

Yeni Medya ve Dezenformasyon: Makine Öğrenmesi ile Sahte Haberlerin Tespitine İletişimsel Bir Yaklaşım

New Media and Disinformation: A Communicative Approach to Detecting Fake News Using Machine Learning

İsnur İnci ARMUTLU¹, Şahide ŞEKER²

Başvuru Tarihi/Submitted: 03.03.2025 Kabul Tarihi/Accepted: 22.11.2025

Makale Türü/Article Type: Araştırma Makalesi/Research Article

Öz

Yeni medya, bilgiye hızlı erişim ve geniş kitlelere anında ulaşım imkânı sunarak iletişim alanında köklü değişimlere yol açmaktadır. Ancak bu dönüşüm, aynı zamanda dezenformasyonun yayılmasını kolaylaştırarak sahte haberlerin artmasına neden olmuştur. Sahte haberlerin tespiti için geliştirilen yöntemler arasında, makine öğrenmesi algoritmaları öne çıkmaktadır. Bu çalışmada, sahte haber tespitinde makine öğrenmesi tekniklerinin etkinliği incelenmiştir. Çalışmada, Kaggle platformunda yer alan "Fake and Real News Dataset" veri seti kullanılarak lojistik regresyon modeli ile bir sınıflandırma gerçekleştirilmiştir. Haber metinleri, TF-IDF yöntemiyle vektörleştirilmiş ve modelin eğitimi bu özellikler üzerinden yapılmıştır. Modelin başarısı doğruluk, kesinlik, duyarlılık ve F1-skoru gibi değerlendirme metrikleriyle ölçülmüştür. Elde edilen bulgular, lojistik regresyonun sahte haberleri tespit etmede yüksek performans gösterdiğini ortaya koymaktadır. Çalışma, dijital medya ekosisteminde sahte haberlerin tespitinde veri odaklı yaklaşımların önemini vurgulamakta ve medya okuryazarlığının artırılmasına yönelik öneriler sunmaktadır.

Anahtar Kelimeler: Yeni medya, iletişim, dezenformasyon, makine öğrenmesi, lojistik regresyon modeli

JEL Kodu: M3, C8

Abstract

New media has brought fundamental changes to the field of communication by providing rapid access to information and enabling instant reach to large audiences. However, this transformation has also facilitated the spread of disinformation, leading to an increase in fake news. Among the methods developed for detecting fake news, machine learning algorithms stand out. This study examines the effectiveness of machine learning techniques in detecting fake news. A classification was performed using a logistic regression model on the Fake and Real News Dataset available on the Kaggle platform. The news texts were vectorized using the TF-IDF method, and the model was trained based on these features. The performance of the model was evaluated using accuracy, precision, recall, and F1-score metrics. The findings indicate that logistic regression demonstrates high performance in detecting fake news. This study highlights the importance of data-driven approaches in identifying fake news within the digital media ecosystem and offers recommendations for enhancing media literacy.

Keywords: New media, communication, disinformation, machine learning, logistic regression model

JEL Code: M3, C8

¹İstanbul Medipol Üniversitesi, Sosyal Bilimler Meslek Yüksek Okulu, Dr. Öğr. Gör., isnur.armutlu@medipol.edu.tr, ORCID: 0000-0003-0351-2493

²İstanbul Nişantaşı Üniversitesi, Lisansüstü Eğitim Enstitüsü, Yapay Zeka Mühendisliği Anabilim Dalı, sahideseker@hotmail.com, ORCID: 0009-0007-3228-3171

Giriş

20. yüzyılın son birkaç on yılında, basın, radyo ve televizyon gibi geleneksel kitle iletişim araçlarına, kendine özgü özelliklere sahip yeni medya ve iletişim araçları eklenmiştir. Riva'ya göre yeni medya, ortak bir terim olarak, küresel ölçekte yayılmış olup toplumsal iletişimde giderek daha büyük bir rol üstlenmekte ve benzeri görülmemiş bir gelişim göstermektedir (Levak, 2021). Günümüzde bile yeni medya ile "geleneksel medya" arasındaki anlam, kapsam ve içerik açısından kesin sınırlar çizmek zordur; özellikle de bazı temel özellikler bakımından bu ayrımı yapmak güçtür (Levak, 2021). Yeni medya, internet devriminin ardından ortaya çıkan geniş bir dijital iletişim platformları ve teknolojileri yelpazesini kapsar. Bu terim genellikle, kullanıcı tarafından oluşturulan içeriğe, sosyal ağlara ve gerçek zamanlı iletişime olanak tanıyan etkileşimli medya biçimlerini ifade eder. Yeni medyanın tanımları disiplinler arasında farklılık gösterse de genellikle sosyal medya, bloglar, podcast'ler ve çevrim içi video platformları gibi unsurları içerir. Yeni medya genellikle, üreticilerden tüketicilere tek yönlü iletişim ile karakterize edilen geleneksel medyanın karşısı olarak tanımlanır. Bu dönüşüm, bireyleri yalnızca içerik tüketicisinden aktif birer üreticiye dönüştürerek iletişim süreçlerinin yönünü köklü biçimde değiştirmiştir. Ayrıca, bilgi üretimi ve dolaşımı üzerindeki merkezî kontrolün zayıflaması, medya okuryazarlığını her zamankinden daha önemli hâle getirmiştir.

Ekawati'ye göre yeni medya, özellikle eğitim bağlamında, kullanıcıların içerik oluşturma ve yayma süreçlerine aktif olarak katılmasına olanak tanıyan iki yönlü bir etkileşimi kolaylaştırır (2022). Bu etkileşim, yeni medyanın temel özelliklerinden biri olup kullanıcıların geleneksel medya formatlarında mümkün olmayan şekillerde içerikle etkileşim kurmasını sağlar. Yeni medyanın dinamik yapısı, bilginin hızla yayılmasına ve çevrim içi toplulukların oluşmasına imkân tanıyarak, kamu söylemini ve bireysel davranışları önemli ölçüde etkileyebilir. Yeni medyanın yükselişi, özellikle influencer pazarlaması aracılığıyla, pazarlama stratejilerini de köklü bir şekilde dönüştürmüştür. Łaszkiewicz, markaların sosyal medya fenomenlerinden faydalanarak görünürlüklerini artırdığını ve hedef kitleleriyle etkileşimi güçlendirdiğini tartışmaktadır (2022). Bu bağlamda, yeni medya yalnızca bilgi üretiminde değil, tüketim alışkanlıklarının biçimlenmesinde de belirleyici bir güç hâline gelmiştir. Kullanıcıların aktif katılımı, markalar, kurumlar ve bireyler arasında çok katmanlı bir iletişim ekosisteminin oluşmasına zemin hazırlamıştır.

Bu değişim, tüketici güveninin geleneksel reklamlardan ziyade akran tavsiyelerine yönelmesi şeklindeki daha geniş bir eğilimi yansıtmaktadır. Literatür, sosyal medya platformlarının markalara daha kişiselleştirilmiş ve etkileşimli pazarlama kampanyaları oluşturma fırsatı sunduğunu ve bunun da daha yüksek tüketici katılımı ve marka sadakati sağladığını göstermektedir (Ersoy, 2021). Ayrıca, gerçekleştirilen araştırmada, sosyal medyanın özellikle gelişmekte olan pazarlarda işletmeler için yeni ağ oluşturma fırsatları sunduğunu ve girişimcilik ortamını yeniden şekillendirdiğini ortaya koymaktadır (Francesca et al., 2017). Yeni medyanın ortaya çıkışı, iletişim alanını, özellikle gazetecilik ve bilgi yayılımı alanlarını köklü bir şekilde dönüştürmüştür. Dijital teknolojilere dayanan yeni medya, anında erişim, etkileşim ve çeşitli formatlar sunarak geleneksel medya paradigmalarını zorlamaktadır. Guo'nun belirttiği gibi, yeni medya teknolojilerinin haber iletişimi endüstrisine entegrasyonu, profesyonellerin uyum sağlaması ve yenilik yapması için elzemdir; böylece haber yayılımının etkinliği artırılabilir (2021). Bu görüşü Tan da paylaşmaktadır; Tan, mikrobloglar ve çevrim içi video gibi yeni medya platformlarının iletişimin sosyoekonomik kalıplarını önemli ölçüde değiştirdiğini vurgulamaktadır. Bu dönüşüm, medya ile kitleler arasında daha dinamik bir etkileşime olanak tanımaktadır (Tan, 2015). Ayrıca, yeni ve geleneksel medyanın birleşimi, içerik üretimi ve dağıtımının yeniden tanımlanmasına yol açmıştır. Yeni medya, 1960'lardaki ilk kavramsallaştırılmasından bu yana evrilerek çeşitli dijital iletişim biçimlerini entegre eden güçlü bir platform hâline gelmiş ve böylece medya manzarasını yeniden şekillendirmiştir (Zhou, 2014). Bu birleşim, haberlerin daha kişiselleştirilmiş ve çeşitlendirilmiş bir şekilde sunulmasını kolaylaştırmaktadır. Yang, bunun, medya kuruluşlarının farklı kitlelerin tercihlerine hitap etmesine olanak tanıırken aynı zamanda haber iletiminde hız ve verimliliği artırdığını vurgulamaktadır (2024). Ancak, bu değişim bazı zorlukları da beraberinde getirmektedir. Geleneksel ve yeni medyanın birleşimi, özellikle kullanıcı tarafından oluşturulan içerikler bağlamında, otantiklik ve güvenilirlik konularında soru işaretleri doğurmaktadır (Chao-Chen, 2013). Dolayısıyla, yeni medya teknolojilerinin sunduğu hız ve erişim olanakları, bilgi doğruluğu, etik sorumluluk ve doğrulama mekanizmalarının yeniden değerlendirilmesini zorunlu kılmaktadır. Dijital ortamlarda bilgi üretimi ve dolaşımı giderek ticarileşirken, yanlış veya yanıltıcı içeriklerin yayılımı da hızlanmakta, bu durum ise medya profesyonelleri için yeni doğrulama standartlarının geliştirilmesini gerektirmektedir. Ayrıca, yapay zekâ destekli algoritmaların içerik önerme süreçlerine dâhil

edilmesi, kullanıcıların bilgiye erişim biçimini şekillendirirken filtre balonları ve yankı odaları gibi olguların etkisini artırmaktadır. Bu bağlamda yeni medya, yalnızca iletişim biçimlerini değil, bireylerin gerçeklik algısını da dönüştüren çok katmanlı bir sistem hâline gelmiştir. Sonuç olarak, dijital çağın sunduğu olanaklarla birlikte, medya okuryazarlığı, etik farkındalık ve eleştirel düşünme becerilerinin güçlendirilmesi, iletişim alanında sürdürülebilir bir denge kurmak için temel bir gereklilik olarak ortaya çıkmaktadır.

Yeni medyada görsel unsurların rolü göz ardı edilemez. Di, özellikle video aracılığıyla gerçekleşen görsel iletişimin, yeni medya ortamında kitleleri etkilemek ve onlarla etkileşim kurmak için hayati bir öneme sahip olduğunu vurgulamaktadır (2024). Bu durum özellikle sosyal medya bağlamında büyük önem taşımaktadır; zira içeriğin hızla tüketildiği bu ortamda, dikkat çekmek için görsel açıdan etkileyici formatlara ihtiyaç duyulmaktadır. Yeni medyanın kamu algısını ve marka kimliğini şekillendirmedeki etkinliği, şirketlerin iletişim stratejilerini, sosyal medya platformlarının benzersiz özelliklerinden yararlanacak şekilde uyarlamaları gerektiğini gösteren araştırmalarla da desteklenmektedir (Zhou, 2023). Yeni medyanın sunduğu avantajlara rağmen, yanlış bilgilendirme ve yankı odaları oluşumu gibi önemli endişeler de bulunmaktadır. Bu konuların özellikle yeni medya ortamında giderek daha kritik hâle geldiği vurgulanmaktadır (Zhang, 2024). Bu karmaşıklıklarla başa çıkabilmek için medya okuryazarlığı hayati bir öneme sahiptir; bu sayede kitleler, karşılaştıkları bilgileri eleştirel bir şekilde değerlendirebilir. Ayrıca, yeni medyanın geleneksel haber uygulamalarına entegrasyonu, gazetecilerin değişen rolü hakkında tartışmaları da beraberinde getirmiştir. Artık gazeteciler, izleyici katılımının hem bir gereklilik hem de gazetecilik etiği açısından potansiyel bir tehdit oluşturduğu bir ortamda yollarını bulmak zorundadır (Carlsson ve Nilsson, 2015). Bu çalışma, makine öğrenmesi yöntemleri kullanarak sahte haberleri tespit etmeyi amaçlamakta ve yeni medya ortamında dezenformasyonun yayılma dinamiklerini analiz etmektedir. Araştırma, dijital çağda bilgi üretimi ile doğruluk arasındaki kırılgan dengeyi yeniden düşünmeyi hedeflemektedir. Görsel yoğunluğun hâkim olduğu bir medya ekosisteminde, yapay zekâ destekli analiz araçlarının yalnızca teknik bir yenilik değil, aynı zamanda toplumsal farkındalık yaratma aracı olarak işlev görebileceği varsayılmaktadır. Böylelikle, çalışmanın sonuçları hem iletişim bilimi hem de dijital etik alanlarına katkı sunmayı amaçlamaktadır.

Bu çalışma, yeni medyanın sunduğu hızlı bilgi dolaşımı olanaklarıyla birlikte artan dezenformasyon sorununa odaklanarak sahte haberlerin tespitinde makine öğrenmesi yöntemlerinin iletişimsel bir çerçevede nasıl değerlendirilebileceğini incelemektedir. Yeni medyanın, demokratik katılımı güçlendiren ancak aynı zamanda bilgi kirliliğini hızla yaygınlaştıran yapısı, bu çalışmayı hem teorik hem de pratik açıdan önemli kılmaktadır. Çalışmanın temel amacı, dijital medya ortamında sahte haberlerin tespiti için geliştirilen lojistik regresyon modelinin etkinliğini değerlendirmek ve yapay zekâ tabanlı sistemlerin iletişim süreçlerine entegrasyonuna yönelik katkılar sunmaktır. Kuramsal olarak, araştırma yeni medya kuramı, dezenformasyon literatürü ve makine öğrenmesi yaklaşımlarının kesişiminden beslenmektedir. Bu çerçevede çalışma, iletişim bilimi ile veri bilimi arasında disiplinler arası bir köprü kurarak teknolojik gelişmelerin bilgi doğruluğu üzerindeki etkilerini analiz etmeyi hedeflemektedir. Araştırmanın sınırlılıkları arasında, yalnızca metin tabanlı haber verilerinin kullanılması ve görsel-işitsel içeriklerin kapsam dışında bırakılması yer almaktadır; ayrıca, modelin yalnızca İngilizce veri setiyle test edilmesi kültürel ve dilsel genelleme gücünü kısıtlamaktadır. Çalışma, giriş bölümünde yeni medya ve dezenformasyonun kavramsal arka planını tartışmakta; yöntem bölümünde veri seti, Python kütüphaneleri ve lojistik regresyon modeli detaylandırılmakta; bulgular ve tartışma bölümünde modelin performans metrikleri analiz edilmekte ve sonuç kısmında gelecekteki araştırmalara yönelik öneriler sunulmaktadır.

Çalışma kapsamında, sahte ve gerçek haberlerden oluşan etiketli bir veri seti kullanılarak makine öğrenmesi modellerinden olan lojistik regresyon (logistic regression) modelinin sahte haber tespitindeki performansı incelenmiştir. Veri ön işleme sürecinde metin madenciliği ve doğal dil işleme (NLP) teknikleri uygulanmış, model doğruluğu değerlendirme metrikleri ile analiz edilmiştir. Çalışmada, lojistik regresyon modeli kullanılmış ve modelin performansı; doğruluk (accuracy), kesinlik (precision), duyarlılık (recall), F1-skoru (F1-Score) ve ROC-AUC (Receiver Operating Characteristic- Area Under the Curve) gibi metrikler ile değerlendirilmiştir. Modelin eğitim sürecinde, haber metinleri üzerinde TF-IDF (Term Frequency- Inverse Document Frequency) vektörizasyonu uygulanarak metinlerin sayısal forma dönüştürülmesi sağlanmış, ardından model eğitilip test edilerek sonuçlar karşılaştırılmıştır. Elde edilen bulgular, makine öğrenmesi algoritmalarının sahte haber tespiti konusunda yüksek başarı sağladığını ortaya koymaktadır. Bu çalışma, yeni medya ortamında sahte haberlerin yayılmasını önlemek için makine öğrenmesi tabanlı yaklaşımların etkinliğini değerlendirerek literatüre önemli katkılar sunmayı hedeflemektedir. Sahte haberlerin tespitine

yönelik geliştirilen modellerin doğruluk oranlarının artırılması, medya okuryazarlığının güçlendirilmesi ve dezenformasyonla mücadelede yapay zeka tabanlı sistemlerin kullanımı, bu alandaki araştırmalar için önemli bir çerçeve sunmaktadır. Çalışmanın sonuçları, gelecekteki araştırmalara ışık tutarak, medya ve iletişim bilimleri ile yapay zeka alanlarının kesişiminde daha kapsamlı analizlerin yapılmasına katkı sağlamayı amaçlamaktadır.

Literatür Taraması

Globalleşme (küreselleşme) dünyanın tek bir mekân hâline gelmesi süreci olarak tanımlanabilmektedir. Immanuel Wallerstein gibi düşünürlerin kırk yıl öncesinde de globalleşme terimini ele aldıkları söylenebilir. Geleceğe dair düşüncelerin oluşumunu tarif ederken bugünkü durumu özetleyen bir kavram olarak globalleşme geçmişinin 500 yıla dayandığı görülmektedir. Dijitalleşme ise globalleşme ile paralel iletişim ve medyanın tasarım, tahayyül ve iletişim tarzını değiştiren bir niteliktedir. Diğer bir deyişle; medya teknolojisi gelişmelerinin ekseninde, globalleşen elektronik medya teknolojisi birincil konuma yükselirken matbuatın hâkim medya teknolojisi ikinci planda kalmıştır (Mutlu, 2016). Medya evriminin geniş kapsamlı dönüşümlerini kavramsallaştırma süreci uzun zamandır akademisyenlerin üzerinde çalıştığı bir konu olmuştur. Bu dönüşüm, iletişimin yalnızca fiziksel sınırları aşmasını değil, kültürel ve bilişsel sınırların da yeniden tanımlanmasını beraberinde getirmiştir. Dijitalleşme, küreselleşmeyi soyut bir ekonomik kavramın ötesine taşıyarak bireylerin, kurumların ve toplumların bilgi üretimi ve paylaşımı süreçlerini birbirine bağımlı hâle getirmiştir. Böylece, medya ekosistemi artık sadece içerik aktarımını değil, kültürel anlam üretimini de küresel ölçekte yeniden şekillendiren bir yapıya dönüşmüştür.

Peters (2009), Terminological Potential and the Problem of New Media bölümünde, medya tarihini anlamlandırmada kullanılan kavramsal çerçevelerin geçiciliğine dikkat çeker ve medya yeniliğinin belirli aşamalardan geçtiğini öne sürer. Joseph Schumpeter (1934), Eric Von Hippel (1990, 2005), Norbert Wiener (1993), Everett Rogers (1995) ve Rudolph Stöber (2004) gibi araştırmacılar da medya yeniliğini belirleyen etkenleri inceleyerek yeni medyanın doğasını şekillendiren faktörleri analiz etmişlerdir. Peters'a göre, yeni medya olarak adlandırılan birçok platform, aslında daha önceki medya biçimlerinden türeyerek ve dönüşerek varlık kazanmaktadır (2021). Bu bağlamda, yeni ve geleneksel medyanın birleşimi üzerine yapılan tartışmalar, medya tarihinin döngüsel doğasını göz önünde bulundurarak ele alınmalıdır. Başka bir deyişle, her yeni medya formu, kendinden önceki teknolojilerin mirasını taşıırken aynı zamanda onları dönüştürerek yeniden tanımlar. Bu perspektiften bakıldığında, dijital medya yalnızca teknik bir yenilik değil, tarihsel sürekliliğin dijital ortamda aldığı yeni biçim olarak da okunabilir. Dolayısıyla, medya tarihindeki ilerleme, eski biçimlerin yerini tamamen almaz; aksine, onları yeniden yorumlayarak çok katmanlı bir iletişim evreni oluşturur. Bu yaklaşım hem medya arkeolojisi hem de dijital kültür çalışmaları açısından yeni medyayı tarihsellikten koparmadan anlamaya olanak tanır.

"Yeni medya" terimi, 20. yüzyılın ikinci yarısında ortaya çıkmıştır. Oxford English Dictionary, bu terimin ilk kullanımını kelime ustası, şovmen ve iletişim bilimci Marshall McLuhan'a atfetmekte ve 1960 yılında Journal of Economic History'de geçtiğini belirtmektedir. Ancak, McLuhan bu terimi, en az 1953 yılı kadar erken bir tarihte, Harold Innis üzerine yazdığı bir Queen's Quarterly makalesinde de kullanmıştır. Bu durum, yeni medya teorisinin kökenlerinin Kanada'ya dayandığını iki katmanlı bir biçimde ortaya koymaktadır (B. Peters, 2009; C. Peters et al., 2021). Temelde, yeni medya, ağ teknolojileri tarafından şekillendirilen ve mesajların kullanıcılar tarafından dijital olarak oluşturulup yayıldığı bir sosyokültürel ortam olarak tanımlanabilir. Bu tanım, yeni medyanın katılımcı doğasını vurgulamakta ve bireylerin merkezî bir otorite yerine kolektif bir şekilde içerik üretmesine ve etkileşimde bulunmasına olanak tanıdığını göstermektedir (Tuğtekin ve Koç, 2019). Bu noktada, yeni medya yalnızca bir iletişim kanalı değil, aynı zamanda bilginin üretilme, dolaşıma girme ve meşrulaştırılma biçimlerini dönüştüren bir paradigma hâline gelmiştir. Kullanıcı merkezli bu yapı, klasik medya hiyerarşilerini çözerek bilgi üretiminde çok sesliliği teşvik etmekte, ancak aynı zamanda doğruluk ve güvenilirlik sınırlarını da tartışmaya açmaktadır. Böylece yeni medya hem demokratik katılımın bir aracı hem de dezenformasyonun çoğalabileceği bir alan olarak ikili bir doğaya sahiptir.

Yeni medyanın doğasında bulunan etkileşim, onu geleneksel medyadan ayıran temel unsurlardan biridir; zira geleneksel medyada iletişim genellikle tek yönlüdür. Bu dönüşüm, kullanıcıların yalnızca tüketici olmanın ötesine geçerek medya ortamına aktif olarak katkıda bulunduğu daha demokratik bir eğilimi de beraberinde getirmektedir (Tuğtekin ve Koç, 2019). Yeni medyanın evrimi, özellikle internetin yükselişi ve dijital iletişim araçlarının gelişimi ile yakından ilişkilidir. Tomin ve diğerleri, yeni medyayı bilgisayar ağları ve dijital depolama sistemlerinin getirdiği sosyokültürel değişimleri yansıtan disiplinler arası bir

kavram olarak tanımlanmaktadır (2020). Bu bakış açısı, farklı iletişim biçimlerinin bir araya gelmesiyle bilginin paylaşım ve tüketim biçimlerinin yeniden yapılandığını vurgulamaktadır. Bennett ise yeni medyayı, mobil telefonlar, akış hizmetleri ve Dünya Çapında Ağ (World Wide Web) gibi yeni teknolojileri kapsayan bir olgu olarak tanımlayarak bu unsurların bilgi yayılımının kalitesini ve erişilebilirliğini artırdığına dikkat çekmektedir (Widyastuti, 2021). Yeni medya, bu çok katmanlı yapısıyla bireyler arası iletişimi yalnızca hızlandırmakla kalmamış, aynı zamanda katılımın doğasını da değiştirmiştir. Kullanıcı merkezli üretim modelleri, medya içeriğini tek yönlü bir aktarım sürecinden çıkarıp kolektif bir anlam inşasına dönüştürmüştür. Böylelikle, bilgi artık yalnızca otoriteler tarafından değil, çevrim içi topluluklar aracılığıyla da şekillendirilen dinamik bir yapıya bürünmüştür.

Bununla birlikte, yeni medya kavramının zaman içinde değiştiğini gösteren çalışmalar da mevcuttur. Méndez ve diğerleri, dijital ve sosyal medyanın her yerde bulunur hâle geldiği günümüz dünyasında "yeni medya" teriminin bir ölçüde eskimeye başladığını belirtmektedir (2019). Bu gözlem, yeni medyanın yalnızca belirli bir teknoloji kategorisi olmadığını, aynı zamanda değişen kültürel pratikleri ve toplumsal normları da yansıttığını göstermektedir. Yeni medyanın günlük yaşama entegrasyonu, iletişim dinamiklerini kökten dönüştürmüş ve bireylerin bu karmaşık ortamda etkili bir şekilde yol alabilmesi için yeni medya okuryazarlığı becerilerini geliştirmelerini zorunlu hâle getirmiştir (Tuğtekin ve Koç, 2019). Yeni medya, dijital teknolojilerin ve internetin sağladığı olanaklarla birlikte, bilgi üretimi ve dağıtımında köklü değişiklikler meydana getirmiştir. Bu ortam, kullanıcıların içerik oluşturma, paylaşma ve etkileşimde bulunma yeteneklerini artırarak iletişim dinamiklerini köklü bir şekilde değiştirmiştir (Guo, 2021). Yeni medya, Web 2.0 özellikleri sayesinde bireyler arasında daha yüksek sosyal etkileşim düzeyleri sağlamakta ve kamu, yerel ve özel sektörler arasında daha etkili iletişim kanalları oluşturmaktadır (Carlsson ve Nilsson, 2015). Sosyal medya platformları, bireylerin içerik üretiminde aktif rol almasına olanak tanıırken bu durum, dezenformasyonun yayılmasını da hızlandırmaktadır. İnternetin yaygınlaşması, web tabanlı iletişim teknolojilerini, sosyal medya platformlarını ve mobil iletişimi başlatan "yeni medya" çağını başlatmıştır. Ancak günümüzde bu çağ, "sürekli yenilenen medya" dönemine evrilmiştir; zira teknolojik gelişmelerin hızı, kavramsal kategorilerin bile geçiciliğini göstermektedir. Artık "yeni medya"dan ziyade, kullanıcı davranışlarını ve bilgi akışını şekillendiren "dijital ekosistemler"den söz etmek daha yerinde görünmektedir. Bu nedenle, akademik tartışmaların yalnızca teknolojik yeniliklere değil, bu yeniliklerin sosyal etkilerine ve etik sonuçlarına da odaklanması gerekmektedir. Bu yaklaşım, dezenformasyon, algoritmik manipülasyon ve veri temelli iletişim gibi çağdaş sorunları anlamada disiplinler arası bir perspektifin önemini vurgulamaktadır.

Yeni medyanın, özellikle genç neslin hayatında belirleyici bir rolü bulunmaktadır (Öztunç, 2020). Yeni medyanın siyasi, toplumsal ve kültürel alanlardaki rolü ve etkisi de son dönemde dikkat çeken konular arasında yer almaktadır. Çeşitli çalışmalar, yeni medya platformlarının kamuoyu oluşturma, gündem belirleme ve siyasal katılımı teşvik etme yönlerini ele almaktadır (Toker et al., 2017). Teknolojik gelişmeler, politik, ekonomik ve toplumsal yapılarda dönüşümlere neden olmuştur. Bu dönüşümlerin bir sonucu olarak yeni medya ortaya çıkmıştır (Gürdağ et al., 2018). Yeni medya; kullanıcı merkezli, interaktif ve çok yönlü iletişimin sağlandığı, gelişen teknolojilerle şekillenmekte olan iletişim biçimidir (Sevim, 2019). Bu dönüşüm sadece iletişim biçimlerini değil, bireylerin dünyayı algılama ve anlamlandırma biçimlerini de yeniden şekillendirmektedir. Özellikle genç kuşak için yeni medya, kimlik inşasının, aidiyetin ve toplumsal katılımın en görünür zemini hâline gelmiştir.

Yeni medya ortamında, görsel unsurların önemi artmaktadır. Flanagin ve Metzger, internetin çağdaş medya ortamındaki rolünü inceleyerek yeni iletişim teknolojilerinin, kullanıcıların iletişim ihtiyaçlarını nasıl karşıladığını vurgulamaktadır (2006). Görsel iletişim, izleyicilerin içerikle daha derin bir bağ kurmasını sağlamakta ve bu da yeni medya ürünlerinin etkisini artırmaktadır. Ancak, bu konuda daha fazla araştırmaya ihtiyaç vardır. Ayrıca, yeni medya, markaların iletişim stratejilerini yeniden şekillendirmekte ve marka tasarımıyla yenilikçi yöntemlerin geliştirilmesine olanak tanımaktadır (Fraga-Lamas ve Fernández-Caramés, 2020). Yeni medya, aynı zamanda organizasyonel iletişimde de önemli bir rol oynamaktadır. Kaplan ve Haenlein'in araştırması, sosyal medyanın organizasyonel iletişim üzerindeki etkilerini incelemekte ve yeni medya araçlarının etkin bir şekilde kullanılmasının, organizasyonların iç iletişimlerini güçlendirdiğini göstermektedir (Guo, 2022). Bu durum, yeni medya kullanımının organizasyonel verimlilik ve etkileşim üzerindeki olumlu etkilerini ortaya koymaktadır. Dezenformasyon, yanlış veya yanıltıcı bilgilerin kasıtlı olarak yayılmasıdır ve bu bağlamda yeni medya, bilgi güvenilirliğini tehdit eden bir ortam yaratmaktadır. Karmila ve arkadaşları, sosyal medyanın siyasi kimlik inşasındaki rolünü inceleyerek bu platformların bireylerin siyasi katılımını nasıl dönüştürdüğünü ve bunun yanı sıra

dezenformasyon ve siyasi kutuplaşma gibi olumsuz sonuçları da beraberinde getirdiğini vurgulamaktadır (2024). Bu durum, yeni medyanın çift yönlü doğasını gözler önüne sermektedir; bir yandan demokratik katılımı ve ifade özgürlüğünü teşvik ederken diğer yandan bilgi kirliliği ve algı yönetimi risklerini artırmaktadır. Bu nedenle, dijital çağda medya etiği ve doğrulama mekanizmalarının güçlendirilmesi kritik bir gereklilik hâline gelmiştir.

Yeni Medya ve Dezenformasyonun Yayılması

Yeni medya, geleneksel medya araçlarının yerini alarak bireylerin bilgiye erişimini kolaylaştırmış, ancak aynı zamanda dezenformasyonun hızla yayılmasına da zemin hazırlamıştır. Pew Research Center tarafından yapılan bir araştırma, sosyal medya kullanıcılarının %64'ünün, sosyal medya üzerinden gördükleri haberlerin çoğunun yanlış bilgi içerdiğini düşündüğünü ortaya koymaktadır (Yamaguchi et al., 2025). Bu durum, sosyal medya platformlarının bilgi akışındaki rolünü sorgulatmaktadır. Yamaguchi ve Tanihara, medya okuryazarlığı ile dezenformasyon yayma davranışı arasındaki ilişkiyi inceleyerek yüksek medya okuryazarlığına sahip bireylerin yanlış bilgi yayma olasılığının daha düşük olduğunu göstermiştir (2023). Dezenformasyonun yayılmasında sosyal medya algoritmalarının etkisi büyüktür. Algoritmalar, kullanıcıların ilgi alanlarına göre içerik önerileri sunarak yanlış bilgilerin daha fazla görünür olmasına neden olmaktadır. Steinfeld, sosyal medya üzerindeki dezenformasyon savaşını ele alarak kullanıcıların yanlış bilgileri yayma ve bu bilgileri düzeltme çabalarını incelemektedir (2022). Bu bağlamda, sosyal medya platformlarının dezenformasyonla mücadele etme konusundaki yetersizlikleri eleştirilmektedir. Dolayısıyla, dijital ekosistemde dezenformasyonla mücadele yalnızca platformlara değil, bireylerin bilinçli medya kullanımına da bağlıdır. Bilgi ekosisteminin sürdürülebilirliği için medya okuryazarlığının eğitim sistemlerine entegre edilmesi ve algoritmik şeffaflığı artırılması büyük önem taşımaktadır.

Yanılıcı ve yanlış bilgi ile haberlerin çeşitli amaçlarla, çoğunlukla propaganda veya manipülasyon amacıyla tasarlanması, üretilmesi ve yayılması, insan iletişimi kadar eski bir olgudur. Bunun kökleri ve başlangıçları, genellikle kasıtlı olarak oluşturulmuş ve kamusal alanda dolaşıma sokulmuş yanlış veya hatalı bilgi ve haberlerin ağırlıklı olarak sözlü olarak yayıldığı antik çağlara kadar uzanmaktadır. Burkhardt'ın (2017) belirttiği gibi, insanlar güç odaklı gruplar hâlinde var olduklarından beri söylentiler ve uydurma hikâyeler de var olmuştur. İlk dönemlerde yeni bilgiler genellikle kişiden kişiye sözlü olarak aktarılıyordu (Levak, 2021). O zamandan günümüze kadar çeşitli kaynaklar ve yazarlar, bu tür bilgi üretimi ve yayılımına dair birçok örnek kaydetmişlerdir. Bu tür bilgiler en yaygın şekilde dezenformasyon veya günümüzde popülerleşen terimiyle "sahte haber" olarak adlandırılmaktadır. Bu fenomenin terminolojik yönleri ve daha ayrıntılı açıklamaları ilerleyen bölümlerde ele alınacaktır. Ancak burada tarih boyunca önemli veya belirleyici olarak işaretlenen bazı vakalara yer verilecektir. Burkhardt'a göre yazının ortaya çıkışı binlerce yıl öncesine dayanmakta olup taş, kil ve papirüs gibi malzemeler üzerine yazılmıştır ve bu bilgiler genellikle grup liderleri, yani imparatorlar, firavunlar, dinî ve askerî liderler gibi seçkin kişilerle sınırlıydı (2017). Bu güç, seçilmiş kişiler tarafından çeşitli şekillerde kullanılmış; bilgi kontrolü, etki yayma ve belirli hedeflere ulaşmak için sahte haber ve yanlış bilgi üretme ve yayma amacıyla değerlendirilmiştir (Levak, 2021). Bu durum, dezenformasyonun yalnızca dijital çağın bir ürünü olmadığını, aksine tarih boyunca iktidar ilişkilerinin bir parçası olarak biçimlendiğini göstermektedir. Günümüzde değişen yalnızca araçlar değil, bu araçların erişim hızı ve etkisinin küresel boyutlara ulaşmasıdır.

Tarihte kaydedilen ilk örneklerden biri, MÖ 6. yüzyıla kadar uzanmaktadır. Bizans tarihçisi Procopius, "Gizli Tarih" adlı eserinde İmparator Justinianus ve eşi hakkında ölümünden sonra yanlış bilgiler yayarak yeni yöneticiye yaranmaya çalışmıştır. Posetti ve Matthews ise dezenformasyon, yanlış bilgi ve propagandanın en az Roma dönemine kadar uzandığını belirtmektedir. Örneğin, MÖ 44 yılında Julius Caesar'ın suikastından sonra Octavian, rakibi Mark Antony'yi itibarsızlaştırmak amacıyla Roma sikkelerine sahte ve yanılıcı sloganlar bastırmıştır. Bu "dezenformasyon savaşı"nı kazanan Octavian, Roma'nın ilk imparatoru Augustus olmuştur (Levak, 2021). Son yıllarda, dijital teknolojilerin ve özellikle internet, sosyal medya ve iletişim platformlarının yükselişiyle birlikte dezenformasyon sorunu tehlikeli boyutlara ulaşmıştır. Özellikle sosyal medya, bireylerin sadece kendi görüşlerine uygun içeriklerle karşılaşmasını sağlayarak bilgi baloncukları yaratmaktadır. Bunun yanı sıra yeni medya ekosisteminin içinde dezenformasyonun tespit edilmesi ve önlenmesi için medya okuryazarlığının geliştirilmesi büyük bir önem taşımaktadır (Zhang, 2024). Bu tarihsel süreklilik, yanlış bilginin biçim değişirse de özünde aynı amaca —iktidar kurma, yönlendirme ve algı yönetimi hizmet ettiğini göstermektedir. Günümüz dijital ortamında bu durum, manipülasyonun daha görünmez ama daha yaygın hâle gelmesine yol açmaktadır.

Yeni medya okuryazarlığı, bireylerin dijital ortamda doğru bilgiye ulaşabilme yeteneklerini geliştirmeleri açısından kritik bir öneme sahiptir. Medya okuryazarlığı, bireylerin medya içeriklerini eleştirel bir bakış açısıyla değerlendirmelerini sağlamakta ve dezenformasyonla başa çıkabilme yeteneklerini artırmaktadır (Soares et al., 2021). Bu bağlamda, eğitim kurumlarının medya okuryazarlığına yönelik programlar geliştirmesi, genç bireylerin dezenformasyonla mücadele etme becerilerini artırmak için önemlidir. Özellikle genç yetişkinler arasında yapılan araştırmalar, medya okuryazarlığı düzeyinin artırılmasının dezenformasyonla mücadelede etkili olduğunu göstermektedir (Su et al., 2021). Yamaguchi, yüksek medya okuryazarlığına sahip bireylerin yanlış bilgi yayma olasılığının daha düşük olduğunu belirtmektedir (2023). Bu nedenle, eğitim politikalarının medya okuryazarlığını teşvik edecek şekilde yeniden yapılandırılması gerekmektedir. Ayrıca, medya okuryazarlığı yalnızca bir bilişsel beceri değil, aynı zamanda dijital etik, eleştirel farkındalık ve sorumlu paylaşım kültürüyle de ilişkilidir. Dijital çağın bireyleri, bilgiyi sadece tüketen değil, aynı zamanda doğruluğunu sorgulayan ve güvenilir bilgi üretimine katkıda bulunan aktörler hâline gelmelidir. Bu dönüşümün sağlanabilmesi için eğitim kurumlarının yanı sıra medya kuruluşları ve devlet politikalarının da ortak bir strateji geliştirmesi gerekmektedir. Böylece dezenformasyonun etkileri yalnızca bireysel düzeyde değil, toplumsal ölçekte de azaltılabilir.

Kuramsal Yaklaşım

Bu çalışma, Yeni Medya Ekolojisi Kuramı çerçevesinde ele alınmaktadır. Yeni Medya Ekolojisi Kuramı, dijital toplumların dinamiklerini inceleyen bir çerçeve sunarak bilgi toplumu anlayışını geliştiren önemli bir araştırma alanıdır. Bu kuram, medyanın ve teknolojinin sosyal, kültürel ve ekonomik etkileşimleri nasıl şekillendirdiğini analiz etmektedir. Özellikle sosyal medya platformlarının etkisi, yeni medya ekolojisinin modern toplumların bilgi üretim ve dağıtım biçimlerine nasıl dönüştürücü bir katkı sağladığını göstermektedir.

Sosyal medya, bireylerin bilgiye erişim yollarını köklü bir şekilde değiştirerek kamuoyunun oluşumunda önemli bir rol oynamaktadır. Gandini ve arkadaşlarının çalışması, dijital teknolojilerin sosyoteknik bileşenler olarak kamuoyunu şekillendirmedeki etkisini vurgularken aynı zamanda algoritmaların bu süreçteki belirleyici rolünü ele almaktadır (Gandini et al., 2022). Lee ve Valenzuela'nın önerdiği yeni çerçeve, sosyal medya ortamında bilgi edinme ve siyasi katılım üzerindeki etkileri incelemekte, klasik iletişim bilimleri kavramlarının yenilikçi bir şekilde revize edilmesi gerekliliğini ortaya koymaktadır (Lee ve Valenzuela, 2024).

Yeni medya ekolojisi, toplumun bilgi üretim süreçlerini demokratikleştirirken bireylere yaratıcı otonomi sağlamaktadır. Mehta'nın çalışması, sosyal medya platformları aracılığıyla içerik oluşturma, geleneksel medyanın kapalı ağlarının ötesine geçtiğini ve yeni yaratıcı ifade biçimlerini mümkün kıldığını belirtmektedir (Mehta, 2019). Bu durum, bilgi toplumunun gelişimi açısından önemli bir dönüşüm olarak değerlendirilmektedir.

Bununla birlikte, kamu kütüphanelerinin yeniden yapılandırılması ve sosyal medya araçlarının haber üretimindeki etkileri, bilgi toplumunun dinamiklerini daha geniş bir bağlamda anlamada önemli rol oynamaktadır. Blewitt'in belirttiği gibi kamusal kütüphaneler, sosyal ve medya alanları olarak sadece bilgi kaynakları değil, aynı zamanda toplumsal dönüşüm için kritik mekânlar haline gelmiştir (Blewitt, 2014). Olesen'in çalışmaları, yeni medya araçları ile protestolar arasında bağlantılar kurarak toplumsal eleştiri ve adalet arayışlarını ele almaktadır (Olesen, 2017, 2020). Dezenformasyonun bu dijital ekosistemde yarattığı tehditler göz önüne alındığında, iletişim bilimleri alanındaki kuramsal temellerin, doğruluk tespiti için kullanılan metodolojik yaklaşımlarla desteklenmesi zorunlu hâle gelmektedir.

İşte bu zorunluluk bağlamında, makine öğrenmesi yaklaşımları arasında lojistik regresyon ve TF-IDF (Term Frequency-Inverse Document Frequency) vektörizasyonu, basitliği, hızı ve yorumlanabilirlik avantajları nedeniyle yaygın olarak kullanılan ve güvenilir bir temel model (baseline) olarak kabul edilen bir yaklaşımdır. Literatürde, bu modelin performansını daha karmaşık algoritmalarla karşılaştıran çalışmalar bulunmaktadır. Örneğin, Aydın et al. (2018), Twitter verileri üzerinde sahte haber tespiti için lojistik regresyon ve destek vektör makineleri (SVM) algoritmalarını karşılaştırmış ve lojistik regresyonun daha iyi sonuçlar verdiğini göstermiştir.

Bu alandaki güncel çalışmalardan biri olan Adeyiga et al. (2023) tarafından yapılan araştırma ise lojistik regresyon modelinin %97,90'lık yüksek bir doğruluk (Accuracy) oranı elde etmesiyle, klasik modellerin dahi bu tür metin sınıflandırma görevlerinde etkili çözümler sunabildiğini kanıtlamaktadır. Bizim çalışmamız da bu literatürdeki eğilimi doğrulamakta ve lojistik regresyon modelinin, karmaşık derin

öğrenme modellerine gerek kalmadan, sahte haberlerin dilsel kalıplarını yüksek doğrulukla (%98,5) tespit edebildiğini göstererek yorumlanabilir ve hesaplama açısından verimli bir temel çözüm sunmaktadır.

Sonuç olarak yeni medya ekolojisi kuramı, bilgi toplumunun gelişimini destekleyen dinamikler sunarken aynı zamanda eleştirel bir bakış açısıyla bilgi edinme süreçlerini de sorgulamaktadır. Bu süreçlerin daha iyi anlaşılması, toplumsal katılımın artırılmasında ve bireylerin bilgiye erişim eşitsizliklerini azaltmada büyük önem taşımaktadır. Gelecekte, dijital teknolojilerin etkisi ve medyanın rolü, bilgi toplumu anlayışımızı derinleştirmeye devam edecektir.

Yöntem

Veri Seti

Bu çalışmada, sahte haberlerin tespitine yönelik bir makine öğrenmesi modeli geliştirmek amacıyla Kaggle platformunda yer alan "Fake and Real News Dataset" veri seti (URL1, 2025) kullanılmıştır. Günümüzde dezenformasyonun yayılması, özellikle dijital medyada önemli bir problem hâline gelmiştir. Sahte haberler, toplum üzerinde yanlış yönlendirme ve manipülasyon gibi ciddi sonuçlara yol açmaktadır. Bu nedenle, sahte haberleri otomatik olarak tespit eden güvenilir sistemlerin oluşturulması büyük önem taşımaktadır. Bu doğrultuda, çalışmada kullanılan veri seti, gerçek ve sahte haberleri içeren geniş bir veri kümesi olması sebebiyle uygun bulunmuştur. Kullanılan veri seti, toplam 44.898 haber metninden oluşmaktadır. Bu haberlerin 23.481'i sahte (fake), 21.417'si gerçek (true) olarak etiketlenmiştir. Veri seti; haber başlığı (title), haber metni (text), konu başlığı (subject) ve tarih (date) olmak üzere dört temel değişken içermektedir. Ancak, çalışmanın amacı haber metinlerinin analizi ile sahte haberleri tespit etmek olduğu için yalnızca title ve text sütunları kullanılmış, subject ve date sütunları veri ön işleme aşamasında çıkarılmıştır. Çalışmada kullanılan veri seti, etiket dengesizliği problemi içermemektedir. Genellikle sahte haber tespitinde kullanılan veri setlerinde dengesiz veri dağılımı (data imbalance) problemi görülebilir. Ancak, bu veri setinde sahte ve gerçek haber sayıları birbirine yakın olduğu için herhangi bir aşırı örnekleme (oversampling) veya eksik örnekleme (undersampling) işlemi uygulanmamıştır. Bununla birlikte veri setinin, haber kaynaklarına dair doğrudan bir bilgi içermemesi, modelin yalnızca metin verisi üzerinden öğrenmesine neden olmuştur. Ancak, sahte haberlerin dil özelliklerinin farklı olabileceği ve belirli kelime kalıplarının dezenformasyon içerikli haberlerde daha sık geçtiği göz önüne alındığında, bu tür modellerin başarılı sonuçlar verebileceği öngörülmektedir. Sonuç olarak bu veri seti, makine öğrenmesi teknikleri ile sahte haber tespit sistemlerinin performansını test etmek için uygun bir yapı sunmaktadır. Sahte ve gerçek haberleri içeren geniş kapsamlı veri kümesi, modelin çeşitli dil kalıplarını öğrenmesine olanak tanımaktadır. Bu çalışma kapsamında kullanılan veri seti, gelecekte daha büyük ve doğrulanmış kaynaklardan elde edilen veri kümeleri ile desteklenerek modelin genelleme kapasitesinin artırılması için önemli bir temel oluşturmaktadır.

Python Kütüphaneleri

Bu çalışmada, sahte haberlerin tespiti için kullanılan veri analizi, metin ön işleme, makine öğrenmesi modeli eğitimi ve değerlendirme süreçlerinde çeşitli Python kütüphanelerinden yararlanılmıştır. Makine öğrenmesi projelerinde doğru kütüphanelerin kullanılması hem model doğruluğunu artırmak hem de işlem süresini optimize etmek açısından kritik bir öneme sahiptir. Bu kapsamda, kullanılan kütüphaneler veri işleme, doğal dil işleme, model eğitimi ve değerlendirme gibi farklı aşamalarda etkin bir şekilde kullanılmıştır. Veri işleme ve manipülasyon süreçlerinde Pandas ve NumPy kütüphaneleri kullanılmıştır. Pandas, veri setinin yüklenmesi, temizlenmesi, eksik verilerin kontrol edilmesi ve dönüştürülmesi gibi işlemler için kullanılmıştır. NumPy ise sayısal işlemleri hızlandırarak veri manipülasyonlarını daha verimli hâle getirmiştir. Metin ön işleme sürecinde NLTK (Natural Language Toolkit) ve RE (Regular Expressions) kütüphanesi kullanılmıştır. NLTK, stopwords (önemsiz kelimeler) temizleme, lemmatization (kelime köklerine ayırma) ve metinlerde gereksiz kelimeleri filtreleme gibi işlemler için tercih edilmiştir. Özellikle önemsiz kelimelerin kaldırılması, modelin gereksiz bilgi ile eğitilmesini engelleyerek daha anlamlı kelimelerin ön plana çıkmasını sağlamıştır. Ayrıca, RE kütüphanesi kullanılarak haber metinlerinden noktalama işaretleri, özel karakterler ve sayılar temizlenmiş, modelin veriyi daha iyi öğrenmesi desteklenmiştir. Makine öğrenmesi modeli oluşturma ve değerlendirme aşamalarında Scikit-learn (sklearn) kütüphanesi kullanılmıştır (URL2, 2025). Scikit-learn, makine öğrenmesi modellerinin oluşturulmasını ve değerlendirilmesini kolaylaştıran güçlü bir kütüphanedir. Bu çalışmada TF-IDF vektörizasyonu, lojistik regresyon modelinin eğitimi ve tahmin yapılması, model performansının ölçülmesi işlemleri için kullanılmıştır. Metin verilerinin sayısal hâle getirilmesi için TF-IDF yöntemi uygulanmıştır. Bu yöntem, bir kelimenin belirli bir belge içindeki önemini hesaplayarak modelin daha bilgilendirici kelimelere daha

fazla ağırlık vermesini sağlamıştır. Model eğitimi için lojistik Regresyon algoritması kullanılmıştır. Lojistik regresyon, metin sınıflandırma problemlerinde yaygın olarak kullanılan ve yüksek başarı sağlayan bir yöntem olduğu için tercih edilmiştir. Modelin değerlendirme sürecinde ise doğruluk, kesinlik, duyarlılık ve F1-Skoru gibi metrikler hesaplanarak model performansı analiz edilmiştir.

Sonuçların görselleştirilmesi için Matplotlib ve Seaborn kütüphaneleri kullanılmıştır. Matplotlib, temel grafiklerin oluşturulmasında kullanılırken Seaborn verilerin daha anlaşılır hâle getirilmesini sağlayan istatistiksel görselleştirme teknikleri için kullanılmıştır. Çalışmada, karmaşıklık matrisi (confusion matrix), ROC eğrisi (Receiver Operating Characteristic Curve) ve Precision-Recall eğrisi gibi kritik model değerlendirme grafikleri oluşturularak modelin güçlü ve zayıf yönleri analiz edilmiştir. Sonuç olarak kullanılan tüm kütüphaneler; veri temizleme, model oluşturma, modelin değerlendirilmesi ve görselleştirme gibi aşamalarda verimli bir şekilde kullanılmış ve modelin başarısını artırmada önemli bir rol oynamıştır. Sahte haber tespiti gibi doğal dil işleme ve makine öğrenmesi gerektiren projelerde, uygun kütüphanelerin seçilmesi hem modelin doğruluğunu artırmakta hem de işlem sürecini optimize etmektedir.

Veri Ön İşleme

Makine öğrenmesi modellerinin başarılı bir şekilde çalışabilmesi için veri setinin özenle hazırlanması ve işlenmesi gerekmektedir. Özellikle doğal dil işleme çalışmalarında, ham metinlerin doğrudan modele verilmesi genellikle düşük performansa yol açmaktadır. Çünkü metin verileri; noktalama işaretleri, gereksiz kelimeler, büyük-küçük harf farklılıkları ve özel karakterler gibi modele doğrudan anlam kazandırmayan unsurlar içerebilmektedir. Bu nedenle, çalışmada kullanılan haber metinleri belirli veri ön işleme adımlarına tabi tutularak makine öğrenmesi modeline uygun hâle getirilmiştir. Veri ön işleme sürecinin ilk aşamasında, haber metinleri üzerinde kapsamlı bir metin temizleme işlemi gerçekleştirilmiştir.

Bu aşamada, haber başlıkları ve metinleri üzerinde yer alan noktalama işaretleri, özel karakterler ve sayılar temizlenmiştir. Noktalama işaretlerinin kaldırılması, kelimeler arasındaki anlam bütünlüğünü bozmadan modelin daha verimli çalışmasını sağlarken özel karakterler ve sayılar, modelin öğrenme sürecine anlamlı bir katkı sunmadığından çıkarılmıştır. Bunun yanı sıra haber metinlerindeki bütün kelimeler küçük harfe dönüştürülerek büyük/küçük harf duyarlılığı ortadan kaldırılmıştır. Böylece, aynı anlama gelen ancak farklı büyük/küçük harf kullanımları nedeniyle farklı kabul edilebilecek kelimelerin modelde tek bir şekilde işlenmesi sağlanmıştır. Metin temizleme aşamasının ardından, önemsiz kelimelerin çıkarılması işlemi gerçekleştirilmiştir. Doğal dil işleme tekniklerinde, bazı kelimeler dil bilgisel açıdan önemli bir değer taşımadığından, modelin öğrenme sürecine olumlu bir katkı sağlamamaktadır.

Bu çalışmada, NLTK kütüphanesinin sağladığı stopwords listesi kullanılarak model için anlamlı bilgi taşımayan kelimeler veri setinden çıkarılmıştır. Stopwords temizleme işlemi, modelin yalnızca anlam içeren kelimelerle çalışmasını sağlayarak modelin öğrenme sürecini hızlandırmış ve gereksiz bilgi yükünü azaltmıştır. Bir sonraki aşamada, haber metinlerindeki kelimelerin köklerini ayırma işlemi uygulanmıştır. Doğal dil işleme süreçlerinde, kelimelerin farklı eklerle kullanılması, modelin aynı köke sahip kelimeleri farklı kelimeler olarak algılamasına neden olabilmektedir. Bu sorunu gidermek için lemmatization yöntemi kullanılmıştır. Lemmatization, kelimenin bağlamını ve dil bilgisel yapısını koruyarak anlamlı bir kök formuna indirgenmesini sağlamaktadır. Bu yöntemin kullanılması, modelin benzer kelimeleri farklı kelimeler olarak değerlendirmesini engelleyerek doğruluk oranının artırılmasına yardımcı olmuştur. Son aşamada, haber metinleri makine öğrenmesi algoritmalarının işleyebileceği sayısal forma dönüştürülmüştür. Metin verileri doğrudan sayısal verilerle çalışmadığı için haber içeriklerinin sayısal temsili gereklidir. Bu çalışmada, TF-IDF yöntemi kullanılarak haber metinleri vektör hâline getirilmiştir. TF-IDF yöntemi, kelimelerin metin içindeki önemini belirleyerek sık kullanılan ancak düşük bilgilendiriciliğe sahip kelimelerin etkisini azaltmakta ve daha anlamlı kelimelere daha yüksek ağırlık vermektedir. TF-IDF yöntemi iki temel bileşenden oluşmaktadır:

- Terim Frekansı (TF): Bir kelimenin belirli bir metin içinde kaç kez geçtiğini ifade eder ve formülü;

$$TF(t, d) = \frac{f_t}{n} \quad (1)$$

şeklinde ele alınmaktadır. f_t Kelimenin metin içinde tekrar edilme sayısını, n ise metindeki toplam kelime sayısını göstermektedir.

- Ters Belge Frekansı (IDF): Bir kelimenin tüm veri kümesi içindeki önemini belirler ve formülü;

$$IDF(t) = \log \left(\frac{N}{df_t} + 1 \right) \quad (2)$$

şeklinde ele alınmaktadır. N Toplam belge sayısını, df_t ise ilgili kelimenin geçtiği belge sayısını temsil etmektedir.

- TF-IDF Skoru: Bir kelimenin hem belirli bir belge içindeki frekansını hem de genel veri kümesindeki önemini hesaba katar ve formülü;

$$TFIDF(t, d) = TF(t, d) \times IDF(t) \quad (3)$$

şeklinde ele alınmaktadır.

Bu yöntem sayesinde, sık kullanılan ancak düşük bilgilendirici kelimeler düşük katkı sağlarken anlam açısından önemli kelimeler modele daha fazla katkı sağlamaktadır. Bu sayede model, yalnızca haber içeriklerinde geçen kelimeleri değil, bu kelimelerin önem derecesini de öğrenerek daha başarılı tahminler yapabilmektedir. Sonuç olarak, veri ön işleme süreci hem haber metinlerini modelin öğrenmesi için uygun hâle getirmiş ve modelin başarısını doğrudan etkilemiştir. Yapılan temizleme, önemsiz kelimelerin çıkarılması, kelime köklerine ayırma ve TF-IDF vektörizasyonu gibi işlemler, modelin hem doğruluğunu artırmış hem de işlem süresini optimize etmiştir. Gelecek çalışmalarda, Word2Vec, BERT ve GloVe gibi derin öğrenme tabanlı vektörleştirme teknikleri kullanılarak modelin bağlamsal analiz yeteneğinin geliştirilmesi önerilmektedir.

Makine Öğrenmesi Modeli

Makine öğrenmesi teknikleri, sahte haberlerin tespit edilmesi ve yeni medyada dezenformasyonun önlenmesi açısından önemli bir yere sahiptir. Bu çalışmada, sahte ve gerçek haberleri ayırabilen bir makine öğrenmesi modeli geliştirilmiş ve modelin etkinliği değerlendirilmiştir. Modelin oluşturulma süreci; veri setinin hazırlanması, metin ön işleme aşamaları, model eğitimi ve değerlendirme metriklerinin hesaplanması gibi temel bileşenleri içermektedir.

Veri Setinin Eğitim ve Test Olarak Ayrılması

Makine öğrenmesi modellerinin genelleme başarısını değerlendirebilmek için veri seti, eğitim ve test olarak ikiye ayrılmıştır. Eğitim verisi modelin öğrenme sürecinde kullanılırken test verisi modelin daha önce görmediği haberleri doğru sınıflandırıp sınıflandıramadığını ölçmek amacıyla kullanılmıştır. Veri seti, %80 eğitim ve %20 test verisi olacak şekilde bölünmüştür. Veri setinin eğitim ve test olarak ayrılması işlemi, Scikit-learn kütüphanesinin `train_test_split` fonksiyonu kullanılarak gerçekleştirilmiştir. Bu yöntem, verilerin rastgele bölünmesini sağlayarak modelin aşırı öğrenme (overfitting) yapmasını engellemekte ve genelleme performansını artırmaktadır.

Lojistik Regresyon

Bu çalışmada sahte haberlerin tespiti için lojistik regresyon modeli (Torres, 2008) kullanılmıştır. Lojistik regresyon, ikili sınıflandırma (binary classification) problemlerinde yaygın olarak kullanılan bir yöntem olup haberlerin sahte veya gerçek olma olasılığını hesaplamak için sigmoid aktivasyon fonksiyonunu kullanmaktadır.

$$z = w_1x_1 + w_2x_2 + \dots + w_nx_n + b \quad (4)$$

Burada, x_i haber metnindeki kelime özelliklerini, x_i modelin öğrendiği ağırlık katsayılarını ve b modelin sabit değerini temsil etmektedir. Model, sigmoid fonksiyonu ile çıktıyı olasılık değerine dönüştürmektedir;

$$\sigma(z) = \frac{1}{1 + e^{-z}} \quad (5)$$

Sigmoid fonksiyonu, modelin tahmin ettiği değeri 0 ile 1 arasında bir olasılığa çevirerek bir haberin sahte olma ihtimalini hesaplamaktadır. Elde edilen olasılık değeri belirli bir eşik değeri ile kıyaslanarak haberin sınıflandırılması gerçekleştirilmiştir. Genellikle 0.5 eşik değeri kullanılmış olup tahmin edilen olasılık 0.5'ten büyükse, haber sahte olarak, 0.5'ten küçükse, haber gerçek olarak etiketlenmiştir.

Lojistik regresyon modelinin bu çalışmada tercih edilme nedenleri; (1) büyük ölçekli metin verileri üzerinde hızlı eğitilebilir ve hesaplama açısından verimli olması, (2) yorumlanabilirlik avantajı sunmasıdır.

Modelin katsayıları sayesinde, hangi kelimelerin veya kelime gruplarının sahte haber tespitinde etkili olduğunu analiz etmek mümkündür. Ancak, bu klasik modelin bazı sınırlamaları bulunmaktadır. Lojistik regresyon, değişkenler arasındaki ilişkinin doğrusal olduğunu varsayar ve bu nedenle doğal dildeki karmaşık, doğrusal olmayan örüntüleri, ilişkileri yakalamakta derin öğrenme modellerine kıyasla kısıtlı kalır. Bu çalışmada, elde edilen yüksek başarımlı metrikleri dikkate alınarak nispeten daha basit ve yorumlanabilir bir temel modelin dahi dezenformasyonun dilsel kalıplarını tespit etmede ne kadar etkili olabileceğini göstermek amacıyla lojistik regresyon tercih edilmiştir. Bu seçim, karmaşık teknolojiye ihtiyaç duymadan da yüksek başarı elde edilebileceği argümanını güçlendirmektedir.

Modelin Eğitilmesi ve Test Edilmesi

Modelin eğitimi, eğitim verileri kullanılarak gerçekleştirilmiş ve lojistik regresyonun parametreleri optimize edilmiştir. Model, TF-IDF ile vektörleştirilmiş haber metinleri üzerinde eğitilmiş ve test edilmiştir. Eğitim sürecinde model, kelime frekanslarına dayalı olarak sahte haberlerin dil yapısını öğrenmiş ve haberlerin içeriğine bağlı olarak bir olasılık tahmininde bulunabilecek şekilde optimize edilmiştir. Modelin eğitim aşaması tamamlandıktan sonra test verisi ile tahminler yapılmış ve modelin başarı oranı çeşitli değerlendirme metrikleri kullanılarak hesaplanmıştır.

Sonuçların Görselleştirilmesi

Makine öğrenmesi modelinin performansını analiz etmek için çeşitli görselleştirme teknikleri kullanılmıştır. Modelin doğruluk oranını ve hata dağılımını incelemek amacıyla karmaşıklık matrisi, ROC eğrisi ve Precision-Recall eğrisi gibi grafikler oluşturulmuştur. Karmaşıklık matrisi, modelin tahmin ettiği doğru ve yanlış sınıflandırmaları görselleştirmek için kullanılmıştır. Bu matris, gerçek ve tahmin edilen etiketler arasındaki ilişkiyi göstermektedir. ROC Eğrisi ve AUC Skoru (Area Under the Curve Score), modelin farklı eşik değerleri altında nasıl performans gösterdiğini incelemek için kullanılmıştır. ROC eğrisi, gerçek pozitif oranı (true positive rate) ile yanlış pozitif oranını (false positive rate) karşılaştırarak modelin doğruluk performansını ölçmektedir. Precision-Recall Eğrisi, modelin sahte haberleri tespit etme performansını daha detaylı bir şekilde incelemek için oluşturulmuştur. Precision, modelin tahmin ettiği sahte haberlerin ne kadarının gerçekten sahte olduğunu ölçerken duyarlılık ise tüm sahte haberler içinde modelin ne kadarını doğru tespit ettiğini belirlemektedir. Görselleştirme sürecinde Matplotlib ve Seaborn kütüphaneleri kullanılarak elde edilen grafikler, modelin güçlü ve zayıf yönlerini anlamaya yardımcı olmuştur.

Değerlendirme Metrikleri

Makine öğrenmesi modellerinin başarısını ölçmek için farklı değerlendirme metrikleri kullanılmaktadır. Özellikle sınıflandırma problemlerinde, modelin doğru tahmin oranlarını, hata oranlarını ve genel performansını değerlendirebilmek adına çeşitli istatistiksel ölçütler hesaplanmaktadır. Bu çalışmada, sahte haberlerin tespit edilmesi için kullanılan modelin performansını belirlemek amacıyla doğruluk, kesinlik, duyarlılık, F1-Skoru ve ROC-AUC metrikleri incelenmiştir. Değerlendirme metriklerinin hesaplanabilmesi için öncelikle modelin tahminlerinin doğru ve yanlış sınıflandırmaları belirlenmiş ve karmaşıklık matrisi üzerinden analiz edilmiştir. Karmaşıklık matrisi, modelin gerçek sınıflarla tahmin edilen sınıfları karşılaştırarak sınıflandırma hatalarını ayrıntılı bir şekilde görmemizi sağlar. Model tahminleri şu dört temel duruma göre değerlendirilmektedir:

- TP (True Positive – Doğru Pozitif): Modelin gerçekten sahte olan bir haberi sahte olarak doğru tahmin etmesi durumudur. Yani, modelin hedeflediği sahte haberleri başarılı bir şekilde tespit ettiğini gösterir.
- TN (True Negative – Doğru Negatif): Modelin gerçekten gerçek olan bir haberi gerçek olarak doğru tahmin etmesi durumudur. Bu, modelin gerçek haberleri yanlış bir şekilde sahte olarak sınıflandırmadığını ve güvenilir haberleri doğru tanıyabildiğini gösterir.
- FP (False Positive – Yanlış Pozitif): Modelin gerçek olan bir haberi yanlışlıkla sahte olarak sınıflandırması durumudur. Bu durumda, model hatalı bir şekilde gerçek bir haberi sahte olarak işaretlemiştir. Yanlış pozitiflerin fazla olması, güvenilir haberlerin yanlış yere sahte olarak etiketlenmesine yol açabilir ve bu durum haber güvenilirliğini olumsuz etkileyebilir.
- FN (False Negative – Yanlış Negatif): Modelin sahte olan bir haberi yanlışlıkla gerçek olarak sınıflandırması durumudur. Bu hata türü, sahte haberlerin tespit edilememesine ve yanlış bilgilendirme

riskinin artmasına sebep olabilir. Yanlış negatif oranının düşük olması, modelin sahte haberleri başarıyla tespit ettiğini gösteren önemli bir kriterdir.

Bu temel sınıflandırma durumları kullanılarak modelin başarısını ölçen metrikler aşağıdaki gibi hesaplanmıştır:

Doğruluk

Doğruluk metriği, modelin tüm sınıfları doğru tahmin etme oranını gösteren temel bir ölçüttür. Modelin tüm veri seti üzerinde ne kadar doğru tahminde bulunduğunu ölçer. Ancak, doğruluk metriği tek başına yeterli bir değerlendirme sağlamayabilir; özellikle dengesiz veri setlerinde yüksek doğruluk oranı yanıltıcı olabilir. Örneğin, eğer model tüm haberleri "gerçek" olarak sınıflandırırsa sahte haberleri hiç tespit etmese bile doğruluk oranı yüksek çıkabilir. Doğruluk formülü;

$$\frac{TP+TN}{TP+FP+TN+FN} \quad (6)$$

şeklinde ele alınmaktadır. Bu formül, modelin doğru tahmin ettiği sahte ve gerçek haberlerin toplamını, tüm tahmin edilen haberlerin toplamına bölerek doğruluk oranını belirler.

Kesinlik

Kesinlik, modelin sahte haber olarak sınıflandırdığı haberlerin gerçekten sahte olup olmadığını ölçen bir metriktir. Yanlış pozitiflerin yüksek olduğu durumlarda, modelin sahte haberleri yanlışlıkla fazla işaretleyip işaretlemediğini anlamak için önemlidir. Kesinlik formülü;

$$\frac{TP}{TP+FP} \quad (7)$$

şeklinde ele alınmaktadır. Kesinlik değeri yüksekse modelin sahte haber olarak tahmin ettiği haberlerin büyük çoğunluğu gerçekten sahte demektir. Yanlış pozitiflerin fazla olması hâlinde, kesinlik skoru düşecek ve modelin sahte haberleri güvenilir bir şekilde ayırt edemediğini gösterecektir.

Duyarlılık

Duyarlılık, modelin gerçek sahte haberleri ne kadar iyi tespit edebildiğini gösteren bir metriktir. Yanlış negatiflerin fazla olduğu durumlarda, modelin sahte haberleri gözden kaçırıp kaçırmadığını anlamak için kullanılır. Eğer bir modelin duyarlılık değeri düşükse sahte haberlerin büyük bir kısmı tespit edilemiyor demektir. Duyarlılık formülü;

$$\frac{TP}{TP+FN} \quad (8)$$

şeklinde ele alınmaktadır. Duyarlılığın yüksek olması, modelin sahte haberleri kaçırmadan tespit ettiğini gösterir. Ancak, duyarlılığı artırmaya çalışırken yanlış pozitiflerin artması riski de göz önünde bulundurulmalıdır.

F1-Skoru

F1-Skoru, kesinlik ve duyarlılığı dengeleyen bir ölçümdür. Tek başına kesinlik ya da duyarlılık kullanıldığında, modelin performansı hakkında yanıltıcı sonuçlar elde edilebilir. Örneğin, bir modelin duyarlılığı yüksek fakat kesinliği düşükse model, sahte haberleri tespit etmekte iyi olabilir ancak çok fazla yanlış pozitif üretiyor olabilir. F1-Skoru, bu iki metriği dengeleyerek genel performansı anlamak için daha güvenilir bir ölçü sağlar. F1-Skoru formülü;

$$2. \frac{\text{Kesinlik} \cdot \text{Duyarlılık}}{\text{Kesinlik} + \text{Duyarlılık}} \quad (9)$$

şeklinde ele alınmaktadır. F1-Skorunun yüksek olması, modelin hem yanlış pozitifleri hem de yanlış negatifleri düşük tutarak iyi bir denge sağladığını gösterir.

ROC-AUC

ROC eğrisi, modelin farklı eşik değerlerine göre nasıl performans gösterdiğini grafiksel olarak sunan bir ölçümdür. ROC eğrisinin altındaki alan, modelin genel başarısını gösteren önemli bir metriktir. AUC skoru 1'e ne kadar yakınsa modelin doğruluk oranı o kadar yüksektir. ROC-AUC formülü;

$$\int_0^1 TPR(FPR) dFPR \quad (10)$$

şeklinde ele alınmaktadır. Bu metrik, özellikle modelin farklı karar eşiklerinde nasıl performans gösterdiğini anlamak için kullanılır.

Karmaşıklık Matrisi

Modelin tahmin performansını daha ayrıntılı bir şekilde değerlendirmek için karmaşıklık matrisi kullanılmıştır. Bu matris, modelin yanlış ve doğru sınıflandırmalarını detaylı bir şekilde analiz etmeye yardımcı olmaktadır. Bu matrisin analizi, modelin hangi sınıfları yanlış ve doğru tahmin ettiğini net bir şekilde ortaya koyarak modelin güçlü ve zayıf yönlerinin belirlenmesine olanak tanımaktadır. Bu bölümde, modelin performansını ölçmek için kullanılan değerlendirme metrikleri açıklanmış ve her bir metriğin nasıl hesaplandığı belirtilmiştir.

Çalışmanın Literatüre Katkısı ve Gelecek Çalışmalar

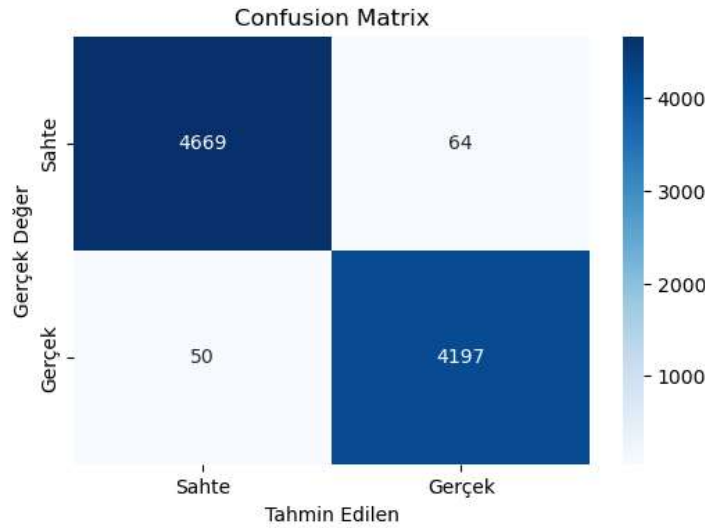
Sahte haber tespiti konusunda makine öğrenmesi tekniklerinin etkinliğini değerlendirerek literatüre önemli katkılar sunmaktadır. Günümüzde dijital medya platformlarında sahte haberlerin paylaşılması, yanlış bilgilendirme ve dezenformasyonun yayılmasında büyük bir sorun teşkil etmektedir. Bu bağlamda, geliştirilen model, haber metinleri üzerinden sahte içerikleri tespit edebilmek için lojistik regresyon tabanlı bir makine öğrenmesi yaklaşımı kullanarak sahte haber tespiti alanında yapılan çalışmalara bir katkı sağlamaktadır. Bu çalışma, TF-IDF vektörizasyonu ile metinleri sayısal hâle getirerek lojistik regresyon modeli kullanımı üzerine odaklanmasıyla farklılık göstermektedir. Çalışmada kullanılan dengeli veri kümesi, modelin, farklı haber türleri arasındaki ayrımı daha sağlıklı yapabilmesini sağlamıştır. Ayrıca, değerlendirme metrikleri detaylı bir şekilde ele alınarak modelin başarısı doğruluk, kesinlik, duyarlılık, F1-Skoru ve ROC-AUC metrikleri ile kapsamlı bir şekilde analiz edilmiştir. Bu sayede, sahte haberlerin belirlenmesinde hangi metriklerin daha güvenilir sonuçlar verdiği tartışmaya açılmıştır. Gelecek çalışmalar açısından, sahte haber tespiti için daha gelişmiş derin öğrenme modellerinin (örneğin BERT, LSTM, Transformer tabanlı modeller) test edilmesi önerilmektedir. Bu tür modeller, kelimelerin bağlam içindeki anlamını daha iyi yakalayarak daha yüksek doğruluk oranları sağlayabilir. Ayrıca, haberlerin kaynak güvenilirliği, dil bilgisi analizi ve yazar profilleri gibi ek faktörlerin modele dâhil edilmesi, sahte haber tespit sistemlerinin daha doğru çalışmasına yardımcı olabilir. Veri setinin genişletilmesi ve farklı dillerdeki haberlerin analiz edilmesi, modelin genelleme kapasitesini artırarak küresel ölçekte daha kapsamlı bir sistem geliştirilmesini sağlayabilir. Bunun yanı sıra hibrit modellerin kullanımı da gelecek araştırmalar için önemli bir potansiyel taşımaktadır. Örneğin, makine öğrenmesi ve derin öğrenme modellerinin bir arada kullanılması, farklı model türlerinin güçlü yönlerinden yararlanılarak daha etkili sahte haber tespit sistemleri geliştirilmesine olanak sağlayabilir. Ayrıca, gerçek zamanlı sahte haber tespiti yapabilen sistemlerin geliştirilmesi ve büyük ölçekli sosyal medya verilerinin analiz edilmesi, dezenformasyonla mücadelede kritik bir adım olacaktır. Sonuç olarak bu çalışma, sahte haberlerin tespit edilmesine yönelik makine öğrenmesi tabanlı yöntemlerin etkili bir şekilde uygulanabileceğini göstermektedir. Ancak modelin, daha fazla veri ile test edilmesi, farklı doğal dil işleme teknikleriyle güçlendirilmesi ve hibrit yöntemlerle desteklenmesi, gelecekte bu alandaki çalışmaların daha ileri seviyelere taşınmasını sağlayacaktır.

Bulgular ve Tartışma

Bu çalışmada, sahte haberlerin tespitine yönelik geliştirilen lojistik regresyon modeli, çeşitli değerlendirme metrikleri kullanılarak analiz edilmiştir. Modelin performansını ölçmek amacıyla doğruluk, kesinlik, duyarlılık ve F1-skora dayalı analizler gerçekleştirilmiştir. Bunun yanı sıra, modelin yanlış sınıflandırma durumları karmaşıklık matrisi kullanılarak görselleştirilmiş, ROC eğrisi ile modelin sınıflandırma gücü incelenmiş ve Precision-Recall eğrisi ile pozitif sınıflamaların doğruluğu değerlendirilmiştir. Modelin genel performansı, doğruluk oranı açısından yüksek bir başarı sergilemiştir. Doğruluk metriği, modelin tüm örnekler üzerinde doğru tahmin yapma oranını göstermektedir ve bu çalışma kapsamında geliştirilen modelin doğruluk değeri oldukça yüksek çıkmıştır. Bununla birlikte, modelin kesinlik ve duyarlılık değerleri de dengeli bir sonuç ortaya koymuştur. Özellikle duyarlılık metriği, modelin sahte haberleri tespit etme başarısını ölçmek açısından kritik bir gösterge olup elde edilen yüksek duyarlılık oranı modelin sahte haberleri büyük ölçüde doğru bir şekilde tespit ettiğini göstermektedir.

Karmaşıklık Matrisi

Aşağıdaki Şekil 1'de, modelin tahmin ettiği ve gerçek sınıflar arasındaki ilişkiyi gösteren karmaşıklık matrisi yer almaktadır. Bu matris, modelin doğru ve yanlış tahminlerini ayrıntılı bir şekilde özetlemektedir.

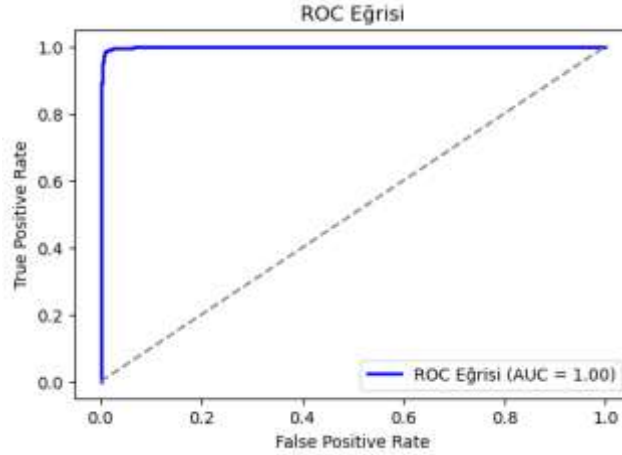


Şekil 1. Karmaşıklık Matrisi

Matris incelendiğinde, modelin 4.669 sahte haberi doğru bir şekilde sahte olarak tahmin ettiği (TP) ve 4.197 gerçek haberi doğru bir şekilde gerçek olarak tahmin ettiği (TN) görülmektedir. Yanlış sınıflandırmalar ise 64 adet yanlış pozitif (FP) ve 50 adet yanlış negatif (FN) olarak gözlemlenmiştir. Yanlış pozitif oranı düşük olsa da yanlış negatiflerin tamamen sıfıra indirilememesi, modelin bazı sahte haberleri gerçek olarak değerlendirdiğini göstermektedir. Ancak genel hatlarıyla bakıldığında, modelin doğru sınıflandırma oranının oldukça yüksek olduğu anlaşılmaktadır. Sınıflandırma hatalarının detaylı analizi; çalışmada gözlemlenen yanlış sınıflandırmalar, modelin hata yapma eğilimleri ve sınırları hakkında önemli bilgiler sunmaktadır. Özellikle dezenformasyonla mücadele bağlamında en kritik hata türü olan 50 adet yanlış negatif (FN) hatanın (gerçekte sahte olan haberin, model tarafından yanlışlıkla gerçek olarak etiketlenmesi) nedenleri derinlemesine incelenmiştir. Bu tür sahte haberlerin, gerçek haberlere benzemek amacıyla genellikle daha resmî ve karmaşık bir dil kullandığı veya duygu yükü yerine nicel veri manipülasyonuna odaklandığı saptanmıştır. Öte yandan, 64 adet yanlış pozitif (FP) hatanın (gerçekte doğru olan haberin, model tarafından yanlışlıkla sahte olarak etiketlenmesi) ise genellikle haber metnindeki yoğun duygu yüklü ifadelerden, polemik içeren dil kullanımından veya aşırı iddialı başlıklardan kaynaklandığı anlaşılmaktadır. Bu durum, kullanılan TF-IDF vektörizasyon tekniğinin yalnızca kelime sıklığına odaklanması ve kelimelerin bağlamsal ilişkisini dikkate almaması nedeniyle, modelin duygusal metinleri sahte haber kalıbıyla karıştırma eğilimini ortaya koymaktadır. Bu detaylı hata analizi, lojistik regresyon modelinin metin içeriğini anlama kapasitesinin sınırlarını netleştirmekte ve gelecekteki bağlamsal analize dayalı çalışmalar için yol göstermektedir.

ROC Eğrisi ve AUC Skoru

Modelin genel performansını değerlendirmek için kullanılan bir diğer ölçüt ROC eğrisi ve AUC skoru olmuştur. ROC eğrisi, modelin farklı sınıflandırma eşiklerinde ne kadar başarılı olduğunu gösteren önemli bir grafikdir. Aşağıdaki Şekil 2'de, ROC eğrisi verilmiştir.

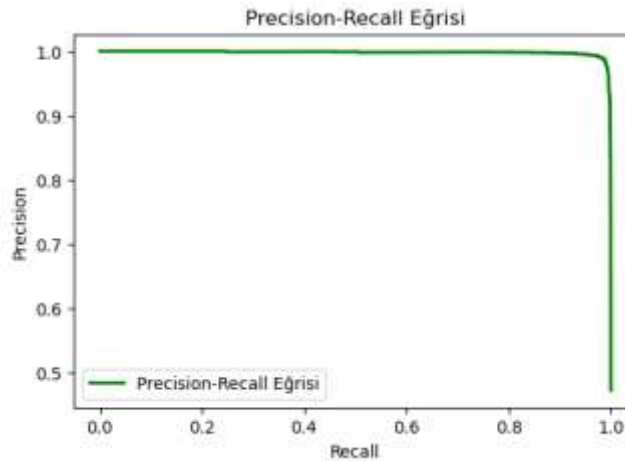


Şekil 1. ROC Eğrisi ve AUC Skoru

ROC eğrisi incelendiğinde, modelin AUC skorunun 1.00'a yakın olduğu görülmektedir. AUC skorunun 1'e yaklaşması, modelin sınıflandırma başarısının oldukça yüksek olduğunu göstermektedir. Bu durum, modelin hem sahte hem de gerçek haberleri yüksek doğrulukla ayırt edebildiğini kanıtlamaktadır.

Precision-Recall Eğrisi

Sahte haber tespitinde, pozitif sınıflamaların güvenilirliğini ölçmek için Precision-Recall eğrisi de analiz edilmiştir. Precision-Recall eğrisi, özellikle dengesiz veri setleriyle çalışırken modelin performansını daha iyi değerlendirmek için kullanılan bir yöntemdir. Aşağıdaki Şekil 3'te, modelin Precision-Recall eğrisi gösterilmektedir.



Şekil 1. Precision-Recall Eğrisi

Eğri incelendiğinde, modelin yüksek bir kesinlik ve duyarlılık oranına sahip olduğu ve ürettiği pozitif tahminlerin büyük bir kısmının doğru olduğu görülmektedir. Bu durum, modelin sahte haberleri tespit etme konusunda tutarlı olduğunu ve yanlış pozitif oranının düşük seviyede kaldığını göstermektedir.

Model Performansının Genel Değerlendirmesi

Modelin doğruluk, kesinlik, duyarlılık ve F1-Skoru değerlendirme sonuçları Tablo 1'de özetlenmiştir;

Tablo 1

Model Performansının Genel Değerlendirmesi

ÖLÇEKLER	Madde Sayısı
Doğruluk	%98,5
Kesinlik	%98,7
Duyarlılık	%98,9
F1-Skoru	%98,8

Bu tabloda (Tablo 1) görüldüğü gibi, modelin performans metrikleri oldukça başarılı sonuçlar vermiştir. Doğruluk oranının yüksek olması, modelin büyük ölçüde doğru tahminler yaptığını göstermektedir. Kesinlik ve duyarlılık değerleri dengeli bir şekilde yüksek olup modelin yanlış pozitif oranının düşük, yanlış negatif oranının ise oldukça az olduğunu ortaya koymaktadır. Bu sonuçlar, geliştirilen lojistik regresyon modelinin sahte haberleri tespit etmek için güvenilir bir yöntem sunduğunu ve modelin genel başarısının yüksek olduğunu göstermektedir. Ancak, modelin yanlış negatif oranını daha da azaltmak ve sahte haberleri daha hassas bir şekilde tespit etmek adına, gelecekte farklı metin vektörleştirme teknikleri veya derin öğrenme modelleri kullanılabilir.

Sonuç

Günümüzde yeni medya, bilgiye hızlı erişim sağlayan, ancak aynı zamanda dezenformasyonun yayılmasını kolaylaştıran bir yapı sunmaktadır. Tuğtekin ve Koç (2019) tarafından belirtildiği gibi, geleneksel medyadan farklı olarak yeni medya, kullanıcıların aktif katılımıyla içeriğin anlık olarak üretildiği ve yayıldığı bir ortamdır. Bu durum, sahte haberlerin kontrolsüz bir şekilde yayılmasına yol açmakta ve toplum üzerinde ciddi etkiler yaratmaktadır. Çalışmamız, sahte haberlerin tespit edilmesine yönelik bir makine öğrenmesi modeli geliştirerek bu soruna çözüm üretmeyi hedeflemektedir. Çalışmamızda kullanılan lojistik regresyon modeli, sahte haberlerin tespitinde başarılı sonuçlar elde etmiştir. McLuhan (1964) ve Peters (2009), yeni medya kavramının bilgi akışındaki değişimlerle birlikte evrildiğini ve dezenformasyonun yayılma hızının arttığını vurgulamaktadır. Bu bağlamda, önerdiğimiz model, yeni medya ortamında yayılan yanlış bilgilerin otomatik olarak tespit edilmesine yönelik bir yaklaşım sunmaktadır. Elde edilen bulgular, sahte haberlerin dil yapılarında belirgin kalıpların bulunduğunu ve bu yapıların makine öğrenmesi algoritmalarıyla tespit edilebileceğini göstermektedir.

Ayrıca Steinfeld (2022), sosyal medya platformlarının algoritmalarının, kullanıcıların ilgi alanlarına uygun içerikleri öne çıkarması nedeniyle yanlış bilgilerin daha fazla yayılmasına neden olabileceğini belirtmiştir. Çalışmamızda kullanılan metin madenciliği teknikleri ve TF-IDF vektörizasyonu, sahte haberlerin yapısal özelliklerini analiz ederek yanlış bilgilendirme içeren haberlerin belirlenmesine katkıda bulunmaktadır. Bu durum, sahte haberlerin yayılımını önlemek amacıyla platformlara entegre edilebilecek otomatik tespit sistemlerinin geliştirilmesi açısından önemlidir. Medya okuryazarlığının dezenformasyonla mücadelede kritik bir faktör olduğu bilinmektedir. Zhang (2024) ve Soares et al. (2021) çalışmalarında, bireylerin dijital medyada maruz kaldıkları bilgileri eleştirel bir bakış açısıyla değerlendirmeleri gerektiğini vurgulamışlardır. Bu bağlamda, çalışmamızın ortaya koyduğu model, medya okuryazarlığı eksikliğinden kaynaklanan yanlış bilgiye maruz kalma riskini azaltmaya yardımcı olabilir. Sahte haberlerin belirli dil kalıplarıyla tespit edilebileceğini gösteren sonuçlarımız, medya okuryazarlığına dayalı eğitim programları için de faydalı olabilecek veriler sunmaktadır. Çalışmamız, yeni medya ortamında sahte haberlerin tespit edilmesi için makine öğrenmesi tekniklerinin uygulanabilirliğini göstermiştir. Ancak, dezenformasyon yalnızca belirli kelime kalıpları veya haber metinlerinin yapısıyla sınırlı değildir. Tomin et al. (2020) tarafından ifade edildiği gibi dijital medyanın karmaşık doğası, sahte haberlerin farklı biçimlerde ortaya çıkmasına neden olmaktadır. Bu nedenle, gelecekteki çalışmalarda, görsel ve video tabanlı sahte haberlerin tespitine yönelik daha gelişmiş modellerin geliştirilmesi önerilmektedir. Sonuç olarak bu çalışma, sahte haberlerin dilsel özelliklerine odaklanarak makine öğrenmesi modellerinin dezenformasyonla mücadelede nasıl kullanılabileceğine dair önemli bir katkı sunmaktadır. Gelecekte, daha büyük ve çeşitli veri setleriyle yapılan çalışmalar, modelin genelleme yeteneğini artırarak farklı medya türlerinde de etkili bir şekilde

uygulanmasını sağlayabilir. Ek olarak, sahte haberlerin yalnızca metin içeriklerinden değil, sosyal medya etkileşimlerinden ve kullanıcı davranışlarından da tespit edilmesine yönelik hibrit yaklaşımlar, dezenformasyonun yayılımını önlemede daha etkili olabilir.

Kaynakça

- Aydin, I., Mehmet, S., ve Salur, M. U. (2018), Detection of Fake Twitter Accounts with Machine Learning Algorithms. 2018 International Conference on Artificial Intelligence and Data Processing (IDAP).
- Blewitt, J. (2014). Public Libraries and the Right to the [Smart] City. *International Journal of Social Ecology and Sustainable Development*, 5(2), 55. <https://doi.org/10.4018/ijsesd.2014040105>
- Carlsson, E., & Nilsson, B. (2015). Technologies of participation: Community news and social media in Northern Sweden. *Journalism* (Vol. 17, Issue 8, p. 1113). SAGE Publishing. <https://doi.org/10.1177/1464884915599948>
- Chao-Chen, L. (2013). Convergence of new and old media: new media representation in traditional news. *Chinese Journal of Communication* (Vol. 6, Issue 2, p. 183). Taylor & Francis. <https://doi.org/10.1080/17544750.2013.785667>
- Di, J. (2024). Principles of AIGC technology and its application in new media micro-video creation. *Applied Mathematics and Nonlinear Sciences* (Vol. 9, Issue 1). De Gruyter. <https://doi.org/10.2478/amns-2024-1393>
- Ekawati, A. D. (2022). Students' Beliefs About Social Media In Efl Classroom: A Review Of Literature. *English Review Journal of English Education* (Vol. 10, Issue 2, p. 587). University of Kuningan and Association of Indonesian Scholars of English Education (AISEE). <https://doi.org/10.25134/erjee.v10i2.6261>
- Ersoy, A. B. (2021). Social Media Shopping as a Driver for Brand Trust and Brand Commitment. *Journal of Business Administration Research* (Vol. 4, Issue 4). <https://doi.org/10.30564/jbar.v4i4.3454>
- Flanagin, A., & Metzger, M. (2006). Internet use in the contemporary media environment. *Human Communication Research* (Vol. 27, Issue 1, p. 153). Wiley. <https://doi.org/10.1111/j.1468-2958.2001.tb00779.x>
- Fraga-Lamas, P., & Fernández-Caramés, T. M. (2020). Fake News, Disinformation, and Deepfakes: Leveraging Distributed Ledger Technologies and Blockchain to Combat Digital Deception and Counterfeit Reality. *IT Professional* (Vol. 22, Issue 2, p. 53). IEEE Computer Society. <https://doi.org/10.1109/mitp.2020.2977589>
- Francesca, M., Demartini, P., & Paoloni, P. (2017). Women in business and social media: Implications for female entrepreneurship in emerging countries. *African Journal of Business Management* (Vol. 11, Issue 14, p. 316). Academic Journals. <https://doi.org/10.5897/ajbm2017.8281>
- Gandini, A., Gerosa, A., Gobbo, B., Keeling, S., Leonini, L., Mosca, L., Orofino, M., Reviglio, U., & Splendore, S. (2022). *The algorithmic public opinion: a literature review* [Review of *The algorithmic public opinion: a literature review*]. <https://doi.org/10.31235/osf.io/m6zn8>
- Guo, H. (2021). The development and application of new media technology in news communication industry. *International Journal of Electrical Engineering Education* (Vol. 60, p. 405). SAGE Publishing. <https://doi.org/10.1177/0020720921996593>
- Guo, J. (2022). Psychological Harm to Underage Children in the Era of New Media. *Advances in Social Science, Education and Humanities Research/Advances in social science, education and humanities research* (p. 2021). https://doi.org/10.2991/978-2-494069-89-3_231
- Gürdağ, B., Akkavak, K. K., & Elibol, G. C. (2018). Tasarım Eğitiminde Alternatif Mekân Olarak Sosyal Medya. *Doaj (Doaj: Directory of Open Access Journals)*. <https://doaj.org/article/5a97be8cd1f3471692cb590ba5f4e08e>
- J. Adeleke Adeyiga, P. Gbounmi Toriola, T. Elizabeth Abioye, and A. Esther Oluwatosin, "Fake News Detection Using a Logistic Regression Model and Natural Language Processing Techniques," pp. 1–18, 2023, [Online]. Available: <https://doi.org/10.21203/rs.3.rs-3156168/v1>
- Karmila, L., Rachmiate, A., K, S. S., Fardiah, D., Ahmadi, D., & Muhtadi, A. S. (2024). The role of social media in the political construction of identity: Implications for political dynamics and democracy in Indonesia. *Journal of Infrastructure Policy and Development* (Vol. 8, Issue 14, p. 9171). <https://doi.org/10.24294/jipd9171>
- Łaszkiiewicz, A. (2022). Areas of influence in influencer marketing. To what extent is the communication under brand control? *Marketing of Scientific and Research Organizations* (Vol. 46, Issue 4, p. 1). Moscow Aviation Institute. <https://doi.org/10.2478/minib-2022-0018>

- Lee, S., & Valenzuela, S. (2024). A Self-Righteous, Not a Virtuous, Circle: Proposing a New Framework for Studying Media Effects on Knowledge and Political Participation in a Social Media Environment. *Social Media + Society*, 10(2). <https://doi.org/10.1177/20563051241257953>
- Levak, T. (2021). Disinformation in the New Media System – Characteristics, Forms, Reasons for its Dissemination and Potential Means of Tackling the Issue. *Medijska istraživanja* (Vol. 26, Issue 2, p. 29). Naklada Medijska istraživanja. <https://doi.org/10.22572/mi.26.2.2>
- McLuhan, M. (1964). “The medium is the message.” In *Understanding Media: The Extensions of Man*, pp. 23–35, 63–7. New York: Signet
- Mehta, S. (2019). Indian intermediaries in new media: Mainstreaming the margins. *Convergence The International Journal of Research into New Media Technologies*, 27(1), 229. <https://doi.org/10.1177/1354856519896167>
- Méndez, M. C., Codina, L., & Salaverría, R. (2019). What is new media? The views of 70 Hispanic experts. <https://doi.org/10.4185/rlds-2019-1396n>
- Mutlu, E. (2016). *Globalleşme, Popüler Kültür ve Medya*. Ütopya Yayınevi.
- Olesen, T. (2017). Memetic protest and the dramatic diffusion of Alan Kurdi. *Media Culture & Society*, 40(5), 656. <https://doi.org/10.1177/0163443717729212>
- Olesen, T. (2020). Greta Thunberg’s iconicity: Performance and co-performance in the social media ecology. *New Media & Society*, 24(6), 1325. <https://doi.org/10.1177/1461444820975416>
- Öztunç, M. (2020). İnternet Haberciliğinde Koronavirüs Salgınının Sunum Biçimi. *DergiPark (Istanbul University)*. Istanbul University. <https://dergipark.org.tr/tr/pub/itobiad/issue/56503/756934>
- Peters, B. (2009). And lead us not into thinking the new is new: a bibliographic case for new media history. *New Media & Society* (Vol. 11, p. 13). SAGE Publishing. <https://doi.org/10.1177/1461444808099572>
- Peters, C., Schröder, K. C., Lehaff, J., & Vulpius, J. (2021). News as They Know It: Young Adults’ Information Repertoires in the Digital Media Landscape. *Digital Journalism* (Vol. 10, Issue 1, p. 62). Taylor & Francis. <https://doi.org/10.1080/21670811.2021.1885986>
- Sevim, Z. (2019). Toplumsal Bir Hafıza Yaratma Gücü: Medya / “Yeni” Medya. *Journal of International Social Research* (Vol. 12, Issue 63, p. 673). The Journal of International Social Research. <https://doi.org/10.17719/jisr.2019.3264>
- Soares, F. B., Recuero, R., Volcan, T., Fagundes, G., & Sodr , G. (2021). Research note: Bolsonaro’s firehose: How Covid-19 disinformation on WhatsApp was used to fight a government political crisis in Brazil. <https://doi.org/10.37016/mr-2020-54>
- Steinfeld, N. (2022). The disinformation warfare: how users use every means possible in the political battlefield on social media. *Online Information Review* (Vol. 46, Issue 7, p. 1313). Emerald Publishing Limited. <https://doi.org/10.1108/oir-05-2020-0197>
- Su, Y., Lee, D. K. L., & Xiao, X. (2021). “I enjoy thinking critically, and I’m in control”: Examining the influences of media literacy factors on misperceptions amidst the COVID-19 infodemic. *Computers in Human Behavior* (Vol. 128, p. 107111). Elsevier BV. <https://doi.org/10.1016/j.chb.2021.107111>
- Tan, X. (2015). Applied-Information Technology with Transmission Features of New Media on News Communication. *Advances in computer science research*. Atlantis Press. <https://doi.org/10.2991/isci-15.2015.152>
- Toker, H., Erdem, S., & Özşarlak, P. (2017). 2015 Haziran ve Kasım Seçimlerinde Siyasal Eğilim: Yeni Bir Kamuoyu Ölçümleme Aracı Olarak Twitter. *Erciyes İletişim Dergisi* (Vol. 5, Issue 1, p. 96). <https://doi.org/10.17680/erciyesakademia.291888>
- Tomin, V. V., Erofeeva, N. E., Borzova, T. V., Lisitzina, T. B., Rubanik, V. E., Aliyev, H. K., & Швайпова, П. Г. (2020). Internet Media as Component of Information and Communication Environment in Electoral Process: Features and Tools. *Online Journal of Communication and Media Technologies* (Vol. 10, Issue 3). <https://doi.org/10.29333/ojcm/7932>
- Torres, Oscar.R., Getting Started in Logit and Ordered Logit Regression”, <https://www.princeton.edu/~otorres/>, Vol.3, No.1, pp.1-15,2008
- Tuğtekin, E. B., & Koç, M. (2019). Understanding the relationship between new media literacy, communication skills, and democratic tendency: Model development and testing. *New Media & Society* (Vol. 22, Issue 10, p. 1922). SAGE Publishing. <https://doi.org/10.1177/1461444819887705>

- URL1, (2025). B. Jikadara, "Fake News Detection Dataset," Kaggle. Erişim: 16 Şubat 2025, <https://www.kaggle.com/datasets/bhavikjikadara/fake-news-detection/data>
- URL2, (2025). "sklearn.linear_model.LogisticRegression — scikit-learn 1.5 documentation." Erişim: 16 Şubat 2025, https://scikit-learn.org/1.5/modules/generated/sklearn.linear_model.LogisticRegression.html
- Widyastuti, D. A. R. (2021). Using New Media and Social Media in Disaster Communication. *Komunikator* (Vol. 13, Issue 2, p. 100). Muhammadiyah University of Yogyakarta. <https://doi.org/10.18196/jkm.12074>
- Yamaguchi, S., & Tanihara, T. (2023). Relationship between misinformation spreading behaviour and true/false judgments and literacy: an empirical analysis of COVID-19 vaccine and political misinformation in Japan. *Global Knowledge Memory and Communication*. Emerald Publishing Limited. <https://doi.org/10.1108/gkmc-12-2022-0287>
- Yamaguchi, S., Oshima, H., Watanabe, T., Osaka, Y., Tanihara, T., Inoue, E., & Tanabe, S. (2025). An analysis of literacy differences related to the identification and dissemination of misinformation in Japan. *Global Knowledge Memory and Communication* (Vol. 74, Issue 11, p. 121). Emerald Publishing Limited. <https://doi.org/10.1108/gkmc-07-2024-0419>
- Yang, Y. (2024). The Trend of Media Convergence and its Impact in Journalism and Communication Studies. *Transactions on Social Science Education and Humanities Research* (Vol. 7, p. 250). <https://doi.org/10.62051/9yd2h686>
- Zhang, Y. (2024). The Impact of New Media on Traditional Media. *Journal of Education Humanities and Social Sciences* (Vol. 28, p. 691). <https://doi.org/10.54097/cfy60b30>
- Zhou, J. (2014). Research on Dynamic Extension of Visual Communication of News Website Information. *Applied Mechanics and Materials* (p. 2043). Trans Tech Publications. <https://doi.org/10.4028/www.scientific.net/amm.687-691.2043>
- Zhou, Y. (2023). On The Influence of Modern New Social Media on Brand Effect. *Lecture Notes in Education Psychology and Public Media* (Vol. 4, Issue 1, p. 249). <https://doi.org/10.54254/2753-7048/4/20220297>

Extended Abstract

Purpose

The spread of fake news in the new media environment negatively affects individuals' access to information and the formation of public opinion, thereby undermining societal trust. This study aims to detect fake news using machine learning methods. Identifying fake news is a critical step in combating disinformation. In this research, a logistic regression model was employed to detect fake news circulating on new media platforms, and the model's performance was evaluated. The study examines the effectiveness of artificial intelligence-based solutions in detecting fake news and aims to contribute to enhancing media literacy and digital reliability. Positioned at the intersection of communication sciences and artificial intelligence, this research offers a scientific perspective on fake news analysis.

Design and Methodology

This study was conducted to develop a machine learning model for detecting fake news. The "Fake and Real News Dataset" from the Kaggle platform (URL1, 2025) was used. The dataset consists of 44,898 news articles, including 23,481 fake and 21,417 real news items. Only the title and text columns were used, while subject and date variables were removed. The balance between fake and real news in the dataset eliminated the need for additional sampling techniques. During the data preprocessing phase, punctuation marks, special characters, and numbers were removed from the news texts, and all words were converted to lowercase for standardization. Stopwords removal and lemmatization were performed using the NLTK library. The texts were then vectorized using TF-IDF (Term Frequency - Inverse Document Frequency) to be processed by machine learning algorithms. A logistic regression model was chosen for fake news detection. The dataset was split into 80% training and 20% testing subsets. The model's performance was analyzed using evaluation metrics such as accuracy, precision, recall, F1-score, and ROC-AUC. To visualize the results, Matplotlib and Seaborn libraries were used to generate the confusion matrix, ROC curve, and Precision-Recall curve. The study assessed the effectiveness of logistic regression in detecting fake news and evaluated its performance. Future research suggests incorporating deep learning models (BERT, LSTM, etc.) to improve model accuracy.

Findings

The logistic regression model developed in this study was analyzed using various evaluation metrics for fake news detection. The model's success was measured using accuracy, precision, recall, and F1-score. Additionally, the confusion matrix, ROC curve, and Precision-Recall curve were used to analyze the model's predictive power.

According to the confusion matrix, the model correctly classified 4,669 fake news and 4,197 real news, with 64 false positives (misclassifying real news as fake) and 50 false negatives (misclassifying fake news as real). The low false-negative rate indicates that the model effectively detects fake news. The ROC curve and AUC score show that the model achieved an AUC value close to 1.00, demonstrating its high accuracy in distinguishing between fake and real news. The Precision-Recall curve indicates that the model consistently detects fake news while maintaining a low false-positive rate. The overall model performance is summarized in Table 2.

Tablo 1
General Evaluation of Model Performance

Metric	Value
Accuracy	%98,5
Precision	%98,7
Recall	%98,9
F1-Score	%98,8

The results show that the developed model performs highly in detecting fake news. However, reducing the false-negative rate further may require deep learning models or alternative text vectorization techniques.

Research Limitations

Although the logistic regression model developed in this study has shown successful results in detecting fake news, it has certain limitations. First, the model only analyzes text-based news. Since fake news today spreads through various formats such as images, videos, and social media posts, the model lacks the capability to detect non-textual misinformation. Additionally, the model was trained using TF-IDF vectorization, which does not deeply analyze the contextual relationships between words. The use of BERT or Transformer-based models in future research could better capture the contextual meaning of news articles. Furthermore, since the model was trained on a single dataset, its ability to generalize across different languages and news sources remains untested. Future studies could improve the model's generalization capacity by incorporating larger and more diverse datasets.

Implication

The spread of fake news in the new media environment increases information pollution, shapes public perception, and facilitates disinformation dissemination. Detecting fake news is crucial for enhancing media literacy and ensuring reliable information flow in digital environments. The study's findings contribute to the development of automated fake news detection systems, providing valuable insights for future research.

Originality and value

The originality of this study lies in its analysis of fake news detectability using a machine learning-based model. The logistic regression model utilized in this study effectively identified linguistic patterns in fake news and achieved high accuracy rates. By employing Natural Language Processing (NLP) and TF-IDF vectorization, news texts were analyzed, and the model's capacity to distinguish fake news was assessed. This research is positioned at the intersection of communication sciences and artificial intelligence, providing a scientific perspective on fake news detection. The study establishes a foundation for developing automated detection systems to combat disinformation in digital platforms.