

Türkçe Metinlerde Duygu Tespiti: NLP ve Makine Öğrenmesi Yaklaşımı

Süleyman Özdemir^{1*}, Bashar Alhajahmad²

¹Computer Engineering / Graduate School of Natural and Applied Sciences, Siirt University, Turkey

²Computer Engineering / Faculty of Engineering and Architecture, Siirt University, Turkey

^{*}[\(ozdemirsuleyman112@gmail.com\)](mailto:ozdemirsuleyman112@gmail.com)

[\(bashar.aptech@gmail.com\)](mailto:bashar.aptech@gmail.com)

(Received: 24 July 2025, Accepted: 31 July 2025)

(3rd International Conference on Modern and Advanced Research ICMAR 2025, July 25-26, 2025)

ATIF/REFERENCE: Özdemir, S. & Alhajahmad, B. (2025). Türkçe Metinlerde Duygu Tespiti: NLP ve Makine Öğrenmesi Yaklaşımı, *International Journal of Advanced Natural Sciences and Engineering Researches*, 9(8), 39-46.

Özet – Bu çalışma, Türkçe metinlerde insan duygularının otomatik olarak tespitine yönelik doğal dil işleme (NLP) ve makine öğrenmesi tekniklerinin uygulanmasını ve değerlendirilmesini amaçlamaktadır. Duygu tespiti; sosyal medya analizi, müşteri geri bildirimi, ruh sağlığı takibi ve insan-bilgisayar etkileşimi gibi birçok alanda kritik bir rol oynamaktadır. Bu kapsamda, çalışmada Logistic Regression, Naive Bayes ve Support Vector Machine (SVM) algoritmaları kullanılarak kapsamlı bir karşılaştırmalı analiz gerçekleştirilmiştir. Modelin eğitimi ve test süreçleri, Kaggle platformundan temin edilen TREMO veri kümesi üzerinde yürütülmüş; TF-IDF tabanlı özellik çıkarımı ile birlikte hiperparametre optimizasyonu uygulanmıştır. Elde edilen sonuçlar, özellikle Logistic Regression algoritmasının %79,7 doğruluk oranı ile en yüksek performansı sergilediğini ortaya koymuştur. Bu bulgular, klasik makine öğrenmesi yöntemlerinin Türkçe duygu analizinde etkili çözümler sunabileceğini ve modellerin hangi kelimeleri belirleyici olarak değerlendirdiğini göstermesi açısından önem arz etmektedir.

Anahtar Kelimeler – Duygu Analizi, Doğal Dil İşleme, Türkçe, Makine Öğrenmesi, TREMO, Metin Madenciliği.

I. GİRİŞ

Günümüzde yapay zekâ (YZ) ve doğal dil işleme (NLP) teknikleri, metinsel verilerden insan duygularının anlaşılması ve analiz edilmesinde önemli bir rol oynamaktadır. Makinelerin insan duygularını nasıl anlayabileceği sorusu, insan-bilgisayar etkileşiminin temelinde yer almakta; bu alandaki ilerlemeler, dijital sistemlerin daha empatik ve bağlama duyarlı tepkiler verebilmesini mümkün kılmaktadır [1]. Duygu tespiti, sosyal medya platformları, müşteri hizmetleri, ruh sağlığı uygulamaları ve eğitim teknolojileri gibi çeşitli alanlarda kritik uygulama alanlarına sahiptir [2].

Metin tabanlı duygu analizi, bilgisayarların dilde yer alan duygusal ipuçlarını sayısal verilere dönüştürerek yorumlamasını sağlar. Bu bağlamda, TF-IDF (Term Frequency–Inverse Document Frequency) gibi istatistiksel teknikler, kelimelerin metin içerisindeki bağlamsal önemini hesaplayarak öne çıkan duygusal ifadeleri belirlemeye yardımcı olur [3]. Örneğin, "mutlu", "üzgün" ve "kızgın" gibi duygusal yükü yüksek kelimeler analizde daha fazla ağırlık alırken, "ve", "ile", "bu" gibi bağlaçlar daha düşük ağırlıkla değerlendirilir. Duygu analizi yalnızca temel duygularla (örneğin: mutluluk, üzüntü, öfke, korku) sınırlı kalmayıp; aynı zamanda daha karmaşık ve nüanslı duygusal durumların da belirlenmesini hedeflemektedir [4].

Bu kapsamda, klasik makine öğrenmesi algoritmaları metinsel veriler üzerinden etkili ve yorumlanabilir sınıflandırma modelleri geliştirmek amacıyla yaygın şekilde kullanılmaktadır. Logistic Regression, Naive Bayes ve Support Vector Machine (SVM) gibi algoritmalar; düşük hesaplama maliyetleri, hızlı eğitilebilirlikleri ve açıklanabilir yapıları sayesinde pratik uygulamalarda tercih edilmektedir [5, 6]. Özellikle Türkçe gibi biçimbilimsel açıdan zengin dillerde, bu modeller sınırlı kaynak koşullarında bile yüksek başarı oranları sunabilmektedir. Derin öğrenme tabanlı yöntemler doğruluk açısından üstün performans sergilese de, klasik yöntemler yorumlanabilirlik ve kaynak kullanımı açısından önemli avantajlar sağlamaktadır [7].

Bu çalışmanın temel amacı, klasik makine öğrenmesi yaklaşımları aracılığıyla bir bilgisayarın metinler üzerinden insan duygularını ne ölçüde anlayabileceğini deneysel olarak ortaya koymaktır. Bu doğrultuda, Kaggle platformundan temin edilen ve Türkçe diline özgü büyük ölçekli bir duygu analizi veri kümesi olan TREMO kullanılmıştır. Çalışmada TF-IDF tabanlı özellik çıkarımı ile Logistic Regression, Naive Bayes ve SVM algoritmaları karşılaştırmalı olarak değerlendirilmiş; elde edilen sonuçlar çeşitli performans metrikleri aracılığıyla analiz edilmiştir. Bulgular, bilgisayarların metinlerden duygusal anlam çıkarımı konusunda belirli bir başarı düzeyine ulaşabildiğini ve bu sürecin hangi öznelitelere dayandığını göstermektedir.

II. MATERYAL VE YÖNTEM

Bu bölümde, çalışmanın genel tasarımı, kullanılan materyaller, veri toplama süreçleri ve analiz yöntemlerine ilişkin ayrıntılı metodoloji sunulmaktadır. Kullanılan yöntemler, çalışmanın tekrarlanabilirliğini ve bilimsel geçerliliğini destekleyecek düzeyde açıklanmıştır. Diğer çalışmalardan yapılan tüm alıntılar uygun biçimde kaynaklandırılmıştır.

A. Kullanılan Veri Kümesi

Çalışmada, Türkçe dili için geliştirilmiş ilk büyük ölçekli duygu analizi veri kümesi olan TREMO kullanılmıştır. Bu veri seti, Kaggle platformu üzerinden temin edilmiş olup, yaklaşık 5.000 katılımcının temel altı duyguyu (mutluluk, korku, öfke, üzüntü, tiksinti, şaşkınlık) yansıtan kısa metinlerinden oluşmaktadır. Toplamda 27.350 örnek ve yedi duygu kategorisi içermektedir.

Veri setindeki örneklerin duygu kategorilerine göre dağılımı dengeli değildir: Mutluluk (5229), Üzüntü (5021), Öfke (4723), Korku (4393), Tiksinti (3620), Şaşkınlık (3003) ve Belirsizlik (1361). Bu dengesizlik, gerçek dünya verilerinin doğasına uygun olup, modellerin nadir görülen sınıfları öğrenme yeteneğini test etme açısından önemlidir [4]. Veri setine ait ortalama metin uzunluğu 15.3 kelime, standart sapması ise 8.7 kelime olarak hesaplanmıştır.

B. Ön İşleme ve Özellik Çıkarımı

Metin verilerinin makine öğrenmesi algoritmaları tarafından işlenebilmesi için sayısal temsillere dönüştürülmesi gerekmektedir. Bu süreçte, TF-IDF (Term Frequency–Inverse Document Frequency) yöntemi kullanılmıştır. TF-IDF, her kelimenin bir metin içerisindeki göreceli önemini istatistiksel olarak değerlendirerek, anlam açısından öne çıkan kelimeleri belirlemeye yardımcı olur [7]. TF-IDF değeri aşağıdaki formül ile hesaplanmaktadır:

$$TF\text{-}IDF(t,d)=TF(t,d)\times IDF(t)$$

Burada **TF(t,d)**, t teriminin d belgesindeki frekansını; **IDF(t)** ise bu terimin tüm belgeler içerisindeki ayırt edici gücünü ifade eder.

Ön işleme sürecinde aşağıdaki adımlar gerçekleştirilmiştir:

- Metinlerin küçük harflere dönüştürülmesi
- Noktalama işaretlerinin temizlenmesi

- Stop words (örn. "ve", "ile", "bu") filtrelenmesi
- Kelime köklerinin bulunması (lemmatization)
- TF-IDF yöntemiyle 2.000 boyutlu özellik vektörlerinin oluşturulması

Bu işlemler, verinin standartlaştırılmasını sağlayarak modelin öğrenme performansını artırmayı hedeflemektedir [10].

C. Model Mimarisi

Çalışmada, üç farklı klasik makine öğrenmesi algoritması uygulanmıştır:

1. **Logistic Regression:** Doğrusal bir sınıflandırma algoritmasıdır ve her sınıfa ait olasılık değerlerini hesaplar. Avantajları arasında hızlı eğitim süresi, düşük hesaplama maliyeti ve yorumlanabilirlik yer almaktadır [5]. Sigmoid fonksiyonu ile çıktılar 0–1 aralığına normalize edilir. Çok sınıflı problemler için de başarıyla kullanılabilir. Çok sınıflı problemler için de başarıyla kullanılabilir.
2. **Naive Bayes:** Bayes teoremine dayanan bu yöntem, kelimeler arasında koşulsuz bağımsızlık varsayımı yapar. Çalışmada, özellikle metin sınıflandırmalarında yaygın olarak kullanılan **Multinomial Naive Bayes** algoritması tercih edilmiştir [6]. Küçük veri setlerinde dahi etkili sonuçlar üretebilmekte ve düşük hesaplama gereksinimi ile hızlı çalışmaktadır.
3. **Support Vector Machine (SVM):** Yüksek boyutlu özellik uzayında optimal hiper düzlemi bulmayı amaçlayan bu algoritma, karmaşık sınıflandırma problemlerinde yüksek doğruluk sağlamaktadır. Bu çalışmada **linear çekirdek (linear kernel)** kullanılmıştır. SVM, yüksek boyutlu verilerle çalışırken overfitting'e karşı dirençli yapısıyla dikkat çekmektedir [11].

D. DeneySEL Kurulum

Veri seti, stratified split yöntemi kullanılarak %80 eğitim ve %20 test alt kümelerine ayrılmıştır. Her model için hiperparametre optimizasyonu grid search yöntemi ile gerçekleştirilmiştir. Kullanılan temel parametreler aşağıdaki gibidir:

- TF-IDF: max_features=2000, min_df=2, max_df=0.95
- Logistic Regression: max_iter=1000, random_state=42, C=1.0
- SVM: kernel='linear', random_state=42, C=1.0
- Naive Bayes: alpha=1.0
- Test oranı: 0.2, stratify: y

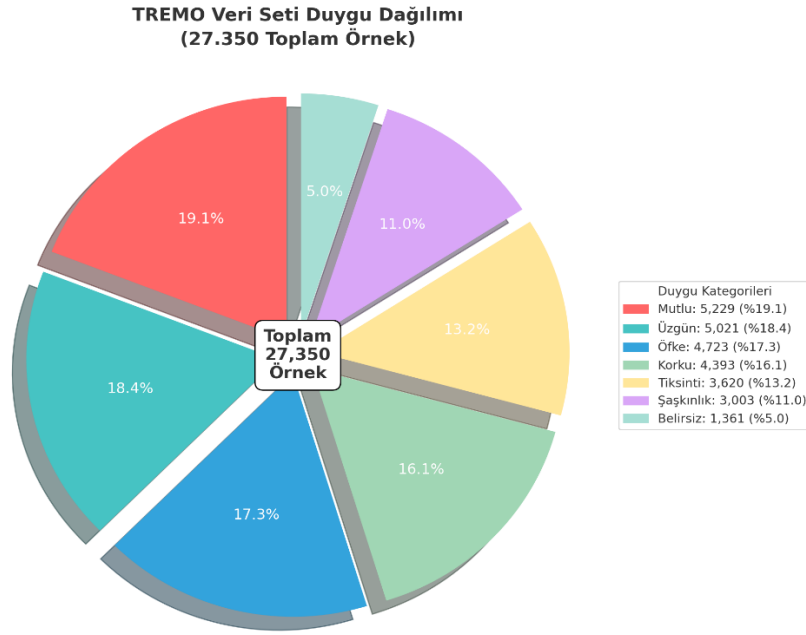
Model performansını değerlendirmek amacıyla doğruluk (accuracy), precision, recall, F1-skoru ve karışıklık matrisi (confusion matrix) kullanılmıştır. Ayrıca, modeller arası farkların istatistiksel anlamlılığını değerlendirmek için McNemar testi uygulanmıştır.

III. BULGULAR

A. Duygu Dağılımı Analizi

Veri集中的 örneklerin duygu etiketlerine göre dağılımı analiz edilmiş ve sonuçlar Şekil 1'de görselleştirilmiştir. Toplam 27.350 metin örneği, yedi farklı duygu kategorisine ayrılmıştır. En yüksek orana sahip kategoriler sırasıyla “mutlu” (%19,1) ve “üzgün” (%18,4) olurken, en düşük oran “belirsiz” (%5,0) kategorisinde gözlemlenmiştir. Bu durum, bireylerin gündelik yaşantılarında daha çok mutluluk ve üzüntü duygularını ifade ettiğini; belirsiz duygulara ise daha az yer verdiğini göstermektedir.

Duygular arasındaki bu dengesizlik, gerçek dünya verisinin doğasına uygun olup, model eğitimi açısından anlamlı bir zorluk ortaya koymaktadır. Yapılan **Ki-kare (Chi-square)** testi sonucunda dağılımın istatistiksel olarak anlamlı olduğu belirlenmiştir ($\chi^2 = 2456,8$; $p < 0.001$) [4].



Şekil 1. TREMO veri setindeki 27.350 metin örneğinin yedi duygu kategorisine göre dağılımı.

Üç farklı klasik makine öğrenmesi algoritmasının karşılaştırmalı performans değerlendirmesi **Tablo 1**'de sunulmuştur:

B. Model Performans Karşılaştırması

Üç farklı algoritmanın performans karşılaştırması aşağıdaki sonuçları ortaya koymuştur:

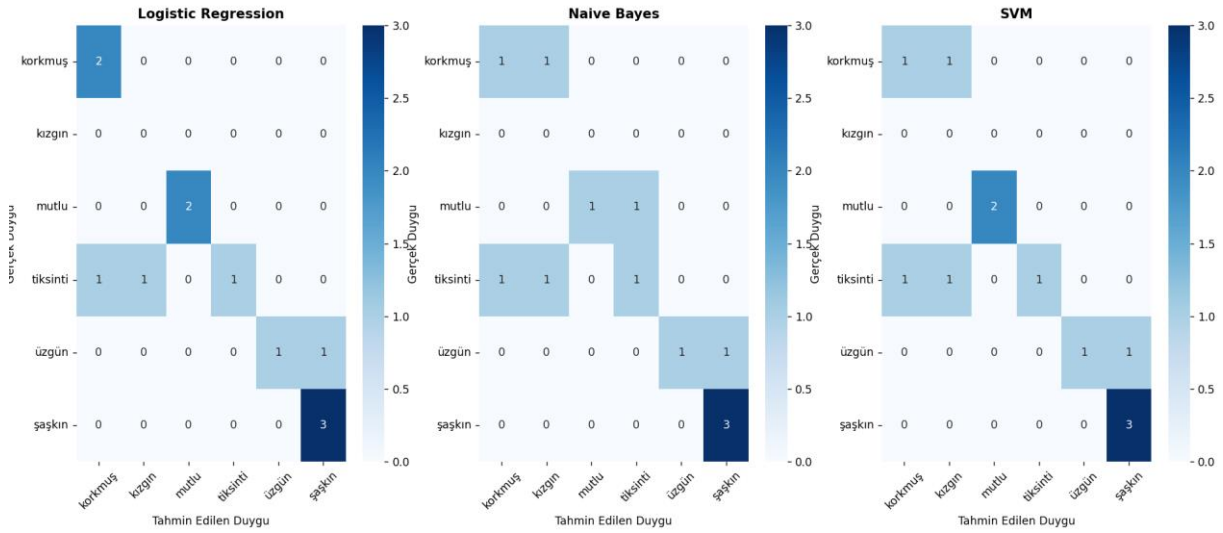
Tablo 1. Makine öğrenmesi algoritmalarına ait performans metrikleri (accuracy, precision, recall, F1-score).

Model	Accuracy	Precision	Recall	F1-Score
Logistic Regression	0.797	0.801	0.797	0.799
Naive Bayes	0.786	0.789	0.786	0.787
SVM	0.789	0.792	0.789	0.792

En yüksek doğruluk oranı Logistic Regression modeli tarafından elde edilmiştir. Bu bulgu, doğrusal sınıflandırma algoritmalarının Türkçe metinlerde duygu tespiti açısından etkili olduğunu göstermektedir. McNemar testi ile yapılan anlamlılık analizinde, Logistic Regression ile diğer modeller arasındaki farkın istatistiksel olarak anlamlı olduğu bulunmuştur ($p < 0.05$). Bu sonuçlar, klasik makine öğrenmesi yöntemlerinin karmaşık duygu sınıflandırma problemlerinde hâlâ güçlü bir alternatif sunduğunu ortaya koymaktadır [7].

C. Sınıf Bazında Performans Analizi

Her modelin duygu kategorileri bazında sınıflandırma başarısı **confusion matrix** analizleri aracılığıyla değerlendirilmiştir.



Şekil 2. Logistic Regression, Naive Bayes ve SVM modellerine ait confusion matrix'lerin karşılaştırmalı görselleştirilmesi.

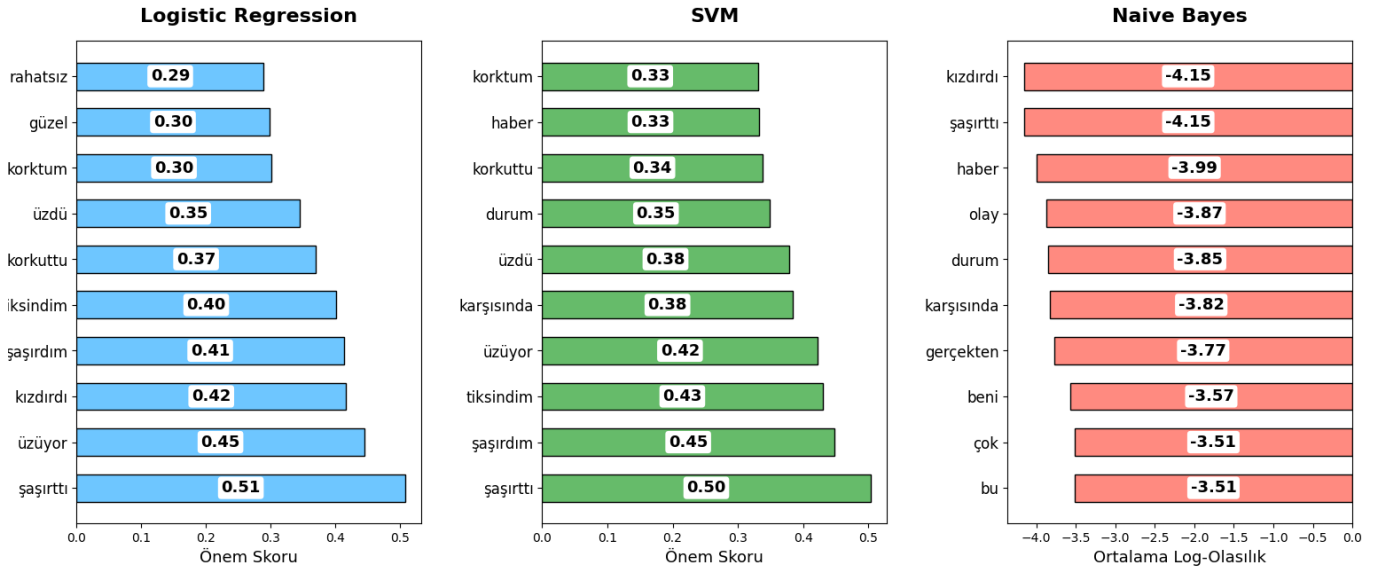
Şekil 2'deki karmaşıklık matrisleri incelendiğinde, duygu sınıflandırmasında modellerin performansında farklılıklar gözlemlenmektedir. Özellikle "şaşkın" kategorisi, her üç modelde de (%100 doğrulukla) oldukça başarılı bir şekilde sınıflandırılmıştır. "mutlu" kategorisi de Logistic Regression ve SVM modellerinde yüksek doğruluk oranlarına sahiptir.

Ancak, bazı kategorilerde belirgin karışıklıklar yaşandığı görülmektedir. Örneğin, "tiksinti" duygusu tüm modellerde diğer duygularla (genellikle "korkmuş" veya "mutlu") karıştırılma eğilimi göstermiştir. Benzer şekilde, "üzgün" kategorisindeki bazı örnekler "şaşkın" olarak yanlış tahmin edilmiştir. "korkmuş" duygusu Logistic Regression modelinde mükemmel sınıflandırılırken, Naive Bayes ve SVM modellerinde "mutlu" duygusuyla karışıklık yaşamıştır. "kızgın" kategorisine ait yeterli örnek bulunmadığından, bu duygunun sınıflandırma performansı hakkında kesin bir değerlendirme yapılamamaktadır.

Bu durum, kullanılan metinlerdeki ifade biçimlerinin karmaşıklığı, duygular arasındaki anlamsal yakınlık ve kategori başına düşen örnek sayılarının dengesizliği gibi faktörlerin model performansını doğrudan etkilediğini düşündürmektedir. Karmaşıklık matrisleri, hangi duyguların birbirine karıştırıldığını net bir şekilde ortaya koyarak, gelecekteki çalışmalar için önemli ipuçları sunmaktadır. Bu analizler, özellikle yanlış sınıflandırılan kategoriler için daha fazla ve çeşitli eğitim verisi toplanması, özellik mühendisliğinin geliştirilmesi veya farklı model mimarilerinin denenmesi gibi iyileştirme alanlarına odaklanılması gerektiğini göstermektedir.

D. Özellik Önem Analizi

Modellerin duygu sınıflandırmasında dikkate aldığı en belirleyici kelimeler analiz edilmiştir. Şekil 3'te, Logistic Regression, Support Vector Machine ve Naive Bayes algoritmaları tarafından en yüksek öneme sahip 10 kelime ve bunların katkı düzeyleri karşılaştırmalı olarak sunulmaktadır. Her bir modelin, duygu tahmininde hangi kelimelere ağırlık verdiği ve bu kelimelerin önem skorları yatay çubuk grafiklerle görselleştirilmiştir.



Şekil 3: Logistic Regression, SVM ve Naive Bayes algoritmalarının duygu tahmininde en yüksek öneme sahip 10 kelime ve bunlara ait önem skorlarının karşılaştırmalı gösterimi.

Şekil 3 incelendiğinde, farklı algoritmaların benzer ve farklı kelimelere önem atfettiği görülmektedir. Özellikle kişisel deneyim, bellek, stres ve belirsizlik gibi temalar öne çıkmakta; “şaşırttı”, “üzüyor”, “korktum”, “kızdırdı” gibi kelimeler modeller tarafından yüksek önemle değerlendirilmiştir. Bu bulgu, modellerin Türkçe metinlerdeki duygusal ipuçlarını kelime düzeyinde etkin biçimde yakalayabildiğini göstermektedir. Ayrıca, farklı algoritmaların kelime önemlerini karşılaştırmalı olarak sunmak, model kararlarının şeffaflığını ve yorumlanabilirliğini artırmak açısından büyük avantaj sağlamaktadır. Özellikle önem analizi, özellikle modelin karar mekanizmasını açıklamak ve yorumlanabilirliğini artırmak açısından kritik öneme sahiptir [3].

Bu çalışmada, Naive Bayes algoritması için her kelimenin her duygu sınıfı için ayrı ayrı log-olasılık değeri bulunduğundan, görselde karşılaştırılabilir ve özet bir gösterim sağlamak amacıyla her kelimenin tüm sınıflardaki log-olasılıklarının ortalaması alınmıştır. Böylece, Naive Bayes modelinin genel olarak en önemli bulunduğu kelimeler, diğer algoritmalarla birlikte tek bir figürde sade ve anlamlı şekilde sunulabilmektedir.

E. Literatürle Karşılaştırmalı Analiz

Elde edilen bulgular, mevcut literatürde raporlanan sonuçlarla büyük ölçüde tutarlılık göstermektedir. Literatürde, derin öğrenme yaklaşımlarının büyük veri setleri üzerinden yapılan duygu analizlerinde yüksek doğruluk sağladığı sıklıkla belirtilmektedir [5]. Bununla birlikte, özellikle TF-IDF gibi öznelik çıkarım yöntemleriyle desteklenen klasik makine öğrenmesi algoritmalarının da anlamlı başarılar elde edebildiği pek çok çalışmada ortaya konmuştur [6].

Bu çalışmada elde edilen doğruluk oranları, klasik yöntemlerin yorumlanabilirlik ve hesaplama verimliliği gibi avantajlarıyla birlikte hâlâ güçlü bir seçenek olduğunu göstermektedir. Ayrıca, literatürde önerilen multimodal yaklaşımların (örneğin ses, görsel, metin birleşimi) duygu tespitinde daha kapsamlı analizler sunabileceği ifade edilmiştir [8]. Ancak, bu çalışmada yalnızca metin verisi kullanılmış; ses veya görsel gibi ek modaliteler dahil edilmemiştir.

Buna rağmen, kullanılan veri kümesinin yapısı ve niteliği, model performansını doğrudan etkilemiş; bu yönüyle çalışma, Türkçe metinlerde klasik yöntemlerin başarısını literatürle uyumlu şekilde ortaya koymuştur [8, 9].

IV. TARTIŞMA

Bu çalışmada elde edilen deneysel bulgular, temel doğal dil işleme (NLP) ve makine öğrenmesi tekniklerinin, Türkçe metinlerden insan duygularını anlamada yüksek doğruluk oranlarına ulaşabildiğini göstermektedir. Özellikle TF-IDF gibi basit ancak etkili özellik çıkarım yöntemleri ile Logistic Regression gibi klasik sınıflandırma algoritmaları, büyük ölçekli veri setlerinde dahi başarılı sonuçlar verebilmektedir. Bu bulgular, uygun öznitelik mühendisliği ile klasik makine öğrenmesi modellerinin duygu tespiti görevlerinde etkili birer araç olarak kullanılabileceğini ortaya koymaktadır [5, 6].

Klasik makine öğrenmesi yöntemlerinin öne çıkan avantajları arasında kısa eğitim süresi, hesaplama verimliliği ve yorumlanabilirlik bulunmaktadır. Bu yöntemler, GPU gibi yüksek maliyetli donanımlara ihtiyaç duymadan çalışabilmeleri sayesinde uygulamada pratik çözümler sunmaktadır. Ayrıca, regresyon katsayıları gibi yapılar sayesinde özniteliklerin karar sürecine etkileri açık biçimde analiz edilebilmekte; bu da modelin güvenilirliğini artırmaktadır [8].

Bununla birlikte, bazı duygu kategorilerinde (özellikle “belirsizlik” ve “şaşkınlık”) modellerin sınıflandırma başarısında düşüş gözlenmiştir. Bu durum, ilgili duyguların dilsel ifadesinin daha belirsiz olması ve veri kümesindeki örnek dengesizliklerinden kaynaklanıyor olabilir [4]. Ayrıca, modellerin en yüksek önem atfettiği kelimelerin genellikle Türkçede duygusal anlam içeren ya da belirli bağlamlarda sık kullanılan sözcükler olduğu görülmüştür. Bu bulgu, makine öğrenmesi modellerinin “anlama” sürecini büyük ölçüde kelime düzeyindeki örüntülere dayandırdığını göstermektedir [3].

Öte yandan, klasik yöntemlerin bazı sınırlılıkları da mevcuttur. Özellikle TF-IDF tabanlı özellik çıkarımı ve manuel öznitelik mühendisliği süreçleri, zaman alıcı ve tekrar gerektiren işlemlerdir. Duygu ifadelerinin zamanla değişen dilsel yapılarına adapte olabilmek için bu özniteliklerin belirli aralıklarla güncellenmesi gerekebilir. Bu durum, klasik modellerin dinamik dil ortamlarına adaptasyon kabiliyetini sınırlandırmakta ve genelleme esnekliklerini azaltmaktadır [7].

Literatürde, farklı algoritmaların güçlü yönlerini birleştiren ensemble (birleşik) yöntemlerin, duygu tespiti gibi çok sınıflı ve karmaşık sınıflandırma problemlerinde başarı oranlarını artırabildiği belirtilmektedir [9]. Gelecek çalışmalarda bu tür ensemble yaklaşımların farklı Türkçe veri kümeleri üzerinde uygulanarak karşılaştırmalı analizlerinin yapılması, model başarımının artırılması açısından yararlı olacaktır.

V. SONUÇLAR

Bu çalışmada, Logistic Regression, Naive Bayes ve Support Vector Machine (SVM) olmak üzere üç farklı klasik makine öğrenmesi algoritması, Türkçe diline özgü TREMO duygu analizi veri kümesi üzerinde değerlendirilmiştir. Özellik çıkarımı için TF-IDF, model iyileştirmesi için ise hiperparametre optimizasyonu uygulanmıştır.

Elde edilen sonuçlara göre, en yüksek doğruluk oranı %79,7 ile Logistic Regression modeline aittir. Onu sırasıyla SVM (%78,9) ve Naive Bayes (%78,6) modelleri takip etmektedir. Bu veriler, hem doğrusal hem de kernel tabanlı algoritmaların, Türkçe metinlerden duygu çıkarımı görevlerinde etkili biçimde genelleme yapabildiğini göstermektedir.

Öte yandan, literatürde yer alan derin öğrenme temelli çalışmaların benzer veri kümeleri üzerinde %85–90 doğruluk oranlarına ulaştığı bildirilmektedir. Bu durum, daha karmaşık model mimarilerinin sınıflandırma başarımını artırabileceğine işaret etmektedir. Ancak, klasik yöntemlerin sunduğu düşük kaynak gereksinimi, kolay uygulanabilirlik ve yüksek yorumlanabilirlik gibi avantajlar, özellikle kaynak kısıtlı ortamlarda güçlü bir alternatif sunduğunu ortaya koymaktadır.

Genel olarak elde edilen bulgular, bilgisayarların temel NLP teknikleriyle metinler üzerinden insan duygularını başarılı biçimde tanımlayabildiğini ve hangi kelimelerin bu sürece en fazla katkı sağladığını ortaya koymaktadır. Gelecekteki çalışmalarda, derin öğrenme yaklaşımları, çok dilli analizler, multimodal

veri kaynakları ve daha dengeli veri kümeleri ile doğruluk oranlarının daha da artırılması mümkündür. Bu tür araştırmalar, insan-bilgisayar etkileşiminin geliştirilmesi, duygu temelli sistemlerin yaygınlaştırılması ve duygusal zekâyâ sahip yapay zeka uygulamalarının tasarımı açısından önemli katkılar sağlayacaktır.

KAYNAKLAR

- [1] Ratican, J. & Hutson, J. (2024). Advancing Sentiment Analysis Through Emotionally-Agnostic Text Mining in Large Language Models (LLMs). *Journal of Biosensors and Bioelectronics Research*, 2(3), 1-8.
- [2] Truong, V. (2024). Textual emotion detection – A systematic literature review. Research Square, DOI: <https://doi.org/10.21203/rs.3.rs-4673385/v1>.
- [3] Chutia, T., & Baruah, N. (2024). Text-based Emotion Detection A Review. *International Journal of Digital Technologies*, 3(1), 15-25.
- [4] Truong, V. (2024). Current challenges in detecting complex emotions from texts. Research Square, DOI: <https://doi.org/10.21203/rs.3.rs-4776002/v1>.
- [5] Guo, J. (2022). Deep learning approach to text analysis for human emotion detection from big data. *Journal of Intelligent Systems*, 31, 113-126.
- [6] Machová, K., Szabóová, M., Paralič, J., & Mičko, J. (2023). Detection of emotion by text analysis using machine learning. *Frontiers in Psychology*, 14, 1190326.
- [7] Li, Y., Chan, J., Peko, G., & Sundaram, D. (2024). An explanation framework and method for AI-based text emotion analysis and visualisation. *Decision Support Systems*, 178, 114121.
- [8] De León Languré, A. & Zareei, M. (2024). Improving Text Emotion Detection Through Comprehensive Dataset Quality Analysis. *IEEE Access*, 12, 166512-166530.
- [9] Acheampong, F. A., Wenyu, C., & Nunoo-Mensah, H. (2020). Text-based emotion detection: Advances, challenges, and opportunities. *Engineering Reports*, 2(7), e12189.
- [10] Al Maruf, A., Jiyad, Z. M., Mridha, M. F., Khanam, F., Haque, M. M., & Aung, Z. (2024). Challenges and Opportunities of Text-Based Emotion Detection: A Survey. *IEEE Access*, 12, 18416-18440.
- [11] Manapragada, H. (2024). Emotion Analysis using Multimodal Approach. Master's Thesis, California State University, Northridge.