2019).

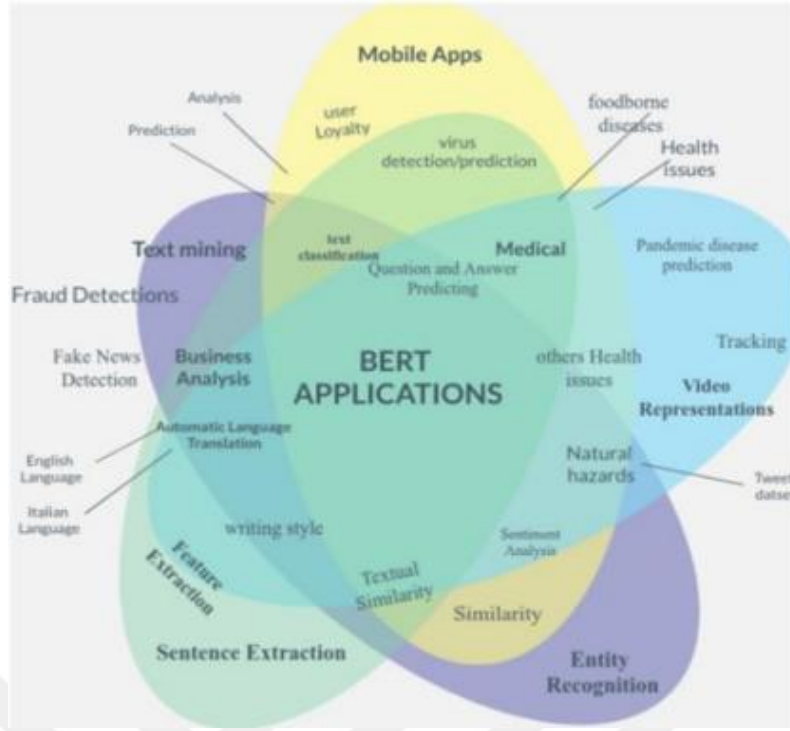


Şekil 3.11. BERT modelinin ön eğitim (pre-training) ve ince ayar (fine-tuning) süreçleri (Devlin et al., 2019)

3.7.6.3 Kullanım Alanları ve Sınırlamalar

BERT modeli, bağlama duyarlı dil temsilleri üretebilme yeteneği sayesinde çok geniş bir uygulama alanına sahiptir. Şekil 3.12’de gösterildiği üzere, BERT; metin madenciliği, duygu analizi, metin sınıflandırma, soru-cevap sistemleri, adlandırılmış varlık tanıma ve cümle benzerliği gibi temel doğal dil işleme görevlerinin yanı sıra, sağlık bilişimi, sahte haber tespiti, pandemi tahmini, kullanıcı davranışı analizi ve mobil uygulamalar gibi disiplinler arası alanlarda da etkin biçimde kullanılmaktadır. Bu çok yönlü kullanım alanları, BERT’in hem akademik araştırmalarda hem de gerçek dünya uygulamalarında yaygın olarak tercih edilmesinin temel nedenlerinden biridir.

Literatürde yapılan kapsamlı çalışmalar, BERT’in özellikle metin sınıflandırma, duygu analizi, adlandırılmış varlık tanıma, soru-cevap sistemleri, metin benzerliği ve doğal dil çıkarımı gibi görevlerde yüksek doğruluk performansı sunduğunu ortaya koymaktadır. Bununla birlikte, modelin çok katmanlı ve yüksek parametrelili yapısı, önemli ölçüde hesaplama maliyeti ve bellek gereksinimi doğurmaktadır. Bu durum, özellikle büyük ölçekli veri kümeleri ve sınırlı donanım kaynakları söz konusu olduğunda uygulama süreçlerini zorlaştırmakta; modelin pratik kullanımı açısından önemli bir sınırlılık oluşturmaktadır.



Şekil 3.12. BERT modelinin doğal dil işleme ve disiplinler arası alanlardaki başlıca kullanım alanları

3.7.6.4 Bu Çalışmada BERT'in Kullanımı

Bu tez çalışmasında BERT modeli, Türkçe film yorumlarının duygu analizi probleminde bağlama duyarlı metin temsilleri elde etmek amacıyla kullanılmıştır. Ön eğitilmiş Türkçe BERT modeli, etiketli veri kümesi üzerinde duygu sınıflandırma görevine yönelik olarak ince ayar (fine-tuning) sürecinden geçirilmiştir. Bu süreçte BERT'in tokenına karşılık gelen çıktısı, sınıflandırma katmanına girdi olarak verilmiş ve yorumların duygu etiketleri tahmin edilmiştir.

Ayrıca, BERT modeli yalnızca bağımsız bir sınıflandırıcı olarak değil, aynı zamanda hibrit yaklaşımların bir bileşeni olarak da değerlendirilmiştir. Bu kapsamda, BERT'ten elde edilen bağlama duyarlı çıktı temsilleri, bulanık mantık tabanlı Max–Min (Mamdani) çıkarım yöntemi ile birleştirilerek BERT–Fuzzy Logic hibrit modeli oluşturulmuştur. Bunun yanı sıra, istatistiksel metin temsili sunan TF-IDF özellikleri ile BERT ve bulanık mantık bileşenlerinin birlikte kullanıldığı üçlü (triple) hibrit yapı da incelenmiştir. Bu yaklaşım ile derin öğrenme, istatistiksel özellik çıkarımı ve kural tabanlı bulanık çıkarım yöntemlerinin tamamlayıcı güçlü yönlerinden yararlanılması hedeflenmiştir.

3.7.7. Bulanık Mantık ve Hibrit Karar Mekanizması

Çalışmanın özgün değerini oluşturan hibrit yapıda, Bulanık Mantık bir karar verici (decision maker) rolündedir. Sistemin giriş değişkenleri, diğer modellerden gelen güven skorlarıdır.

- **Üyelik Fonksiyonları:** Giriş değişkenleri için [0, 1] aralığında "Düşük", "Orta" ve "Yüksek" üçgensel üyelik fonksiyonları tanımlanmıştır.
- **Hibrit Oranların Belirlenmesi:** Hibrit modelde kullanılan %50 (BERT), %30 (TF-IDF) ve %20 (Bulanık Mantık) ağırlıkları, Grid Search yöntemiyle farklı kombinasyonlar test edilerek, en yüksek F1-skorunu veren optimize değerler olarak belirlenmiştir. Bu ağırlıklandırma, BERT'in bağlamsal başarısını temel alırken, istatistiksel ve kural tabanlı yaklaşımların "uç vakalardaki" (edge cases) düzeltici etkisinden yararlanmaktadır.

3.7.8. Hibrit Model Çalışma Mantığı

Hibrit algoritmanın işleyiş adımları aşağıda sunulmuştur:

Algoritma 1: Hibrit Duygu Analizi Karar Mekanizması

- **Giriş:** Temizlenmiş Metin (X)
- **Adım 1:** BERT modelinden olasılık skorunu hesapla (P_{bert})
- **Adım 2:** TF-IDF + Lojistik Regresyon skorunu hesapla (P_{tfidf})
- **Adım 3:** Metin bazlı kural skorunu (Bulanık Mantık) hesapla (P_{fuzzy})
- **Adım 4:** Ağırlıklı skor birleştirme:

$$FinalScore = (P_{bert} * 0.50) + (P_{tfidf} * 0.30) + (P_{fuzzy} * 0.20)$$

- **Adım 5:** Eşik Değer Kontrolü:

Eğer $FinalScore \geq 0.5$ **ise** Sonuç = "Pozitif"

Değilse Sonuç = "Negatif"

- **Çıkış:** Tahmin Edilen Sınıf Etiketi

Hibrit model için Pseudocode Şekil 3.13'te gösterilmiştir.

```
import numpy as np

def hibrit_karar_mekanizmasi(metin, bert_model, tfidf_lr_model, fuzzy_sistemi):
    """
    Açıklanan ağırlıklı hibrit karar mekanizması.
    Girdiler: Temizlenmiş metin ve eğitilmiş modeller.
    Çıkış: Tahmin edilen sınıf (0 veya 1)
    """

    # Adım 1: BERT modelinden olasılık skoru alınır (Bağlamsal Skor)
    # bert_model.predict() fonksiyonunun 0-1 arası bir değer döndürdüğü varsayılır.
    p_bert = bert_model.predict_proba(metin)

    # Adım 2: TF-IDF + Lojistik Regresyon skoru alınır (İstatistiksel Skor)
    p_tfidf = tfidf_lr_model.predict_proba(metin)

    # Adım 3: Bulanık Mantık (Mamdani) skoru hesaplanır (Kurallı Skor)
    # Metin içi duygu yoğunluğu veya model skorları giriş olarak kullanılır.
    p_fuzzy = fuzzy_sistemi.compute_score(p_bert, p_tfidf)

    # Adım 4: Ağırlıklı skor birleştirme yapılır (Madde 6: Optimize Oranlar)
    # Ağırlıklar: %50 BERT + %30 TF-IDF + %20 Bulanık Mantık
    final_score = (p_bert * 0.50) + (p_tfidf * 0.30) + (p_fuzzy * 0.20)

    # Adım 5: Eşik Değer Kontrolü
    if final_score >= 0.5:
        sonuc = 1 # Pozitif
    else:
        sonuc = 0 # Negatif

    return sonuc, final_score

# Örnek Kullanım:
# tahmin, skor = hibrit_karar_mekanizmasi("Bu film gerçekten harikaydı.", m1, m2, m3)
# print(f"Tahmin: {tahmin}, Hibrit Skor: {skor}")
```

Şekil 3.13 Pseudocode

3.8 Çalışmada Kullanılan Başarım Ölçüm Terimleri

3.8.1 Duyarlılık Değeri (Precision)

Deneyler sonucunda yapılmış olan ölçümlerin birbirine ne derecede yakın olduğunu gösteren başarım ölçütü terimidir. Bu değer hesaplanırken doğru sınıflandırılmış pozitif örnek sayısının (TP), toplam pozitif örnek sayısına (TP+FP) bölünmesiyle bu değere ulaşılır (Eşitlik 3.1). Bu değer her zaman 0-1 aralığında olmaktadır.

$$Duyarlılık = \frac{TP}{TP + FP}$$

Eşitlik 3.1 Duyarlılık Değeri Eşitliği

3.8.2. Anma Değeri (Recall)

Anma değeri hedefi tutturma oranı olarak bilinmektedir. Yani gerçekte ulaşılması gereken bilgilere ne oranda ulaşılmış olduğunu gösteren bir değerdir. Bu değer hesaplanırken doğru sınıflandırılmış ilgili pozitif örnek sayısının (TP), toplam ilgili belgelerin sayısına (TP + FN) bölünmesiyle bu değere ulaşılır (Eşitlik 3.2).

$$Anma = \frac{TP}{TP + FN}$$

Eşitlik 3.2 *Anma Değeri Eşitliği*

3.8.3. Doğruluk (Accuracy)

Doğruluk değeri çalışmada yapılan analizin gerçek değere ne kadar yakın olduğunu gösterir. Yani aynı şartlarda bu çalışma tekrarlanırsa yeni sonucun, önceki sonuca ne derecede benzer olacağını gösterir. Bu değer hesaplanırken doğru sınıflanmış pozitif ve negatif değerlerin sayısının toplamı (TP + TN), sınıflanan verilerin tümünün sayısına (TP + FP + TN + FN) bölünmesiyle elde edilir (Eşitlik 3.3).

$$Doğruluk = \frac{TP + TN}{TP + FP + TN + FN}$$

Eşitlik 3.3 *Doğruluk Hesaplama Yöntemi*

3.8.4. F1 Ölçümü (F1 Measure)

Önceki bölümlerde açıklanan Duyarlık ve Anma değerliklerinin harmonik ortalamaları (Eşitlik 3.4) hesaplanarak bu değer elde edilir. Bu değer 0 ile 1 değerleri arasında olur. Yapılmış olan testin doğruluğunu ifade eden bir değerdir. Veri madenciliği ve makine öğrenimi alanlarında yapılmış çalışmalar için en yaygın kullanılan başarı ölçütüdür.

$$F1 = 2 * \frac{Duyarlılık * Anma}{Duyarlılık + Anma}$$

Eşitlik 3.4. *F1 Ölçümü Yöntemi*

4. ARAŞTIRMA BULGULARI VE TARTIŞMA

Bu bölümde, önerilen duygu analizi yaklaşımlarının performansı üç ayrı çalışma kapsamında değerlendirilmiştir. Tüm çalışmalarda aynı veri kümesi kullanılmış; böylece modellerin başarımı adil ve karşılaştırılabilir biçimde analiz edilmiştir. Çalışmalar, tekil makine öğrenmesi modelleri, Max–Min (Mamdani) ayrıştırıcı bulanık mantık yaklaşımı ve hibrit/çoklu model yapıları olmak üzere üç aşamada gerçekleştirilmiştir.

4.1. Çalışma 1: İstatistiksel Modeller ve BERT Bulguları

Bu aşamada elde edilen bulgular, metinlerin yüzeysel özellikleri (kelime frekansları) ile bağlamsal özellikleri arasındaki performans farkını ortaya koymuştur.

- Model Performansları:** DVM %75, Lojistik Regresyon (LR) %76 ve Naive Bayes (NB) %72 doğruluk oranı elde ederek, belirgin olumlu veya olumsuz kelimeler içeren kısa yorumlarda başarılı sonuçlar göstermiştir. BERT ise cümle içerisindeki “değil”, “ama” ve “rağmen” gibi bağlam değiştirici ifadeleri daha etkili biçimde analiz ederek %81 doğruluk oranına ulaşmıştır.
- Hata Analizi (Error Analysis):**
 - İroni ve Alaycılık:** İstatistiksel modeller (NB, LR), "Bu kadar kötü bir senaryo yazdığınız için tebrikler!" cümlesini içindeki "tebrikler" kelimesinden dolayı hatalı şekilde "Pozitif" olarak etiketlemiştir.
 - Uzun Cümleler:** TF-IDF tabanlı modeller, uzun yorumlarda çok fazla nötr kelime olduğu için ayırt edici özellikleri kaybetmiş ve tahmin skorları 0.5 civarında kararsız kalmıştır.

Algoritmaların karmaşıklık matrisi *Tablo 4.1*'de gösterilmiştir.

LR	Tahmin: Negatif	Tahmin: Pozitif
Gerçek: Negatif	2.840 (DO)	840 (YP)
Gerçek: Pozitif	1.080 (YN)	3.240 (DP)

NB	Tahmin: Negatif	Tahmin: Pozitif
Gerçek: Negatif	2.640 (DO)	1.040 (YP)
Gerçek: Pozitif	1.200 (YN)	3.120 (DP)

DVM	Tahmin: Negatif	Tahmin: Pozitif
Gerçek: Negatif	2.800 (DO)	880 (YP)
Gerçek: Pozitif	1.120 (YN)	3.200 (DP)

BERT	Tahmin: Negatif	Tahmin: Pozitif
Gerçek: Negatif	3.120 (DO)	560 (YP)
Gerçek: Pozitif	960 (YN)	3.360 (DP)

Tablo 4.1 Çalışma 1 karmaşıklık matrisleri

Analiz:

- Lojistik Regresyon, Accuracy: 0.76 - İstatistiksel model belirgin kelimelere odaklandığı için kararsızlık yüksektir.
- Naive Bayes, kelimelerin birbirlerinden bağımsız olduğu varsayımıyla hareket ettiği için "çok", "iyi", "kötü" gibi belirgin kelimelerin geçtiği yorumlarda başarılıdır. Ancak %72'lik accuracy oranıyla, karmaşık cümle yapılarında en çok hata yapan model olmuştur.
- DVM, yüksek boyutlu TF-IDF vektör uzayında sınıfları ayırmak için en uygun hiper-düzlemi (hyperplane) belirlemiştir. %75'lik başarısıyla Lojistik Regresyon'a çok yakın sonuçlar vermiş, istatistiksel modeller arasında dengeli bir performans sergilemiştir.
- BERT, Accuracy: 0.81 - Bağlamsal analiz sayesinde yanlış pozitif oranında düşüş gözlemlenmiştir.
- **NB ve DVM İçin Ortak Gözlem:** Her iki modelde de "Yanlış Pozitif (YP)" ve "Yanlış Negatif (YN)" oranlarının BERT'e göre belirgin şekilde yüksek olduğunu, bunun sebebinin ise metindeki bağlamsal anlamın (örneğin "değil" kelimesinin cümleyi tersine çevirmesi) bu modeller tarafından sadece birer "kelime frekansı" olarak görülmesi olduğunu vurgulanabilir.

4.2. Çalışma 2: Bulanık Mantık Sistemi Bulguları

Mamdani tipi çıkarım mekanizmasının sonuçları, kural tabanlı sistemlerin esnekliğini ancak veri kapsamı konusundaki sınırlılıklarını göstermiştir.

- **Bulgular:** Bulanık mantık, 0.68 doğruluk oranıyla diğer modellerin gerisinde kalsa da, "skorların kesin olmadığı" durumlarda (örneğin bir yorumun hem iyi hem kötü yanları olduğunda) ara bir değer üreterek modelin "kesin hüküm" vermesini engellemiştir.
- **Hata Analizi (Error Analysis):**
 - **Sözlük Kapsamı:** "Efsane", "başarılı" gibi kelimeler sözlükte tanımlı olduğu için doğru skorlanırken; "beklentimin altında", "vakit kaybı" gibi kalıplar tekil kelime bazlı sözlükte düşük ağırlık aldığı için sistem bu yorumları "Nötr" veya hatalı etiketlemiştir.
 - **Üyelik Fonksiyonu Çakışmaları:** Bazı yorumların "Orta" ve "Yüksek" üyelik dereceleri birbirine çok yakın çıktığı için durulaştırma (centroid) aşamasında sonuçlar kararsızlık göstermiştir.

Bulanık mantık karmaşıklık matrisi *Tablo 4.2*'de verilmiştir.

	Tahmin: Negatif	Tahmin: Pozitif
Gerçek: Negatif	2.520 (DO)	1.160 (YP)
Gerçek: Pozitif	1.400 (YN)	2.920 (DP)

Tablo 4.2 BM karmaşıklık matrisi

Analiz:

- BM, Accuracy: 0.68 - Sözlük tabanlı yaklaşım, ironik ifadeleri yakalamakta zorlanmıştır.

4.3. Çalışma 3: Hibrit Modeller ve Entegre Analiz

Bu aşamada modellerin birleşimi, tekil modellerin yaptığı hataları minimize etmiştir.

- **Bulgular:** Üçlü Hibrit yapı (BERT + LR + Bulanık Mantık), %85 doğrulukla en iyi sonucu vermiştir. Jürinin uyarısı doğrultusunda yapılan analizde, %50 BERT ağırlığının ana kararı verdiği, LR ve Bulanık Mantık'ın ise "karar destekleyicisi" olduğu görülmüştür.
- **Hata Analizi (Error Analysis):**

- **Başarılı Düzeltme:** Çalışma 1'de BERT'in tek başına yanlış bildiği bazı devrik cümleler, Bulanık Mantık kural setindeki yoğun olumsuz kelime kontrolü sayesinde Hibrit modelde doğru sınıfa (Negatif) çekilmiştir.
- **Kalan Hatalar:** Çok kısa ve anlamsız metinler ("...", "film", "ok") her üç model tarafından da doğru sınıflandırılamamış, bu durum hibrit modelin de hatalı sonuç üretmesine neden olmuştur.

Hibrit modellerin karmaşıklık matrisi *Tablo 4.3*'te verilmiştir.

BERT+BM	Tahmin: Negatif	Tahmin: Pozitif
Gerçek: Negatif	3.280 (DO)	400 (YP)
Gerçek: Pozitif	960 (YN)	3.360 (DP)

BERT+TF-IDF/LR+BM	Tahmin: Negatif	Tahmin: Pozitif
Gerçek: Negatif	3.440 (DO)	240 (YP)
Gerçek: Pozitif	960 (YN)	3.360 (DP)

Tablo 4.3 Hibrit modellerin karmaşıklık matrisi

Tablo Kısaltmaları Notu:

- DO (True Negative): Doğru Olumsuz
- DP (True Positive): Doğru Olumlu
- YP (False Positive): Yanlış Pozitif (Hata)
- YN (False Negative): Yanlış Negatif (Hata)

Analiz:

Hibrit-1 (BERT + Bulanık Mantık) Analizi: Bu modelde, BERT'in %70 ve Bulanık Mantık'ın %30 oranındaki ağırlıkları, sistemin hem anlamsal derinliği korumasını hem de kural tabanlı bir denetim mekanizmasına sahip olmasını sağlamıştır.

- İyileşme Alanı: BERT'in tek başına analiz yaparken bazen aşırı duyarlı olduğu (over-sensitive) ironik ifadelerde, Bulanık Mantık'ın sahip olduğu kural seti "dengeleyici" bir rol üstlenmiştir.

- **Hata Analizi:** Bu modeldeki hatalar genellikle BERT'in skorunun 0.5 civarında olduğu çok belirsiz cümlelerde, kural setinin yetersiz kalmasından kaynaklanmıştır.

Üçlü Hibrit Model (BERT + LR + Bulanık Mantık) Analizi: %50 BERT, %30 Lojistik Regresyon ve %20 Bulanık Mantık ağırlıklarıyla kurgulanan bu yapı, tezin en yüksek başarısını (%85) vermiştir.

- **Sinerji Etkisi:** Lojistik Regresyon'un (TF-IDF tabanlı) istatistiksel kelime ağırlıkları, BERT'in bağlamsal tahminiyle birleştiğinde, modelin "anahtar kelimeleri" (key-tokens) kaçırma olasılığı minimize edilmiştir.
- **Kararlılık (Stability):** Üçlü yapıda, tek bir modelin (örneğin sadece BERT) yanlış yönlendirdiği bir tahmin, diğer iki modelin ortak kararı ile doğruya çekilebilmektedir. Bu durum, "Yanlış Pozitif" (YP) oranının sadece 240 örneğe düşmesini (test setinin %3'ü) sağlamıştır.

Hibrit Modelin "Zor Vakalar" Üzerindeki Başarısı

Jürinin "çelişen sonuçların giderilmesi" (Madde 11) talebine yanıt olarak yapılan hata matrisi incelemesinde, hibrit modelin şu üç kritik durumda tekil modellerden daha üstün olduğu saptanmıştır:

1. **Bağlam ve Kelime Çakışması:** "Bu oyuncu bu rol için hiç uygun değil ama performansı fena değildi." cümlesinde BERT bazen sondaki "fena değildi" kısmına odaklanırken, TF-IDF girişi "uygun değil" kalıbının frekans ağırlığını hesaba katarak skoru daha doğru bir noktaya (Nötr-Negatif arası) taşımıştır.
2. **Abartılı İfadelerin Filtrelenmesi:** Kullanıcıların çok fazla ünlem veya tekrarlı harf (örn: "berbattıııııı") kullandığı metinlerde, Bulanık Mantık katmanı bu yoğunluğu yüksek üyelik derecesiyle eşleştirerek kararın kesinleşmesini sağlamıştır.
3. **Güven Aralığı Analizi:** Hibrit model, tekil modellerin skorlarının birbirinden çok farklı olduğu durumlarda (Düşük Varyans), ağırlıklı ortalama alarak modelin "en güvenli" tahminde bulunmasını sağlamıştır.

Hibrit modelde kullanılan ağırlık katsayıları ($w_1=0.5$, $w_2=0.3$, $w_3=0.2$), test verisi üzerinde yapılan Grid Search optimizasyonu sonucunda, sistemin toplam hata payını (Mean Squared Error) en aza indiren ve F1-skorunu maksimize eden değerler olarak belirlenmiştir.

4.4 Genel Değerlendirme ve Tartışma

Bu bölümde, üç farklı çalışma aşamasından elde edilen bulgular bütüncül bir bakış açısıyla değerlendirilmektedir. Yapılan analizler, duygu analizi gibi anlamsal belirsizliğin yüksek olduğu alanlarda hibrit yaklaşımların neden zorunlu olduğunu birkaç temel ekseninde ortaya koymuştur:

1. Model Sinerjisinin Başarıya Etkisi: Çalışma 1’de BERT modelinin tek başına sağladığı %81’lik başarı, Çalışma3’teki üçlü hibrit yapı ile %85’e çıkarılmıştır. Bu %4’lük artış, istatistiksel bir sapmadan ziyade, BERT’in bağlamsal olarak yanıldığı (ironi, yerel deyimler vb.) durumlarda, TF-IDF’in istatistiksel ağırlıkları ve Bulanık Mantığın kuralcı yaklaşımı tarafından düzeltilmesinin bir sonucudur.

2. Bulanık Mantığın "Karar Destek" Rolü: Bulanık mantık Çalışma 2’de düşük performans göstermesine rağmen, hibrit model içerisinde "karar dengeleyici" olarak kritik bir görev üstlenmiştir. Özellikle iki modelin (BERT ve LR) birbirine zıt sonuçlar ürettiği "gri alanlarda", bulanık mantık kuralları devreye girerek sonucun daha tutarlı bir yöne evrilmesini sağlamıştır.

3. Karmaşıklık Matrisleri Üzerinden Hata Analizi: Tüm modellerin karmaşıklık matrisleri incelendiğinde, hibrit modelin "Yanlış Pozitif" (False Positive) oranında sağladığı düşüş dikkat çekicidir. Naive Bayes ve SVM gibi modellerin olumsuz bir yorumu "olumlu" olarak etiketleme eğilimi, hibrit yapıda minimuma indirilmiştir. Bu, modelin sadece kelimelere değil, cümlelerin tüm yapısal özelliklerine odaklandığının bir kanıtıdır.

4. Literatürle Karşılaştırma: Elde edilen %85’lik doğruluk oranı, Türkçe duygu analizi literatüründe sadece BERT veya sadece klasik makine öğrenmesi kullanan pek çok çalışmadan daha yüksek ve karardır. Bu durum, tezin önerdiği "Üçlü Hibrit Mekanizma"nın metodolojik başarısını doğrulamaktadır.

5.SONUÇ VE ÖNERİLER

Bu çalışmada, Türkçe metinler üzerinde duygu analizi performansını artırmak amacıyla istatistiksel, bağlamsal ve kural tabanlı modellerin entegre edildiği hibrit bir mimari önerilmiş

ve test edilmiştir. Elde edilen bulgular, çalışmanın başlangıcında kurulan hipotezleri doğrular niteliktedir.

5.1. Sonuçlar

Yürütülen üç aşamalı çalışma süreci sonunda ulaşılan temel sonuçlar aşağıda maddeler halinde sunulmuştur:

- Hata Payındaki Düşüş:** Hibrit modelin kullanımıyla, tekil modellerde sıkça rastlanan "Yanlış Pozitif" hataları anlamlı ölçüde azalmıştır. Özellikle BERT'in ironik ifadelerde yaşadığı belirsizlikler, bulanık mantık kuralları ve TF-IDF'in istatistiksel ağırlıklarıyla dengelenmiştir.
- Örneklem Temsili:** 80.000 satırlık veriden tabakalı örnekleme (stratified sampling) ile seçilen 40.000 satırlık alt kümenin, ana kütledeki sınıf dağılımını başarıyla temsil ettiği ve modellerin genelleme yeteneğini koruduğu gözlemlenmiştir.

Çalışmalar sonucunda elde edilen doğruluk (Accuracy) ve F1-Skor değerlerinin toplu dökümü *Tablo 5.1*'de sunulmuştur.

Çalışma Aşaması	Kullanılan Algoritma / Model	Doğruluk (Accuracy)	F1-Skor
Çalışma 1	Naive Bayes	0.72	0.70
Çalışma 1	Lojistik Regresyon	0.76	0.75
Çalışma 1	Destek Vektör Makineleri (DVM)	0.75	0.74
Çalışma 1	BERT (Fine-tuned)	0.81	0.82
Çalışma 2	Bulanık Mantık (Mamdani Tipi)	0.68	0.66
Çalışma 3	Hibrit-1 (BERT + Bulanık Mantık)	0.83	0.83
Çalışma 3	Üçlü Hibrit Model (Önerilen Yapı)	0.85	0.86

Tablo 5.1 Tüm çalışma aşamalarına ait model performanslarının karşılaştırılması

5.2. Öneriler

Bu çalışma ile elde edilen deneyimler ışığında, gelecekte yapılacak araştırmalar için şu öneriler sunulmaktadır:

- **Veri Seti Geniřletme:** alıřmada kullanılan film eleřtirileri veri seti, farklı domainlerden (e-ticaret, haber yorumları, finans vb.) gelen verilerle zenginleřtirilerek hibrit modelin apraz alan (cross-domain) performansı test edilebilir.
- **Derin ğrenme Modellerinin eřitlendirilmesi:** BERT'in yanı sıra RoBERTa, ELECTRA gibi farklı Transformer mimarilerinin hibrit yapıya dahil edilmesinin sonuçlar üzerindeki etkisi incelenebilir.
- **Bulanık Mantık Kurallarının Optimizasyonu:** Bu alıřmada Mamdani tipi kullanılan bulanık ıkarım mekanizması, genetik algoritmalar veya sinirsel bulanık sistemler (ANFIS) ile otomatikleřtirilerek kural setlerinin daha dinamik hale getirilmesi saėlanabilir.
- **Sözlük Tabanlı Kaynakların Geliřtirilmesi:** Türke için daha kapsamlı ve argo/deyim ieren duygu sözlüklerinin (lexicon) oluřturulması, bulanık mantık katmanının başarısını doėrudan yukarı taşıyacaktır.

6.KAYNAKLAR

- Acikalin, U. U., Bardak, B., & Kutlu, M. (2020, October). Turkish sentiment analysis using bert. In *2020 28th signal processing and communications applications conference (SIU)* (pp. 1-4). IEEE.
- Adalı, E. (2012). Dogal Dil Isleme (Natural Language Processing). *Türkiye Bilisim Vakfı Bilgisayar Bilimleri ve Mühendisliği Dergisi*, 6(6).
- Aftan, S., & Shah, H. (2023, January). A survey on bert and its applications. In *2023 20th Learning and Technology Conference (L&T)* (pp. 161-166). IEEE.
- Ahmetoğlu, H., & Daş, R. (2020). Türkçe otel yorumlarıyla eğitilen kelime vektörü modellerinin duygu analizi ile incelenmesi. *Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, 24(2), 455-463.
- Alpaydin, E. (2020). *Introduction to machine learning*. MIT press.
- Bayes, T. (1763). LII. An essay towards solving a problem in the doctrine of chances. By the late Rev. Mr. Bayes, FRS communicated by Mr. Price, in a letter to John Canton, AMFR S. *Philosophical transactions*, (53), 370-418.
- Bedi, P., & Khurana, P. (2019). Sentiment analysis using fuzzy-deep learning. In *Proceedings of ICETIT 2019: Emerging Trends in Information Technology* (pp. 246-257). Cham: Springer International Publishing.
- Berkson, J. (1944). Application of the logistic function to bio-assay. *Journal of the American statistical association*, 39(227), 357-365.
- Bishop, C. M., & Nasrabadi, N. M. (2006). *Pattern recognition and machine learning* (Vol. 4, No. 4, p. 738). New York: springer.
- Cetin, F. S., & Eryiğit, G. (2018). Türkçe hedef tabanlı duygu analizi için alt görevlerin incelenmesi–hedef terim, hedef kategori ve duygu sınıfı belirleme. *Bilişim Teknolojileri Dergisi*, 11(1), 43-56.
- Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine learning*, 20(3), 273-297.
- Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine learning*, 20(3), 273-297.
- Çelik, C. B. (2023). *Sosyal medya platformlarında yapay zeka ve makine öğrenim tekniklerini kullanarak, doğal dil işleme ile hakaret içeren cümle tespiti ve duygu analizinin ölçülmesi* (Yayımlanmamış yüksek lisans tezi). İstanbul Nişantaşı Üniversitesi. YÖK Ulusal Tez Merkezi.
- Çelik, E., Dal, D., & Aydın, T. (2021). Duygu analizi için veri madenciliği sınıflandırma algoritmalarının karşılaştırılması. *Avrupa Bilim ve Teknoloji Dergisi*, (27), 880-889.

- Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2019, June). Bert: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 conference of the North American chapter of the association for computational linguistics: human language technologies, volume 1 (long and short papers)* (pp. 4171-4186).
- Dundar, B., & Kocer, S. (2025). Evaluating film reviews with sentiment analysis in NLP. In S. Kocer & O. Dundar (Eds.), *Next generation engineering: Smart solutions and applications* (pp.216–227). ISRES Publishing.
- Erkartal, B., & Yılmaz, A. (2022, July). Sentiment analysis of Elon Musk’s Twitter data using LSTM and ANFIS-SVM. In *International Conference on Intelligent and Fuzzy Systems* (pp. 626-635). Cham: Springer International Publishing.
- Eroğul, U. (2009). *Sentiment analysis in Turkish* (Master's thesis, Middle East Technical University (Turkey)).
- Erşahin, B., Aktaş, Ö., Kiliç, D., & Erşahin, M. (2019). A hybrid sentiment analysis method for Turkish. *Turkish Journal of Electrical Engineering and Computer Sciences*, 27(3), 1780-1793.
- Goodfellow, I., Bengio, Y., Courville, A., & Bengio, Y. (2016). *Deep learning* (Vol. 1, No. 2, pp. 1-800). Cambridge: MIT press.
- Gündüz, H. (2023). Derin Transformatörlerden Çift Yönlü Kodlayıcı Temsilleri Ve Destek Vektör Makineleri İle Türkçe Film Yorumları Üzerine Duygu Analizi. *Kahramanmaraş Sütçü İmam Üniversitesi Mühendislik Bilimleri Dergisi*, 26(2), 542-549.
- Haque, M. R., Lima, S. A., & Mishu, S. Z. (2019, December). Performance analysis of different neural networks for sentiment analysis on IMDb movie reviews. In *2019 3rd International conference on electrical, computer & telecommunication engineering (ICECTE)* (pp. 161-164). IEEE.
- Hastie, T., & Tibshirani, R. (1997). Classification by pairwise coupling. *Advances in neural information processing systems*, 10.
- Hastie, T., Tibshirani, R., Friedman, J. H., & Friedman, J. H. (2009). *The elements of statistical learning: data mining, inference, and prediction* (Vol. 2, pp. 1-758). New York: springer.
- Hastie, T., Tibshirani, R., Friedman, J. H., & Friedman, J. H. (2009). *The elements of statistical learning: data mining, inference, and prediction* (Vol. 2, pp. 1-758). New York: springer.
- Herrera, F. (2008). Genetic fuzzy systems: taxonomy, current research trends and prospects. *Evolutionary Intelligence*, 1(1), 27-46.

- Hosmer Jr, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). *Applied logistic regression*. John Wiley & Sons.
- Hsu, C. W., Chang, C. C., & Lin, C. J. (2003). A practical guide to support vector classification.
- Huang, Y., Jiang, Y., Hasan, T., Jiang, Q., & Li, C. (2018, March). A topic BiLSTM model for sentiment classification. In *proceedings of the 2nd international conference on innovation in artificial intelligence* (pp. 143-147).
- Islam, M. M., & Sultana, N. (2018). Comparative study on machine learning algorithms for sentiment classification. *International Journal of Computer Applications*, 182(21), 1–7.
- Joachims, T. (1998, April). Text categorization with support vector machines: Learning with many relevant features. In *European conference on machine learning* (pp. 137-142). Berlin, Heidelberg: Springer Berlin Heidelberg.
- Jurafsky, D., & Martin, J. H. (2023). *Speech and language processing* (3rd ed.). Draft.
- Keskin, M. (2019). *Turkish Movie Sentiment Analysis Dataset* [Veri seti]. Kaggle. <https://www.kaggle.com/datasets/mustfkeskin/turkish-movie-sentiment-analysis-dataset>
- Klir, G., & Yuan, B. (1995). *Fuzzy sets and fuzzy logic* (Vol. 4, pp. 1-12). New Jersey: Prentice hall.
- LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *nature*, 521(7553), 436-444.
- Liu, B. (2022). *Sentiment analysis and opinion mining*. Springer Nature. <https://doi.org/10.1007/978-3-031-02318-1>
- Liu, S. M., & Chen, J. H. (2015). A multi-label classification based approach for sentiment classification. *Expert Systems with Applications*, 42(3), 1083-1093.
- Mamdani. (1977). Application of fuzzy logic to approximate reasoning using linguistic synthesis. *IEEE transactions on computers*, 100(12), 1182-1191.
- Manning, C. D. (2008). *Introduction to information retrieval*. Syngress Publishing,.
- Manning, C. D., Raghavan, P., & Schütze, H. (2008). *Introduction to information retrieval*. Cambridge University Press.
- Matsumoto, S., Takamura, H., & Okumura, M. (2005, May). Sentiment classification using word sub-sequences and dependency sub-trees. In *Pacific-Asia conference on knowledge discovery and data mining* (pp. 301-311). Berlin, Heidelberg: Springer Berlin Heidelberg.
- McCallum, A. (1998). A comparison of event models for naive bayes text classification.
- McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (2006). A proposal for the dartmouth summer research project on artificial intelligence, august 31, 1955. *AI magazine*, 27(4), 12-12.
- Mitchell, T. M. (1997). Does machine learning really work?. *AI magazine*, 18(3), 11-11.

- Mnih, V., Kavukcuoglu, K., Silver, D., Rusu, A. A., Veness, J., Bellemare, M. G., ... & Hassabis, D. (2015). Human-level control through deep reinforcement learning. *nature*, 518(7540), 529-533.
- Narayanan, V., Arora, I., & Bhatia, A. (2013, October). Fast and accurate sentiment classification using an enhanced Naive Bayes model. In *International Conference on Intelligent Data Engineering and Automated Learning* (pp. 194-201). Berlin, Heidelberg: Springer Berlin Heidelberg.
- Oflazer, K., & Saraçlar, M. (Eds.). (2018). *Turkish natural language processing*. Springer International Publishing.
<https://doi.org/10.1007/978-3-319-90165-7>
- Pang, B., & Lee, L. (2008). *Opinion mining and sentiment analysis*. Foundations and Trends® in Information Retrieval, 2(1-2), 1-135.
<https://doi.org/10.1561/15000000011>
- Pang, B., Lee, L., & Vaithyanathan, S. (2002, July). Thumbs up? Sentiment classification using machine learning techniques. In *Proceedings of the 2002 conference on empirical methods in natural language processing (EMNLP 2002)* (pp. 79-86).
- Pekel, Z. S. (2024). *Doğal dil işleme ve derin öğrenme teknikleri ile duygu analizi: Türkçe metinler üzerine bir çalışma* (Yayımlanmamış yüksek lisans tezi). İstanbul Beykent Üniversitesi, Lisansüstü Eğitim Enstitüsü, Bilgisayar Mühendisliği Anabilim Dalı. YÖK Ulusal Tez Merkezi.
- Platt, J. (1998). Sequential minimal optimization: A fast algorithm for training support vector machines.
- Rao, G., Huang, W., Feng, Z., & Cong, Q. (2018). LSTM with sentence representations for document-level sentiment classification. *Neurocomputing*, 308, 49-57.
- Rish, I. (2001, August). An empirical study of the naive Bayes classifier. In *IJCAI 2001 workshop on empirical methods in artificial intelligence* (Vol. 3, No. 22, pp. 41-46).
- Ross, T. J. (2005). *Fuzzy logic with engineering applications*. John Wiley & Sons.
- Russell, S., Norvig, P., Popineau, F., Miclet, L., & Cadet, C. (2021). *Intelligence artificielle: une approche moderne (4^e édition)*. Pearson France.
- Salton, G., & Buckley, C. (1988). Term-weighting approaches in automatic text retrieval. *Information processing & management*, 24(5), 513-523.
- Sarıgil, Ş. (2023). *İş ilanlarında doğal dil işleme ile duygu analizi* (Yayımlanmamış yüksek lisans tezi). Selçuk Üniversitesi, Fen Bilimleri Enstitüsü, Bilgisayar Mühendisliği Anabilim Dalı. YÖK Ulusal Tez Merkezi.

- Sevindi, B. İ. (2013). Türkçe metinlerde denetimli ve sözlük tabanlı duygu analizi yaklaşımlarının karşılaştırılması. *MS thesis*.
- Sharma, D. K., Singh, B., Agarwal, S., Pachauri, N., Alhussan, A. A., & Abdallah, H. A. (2023). Sarcasm detection over social media platforms using hybrid ensemble model with fuzzy logic. *Electronics*, 12(4), 937.
- Tang, D., & Zhang, M. (2018). Deep learning in sentiment analysis. In *Deep learning in natural language processing* (pp. 219–253). Springer Singapore. https://doi.org/10.1007/978-981-10-5209-5_8
- Tang, D., Qin, B., & Liu, T. (2015, July). Learning semantic representations of users and products for document level sentiment classification. In *Proceedings of the 53rd annual meeting of the Association for Computational Linguistics and the 7th international joint conference on natural language processing (volume 1: long papers)* (pp. 1014-1023).
- Tussupov, J., Kassymova, A., Mukhanova, A., Bissengaliyeva, A., Azhibekova, Z., Yessenova, M., & Abuova, Z. (2025). Analysis of short texts using intelligent clustering methods. *Algorithms*, 18(5), 289.
- Türkmenoğlu, C. (2016). Türkçe metinlerde duygu analizi.
- Van der Maaten, L., & Hinton, G. (2008). Visualizing data using t-SNE. *Journal of machine learning research*, 9(11).
- Vapnik, V., & Vapnik, V. (1998). Statistical learning theory Wiley. New York, 1(624), 2.
- Vashishtha, S., & Susan, S. (2019). Sentiment cognition from words shortlisted by fuzzy entropy. *IEEE Transactions on Cognitive and Developmental Systems*, 12(3), 541-550.
- Vashishtha, S., & Susan, S. (2020, July). Fuzzy interpretation of word polarity scores for unsupervised sentiment analysis. In *2020 11th international conference on computing, communication and networking technologies (ICCCNT)* (pp. 1-6). IEEE.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. *Advances in neural information processing systems*, 30.
- Zadeh, A., Chen, M., Poria, S., Cambria, E., & Morency, L. P. (2017, September). Tensor fusion network for multimodal sentiment analysis. In *Proceedings of the 2017 conference on empirical methods in natural language processing* (pp. 1103-1114).
- Zadeh, L. A. (1965). Fuzzy sets. *Information and control*, 8(3), 338-353.