



Şüpheli Haberlerin Tespitine Yönelik Derin Öğrenme Tabanlı Haber Teyit Sisteminin Gerçekleştirilmesi

Development of a Deep Learning Based News Verification System for The Detection of Suspicious News

Nurcan Yardımcı^{1*}, Burhan Ergen², Hatice Karaoğlu³

¹Fırat Üniversitesi, Bilgisayar Mühendisliği Bölümü, nurcann.yardimci@gmail.com, ORCID: 0009-0002-0476-9856

² Fırat Üniversitesi, Bilgisayar Mühendisliği Bölümü, bergen@firat.edu.tr, ORCID: 0000-0003-3244-2615

³ Fırat Üniversitesi, Bilgisayar Mühendisliği Bölümü, haticekaraoglu02@gmail.com, ORCID: 0009-0006-3750-1935

MAKALE BİLGİLERİ

Makale Geçmişi:

Geliş 4 Mart 2025
Revizyon 6 Nisan 2025
Kabul 2 Haziran 2025
Online 30 Haziran 2025

Anahtar Kelimeler:

Anlamsal benzerlik, doğal dil işleme, metin özetleme, şüpheli haber tespiti

ÖZ

Dijital medyanın yaygınlaşması, sahte haberlerin toplumu olumsuz etkilemesine yol açmaktadır. Özellikle sosyal medya platformları, yanlış bilgilerin hızla yayılmasına sebep olmaktadır. Bu durum bireylerin yanlış yönlendirilmesine, toplumsal kutuplaşmaya ve güven kaybına yol açmaktadır. Bu nedenle, sahte haberlerin hızlı ve güvenilir bir şekilde tespiti çağımızın en önemli sorunlarından biri haline gelmiştir. Bu çalışmada, derin öğrenme ve doğal dil işleme (NLP) tabanlı yenilikçi bir haber teyit sistemi geliştirilmiştir. Sistem, dijital haber içeriklerinin doğruluğunu analiz ederek, kullanıcıları yanlış bilgilendirmeden korumayı hedeflemektedir. Geliştirilen model, haber metinlerini Zemberek kütüphanesiyle işleyerek optimize etmekte, ardından BERT tabanlı özetleme yöntemleriyle metinlerin anlamını zenginleştirmektedir. Zemberek, Türkçe metinlerin doğal dil işleme süreçlerini kolaylaştırırken, SBERT modeli ise haber içeriklerinin daha derin anlamsal analizini gerçekleştirmektedir. Ayrıca YAKE algoritması haber içeriklerinden anahtar kelimeler çıkararak güvenilir kaynaklarla karşılaştırılmasını kolaylaştırmaktadır. Metinler arasındaki anlamsal benzerlik SBERT modeli kullanılarak ölçülmektedir. Kelime bazlı benzerlikler ise N-gram yöntemiyle ölçülmektedir. Bu iki yöntemden elde edilen sonuçlar birleştirilerek, haberin doğruluk skoru hesaplanmakta ve kullanıcıya sunulmaktadır. Önerilen sistem, sahte haberlerin tespitinde yüksek doğruluk oranları sunmakta ve gerçek zamanlı veri analizi sayesinde güncel olaylardaki yanlış bilgileri hızla belirlemektedir. Çok dilli ve çok modaliteli bir yapıya kavuşturulması hedeflenen sistemin uluslararası sahte haber tespiti için de etkili bir çözüm sunması hedeflenmektedir.

ARTICLE INFO

Article History:

Received 4 May 2025
Received in revised form 6 April 2025
Accepted 2 June 2025
Available online 30 June 2025

Keywords:

Semantic similarity, natural language processing, text summarization, suspicious news detection

DOI: 10.24012/dumf.1651348

* Sorumlu Yazar

ABSTRACT

The proliferation of digital media causes fake news to negatively affect society. Social media platforms, in particular, cause false information to spread rapidly. This situation leads to misdirection of individuals, social polarization and loss of trust. Therefore, rapid and reliable detection of fake news has become one of the most important problems of our time. In this study, an innovative news verification system based on deep learning and natural language processing (NLP) has been developed. The system aims to protect users from misinformation by analyzing the accuracy of digital news content. The developed model optimizes news texts by processing them with the Zemberek library, and then enriches the meaning of the texts with BERT-based summarization methods. While Zemberek facilitates the natural language processing of Turkish texts, the SBERT model performs deeper semantic analysis of news content. Moreover, the YAKE algorithm extracts keywords from news content, thereby facilitating its comparison with reliable sources. Semantic similarity between texts is measured using the SBERT model. Word-based similarities are measured with the N-gram method. By combining the results obtained from these two methods, the accuracy score of the news is calculated and presented to the user. The proposed system offers high accuracy rates in detecting fake news and quickly identifies misleading information in current events thanks to real-time data analysis. The system, which is intended to be enhanced with a multilingual and multimodal structure, is also expected to provide an effective solution for international fake news detection.

Giriş

Son zamanlarda dijital medyanın hızlı bir şekilde popüler olmasıyla beraber bireylerin ve toplulukların da bilgiye erişimi artmıştır. Bu erişimin artmasıyla beraber sahte haberler de artarak kolay bir şekilde yayılmıştır. Her birey haber üreticisi veya paylaşımcısı haline gelmiştir. Bununla beraber, haberlerin doğruluk ve güvenilirlik açısından denetlenme ihtiyacı da artmıştır. Bilgilerin hızla yayılıyor olması, özellikle doğruluğu teyit edilmemiş ya da kasıtlı olarak oluşturulmuş sahte haberlerin toplum üzerinde ciddi etkiler bırakmasına sebep olmaktadır. Özellikle ulusal güvenlik, kişisel haklar, sosyal düzen ve kamu sağlığı gibi kritik konulardaki sahte haberler; bireylerin yanlış yönlendirilmesine, toplumun manipülasyonuna ve yanlış bilgilendirmeye yol açarak ciddi sosyal, ekonomik ve politik sonuçlara sebep olmaktadır. Bu durumlar dikkate alınırca güvenilir bilgiye ulaşmak giderek daha zor ve önemli hale gelmiştir [1,2]. Sahte haberlerin tespiti için yapılmış olan çalışmalar, bilgi yoğunluğunun fazla olduğu bu dönemde yetersiz kalmaktadır. Özellikle doğrulama sürecinin hız ve güvenilirlik açısından gereksinimleri dikkate alındığı zaman, insan gücüne dayalı yöntemler ölçeklenebilirlik sorunlarına sebep olmaktadır. Bu yüzden, derin öğrenme ve doğal dil işleme (NLP) teknikleri kullanılarak geliştirilmiş otomatik haber teyit sistemleri önem kazanmaktadır [3]. Geliştirilmiş haber teyit sistemlerinin, sosyal medya platformları ve dijital haber siteleriyle entegrasyonu sağlandığında, kullanıcıların yanlış bilgilendirilmesinin önüne geçilmesinde kritik bir rol oynayacaktır [4].

Haber metinlerini anlamlı olacak şekilde analiz edebilmek için, Zemberek gibi kütüphaneler kullanılarak haberlerin köklerine ayrılması, gereksiz kelimelerin temizlenmesi ve metnin analize hazır hale getirilmesi sağlanmaktadır. Sahte haber tespitinin geliştirilmesi noktasında, BERT ve SBERT gibi derin öğrenme tabanlı dil modellerinin imkan sağladığı yüksek doğruluk oranları, özellikle cümle temsillerinin zenginleştirilmesinde etkili bir yöntem olmaktadır [5]. Bu modeller, metinlerin anlamsal düzeyde kıyaslanmasına olanak tanıırken, Chu vd. [6] önermiş olduğu R-SBERT gibi geliştirilen yaklaşımlar, sahte haber tespitinin doğruluğunu daha da arttırmaktadır. Metinlerin yüzeysel benzerliklerini incelemek için N-gram ve Kosinüs Benzerliği gibi ölçütler de kullanılmaktadır. Lopez-Gazpio vd. [7] doğal dil anlamada yaygın olarak kullanılan iki görev olan anlamsal metin benzerliği ve doğal dil çıkarımı üzerinde N-gram tabanlı model geliştirmişlerdir. Model, cümlelerde farklı uzunluktaki N-gramlar arasındaki ilişkileri modellemek için kullanılmıştır ve kelime tabanlı modellere kıyasla önemli iyileştirmeler sağlanmıştır. Sosyal medya üzerindeki sahte haberlerin tespitiyle alakalı başka bir çalışmada da SVM ve LSTM modelleri kullanılmış, özellikle yüksek doğruluk oranları ile sahte haberlerin hızlı bir şekilde tespitinin önemi vurgulanmıştır [8].

Haber teyit sisteminde, metinlerdeki en belirgin bilgilerin öne çıkarılması için YAKE gibi anahtar kelime çıkarma teknikleri de kullanılmaktadır. Campos vd. [9] YAKE

üzerine yaptıkları çalışmada, modelin dil ve alan bağımsız olarak belgelerden anahtar kelime çıkarımında yüksek performans gösterdiği belirtilmiştir. YAKE modeli, haberlerin önemli kelimelerini öne çıkararak incelenecek içeriklerin güvenilir kaynaklarla karşılaştırılmasını kolaylaştırmaktadır. Diğer taraftan TF-IDF ve TextRank algoritmaları birleştirilerek sunulan yöntemde, metin madenciliğindeki anahtar kelime çıkarımının doğruluğunun artırılması hedeflenmiştir [10]. Veri toplama aşamasında, haber içerikleri web scraping ve RSS ile toplanmaktadır. RSS protokolü, belirlenmiş olan sitelerden güncel haber içeriklerin otomatik olarak alınmasını sağlayarak güncellenen içeriklere hızlı erişim sunmaktadır. Web scraping işlemlerinde Selenium ve BeautifulSoup gibi araçlar kullanılarak her haber sitesine özgü HTML yapılarına göre veriler toplanmaktadır [11]. Nair vd. [12], sosyal medya tabanlı sahte haber tespit çalışmalarında da vurguladığı gibi veri madenciliği süreçlerinde veri doğruluğu artırılırken doğrulama sürecinin etkinliği de artırılmaktadır. Derin öğrenme tabanlı modellerin sunmuş olduğu güçlü anlam çıkarımı yetenekleri geleneksel anahtar kelime ve yüzeysel benzerlik ölçütleriyle birleştirilerek, sahte haber tespitinde daha detaylı bir çözüm sunmaktadır. TOP-Rank algoritması gibi konu bazlı sıralama yöntemler ise, anahtar kelime çıkarımının doğruluğunu artırmak için kullanılan diğer bir tekniktir [13]. Ayrıca, Dinçer vd. [14] siber zorbalık tespiti çalışmasında kullandığı makine öğrenmesi ve veri madenciliği yöntemleri de, önerilen haber teyit sisteminin altyapısında kullanılabilecek nitelikte teknikler sunmaktadır. Syed vd. [15], sunduğu hibrit öğrenme modeli, zayıf gözetimli öğrenme teknikleriyle entegre edilerek sahte haberlerin doğruluk oranlarını artırmaktadır. Böylece topluma doğru bilgi aktarma hedefine katkı sunmaktadırlar. Bu doğrulama sistemi, dijital medyada sahte haberlerin hızla tespit edilmesi ve toplumun güvenilir bilgiye erişiminin sağlanması noktasında güçlü bir araç olarak öne çıkmaktadır.

Koru vd. [16] Twitter üzerinden paylaşılan Türkçe haberlerin sahte olup olmadığını tespit etmek amacıyla yapay zekâ temelli çalışma yürütmüştür. Araştırmada, doğrulama platformlarından elde edilen sahte haberler ve ana akım medya hesaplarından alınan gerçek haberlerle oluşturulan TR_FaRe_News adlı Türkçe veri kümesi geliştirilmiştir. Bu veri kümesi, Zemberek kütüphanesi ile ön işleme tabi tutulmuştur. Ardından BoW, TF-IDF, Word2Vec gibi vektörleştirme yöntemleri ve çeşitli makine öğrenimi algoritmaları (SVM, RF, NB vb.) ile sınıflandırılmıştır. Daha sonra, önceden eğitilmiş BERT ve BERTurk modelleri Bi-LSTM ve CNN katmanları ile genişletilerek hem dondurulmuş hem de serbest parametre yaklaşımlarıyla incelemeye alınmıştır.

Özbay vd. [17] sosyal medyada yayılan sahte haberlerin tespitine yönelik iki aşamalı bir model önermiştir. İlk aşamada, yapılandırılmamış haber metinleri metin madenciliği teknikleri ile işlenerek sayısal ve dilsel gürültülerden arındırılmıştır. Ardından Vektör Uzay Modeli (VUM) kullanılarak yapısal formata dönüştürülmüştür.

İkinci aşamada ise, ISOT sahte haber veri kümesi üzerinde toplam on farklı denetimli yapay zekâ algoritması (Naive Bayes, JRip, J48, Random Forest, SGD, LWL, DTNB, Bagging, CvR, OneR) dört farklı eğitim-test senaryosu ile değerlendirilmiştir.

Kulaksız [18], Türkçe ana akım medyada yer alan haberlerin sahte mi gerçek mi olduğunu tespit edebilen bir sistem geliştirmiştir. TRT Gündem ve Teyit.org kaynaklarından toplanan haberler ön işleme adımlarından geçirilerek etiketlenmiştir. Özellik çıkarımı için TF-IDF, CountVectorizer ve Word2Vec yöntemlerini kullanmıştır. Ardından SVM, Lojistik Regresyon, Naive Bayes, K-NN, Karar Ağacı, Rastgele Orman, LSTM ve BERTurk algoritmaları ile haberleri sınıflandırmıştır.

Önceki çalışmalarda Türkçe sahte haber tespiti için farklı yöntemler önerilmiştir. Örneğin, Kuru vd. [16] çalışmalarında BERT ve BERTurk modellerini, Bi-LSTM ve CNN katmanlarıyla genişleterek Twitter üzerinde paylaşılan Türkçe haberlerin sınıflandırılmasını hedeflemiştir. Ancak bu yaklaşım, sabit veri kümesi ile sınırlı kalmış ve dinamik içerik akışına odaklanmamıştır. Özbay vd. [17] ise sosyal medya verilerine dayalı olarak iki aşamalı bir model geliştirmiş ve farklı denetimli makine öğrenmesi algoritmalarını değerlendirmiştir. Ancak gerçek zamanlı veri toplama, özetleme ve çok katmanlı benzerlik analizine yer verilmemiştir. Kulaksız [18] tarafından geliştirilen çalışmada ise, TF-IDF, Word2Vec ve CountVectorizer gibi geleneksel özellik çıkarım yöntemleriyle etiketlenmiş Türkçe haberler sınıflandırılmıştır. Ancak bu yaklaşım da özetleme ya da semantik benzerlik temelli bütüncül analiz yapısından uzaktır. Gerçek zamanlı veri akışını kullanan, SBERT, BERTSum ve YAKE modellerini entegre şekilde birleştirilerek Türkçe haber doğruluğunu analiz eden bir sistem literatürde sınırlı sayıdadır. Bu çalışmada önerilen sistem, gerçek zamanlı haber akışı ile veri toplama, BERTSum ile metin özetleme, YAKE algoritması ile anahtar kelime çıkarımı, SBERT ile anlamsal ve N-gram ile yüzeysel benzerlik analizini birleştiren çok katmanlı bir doğrulama süreci sunarak Türkçe içeriklere özel olarak önemli bir katkı sunmaktadır.

Çalışmada önerilen yöntem; haber metinlerinin işlenmesi, özetlenmesi, anahtar kelime çıkarımı ve benzerlik analizini bir arada kullanan bütüncül bir yaklaşım sunmaktadır. Sistem, metinlerin güvenilir kaynaklardan alınan içeriklerle kıyaslanarak doğruluğunun analiz edilmesi temeline dayanmaktadır. İlk aşamada, dijital ortamdaki haber içerikleri Zemberek kütüphanesi kullanılarak dil ön işleme sürecine tabi tutulmaktadır. Sonrasında BERT tabanlı çıkarımsal özetleme yöntemi olan BERTSum ile metinler özetlenmektedir. Özetlenen metinlerden YAKE algoritması aracılığıyla anahtar kelimeler çıkarılmakta ve bu kelimeler güvenilir haber kaynaklarıyla eşleştirilmektedir. Son olarak, metinler arasındaki benzerlikler SBERT modeli ile anlamsal düzeyde ve N-gram yöntemi ile kelime düzeyinde analiz

edilmektedir. Elde edilen sonuçlar Ortak Skor hesabı ile birleştirilerek haberin doğruluk derecesi hesaplanmaktadır. Bu çok aşamalı sistem sayesinde, sahte haberlerin tespiti hem içerik derinliği hem de yüzeysel benzerlik açısından kapsamlı bir şekilde gerçekleştirilmektedir.

Mevcut literatürde sahte haber tespiti üzerine yapılan çalışmaların büyük bir kısmı sabit, etiketlenmiş veri kümeleri üzerinde gerçekleştirilmiştir. Gerçek zamanlı, dinamik haber akışlarıyla çalışan sistemler oldukça sınırlı kalmıştır. Bu çalışmanın literatüre katkısı, sahte haberlerin tespitine yönelik geliştirilen sistemin gerçek zamanlı veriyle çalışabilmesi, hem anlamsal hem yüzeysel benzerlikleri birleştiren özgün Ortak Skor algoritmasını içermesi ve Türkçe haber içerikleri üzerinde yüksek başarı oranları sunabilmesidir. Literatürde SBERT, YAKE ve BERTSum gibi bileşenlerin birlikte kullanıldığı, özellikle Türkçe içerikler üzerinde çalışan ve güvenilir kaynaklara dayalı otomatik referans eşlemesi yapan sistem sayısı oldukça sınırlıdır. Ayrıca, geliştirilen "Ortak Skor" algoritması ile birden fazla benzerlik yönteminin ağırlıklı kombinasyonu gerçekleştirilerek daha doğruluklu bir tespit modeli sunulmuştur. Gerçek zamanlı veri kullanımı, verilerin işlenebilmesi, çok aşamalı analiz yapısı ve yerli kaynaklardan otomatik içerik çekme mekanizması ile bu çalışma, Türkçe haber doğrulama sistemlerine hem metodolojik hem de uygulamalı anlamda özgün bir katkı sunmaktadır. Özellikle kriz dönemlerinde ortaya çıkan bilgi kirliliğiyle başa çıkmak için sistemin uygulama alanını genişletmekte ve bu yönüyle literatüre katkı sağlamaktadır.

Materyal ve Metot

Veri Seti Kullanımı ve Veri Toplama Süreci

Bu çalışmada, sabit ve önceden etiketlenmiş bir eğitim veri seti kullanılmasından ziyade, sistemin gerçek zamanlı ve güncel haber içerikleri üzerinde çalışması hedeflenmiştir. Önerilen sistem gerçek zamanlı haber akışları üzerinde çalışmak üzere tasarlanmıştır. Veri toplama süreci kapsamında, Türkiye’de yer alan haber kaynaklarından RSS protokolü ve web scraping (Selenium ve BeautifulSoup) teknikleriyle anlık haber içerikleri toplanmaktadır. Çekilen haberler başlık, içerik, yayın tarihi ve haber bağlantısı (URL) bilgilerini içermektedir. Örneklem sürecinde, farklı temalarda (gündem, sağlık, siyaset, kriz anları vb.) yaklaşık 300 adet haber metni üzerinde deneme yapılmıştır. Böylece, sistemin yalnızca geçmiş verilere değil, güncel ve dinamik içeriklere dayalı analiz yapabilmesi sağlanmıştır. Bu yaklaşım, sahte haber tespit sürecini daha esnek ve güncel hale getirerek, özellikle kriz dönemleri gibi bilgi akışının hızlı ve değişken olduğu durumlarda daha etkili sonuçlar elde edilmesine olanak tanımaktadır. Elde edilen veriler, ön işleme adımlarının ardından özetleme, anahtar kelime çıkarımı ve benzerlik analizi süreçlerine dahil edilmekte ve sistem bu veriler üzerinden haberlerin doğruluğunu değerlendirmektedir.

Metin Özetleme Yöntemleri

Metin özetleme, bir veya birden fazla metindeki bilgileri analiz ederek önemli noktaları seçip daha kısa ve anlamlı bir

biçimde sunan doğal dil işleme yöntemidir. Bu süreç, metindeki temel fikirleri ve mesajları kaybetmeden kullanıcıya bilgi yükünü azaltarak hızlı ve etkili bir şekilde sunmayı amaçlar [19].

Çıkarımsal metin özetleme, orijinal metindeki önemli cümle veya ifadeleri seçerek oluşturulan, doğal dil işleme (NLP) alanında sıkça kullanılan bir metin özetleme yöntemidir. Bu yöntem, metni yeniden yazmak yerine, doğrudan metindeki önemli bölümleri seçerek özet oluşturmaktadır. Çıkarımsal metin özetlemede, her cümle veya ifadenin önem derecesi, kelime sıklığı, cümle uzunluğu, metindeki konumu ve bağlamsal ilişkiler gibi ölçütlerle belirlenmektedir. Çıkarımsal özetleme yöntemleri, dilbilgisi hatası yapmaması ve bağlamsal anlamın korunmasını sağlaması sebebiyle avantajlıdır [20]. Ayrıca, büyük veri kümelerinde hızlı ve etkili bir şekilde uygulanabilmektedirler. Ancak bu yöntem, metindeki cümleleri olduğu gibi kullandığı için özetin akıcı olmaması gibi dezavantajlara sahiptir [21]. Çıkarımsal özetleme, haberlerin, akademik makalelerin ve yasal belgelerin özetlenmesi gibi pek çok alanda yaygın olarak kullanılmaktadır. Çıkarımsal metin özetlemede istatistiksel yöntemlere bakıldığı zaman kelime frekans analizi, metindeki en sık kullanılan kelimeleri belirleyerek özetleme sürecinde önemli cümlelerin seçilmesini sağlar. TF-IDF yöntemi ise bir terimin belirli bir belgede ne kadar önemli olduğunu belirlemek için kullanılır [22].

Graf tabanlı yöntemlerden TextRank, graf üzerinde derecelendirme yaparak en önemli cümleleri belirleyerek bir özet oluşturur. LexRank ise PageRank benzeri bir algoritma kullanarak cümlelerin önem derecesini hesaplar. Her iki yöntem de özetleme sürecinde en anlamlı cümleleri seçmek için kullanılmaktadır [23]. Kümeleme tabanlı yöntemler, cümleleri benzerliklerine göre gruplandırarak özet oluşturmaktadır. K-Means ise cümleleri özelliklerine göre belirli sayıda kümeye ayırır ve her kümeden merkezi cümleleri seçerek özet oluşturur [24].

Semantik analiz yöntemleri, metindeki anlamsal ilişkileri inceleyerek özetleme yapar. LSA (Latent Semantic Analysis), tekil değer ayrışımı kullanarak cümleleri vektörlere dönüştürüp bu vektörler arasındaki benzerliklere göre metnin derin anlamsal yapısını analiz eder. Word2Vec/GloVe ise kelimeleri vektör uzayında temsil ederek kelimeler arasındaki bağlamsal ilişkileri öğrenir ve bu bilgiyi özetleme sürecinde kullanır [25]. Derin öğrenme tabanlı yöntemler, cümlelerin önemini belirlemek ve özetleme yapmak için farklı sinir ağı yapılarından faydalanır. DNN, CNN, LSTM ve GRU, BERT tabanlı yöntemler örnek olarak verilebilir. BERTSum ise, BERT tabanlı bir çıkarımsal özetleme yöntemidir ve metin özetleme görevleri için özelleştirilmiştir. Bu model, BERT'in cümle düzeyinde anlam yakalama yeteneğini kullanarak metindeki önemli cümleleri seçer. BERTSum, metni bir dizi cümle olarak işler ve her bir cümle için özetleme için önemini belirler. Önem derecelerine göre sıralanmış olan cümlelerden özet oluşturur. BERTSum, diğer çıkarımsal yöntemlere göre bağlamsal

anlamı daha iyi koruması ve yüksek doğruluk sağlamasıyla öne çıkmaktadır [26].

Anahtar Kelime Çıkarma

Anahtar kelime çıkarma, bir metin belgesindeki cümle öbeklerinin veya kelime öbeklerinin çıkarılması için kullanılmaktadır. Özellikle bilgi arama, belge sınıflandırma ve özetleme gibi alanlarda önemli olup, metinle ilgili anahtar bilgileri ayıklayarak verimli analiz yapılmasına olanak tanır. Kullanılan teknikler arasında istatistiksel analiz, NLP algoritmaları ve makine öğrenimi bulunur. Literatür çalışmalarında yer alan TF-IDF, RAKE, YAKE ve KeyBERT modelleri incelenmiştir. RAKE algoritması, metinlerdeki anahtar kelimeleri belirlemek için kelime frekansı ve kelime gruplarının derecelerini kullanarak etkili bir şekilde çalışır ve özellikle stop-word'leri filtreleyerek önemli kelime öbeklerini çıkarır. TF-IDF yöntemi ise, bir terimin bir metin içinde sık görülme sıklığını (TF) ve diğer metinlerde ne kadar nadir geçtiğini (IDF) değerlendirerek, kelimenin önemini istatistiksel olarak belirlemektedir [27].

Bu çalışma için metindeki cümle veya kelime öbeklerinden bağlamsal vektör temsillerinden koşullu benzerlikte en yüksek puanı alan cümle ve öbekleri anahtar kelime olarak veren KeyBERT ve metin özelliklerini kullanarak otomatik anahtar kelime çıkarımına yönelik denetimsiz bir yaklaşım kullanan YAKE algoritması kullanılmıştır.

YAKE (Yet Another Keyword Extractor), herhangi bir eğitim veri seti veya dış kaynak ihtiyacı olmadan tekil dokümanlardan yerel istatistiksel özelliklere dayalı olarak otomatik anahtar kelime çıkarımı yapan hafif ve denetimsiz bir algoritmadır. Dil ve alan bağımsızlığı ile eğitim veri setine ihtiyaç duymaması, özellikle veri eksikliği olan durumlarda tercih edilmesini sağlar. İncelenen çalışmada, YAKE, TF-IDF, RAKE ve TextRank gibi yöntemlere kıyasla daha yüksek doğruluk ve esneklik sunmuştur. YAKE modeli, beş ana adımdan oluşan bir iş akışı ile çalışmaktadır. İlk adımda metin ön işleme gerçekleştirilir ve aday terimler tanımlanır. Daha sonra, aday terimlerin istatistiksel özellikleri çıkarılır ve bu özellikler kullanılarak terimlerin önem derecesi hesaplanır. Üçüncü adımda, terimler birleştirilerek N-gram ifadeler oluşturulur ve bu ifadelerin aday anahtar kelime puanları hesaplanır. Son olarak, benzer anahtar kelimeler veri temizleme aşamasında filtrelenir ve aday anahtar kelimeler önem derecesine göre sıralanır. Bu süreç, YAKE'nin metinden bağlam bağımsız olarak en uygun anahtar kelimeleri çıkarma kabiliyetini desteklemektedir [9].

Benzerlik Değerlendirme Yöntemleri

Kelime ve Karakter Düzeyinde Benzerlik Yöntemleri

N-gram tabanlı benzerlik, istatistiksel dil modellerinde kullanılan yaygın bir uygulamadır. Yakın zamanlarda kullanılan cümle yerleştirme modelleri, unigram (kelime) yerleştirmeleri bigram yerleştirmelerle tamamlamıştır. Her iki cümle çifti için tek bir kelimeden maksimum N kelimeye kadar değişen bitişik kelime dizilerini çıkarılır ve tüm N-

gram çiftleri için bir matris oluşturulur. Oluşturulan matrisin köşegeni 1 gramı (tek kelime) temsil ederken, sağ tarafta yer alan kareler daha uzun N-gramları temsil etmektedir. Her adımda, önceki N-grama bir kelime daha eklenir. Dikkate alınması gereken maksimum N-gram boyutu bir hiperparametredir (N) ve $y-x < N$ koşulunu sağlamaktadır. Böylece, S cümlesi için model tarafından işlenen toplam N-gram sayısı Denklem 1 gibi hesaplanmaktadır.

$$|S|.N - \sum_{i=1}^{N-1} i \quad (1)$$

$N = 1$ olduğunda, N-gram sayısı $|S|$ 'ye eşit olur. Yani köşegendeki eleman sayısı ($N = |S|$ sayısı) bir üst üçgendeki eleman sayısına eşittir. Tanımlanan boyut $|S|$ matrisidir ve Denklem 2'de olduğu gibi hesaplanmaktadır.

$$\sum_{i=1}^{|S|} i \quad (2)$$

N-gram modeli, bağlam bilgisini kelime tabanlı bir dikkat modeline aşlamak açısından alternatif temellerden (RNN'ler veya CNN'ler) daha üstün sonuçlar göstermiştir [7].

Benzerlik, bazı metotlarda mesafe ölçümü ilkesine dayanırken bazı metotlarda ilişki seviyesi belirlenmesi ilkesine dayanır. Yaygın olarak kullanılan bir diğer benzerlik ölçme yöntemi de Kosinüs benzerliğidir ve bu çalışmada kullanılmıştır. Kosinüs benzerliği iki vektör arasındaki kosinüs açısını çok boyutlu bir uzayda hesaplayarak vektörler arasındaki benzerliği ölçme metodudur [28]. Denklem 3' te Kosinüs benzerliği hesaplama fonksiyonu görülmektedir.

$$\text{benzerlik} = \cos(\theta) = \frac{A \cdot B}{\|A\| \|B\|} = \frac{\sum_{i=1}^n A_i \cdot B_i}{\sqrt{\sum_{i=1}^n A_i^2} \sqrt{\sum_{i=1}^n B_i^2}} \quad (3)$$

Anlamsal Benzerlik Yöntemleri

Anlamsal metin benzerliği, iki metnin ne kadar benzer olduğunu belirlemektedir. Anlamsal benzerlik, makine çevirisi, özetleme, metin üretme, soru yanıtlama, diyalog ve konuşma sistemleri gibi pek çok uygulama alanında değerlendirme amacıyla kullanılmaktadır. Anlamsal metin benzerliği çalışmaları İngilizce'de oldukça yaygındır ve bu çalışmalar insanlar tarafından işaretleme yapılarak benzerlik ölçüt değerleri verilmiş veri kümelerine dayanmaktadır. Ancak, işaretleme yapmak maliyetli ve zaman alan bir iştir. Son zamanlarda, makine çevirilerindeki başarının artması ve çoklu dil modellerinin geliştirilmesiyle birlikte veri kümelerinin bir dilden diğerine çevrilerek kullanılması mümkün hale gelmiştir [29].

Sentence-BERT (SBERT) ile metin benzerliği, siyam ve üçlü ağ yapılarını kullanan önceden eğitilmiş BERT ağıının (veya diğer transformer modellerinin) bir modifikasyonudur. Model, anlamsal olarak benzer cümleler için vektör uzayında birbirine yakın sabit boyutlu cümle gömmeleri türetmektedir. Cümleler arasındaki anlamsal benzerliği ölçmek için kullanılan BERT tabanlı bir modeldir [5]. Metinlerin anlam benzerliğini analiz etmek amacıyla tasarlanmış olan SBERT, cümle gömülülerini hesaplamaktadır (Denklem 4).

$$SBERT - göm_k = BERTModeli(Cümle_k) \quad (4)$$

Ardından bu gömülüler arasındaki kosinüs benzerliğini değerlendirmektedir (Denklem 5).

$$SBERT - Benzerlik = \frac{SBERT - göm_i \cdot SBERT - göm_j}{\|SBERT - göm_i\| \cdot \|SBERT - göm_j\|} \quad (5)$$

Bu cümle benzerlikleri üzerinden referans özet cümleleri ile aday özet cümleleri arasındaki anlamsal benzerlikler de ölçülebilir.

Doğal Dil İşlemede Değerlendirme Metrikleri

Metin özetleme ve makine çevirisi gibi doğal dil işleme (NLP) uygulamalarında, özetleme performansını değerlendirmek için çeşitli otomatik metrikler kullanılmaktadır. Bu metrikler, üretilen metinlerin referans metinlerle ne kadar uyumlu olduğunu ölçerek modellerin doğruluğunu ve etkinliğini belirlemeye yardımcı olur.

Çalışmada da kullanmış olduğumuz ROUGE metriği, özellikle otomatik metin özetleme sistemlerinin değerlendirilmesinde kullanılan bir metriktir. N-gram, kelime dizisi ve en uzun ortak alt dizilim (LCS) gibi ölçümlerle, sistem tarafından üretilen özetlerin referans özetlerle örtüşme derecesini değerlendirir. ROUGE'un çeşitli varyantları bulunmaktadır [30]. Bunlar şu şekildedir;

- ROUGE-N: Sistem ve referans özetler arasındaki N-gram örtüşmelerini hesaplar.
- ROUGE-L: En uzun ortak alt dizilimleri kullanarak cümle düzeyindeki benzerlikleri değerlendirir.
- ROUGE-S: Atlanmış bigramları (skip-bigram) dikkate alarak esnek kelime çiftleri arasındaki örtüşmeleri ölçer.
- ROUGE-W: Ağırlıklı LCS tabanlı istatistikleri kullanarak ardışık LCS'leri daha fazla önemser. Bu metrik, ardışık kelime dizilimlerinin önemini vurgulayarak, özetlerin referans metinle daha tutarlı olmasını sağlar [30].

ROUGE, özellikle özetleme sistemlerinin referans özetlerdeki önemli bilgileri ne kadar iyi kapsadığını belirlemede etkilidir.

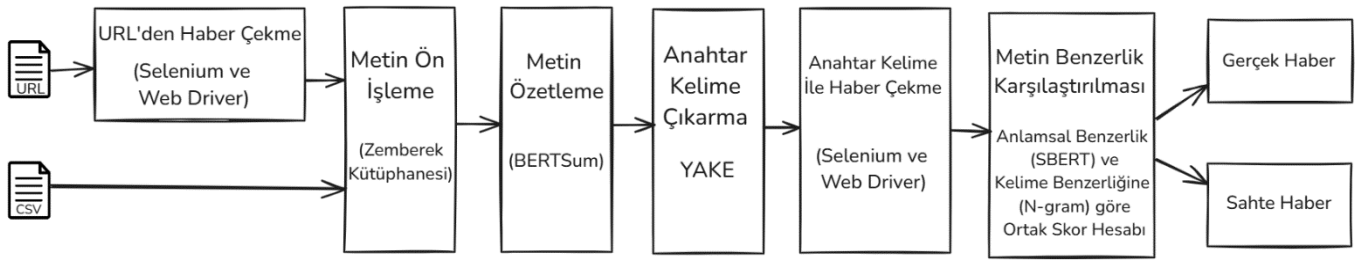
Diğer metriklere bakıldığı zaman BLEU, makine çevirisi çıktılarının kalitesini değerlendirmek için kullanılan bir metriktir ve paylaştığı N-gramlar üzerinden doğruluk skorunu hesaplayan ve çevirilerin insan çevirilerine benzerliğini ölçen bir metriktir [31]. METEOR, kelime kökleri, eş anlamlılar ve kelime sırasını dikkate alarak makine çevirisi çıktılarının referans çevirilere uyumunu ölçen ve insan yargılarıyla daha yüksek korelasyon sağlamayı amaçlayan bir değerlendirme metriğidir [32]. BERTScore, BERT gibi önceden eğitilmiş dil modellerinin bağlamsal kelime gömme temsillerini kullanarak iki metin arasındaki anlamsal benzerliği ölçen bir değerlendirme metriğidir. Bu metrik, üretilen ve referans metinler arasındaki her bir token için gömümleri karşılaştırarak benzerlik skorları hesaplar [33]. MoverScore, bağlamsal kelime gömmeleri ve EMD (Earth Mover's Distance)

yöntemini kullanarak metinler arasındaki benzerliği en düşük taşıma maliyeti üzerinden ölçen ve insan yargılarıyla yüksek korelasyon sağlayan bir değerlendirme metriğidir [34].

Bulgular ve Tartışma

Yapılan çalışma, haber sitelerinde yer alan haberleri çekip analiz ederek onların gerçek mi yoksa şüpheli mi olduğunu belirlemeye yönelik bir süreçten oluşmaktadır. Bu çalışmada sabit bir veri setine bağlı kalınmadan, güncel haber içeriklerinden yararlanılarak bir sistem tasarlanmıştır.

İncelemek istediğimiz haber, herhangi bir haber sitesinden veya sosyal medyadan yayınlanan bir haber olabilir. İncelenecek olan haber otomatik olarak çekilmekte ve işlenmektedir. Gerekli işlemler yapılan haber içeriklerinden anahtar kelimeler çıkarılmaktadır. Çıkarılan anahtar kelimelerle, güvenilir haber kaynakları olarak belirlenen Anadolu Ajansı ve Teyit.org siteleri üzerinden eşleşen içerikler aranmaktadır ve eşleşen haberler otomatik olarak çekilmektedir. Ardından benzerlik analizi gerçekleştirilmektedir.



Şekil 1. Şüpheli haberlerin tespiti için önerilen sistem modeli

Şüpheli haberlerin tespiti için önerilen modelin ilk aşamasında, haberlerin bulunduğu URL'lerden veya CSV dosyasına kaydedilmiş olan haber içeriklerinden oluşan bir haber verisi çekilmektedir.

CSV dosyasına kaydedilmiş olan haber içerikleri başlık, yazar, tarih, kaynak ve metin bilgilerini içermektedir. Aşağıda, bu verilerden bir haber örneğine ait içerik telif hakkı ve gizlilik nedeniyle sınırlı bir şekilde sunulmuştur:

"Gazeteciler 1 Nisan şakası yaparsa ne olur ?

Gaye Güllü / 31.03.2025

İstanbul

1 Nisan, dünya genelinde şaka günü olarak kutlanıyor. Bazı kaynaklar bu şaka geleneğinin 16. yy.'da Fransa'da ..."

Tam haber içerikleri sistem tarafından özetleme ve anahtar kelime çıkarma işlemleri için kullanılmıştır. Yukarıda sadece temsil edici örnek gösterilmiştir.

URL'den çekilecek haberler için Selenium ve Web Driver kullanılarak tarayıcı otomasyonu sağlanmakta ve haber metinleri otomatik olarak alınmaktadır. Çekilen haber

metinleri daha sonra Zemberek kütüphanesi ile ön işleme tabi tutulmaktadır. Ön işlemlerde gereksiz kelimeler çıkartılır ve metin doğal dil işleme için sadeleştirilerek optimize edilmiş olur.

Metin ön işlemeden geçen haberler, BERT tabanlı bir özetleme yöntemi olan BERTSum ile özetlenmektedir. Bu yöntem, metnin ana fikirlerini seçerek daha kısa bir özet oluşturmaktadır. Böylece sonraki analizlerin hızını artırmaya yardımcı olmaktadır. Aşağıda yer alan Tablo 1'de, denenmiş olan farklı çıkarımsal özetleme yöntemlerinin ROUGE metriklerine göre karşılaştırması gösterilmektedir. Sunulan verilere göre, ROUGE-1, ROUGE-2 ve ROUGE-L metrikleri, özetlerin orijinal metinle örtüşme düzeyini farklı açılardan değerlendirmektedir. ROUGE-1, özetlerin tekil kelime seviyesindeki örtüşmesini ölçerken, ROUGE-2 daha bağlamsal detayları ve cümle yapısını analiz etmektedir. ROUGE-L ise metin uyumunu ve dil bilgisel doğruluğu değerlendirerek en uzun ortak kelime dizilerini ölçmektedir. Bu üç metrik birlikte değerlendirildiğinde, özetleme yöntemlerinin yüzeysel doğruluğunu ve bağlamsal derinliğini anlamak için kapsamlı bir araç sunulmaktadır [30].

Tablo 1. Çıkarımsal özetleme yöntemlerinin ROUGE değerlerine göre karşılaştırılması

Çıkarımsal Özetleme Yöntemleri	ROUGE-1	ROUGE-2	ROUGE-L
BERTSum	41.81	27.77	32.72
K-Means	40.23	17.96	24.85
TextRank	39.54	17.14	19.20
LexRank	39.75	12.57	26.08
TF-IDF	39.75	20.12	26.08
LSA	35.52	12.86	20.80

BERTSum'un özellikle ROUGE-2 ve ROUGE-L metriklerindeki üstün performansı, bağlamsal ilişkileri ve metin detaylarını etkili bir şekilde kavradığını göstermektedir. TF-IDF ve LexRank yöntemleri, ROUGE-1 skorlarında iyi performans sergilemesine rağmen, ROUGE-2 skorlarında daha düşük sonuçlar elde ederek metin bağlamını derinlemesine anlamakta sınırlı olduklarını göstermektedir. LSA yöntemi, özellikle ROUGE-2 ve ROUGE-L metriklerinde en düşük performansı göstermiştir, bu da anlamsal bağlamı yeterince yakalayamamasından kaynaklanmaktadır. TextRank ve K-Means yöntemleri, genel anlamı özetlemede başarılı olurken, bağlamsal ve ayrıntılı analiz gerektiren ROUGE-2 metriklerinde daha sınırlı kalmıştır. Diğer yöntemlerin aksine, BERTSum tüm metriklerde en yüksek performansı göstermiş ve özellikle ROUGE-2 ve ROUGE-L skorlarındaki üstünlüğüyle bağlamsal anlamayı ve metin detaylarını etkili bir şekilde kavradığını kanıtlamıştır. BERTSum'un derin öğrenme tabanlı yapısı ve önceden eğitilmiş güçlü BERT modelini kullanması, metinlerin bağlamını ve dil bilgisel doğruluğunu olağanüstü bir şekilde anlamasına olanak tanımaktadır.

Özetlenen metinlerden anahtar kelimeler YAKE algoritması kullanılarak çıkarılmaktadır. Elde edilen anahtar kelimeler kullanılarak, güvenilir olarak belirlenen haber sitelerinden ilgili haberler tekrar çekilmektedir. Bu sayede model, benzer içerikleri karşılaştırmak için ilgili haberleri otomatik olarak bulmaktadır. Özetlenmiş olan metinlerin benzerlik hesaplamalarında kullanılan yöntemler arasında farklı performans seviyeleri bulunmaktadır. N-gram yöntemi kelime bazlı benzerlik analizi yaparak daha hızlı sonuçlar sunmaktadır. Fakat anlamsal bağlamı değerlendirmemektedir. SBERT yöntemi, anlamsal benzerlikleri bağlamsal olarak yakalayarak yüksek doğruluk sağlar ve özellikle uzun metinlerde güçlü bir performans gösterir. Bu durum, her yöntemin güçlü ve zayıf yönlerinin olduğunu ve tek bir yöntemle her türlü metin benzerliğinin doğru bir şekilde değerlendirilmesinin zor olduğunu göstermektedir. Çalışmada önerilen Ortak Skor hesabıyla, SBERT, N-gram ve Kosinüs yöntemlerini birleştirerek daha dengeli bir benzerlik ölçümü sunulmuştur. Tablo 2' de çeşitli yöntemlerin sonuçlarına dayanarak her birinin farklı avantajlar sunduğu görülmektedir.

Tablo 2. Metin benzerliği ölçüm yöntemlerinin başarı metriklerine göre karşılaştırması

Model	Kesinlik	Duyarlılık	F1 Skor
Ortak_score	0.90	0.85	0.87
BERT_score	0.85	0.80	0.82
Word2Vec	0.80	0.75	0.77
N-Gram	0.78	0.70	0.74
SBERT	0.88	0.82	0.85
TD-IDF	0.75	0.68	0.71

İncelenen haber metniyle anlamsal olarak benzer, farklı ve alakasız haber metinleri seçilerek derlenmiş ve modeller üzerinde skor hesaplaması yapmak için test edilmiştir. Metin benzerlik yöntemlerinin başarımını değerlendirmek amacıyla; Kesinlik, Duyarlılık ve F1 Skor gibi yaygın olarak kullanılan metrikler kullanılmıştır.

Kesinlik, gerçek değeri pozitif olup pozitif olarak sınıflandırılan tahmin sayısının tahmin değeri pozitif olan tüm gözlemlere oranıdır. En kötü kesinlik değeri 0 iken, en iyi kesinlik değeri 1 değerine sahiptir. Duyarlılık, gerçek değeri pozitif olan değerlerin ne kadarının pozitif olarak etiketlendiğini gösterir. Gerçek pozitif sayısı ne kadar fazla

olursa duyarlılık değeri artarken yanlış negatif sayısı ne kadar fazla olursa bu sefer de duyarlılık değeri azalır. F1 skoru, modelin kesinlik ve duyarlılık değerlerinin harmonik ortalaması alınarak hesaplanır. F1 skoru 0 ile 1 arasında değer almaktadır. F1 skoru değeri 1'e yaklaştıkça kategorizasyon sonuçları da o kadar doğru olmaktadır [34]. F1 skorunun harmonik ortalamasının hesaplanması ile ortak skorun sınır değeri belirlenmiştir ve denklem 6' da verilmiştir.

$$F1 = 2 \times \frac{\text{Kesinlik} \times \text{Duyarlılık}}{\text{Kesinlik} + \text{Duyarlılık}} \quad (6)$$

Tablo 2'de hesaplanan değerlere göre Ortak Skor değerlerinin diğer modellere göre daha iyi başarımlar gösterdiği gözlenmiştir. Ortak Skor hesabı, SBERT'in bağlamsal anlam yakalama gücünü, N-gram'ın yüzeysel yapıları değerlendirme becerisini ve Kosinüs benzerliğinin matematiksel sağlamlığını bir araya getirmektedir. Böylece, hem metnin semantik anlamı hem de yüzeysel benzerliği değerlendirilerek daha güvenilir ve kapsamlı bir sonuç elde edilmiştir. Özellikle bağlamsal ve yüzeysel benzerliğin aynı anda önemli olduğu uygulamalarda, önerilen yöntemin daha üstün performans göstermesi beklenmektedir. Bu nedenle Ortak Skor yöntemi, yalnızca doğruluk açısından değil, aynı zamanda farklı yöntemlerin güçlü yönlerini birleştirerek güvenilirlik ve uygulama kolaylığı açısından da tercih edilmelidir.

Haberlerin doğruluğunu analiz etmek için, orijinal haber ile çekilen haberler arasındaki benzerlikler karşılaştırılmaktadır. Bu adımda iki farklı yöntem kullanılmaktadır. Anlamsal benzerlik için SBERT ve kelime benzerliği N-gram yöntemleri kullanılmaktadır. SBERT modeli ile metinlerin anlamsal benzerliği karşılaştırılmıştır.

Metindeki kelime benzerliğini hesaplamak için de N-gram benzerlik ölçütü kullanılmıştır. Bu iki yöntemden elde edilen sonuçlara göre bir ortak skor hesaplanmaktadır. Ortak Skor hesabı Denklem 7'de verilmiştir.

$$\text{Ortak Skor} = (K.S. \times K.A.) + (A.S. \times A.A.) \quad (7)$$

Denklem 7, çalışmaya özgü olarak geliştirilen Ortak Skor tabanlı karar mekanizmasına aittir. Literatürdeki benzerlik ölçütlerinden farklı olarak, kelime düzeyinde ve anlamsal düzeyde elde edilen sonuçları ayrı ayrı değerlendirip bunları belirli ağırlıklarla birleştirerek nihai doğruluk kararını üretmektedir. Ortak Skor, Kelime Skor (K.S.) ve Kelime Ağırlığının (K.A.) çarpımı ile Anlamsal Skor (A.S.) ve Anlamsal Ağırlığın (A.A.) çarpımlarının toplamı ile hesaplanmaktadır. Kelime Ağırlığı ve Anlamsal Ağırlık değerleri, metnin anlamsal ve yapısal bağlamının önemine göre değiştirilebilmektedir. Bu projede, haber metni türleri için anlamsal bağlamın daha önemli olduğu göz önünde bulundurularak Anlamsal Ağırlık için 0.7 ağırlığı ve Kelime Ağırlığı için ise 0.3 ağırlığı kullanılmıştır. Son aşamada SBERT, N-gram ve Kosinüs benzerliği yöntemleri kullanılarak tanımlanan Ortak Skor değeri, denklem 6' da belirlenen F1 skorunun harmonik ortalamasına göre hesaplanan değerden büyükse metinlerin benzer olduğu ve haberin gerçek olduğu, ortak skor değeri hesaplanan değerden küçükse metinlerin benzer olmadığı ve haberin sahte olduğu sonucuna varılmaktadır. Böylece önerilen model, haberlerin doğruluğunu belirlemek için metin işleme, özetleme, anahtar kelime çıkarma ve benzerlik analizini bir arada kullanarak otomatik bir doğrulama süreci sunmaktadır. Güncel olaylarla ilgili yanıltıcı içeriklerin tespitini hızlandırarak, kamuoyunun doğru bilgiye ulaşmasını ve ulusal güvenlik açısından risk taşıyan manipülatif girişimlerin önlenmesini sağlamaktadır.

Tablo 3. Literatürdeki çalışmalarla karşılaştırma sonuçları

Kaynak	Veri Seti	Yöntem	F1 Skoru
12	Twitter-SPO-KnowledgeBase	BERT, LSTM, RNN	84%
15	Etiketsiz Sosyal Medya Verisi	Weakly Supervised SVM, Bi-LSTM, TF-IDF	85%
2	Twitter Türkçe Haberler	KNN, SVM, K-Means, LDA	86%
Önerilen Model	Gerçek Zamanlı Veri (Anadolu Ajansı, Teyit.org)	BERTSum, YAKE, SBERT, N-Gram (Ortak Skor Yöntemi)	87%

Önerilen sistemin performansı, literatürdeki benzer çalışmalarda sunulan yaklaşımlarla karşılaştırılmıştır. Yukarıda verilen Tablo 3 incelendiğinde, sabit veri kümeleri

ile eğitilen geleneksel makine öğrenmesi ve derin öğrenme modellerine kıyasla önerilen modelin daha yüksek F1 Skor değerine ulaştığı görülmektedir. Özellikle, gerçek zamanlı

ve dinamik veri akışı üzerinden çalışabilmesi sayesinde, önerilen sistem yalnızca doğruluk açısından değil, aynı zamanda uygulanabilirlik ve güncellik açısından da avantaj sunmaktadır. Diğer çalışmalarda genellikle sabit ve etiketlenmiş veri setleriyle çalışıldığı görülürken, bu çalışmada haber içeriklerinin doğrudan kaynaklardan çekilmesi ve otomatik analiz edilmesi, sistemi pratikte daha esnek ve etkili kılmaktadır. Ayrıca önerilen modelin Anadolu Ajansı ve Teyit.org gibi güvenilir kaynaklardan elde edilen güncel verilerle desteklenmesi, sahte haber tespitinde gerçek zamanlı karar destek sistemleri açısından önemli bir avantaj sunmaktadır.

Sonuçlar

Bu çalışmada, dijital medyada hızla yayılan sahte haberlerin etkili bir şekilde tespit edilmesi amacıyla derin öğrenme ve doğal dil işleme (NLP) tabanlı yenilikçi bir haber teyit sistemi geliştirilmiştir. Geliştirilen sistem, dijital haber içeriklerinin güvenilirliğini analiz ederek kullanıcıları yanlış bilgilendirme riski taşıyan içeriklere karşı korumayı hedeflemiştir. Çalışmada sabit bir veri seti kullanılmasının yerine, sistemin güncel ve gerçek zamanlı veriyle çalışması hedeflenmiştir. Anadolu Ajansı ve Teyit.org gibi güvenilir haber kaynaklarından otomatik olarak çekilen haber içerikleri doğrultusunda değerlendirmeler yapılmıştır. Özellikle anlık sahte haber tespiti hedeflendiğinden, sistemin performansı dinamik veriler üzerinden, güncel haber akışlarıyla değerlendirilmiştir.

Sistem, haber metinlerini hem yüzeysel hem de anlamsal düzeyde analiz ederek şüpheli içeriklerin belirlenmesine olanak sağlamıştır. Bu kapsamda, SBERT ve YAKE gibi modellerin kullanılması, metinlerin derin anlamsal özelliklerinin çıkarılmasına ve anahtar kelimelerle yüzeysel benzerliklerin etkili bir şekilde ölçülmesine olanak tanımıştır. SBERT modeli, metinler arasındaki derin anlamsal ilişkileri anlamada güçlü bir araç olarak öne çıkmıştır. SBERT'in bağlamsal anlamı yüksek doğrulukla yakalayabilmesi, sahte haberlerin tespitinde yanlış pozitif oranlarını düşürmüştür. YAKE algoritmasının eğitim veri setine ihtiyaç duymadan çalışabilmesi, sistemin daha esnek ve uygulanabilir olmasını sağlamıştır. N-gram benzerliği gibi yöntemlerle derin öğrenme modellerinin bir arada kullanılması, sahte haber tespitinin doğruluk oranını önemli ölçüde artırmıştır.

Kullanılan yöntemin bir diğer önemli avantajı, belirli veri setleriyle sınırlı kalmayıp güncel haberler üzerinde çalışabilmesidir. Bu sistem, anlık veri çekme özelliği sayesinde haberlerin en hızlı şekilde doğrulanmasını sağlamaktadır. Güncel olaylara yönelik sahte haberlerin gerçek bilgilerle kıyaslanarak hızlı tespit edilmesi, özellikle kriz anlarında yanlış bilgilendirmenin önüne geçmek açısından kritik bir öneme sahiptir. Bu özellik, sistemin yalnızca tarihsel veri kümeleriyle değil, dinamik ve sürekli değişen haber içerikleriyle çalışmasını mümkün kılmakta, böylece sahte haber tespit süreçlerinin etkili bir şekilde gerçek zamanlı olarak yürütülmesini sağlamaktadır.

Çalışmanın sonuçları, önerilen sistemin sahte haberlerin tespiti konusunda yüksek başarı elde ettiğini göstermiştir. Metin işleme, özetleme ve benzerlik analizi gibi birden fazla yöntemi entegre eden bu sistem, hem bireysel kullanıcılar hem de toplumsal düzeyde fayda sağlamayı hedeflemiştir. Modelin, hem büyük veri kümelerinde çalışabilmesi hem de süreçlerin otomatikleştirilmesi, manuel doğrulama yöntemlerinin zaman ve ölçkleme sorunlarına etkili bir çözüm sunmuştur. Bununla birlikte, güvenilir kaynaklardan alınan bilgilerle desteklenen çok katmanlı bir doğrulama süreci, sistemin sağlamlığını ve güvenilirliğini artırmıştır.

Bu çalışmanın ilerleyen aşamalarında, geliştirilen sistemin çok dilli bir yapıya uyarlanarak farklı dillerdeki haberlerin doğrulanmasında kullanılabilecek şekilde genişletilmesi hedeflenmektedir. Özellikle uluslararası medya kaynaklarındaki sahte haberlerin tespiti için, çok dilli derin öğrenme modellerinin (mesela, mBERT) entegrasyonu önerilebilir. Modelin farklı sahte haber veri setleriyle karşılaştırmalı olarak test edilmesi de planlanmaktadır. Böylece önerilen yapının farklı içerik türlerinde ve çok dilli ortamlarda da geçerliliği değerlendirilebilecektir. Ek olarak, sistemin sadece metin analizine dayalı bir yapıdan çıkarılarak, görsel ve video içeriklerin de analiz edilebileceği çok modaliteli bir doğrulama sistemi haline getirilmesi mümkündür. Bu tür bir geliştirme, özellikle sosyal medya platformlarında yaygınlaşan manipülatif görsel ve video içeriklerin tespitinde etkili bir çözüm sunabilir.

Sonuç olarak, geliştirilen haber teyit sistemi, toplumun doğru bilgiye erişimini kolaylaştıran, ulusal güvenlik risklerini azaltan ve sahte haberlerin etkilerini en aza indiren yenilikçi bir çözüm olarak öne çıkmaktadır. Dijital medyanın giderek daha fazla önem kazandığı günümüzde, bu tür sistemlerin geliştirilmesi, bireylerin ve toplumların bilgi güvenliği açısından büyük bir öneme sahiptir.

Etik kurul onayı ve çıkar çatışması beyanı

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

Yazar Katkıları

Nurcan Yardımcı: Model tasarımı, deneysel çalışma, veri toplama, metodoloji, taslak hazırlama ve revizyon.

Hatice Karaoğlu: Veri ön işleme, deneysel çalışma, veri toplama, metodoloji ve taslak hazırlama.

Burhan Ergen: Denetleme, teorik katkı, düzenleme ve metodoloji.

Teşekkür

Bu çalışma TÜBİTAK 1002-A Hızlı Destek Modülü, 124E844 numaralı proje kapsamında desteklenmektedir.

Kaynaklar

- [1] M. F. Çömlekçi, "Sosyal medyada dezenformasyon ve haber doğrulama platformlarının pratikleri," *Gümüşhane Üniversitesi İletişim Fakültesi*

- Elektronik Dergisi*, vol. 7, no. 3, pp. 1549-1563, Dec. 2019.
- [2] S. G. Taşkın, E. U. Küçükşille, and K. Topal, "Twitter üzerinde Türkçe sahte haber tespiti," *Balıkesir Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, vol. 23, no. 1, pp. 151-172, 2021, DOI:10.25092/baunfbed.843909
- [3] A. M. Kırık, and N. Samadli, "Fake News Spread in Online Media During The Second Karabakh War" *RESS Journal*, vol. 11, no. 5, pp. 97-108, Sep. 2024, DOI: 10.17121/ressjournal.3571.
- [4] S. S. Akyüz, and M. Özkan, "Kriz Dönemlerinde Enformasyon Süreçleri: Ukrayna-Rusya Savaşında Dolaşıma Giren Sahte Haberlerin Analizi" *Uluslararası Kültürel ve Sosyal Araştırmalar Dergisi*, vol. 8, no. 2, pp. 66-82, Dec. 2022, DOI: 10.46442/intjcss.1213993.
- [5] N. Reimers, and I. Gurevych "Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks" *arXiv preprint arXiv vol. 1908*, pp. 10084, Aug. 2019.
- [6] Y. Chu *et al.*, "Refined SBERT: Representing sentence BERT in manifold space," *Neurocomputing*, vol. 555, pp. 126453, Jul. 2023.
- [7] I. Lopez-Gazpio *et al.*, "Word n-gram attention models for sentence similarity and inference," *Expert Systems with Applications*, vol. 132, pp. 1-11, Apr. 2019.
- [8] M. Sudhakar, and K. P. Kaliyamurthi, "Detection of fake news from social media using support vector machine learning algorithms," *Measurement: Sensors*, vol. 32, pp. 101028, Feb. 2024.
- [9] R. Campos, *et al.*, "YAKE! Keyword extraction from single documents using multiple local features," *Information Sciences*, vol. 509, pp. 257-289, Sep. 2019.
- [10] Z. Xu, and J. Zhang, "Extracting keywords from texts based on word frequency and association features," *Procedia Computer Science*, vol. 187, pp. 77-82, 2021.
- [11] G. Yavuz, "Web kazıma ve duygu analizi temelli ürün analiz sistemi/Web scraping and sentiment analysis based product analysis system," M.S. thesis, Aksaray Üniversitesi Sosyal Bilimler Enstitüsü, 2023.
- [12] V. Nair, J. Pareek, and S. Bhatt, "A Knowledge-Based Deep Learning Approach for Automatic Fake News Detection using BERT on Twitter," *Procedia Computer Science*, vol. 235, pp. 1870-1882, 2024.
- [13] M. N. Awan, and M. O. Beg, "Top-rank: a topicalpositionrank for extraction and classification of keyphrases in text," *Computer Speech & Language*, vol. 65, pp. 101116, Jun. 2020.
- [14] E. Ş. Dinçer, D. Kayaoğlu, and S. Sarı, "Metin madenciliği ve duygu analizi ile siber zorbalık tespiti," *Eskişehir Türk Dünyası Uygulama ve Araştırma Merkezi Bilişim Dergisi*, vol. 3, no. 2, pp. 38-45, Apr. 2022.
- [15] L. Syed, *et al.*, "Hybrid weakly supervised learning with deep learning technique for detection of fake news from cyber propaganda," *Array*, vol. 19, pp. 100309, Jul. 2023.
- [16] G. K. Koru and Ç. Uluyol, "Detection of Turkish fake news from tweets with BERT models," *IEEE Access*, vol. 12, pp. 14918-14931, 2024, DOI: 10.1109/ACCESS.2024.3354165
- [17] F. A. Özbay and B. Alataş, "Çevrimiçi sosyal medyada sahte haber tespiti," *Dicle Üniversitesi Mühendislik Fakültesi Mühendislik Dergisi*, vol. 11, no. 1, pp. 91-103, 2020, DOI: 10.24012/dumf.629368
- [18] İ. Kulaksız, "Doğal Dil İşleme Kullanılarak Ana Akım Medyada Türkçe Sahte Haber Tespiti" Yüksek Lisans Tezi, Bilgisayar Mühendisliği ABD, Atatürk Üniversitesi Fen Bilimleri Enstitüsü, Erzurum, Türkiye, 2025.
- [19] N. Z. KAYALI and S. İ. OMURCA, "Candidate Summary Elimination Approach for Turkish Hybrid Text Summarization," *Innovations in Intelligent Systems and Applications Conference (ASYU)*, pp. 1-6, 2024, DOI: 10.1109/ASYU62119.2024.10756999.
- [20] T. Uçkan, and K. Karabulut, "The Effectiveness of Machine Learning Algorithms in Extractive Text Summarization: A Comparative Analysis of K-Means, Random Forest, GBM, Logistic Regression, and SVM," *Doğu Fen Bilimleri Dergisi*, vol. 7 no. 2, pp. 77-91, 2024, DOI: 10.57244/dfbd.1538959.
- [21] A. Nenkov, and K. McKeown, "A survey of text summarization techniques," *Mining text data*, pp. 43-76, Jan. 2012.
- [22] A. C. Akdeniz, "Makine öğrenmesi yöntemleri kullanılarak uzaktan eğitim konulu Türkçe tweetlerin duygu analizi," M.S. thesis, Konya Teknik Üniversitesi, Apr. 2022.
- [23] C. Hark, *et al.*, "Metin özetlemesi için düğüm merkezliklerine dayalı denetimsiz bir yaklaşım," *Bitlis Eren Üniversitesi Fen Bilimleri Dergisi*, vol. 8, no.3, pp. 1109-1118, 2019.
- [24] A. İ. Kısayol, and M. Turan, "Paragraf Tabanlı Çıkarımsal Özetlemede Öbekleme Kullanan İki Yeni Yöntemin Kıyaslanması," *Düzce Üniversitesi Bilim ve Teknoloji Dergisi*, vol.6, no. 4, pp. 1047-1057, 2018.
- [25] G. Soğancıoğlu, "A semantic sentence similarity estimation approach for the biomedical domain," Doctoral dissertation, Boğaziçi University, 2013.
- [26] Y. Liu, and M. Lapata, "Text summarization with pretrained encoders," *arXiv preprint arXiv 1908.08345*, Sep. 2019.
- [27] A. S. Birdevrim, A. Boyacı, and D. A. S. Al Thani, "İYİLEŞTİRİLMİŞ OTOMATİK ANAHTAR KELİME ÇIKARIMI BRAKE," *İstanbul Ticaret Üniversitesi Teknoloji ve Uygulamalı Bilimler Dergisi*, vol. 1, no. 1, pp. 11-19, 2018.
- [28] A. Çelikten, H. Bulut, and A. Onan, "A semantic Similarity-Based approach to extract respiratory disease-symptom relations from biomedical literature," *Journal of the Faculty of Engineering and Architecture of Gazi University*, vol. 40, no. 1, pp. 121-134, 2025.
- [29] L. Soykan, C. Karsak and İ. D. El-Kahlout, "Adult Content Detection in Search Engine Queries," *2022 30th Signal Processing and Communications Applications Conference (SIU)*, pp. 1-4, 2022, DOI: 10.1109/SIU55565.2022.9864759.
- [30] A. Karaca, and Ö. Aydın, "Transformatör mimarisi tabanlı derin öğrenme yöntemi ile Türkçe haber metinlerine başlık üretme," *Gazi Üniversitesi*

- Mühendislik Mimarlık Fakültesi Dergisi*, vol. 39, no. 1, pp. 485-496, 2024, DOI: 10.17341/gazimmfd.963240
- [31] E. Erdağı, and V. Tunalı, "Comparison of feature-based sentence ranking methods for extractive summarization of turkish news texts," *Sigma Journal of Engineering and Natural Sciences*, vol. 42, no. 2, pp. 321-334, Apr. 2024, DOI: 10.14744/sigma.2023.00076.
- [32] S. D. Myla, E. R. Saini and E. N. Kapoor, "Auto Text Summarization in Natural Language Processing: Review," *2nd International Conference on Intelligent Data Communication Technologies and Internet of Things (IDCIoT)*, pp. 1258-1267, Jan. 2024, DOI: 10.1109/IDCIoT59759.2024.10467376.
- [33] T. Zhang, *et al.*, "Bertscore: Evaluating text generation with bert," *arXiv preprint arXiv:1904.09675*, 2019.
- [34] W. Zhao, *et al.*, "MoverScore: Text generation evaluating with contextualized embeddings and earth mover distance," *arXiv preprint arXiv:1909.02622*, Sep. 2019.
- [35] D. Demirbilek and M. Ersoy, "Sosyal Medya Paylaşımlarında Karar Mekanizmalarının Öğrenme Algoritmalarıyla Karşılaştırmalı Analizi," *Uluslararası Sürdürülebilir Mühendislik ve Teknoloji Dergisi*, vol. 8, no. 1, pp. 8–25, 2024, DOI: 10.62301/usmtd.1462808.