
BİLGİSAYAR BİLİMLERİ VE MÜHENDİSLİĞİ

Editör: Doç. Dr. Bekir PARLAK

BİLGİSAYAR BİLİMLERİ VE MÜHENDİSLİĞİ

Editör

Doç. Dr. Bekir PARLAK

yaz
yayınları

2024

BİLGİSAYAR BİLİMLERİ VE MÜHENDİSLİĞİ

Editör: Doç. Dr. Bekir PARLAK

© YAZ Yayınları

Bu kitabın her türlü yayın hakkı Yaz Yayınları'na aittir, tüm hakları saklıdır. Kitabın tamamı ya da bir kısmı 5846 sayılı Kanun'un hükümlerine göre, kitabı yayınlayan firmanın önceden izni alınmaksızın elektronik, mekanik, fotokopi ya da herhangi bir kayıt sistemiyle çoğaltılamaz, yayınlanamaz, depolanamaz.

E_ISBN 978-625-6104-59-4

Ekim 2024 – Afyonkarahisar

Dizgi/Mizanpaj: YAZ Yayınları

Kapak Tasarım: YAZ Yayınları

YAZ Yayınları. Yayıncı Sertifika No: 73086

M.İhtisas OSB Mah. 4A Cad. No:3/3
İscehisar/AFYONKARAHİSAR

www.yazyayinlari.com

yazyayinlari@gmail.com

info@yazyayinlari.com

İÇİNDEKİLER

İş Kazalarında YSA & ANFIS Faktörü	1
<i>Ömer Faruk ÜTÜK, Ömer ASAL, Yunus KAYIR, Hakan DİLİPAK</i>	
İşbirlikçi Filtreleme ile Film Öneri Sistemi	19
<i>Bekir PARLAK, Fadimana DİKİCİ, Esmâ AKSAN</i>	
Türkçe ve İngilizce Haber Verilerinin Sınıflandırılmasında Öznitelik Seçimi ve Globalleştirme Tekniklerinin Rolü.....	34
<i>Bekir PARLAK</i>	
New Opportunities and Challenges in Quantum Natural Language Processing.....	55
<i>Ramazan KATIRCI, Taha OĞUZ</i>	
Transformer Mimarisi.....	73
<i>Ramazan KATIRCI, Hilal ÇELİK</i>	
Powerlifting (Güç Kaldırma) Sporcularının Yarışma Kategorilerinin Geçmiş Sporcu Performanslarına Göre Belirlenmesi.....	95
<i>Ege Kaan SUCU, Ümit YILMAZ, Atınç YILMAZ</i>	

"Bu kitapta yer alan bölümlerde kullanılan kaynakların, görüşlerin, bulguların, sonuçların, tablo, şekil, resim ve her türlü içeriğin sorumluluğu yazar veya yazarlarına ait olup ulusal ve uluslararası telif haklarına konu olabilecek mali ve hukuki sorumluluk da yazarlara aittir."

İŞ KAZALARINDA YSA&ANFIS FAKTÖRÜ

Ömer Faruk ÜTÜK¹

Ömer ASAL²

Yunus KAYIR³

Hakan DİLİPAK⁴

1. GİRİŞ

Önceden planlanmamış, güvenliksiz eylem ve şartlardan doğan, çalışanların güvenliğini tehdit eden, kullanılan teçhizatın zarar görmesine neden olan veya üretimi sekteye uğratan olaylara iş kazası denilmektedir [1].

İş kazalarının önlenmesinde kullanılan geleneksel yöntemleri ele aldığımızda; işçilere güvenli çalışma kuralları hakkında düzenli eğitimler vermek, iş yerindeki potansiyel tehlikeleri belirleyerek bunların nasıl azaltılabileceğine yönelik bir plan yaparak risk analizi oluşturmak, işçilerin güvenliğini sağlama amaçlı kişisel koruyucu donanım kullanımını sağlamak, yasal düzenlemelere uyum sağlamak, düzenli denetimler yaparak potansiyel tehlikeleri belirlemek gibi klasik yöntemlerden bahsetmek mümkündür [2].

¹ Doktora Öğrencisi, Gazi Üniversitesi, Fen Bilimleri Enstitüsü, Kazaların Çevresel ve Teknik Araştırması ABD, Ankara, Türkiye, omerutuk@gmail.com, ORCID: 0000-0003-4658-6652.

² Doç. Dr., Gazi Üniversitesi, Teknoloji Fakültesi, İmalat Mühendisliği Bölümü, Ankara, Türkiye, omerasal@gazi.edu.tr, ORCID: 0000-0002-6339-9202.

³ Prof. Dr., Gazi Üniversitesi, Teknoloji Fakültesi, İmalat Mühendisliği Bölümü, Ankara, Türkiye, ykayir@gazi.edu.tr, ORCID: 0000-0001-6793-7103.

⁴ Prof. Dr., Gazi Üniversitesi, Teknoloji Fakültesi, İmalat Mühendisliği Bölümü, Ankara, Türkiye, hdilipak@gazi.edu.tr, ORCID: 0000-0003-3796-8181.

Maalesef ki geçmişten günümüze kadar kullanılan ve halen geleneksel olarak kullanımına devam edilen klasik kaza önleme metotları kazaları istenilen seviyede önleyememektedir. Tüm bu yöntemler İSG' nin sağlanmasında ve uygulanmasında iş kazası sonucu ölüme ve yaralanmaya engel olamamaktadır. Örneğin, tüm sektörlerde kullanılan kişisel koruyucu donanımların daha güvenli bir çalışma ortamı sağlamak adına yetersiz kaldığı gözlemlenmiştir. Bu gibi sebeplerden dolayı gerek araştırmacılar gerekse de şirketler, iş sahalarında iş yapış şekillerini değiştiren ya da yapay zekâ tekniklerinin kullanıldığı teknolojik yöntemlerin, iş kazalarını azaltmada önemli ölçüde katkı sağlayacağını öngörmüşlerdir. Bu yapay zekâ tekniklerinden Yapay Sinir Ağları ve ANFIS ön plana çıkmaktadır.

Yapay sinir ağları, yinelenebilir basit elemanların yoğun bir şekilde paralel bağlanmasıyla oluşan ağlardır ve sinir hücrelerinin oluşturduğu sinir sisteminin bir benzeridir [3]. Yapay sinir ağları, insan beyninden ilham alınarak geliştirilen, öğrenme aşamasının matematiksel modellemeye dayandığı bir sistemdir [4]. ANFIS ise Bulanık Mantık ve Yapay Sinir Ağı yöntemlerini birleştiren bir yapıya sahiptir [5]. Bu iki yöntem arasındaki en yakın benzerlik, girdi ve çıktı değişkenleri arasındaki durumu kavramaya çalışmalarıdır. Bu maksatla her yöntemde modeller oluşturmak için kullanılabilirler [6].

Bu çalışmada, iş kazalarının önlenmesinde yapay zekâ tabanlı YSA ve ANFIS teknikleri kullanan çalışmalar literatür taraması yapılarak detaylıca incelenmiştir. İş kazalarını minimize etmede geleneksel yöntemlerin tek başına yeterli olmadığı, yapay zekâ tabanlı tekniklerden yapay sinir ağlarının daha etkin ve elverişli olduğu kanısına varılmıştır.

2. YSA & ANFIS

Yapay sinir ağı, yapay zekâ türleri arasında yer alan bir teknolojidir. İnsan beynine yapısal benzerliği olan bu sistem, öğrenebilir ve tecrübe kazanabilir özelliklere sahiptir. Özellikle problemleri çözme konusunda oldukça etkilidir. Öğrenme işlemini örneklerle gerçekleştiren yapay sinir ağları, matematiksel olarak çözümü zor olan problemleri birçok veri üzerinden örneklem yaparak doğru sonuca ulaşmayı öğrenirler [7]. Bu sayede, geleneksel yöntemlerle çözümü oldukça zor olan problemleri yüksek doğruluk derecesiyle çözebilme yeteneğine sahiptirler.

Bulanık mantık, Lotfi Aliasker Zadeh tarafından 1965 yılında klasik mantıktan dilbilimsel ifade kavramına geçiş üzerine yayınlanan makaleye dayanmaktadır. Zadeh'e göre, bir varlığın bir küme içinde yer alma durumu en az üç farklı önerme ile ifade edilebilir [8]. Matematiksel olarak söylemek gerekirse, klasik mantığa göre bir varlık ya bir kümenin içinde yer alır ya da yer almaz. Ancak bulanık mantıkta her varlığın bir kümede yer alma kademesi mevcuttur. Klasik mantıkta “0” veya “1” olarak gösterilenler, bulanık mantıkta “0” ile “1” aralığında bir kademe derecesi olarak gösterilmektedir [9].

ANFIS yöntemi, bulanık mantık ve yapay sinir ağı yöntemlerinin avantajlarını birleştirerek etkili bir çözüm sunmaktadır. Bu yöntem, çeşitli uygulamalarda kullanılabilir. Örneğin, trafik akışını kontrol etmek veya yatırım kararları vermek gibi alanlarda kullanılabilir. ANFIS yöntemi ayrıca finansal verilerin analiz edilmesiyle gelecekteki finansal trendlerin tahmin edilmesinde de kullanılabilir. Bulanık mantık yöntemi, belirsizlik içeren problemlerin çözümünde kullanılan bir yöntemdir. Bu yöntem, insanların karar verme sürecinde kullandığı düşünme yöntemine benzer şekilde çalışır. Yapay sinir ağı yöntemi ise, insan beyninin işleyişini

modelleyen bir yöntemdir. Bu yöntem, öğrenme yeteneği sayesinde verilerden trendleri ve ilişkileri çıkarabilir. ANFIS yöntemi, bu iki yöntemin avantajlarını birleştirerek dezavantajlarını azaltır. Bu sayede, daha etkili bir sonuç elde edilir. Bu metot özellikle karmaşık problemlerin çözümünde tercih edilmektedir. Uygulama alanlarına bakıldığında, bu yöntemi trafik akışını kontrol etmek için kullanabilir, bu sayede trafik sıkışıklığı azaltılabilir ve trafik akışı daha düzenli hale getirilebilmektedir. Aynı zamanda bu yöntem ile finansal veriler analiz edilerek gelecekteki finansal trendler tahmin edilebilir. Bu sayede, şirketler veya yatırımcılar gelecekteki piyasa koşullarına göre stratejilerini belirleyebilirler. Literatür taramasına konu olan iş kazalarını önleme ya da minimize etme amaçlı ANFIS yöntemi kullanıldığında olumlu geri bildirimler alındığı çalışmalarda gösterilmektedir. Sonuç olarak, ANFIS yöntemi bulanık mantık ve yapay sinir ağı yöntemlerinin avantajlarını birleştirerek etkili bir çözüm sunar. Bu yöntem, birçok alanda kullanılabilir ve gelecekte daha fazla uygulama alanı bulabileceği öngörülmektedir. Bu yöntem sayesinde, karmaşık problemler daha etkili bir şekilde çözülerek bilinçli kararlar alınabilmektedir.

Literatürde iş kazalarını önlemeye yönelik YSA&ANFIS tabanlı birçok çalışma karşımıza çıkmaktadır.

İnşaat sektöründe, yüksekten düşme ve iskele ile ilgili kazalar sıkça yaşanmaktadır. Güvenlik denetimleri yapılsa da, iskele kazaları hala meydana gelmektedir. Ancak, mevcut güvenlik izleme uygulamaları pratik değil ve gerçekleştirilemezdir. Bu nedenle Khan ve ark., mobil iskele güvenliği izlemesi ve işçinin güvensiz davranışlarını tespit etmek için ilişki tabanlı bir yaklaşım önermişlerdir. Bu yaklaşım, Mask R-CNN derin sinir ağı ve nesne ilişkisi tespiti modülü kullanılarak işçinin güvenli ve güvensiz davranışlarının tespit edilmesini sağlamaktadır. Test sonuçları, yaklaşımın %86 doğruluk oranı ile işçinin güvensiz davranışlarını etkili bir şekilde

tespit edebildiğini göstermektedir. Bu nedenle, Mask R-CNN tabanlı OCD modülü inşaat ortamında kullanılabilecek bir çözüm olarak öne çıkmaktadır [10].

İnşaat sektöründe yapılan bir çalışmada, keskin bir cisim ile temas sonucu meydana gelen yaralanma kazalarının analizi yapılmıştır. Bu analizde analitik hiyerarşi prosesi ve yapay sinir ağları kullanılarak bir model oluşturulmuştur. Oluşturulan bu model, kaza tehlikelerini %78 doğruluk oranıyla tahmin etmeyi başarmıştır. Keskin cisimlerle temas sonucu meydana gelen yaralanma kazaları, çalışma ortamlarında sıkça karşılaşılan risklerdendir. Bu kazaların önlenmesi ve iş güvenliğinin sağlanması için etkili bir risk değerlendirme yöntemi kullanılması büyük önem taşımaktadır. Analitik hiyerarşi prosesi ve yapay sinir ağları, karmaşık veri setlerinin analiz edilmesi ve sonuçların tahmin edilmesi için yaygın olarak kullanılan yöntemlerdir. Bu çalışmada da bu yöntemler kullanılarak, keskin cisimlerle temas sonucu meydana gelen yaralanma kazalarının nedenleri ve risk faktörleri analiz edilmiştir. Elde edilen veriler kullanılarak bir model oluşturulmuş ve bu model sayesinde kaza tehlikeleri tahmin edilebilmiştir. Keskin cisimlerle temas sonucu meydana gelen yaralanma kazalarının analiz edilmesi ve kaza tehlikelerinin tahmin edilmesi konusunda önemli bir katkı sağlanan bu çalışmada elde edilen bulgular, iş sağlığı ve güvenliği alanında daha iyi politikaların oluşturulmasına ve işyerlerindeki kaza risklerinin azaltılmasına yardımcı olacaktır [11].

Bir çalışma, metal sektöründe meydana gelen iş kazalarını inceleyerek bir kaza veri seti oluşturmuştur [12]. Bu veri seti, çok değişkenli veri analizi yöntemlerini kullanarak çeşitli çıkarımlar elde etmiştir. Bununla birlikte veri kümesinde değişkenlerin indirgenmesi yapılmıştır. Yapılan çalışmalar sonucunda, en iyi tahmin modelini üreten algoritmanın yapay sinir ağları olduğu tespit edilmiştir. Çift katmanlı ileri beslemeli bir yapay sinir ağı kullanılarak bir kaza tahmin modeli oluşturulmuş ve örnek metal

sektörü işyeri verileri üzerinde bu model denenerek, işyerlerindeki olası kaza riskleri değerlendirilmiştir. Bu çalışma, metal sektöründe faaliyet gösteren işyerlerinin güvenliği açısından önemli bir adım olmuştur ve benzer sektörlerde de benzer çalışmaların yapılması önerilmektedir. Veri madenciliği ve makine öğrenimi teknolojilerinin kullanımıyla, iş güvenliği konusunda daha akılcı ve etkili çözümler üretilebileceği görülmüştür. Bu sayede, iş kazalarının önlenmesi ve çalışanların güvenliği sağlanabilmektedir.

Bir yapay sinir ağı kullanarak çalışanların güvenli çalışma davranışını tahmin etmek için bir model geliştirilmiştir. Model, literatür taramasıyla belirlenen 10 güvenlik iklimi yapısını giriş olarak kullanırken, güvenli çalışma davranışını çıkış olarak kullanmaktadır. Hindistan genelindeki çeşitli inşaat projelerinden 222 yanıt bir anket aracılığıyla toplandı. Bu modeli oluşturmak için üç katmanlı bir yapay sinir ağı kullanılmıştır ve yeterli veri setleriyle eğitilmiş, doğrulanmış ve test edilmiştir. Model, çalışanların güvenli çalışma davranışını oldukça iyi tahmin etmektedir. Ayrıca, her bir yapı unsurunun çalışanların güvenli çalışma davranışı üzerindeki etkisini incelemek için bir duyarlılık analizi yapılmıştır. Sonuç olarak, denetim ortamı, iş baskısı, çalışanların katılımı, riskin kişisel takdiri ve destekleyici ortam gibi güvenlik iklimi yapısı unsurlarının çalışanların güvenli çalışma davranışı ile önemli ölçüde ilişkili olduğu bulunmuştur. Bu model, müteahhitlere ve müşterilere güvenli çalışma davranışını teşvik etmede ve inşaat projelerinde çalışanların güvenliğinin verimli bir şekilde yönetilmesine yardımcı olma potansiyeline sahiptir [13].

İnşaat sektöründe meydana gelen ölümlerin en önemli nedeni düşmelerdir ve bu oran %30-35 civarındadır. Ancak, işçilerin düşmelerini algılama yöntemleri bazı sınırlılıklara sahiptir. Bu nedenle, giyilebilir sensörler yerine iskeledeki ivmeleri kullanarak işçilerin hareketlerindeki güvensiz eylem

desenlerini öğrenmek için Lee ve ark., evrişimli sinir ağı (CNN) önermişlerdir. Yaygın olarak kullanılan algoritmaların tanımlanamayan sınıfları sınıflandırmak için uygun olmaması gibi sorunlarla karşılaşilmektedir. Dört modelleme stratejisi kullanılarak tanımlanmış öncü belirteçlerinin algılanması ve sınıflandırılması gerçekleştirilmiştir. Bu sonuçlar, iskelelerdeki ivmelerin işçilerin hareketleri hakkında yeterli bilgi içerdiğini ve modelleme stratejilerinin tanımlanmış öncüleri sorunsuz bir şekilde sınıflandırabileceğini göstermektedir. Bu yaklaşım, yüksek yerlerde çalışan işçilerin sürekli olarak izlenerek düşme ile ilgili olayların azaltılmasında yardımcı olabileceğini göstermişlerdir. Giyilebilir sensörlerin fiziksel müdahale ve gizlilik sorunlarına neden olması da bu yöntemin tercih edilmesinde etkili olmuştur [14].

Petrokimya endüstrisinde sağlık, çevre, güvenlik ve ergonomi (HSEE) tabanlı iş güvenliğini ölçen, değerlendiren ve araştıran uyumlaştırılmış akıllı algoritmalar geliştirmek amaçlanmıştır. Azadeh ve arkadaşları sağlık, güvenlik ve çevre (HSE) ile ergonomi açısından iş güvenliği değerlendirmesi için bir algoritma sunmaktadır [15]. Ayrıca, insan performansının iş güvenliği ve HSEE faktörleri açısından ölçülmesi için ANFIS tabanlı bir parametrik olmayan verimlilik sınırı analizi yöntemi önerilmektedir. Önerilen yaklaşım, büyük bir petrokimya tesisindeki bir dizi operatöre uygulanarak uygulanabilirliği ve üstünlüğü gösterilmektedir. Bu çalışmanın amaçlarına ulaşmak için, iş güvenliği ve HSEE ile ilgili standart anketler operatörler tarafından tamamlanmaktadır. HSEE'nin her kategorisi için ortalama sonuçlar girdi olarak kullanılırken, iş güvenliği çıktı değişkeni olarak kullanılmaktadır. Ayrıca bu algoritma, iş güvenliği ve HSEE açısından operatörleri sıralamaktadır. Son olarak, bu algoritma verimsiz operatörleri belirlemektedir. Bu çalışmanın üstünlüğü ve avantajları da gösterilmektedir. HSEE ve iş güvenliği açısından insan performansının değerlendirilmesi ve

iyileştirilmesi için bir zekâ algoritması sunan ilk çalışma olduğu görülmüştür.

İnşaat işçilerinin eğitim almalarına rağmen, inşaat ortamlarındaki önemli bir kısım tehlikeyi tanıyamadıkları görülmektedir. Bu eksikliği gidermek için Jeelani ve arkadaşları karmaşık ve dinamik inşaat ortamlarında çalışanları ve güvenlik yöneticilerini tehlikeleri tanıma konusunda destekleyen teknolojilerin geliştirilmesini amaçlamışlardır. Bu çalışmada, gerçek zamanlı olarak tehlikeli koşulları ve nesneleri tespit eden otomatik bir sistem için bir çerçeve oluşturulmuştur. Çerçeve, işçilerin yerini belirleme, işçilerin çevresindeki görsel sahnenin anlamsal segmentasyonu ve statik ve dinamik tehlikelerin tespiti için üç bağımsız işlem hattından oluşmaktadır. Çerçeve, inşaat işyerlerinde çalışanların ve güvenlik yöneticilerinin tehlike tespit yeteneklerini otomatikleştirmek ve artırmak için kullanılabilmektedir. Ayrıca, çerçeve, işlemci katkılarına ek olarak, geliştirilmiş gerçek zamanlı işçi yerleştirme yöntemi ve varlık yerleştirme ve nesne tespiti için işlem hatlarının entegrasyonu için verimli bir mimari sunmaktadır. Önerilen çerçeve üzerine geliştirilen derin öğrenme tabanlı bir sistemle, iç ve dış mekân inşaat ortamlarında uygulama yapılarak %93'ten fazla doğruluk elde edilmiştir [16].

Bir metal sektöründe meydana gelen iş kazalarının incelenmesi sonucunda, Yapay Sinir Ağı kullanılarak oluşacak iş gücü kaybı tahmin edilmiştir [17]. Bu çalışma sayesinde, aday çalışanların yıl bazında kaza geçirme riskleri tahmin edilerek üst yönetim tarafından güvenilir kararlar alınabilmektedir. Yapay Sinir Ağı ile geliştirilen model sayesinde, olası kazaların sıklığı önceden tahmin edilebilmekte ve işletmeler iş sağlığı ve güvenliği konusunda daha bilinçli hareket edebilmektedir. Bu çalışma, endüstriyel kuruluşların iş sağlığı ve güvenliği konusunda daha etkili bir şekilde hareket etmesine yardımcı olacak veriler sağlamaktadır. Sonuç olarak, Yapay Sinir Ağı ile

yapılan bu çalışma, iş kazalarının önlenmesi ve iş gücü kaybının minimize edilmesi için önemli bir araçtır.

Bu çalışmada, inşaat sektöründe mesleki risklerin sistematik bir şekilde değerlendirilmesi için uzmanların bilgisi kullanılarak bir ANFIS modeli geliştirilmiştir [18]. İşyeri risk değerlendirmesi, sektörde güvenliği sağlamak için önemli bir önlemdir ancak yetersiz veri veya kesin olmayan bilgiler nedeniyle zorlu bir süreçtir. Bu nedenle, geleneksel nicel yaklaşımlar sıklıkla başarısız olmaktadır. Bu makalede, bu eksiklikleri aşmak için bir Takagi-Sugeno tipi bulanık çıkarım sistemi geliştirilmiştir. Model formülasyon sürecinde, kazaların nedenleri ve kontrol faktörleri girdi parametreleri olarak ele alınmıştır. Her tür yaralanma eğilimli vücut bölgesinin güvenlik seviyeleri analitik hiyerarşi süreci kullanılarak değerlendirilir. Sayısal kümelenme tekniği, kuralların sayısını azaltmak için kullanılmış ve böylece bir başlangıç bulanık çıkarım sistemi oluşturulmuştur. Son olarak, girdi değişkenlerine karşılık gelen tüm parametreleri ayarlayarak hibrit bir öğrenme süreci kullanılarak başlangıç modeli güncellenir. Geliştirilen yöntem, Hindistan'daki birkaç seçilmiş inşaat alanında uygulanmış ve sonuçlar, modelin inşaat sitelerinde riskleri değerlendirmek için uygulanabilir olduğunu ve mevcut güvenlik stratejilerinin ilerlemesini belirlediğini doğrulamaktadır.

İş sağlığı ve güvenliği, 6331 Sayılı İş Sağlığı ve Güvenliği Kanunu ile tüm çalışanları kapsayarak daha da önemli hale gelmiştir. Kanun, iş verenlerin iş yerlerinde iş sağlığı ve güvenliği uzmanı, iş yeri hekimi ve diğer sağlık personeli bulundurmasını zorunlu kılmaktadır. Ancak, uzmanların da hatalar yapabileceği ve koruyucusuz makine ve cihazlar gibi durumlarda yetersiz kalabileceği hatırlatılmaktadır. Bu nedenle, İSG profesyonelleri tarafından alınacak tedbirlerle kazaların önüne geçilebilmesi için bir karar destek sistemi geliştirilmesi önerilmektedir [19]. Yapay sinir ağı tabanlı bu sistem, uzman kişinin kararına yardımcı olarak

muhtemel riskleri ve alınması gereken tedbirleri saptayacak ve geçmişte yapılmış risk analizleri sonucu ortaya çıkan durumları öğrenerek daha kaliteli kararlar alınmasını sağlayacaktır. İş sağlığı ve güvenliği konusunda alınacak tedbirlerin önemi ve bu alanda kullanılacak teknolojilerin gelişimi ile birlikte, çalışanların sağlığı ve güvenliği için daha güçlü adımlar atılacaktır.

İnşaat sektöründeki bu çalışmada, yapı üretim sürecinde meydana gelen iş kazalarının önlenmesi için bütünlük bir model geliştirilmiştir. Geçmiş kaza verileri kullanılarak oluşturulan model, AHP ve YSA metodlarının birbirine entegre edilmesiyle oluşturulmuştur. Model, gerçek verilerle test edilerek anlamlılığı kanıtlanmıştır. Örneklem için seçilen 4 tür iş kazasında, toplanan verilerle risk azaltıcı önlemler ile kaza şiddeti arasındaki ilişki %90 oranında anlamlı bulunmuştur. Bu çalışma, iş sağlığı ve güvenliği uzmanlarına, yapı üretim sürecinde iş kazalarının önlenmesinde kullanılabilecek etkili bir model sunmaktadır. Çalışma, inşaat sektöründe faaliyet gösteren tüm kurumların iş sağlığı ve güvenliği konularına daha fazla önem vermeleri gerektiğini vurgulamaktadır. Bu sayede, iş kazalarının önlenmesi ve çalışanların güvenliği sağlanarak sektörde daha sağlıklı bir çalışma ortamı oluşturulması hedeflenmektedir [20].

Yaşlı yetişkinlerde düşme riskinin azaltılması için doğru bir şekilde denge eksikliklerinin tespit edilmesi oldukça önemlidir. Bu amaçla yapılan bir çalışmada, insan denge desenlerinin tanımlanması ve düşme riskinin tahmin edilmesi amacıyla derin sinir ağlarının kapasitesi araştırılmaktadır. İnsan denge yeteneği, kuvvet plakası zaman serisinden elde edilen denge ölçütleri gibi denge metriklerine dayanarak karakterize edilebilir. Bu çalışmada, düşme riskinin düşük, orta ve yüksek risk olarak karakterize edilebileceği ve derin öğrenme algoritmaları ile birlikte kuvvet plakası zaman serisinden elde edilen denge metrikleri ile tahmin edilebileceği hipotezi test edilmiştir. Ayrıca, önerilen One-One-One Derin Sinir Ağı

algoritmasının diğer algoritmalara göre önemli bir performans artışı sağladığı gözlenmiştir. Yapılan çalışmada, geniş bir demografik insandan (18-86 yaşında 163 kadın ve erkek) çeşitli duruş koşulları için insan dengesini ölçen açık kaynaklı bir kuvvet plakası veri seti kullanılmıştır. Sonuçlar, önerilen One-One-One Derin Sinir Ağı modelinin kuvvet plakası zaman serisi sinyali kullanarak insan düşme riskini en etkili şekilde tahmin eden en verimli sinir ağı olduğunu göstermiştir. Bu çalışma, insan düşme riskinin doğru bir şekilde tahmin edilmesi için yeni bir metodoloji sunmaktadır [21].

3. YSA VE ANFIS'IN İŞ KAZALARINDAKİ ETKİNLİĞİ

Literatür taramasına genel olarak bakıldığında iş kazalarını önlemede veya minimize etmede geleneksel yöntemlerin yanında yapay zekâ teknolojilerinden yararlanmanın kaçınılmaz olduğu varsayımıyla incelemeler yapılmıştır. Yapay zekâ tekniklerinden YSA ve ANFIS özeline inildiğinde yapay sinir ağlarının iş kazalarını önlemede daha etkin olduğu açıkça görülmüştür. İş kazalarını önleme açısından, ANFIS yönteminin bulanık mantık tabanlı olmasından ötürü YSA'ya göre daha az tercih edildiği kanısına varılmış olursa da bazı çalışmalarda bu yöntemin başarılı olduğu tespit edilmiştir. Nitekim bir inşaat sektörünün mesleki risklerinin sistematik bir şekilde değerlendirilmesi amacıyla uzmanların bilgilerinden yararlanılarak bir ANFIS modeli geliştirilmiş ve uygulanmıştır [18]. Bu mimarinin temel avantajının girdi-çıkı davranışını tanımlamak için alan uzmanlarından elde edilen bulanık kuralları hassaslaştırmasıdır. Geliştirilen metodolojinin, şantiyelerdeki mesleki risklerin FIS kullanılarak değerlendirilmesinde olumlu sonuçlar verebileceği öngörülmektedir.

Bir petrokimya tesisinde uygulanan çalışmada iş güvenliği, sağlık, güvenlik ve çevre (HSEE) ile ergonomi açısından operatörlerin performansını ölçmek için bir algoritma sunulmuştur [15]. Standart anketler kullanarak operatörler tarafından iş güvenliği ve HSEE hakkında veriler toplanmış ve bu veriler, ANFIS kullanılarak analiz edilmiştir. Petrokimya tesisi operatörleri üzerinde uygulanan bu yöntemin, diğer benzer çalışmalara göre daha üstün sonuçlar verdiği tespit edilmiş olsa da ANFIS' ın teknik verimliliğin ölçülmesinde potansiyel bir alternatif olabileceğine ve üretim sürecinin bilinmediği durumlarda diğer tekniklere göre daha iyi performans gösterebileceğine inanılmasına rağmen, verimlilik analizi ve dolayısıyla optimizasyon analizi konusunda hala hem teorik hem de ampirik çalışmaların eksikliği söz konusudur. Bununla birlikte, verimlilik ve optimizasyon analizinde ANFIS ile gelecekte yapılacak araştırmalar önerilmektedir. Gelecekte daha fazla çıktı veya girdiye sahip bazı durumlar incelenebilir ve ayrıca önerilen algoritma veri eksikliğiyle karşılaşılarak genişletilebilir.

Sektörel bazda iş kazalarının çokça yaşandığı inşaat, imalat ve metal sektörlerinde yapılan çalışmalar bir bütün olarak değerlendirildiğinde, yapay sinir ağlarıyla ilgili yapılan çalışmaların iş kazalarını önlemede ANFIS faktörüne göre daha üstün olduğu açıkça görülmüştür. Nitekim bir inşaat sektöründe yapılan çalışmada, işçilerin geleneksel yöntemlerle birlikte İSG eğitimi almalarına rağmen inşaat ortamlarındaki önemli bir kısım tehlikeyi tanıyamadıkları görülmüştür. Bu eksikliği giderme adına yapılan bu çalışmada gerçek zamanlı olarak tehlikeli koşulları ve nesneleri tespit eden otomatik bir sistem için bir çerçeve oluşturulmuştur. Önerilen çerçeve üzerine geliştirilen YSA tabanlı bir sistemle, iç ve dış mekân inşaat ortamlarında uygulama yapılarak %94 civarlarında bir doğrulukla sistem kullanılabilirliği test edilmiştir. Bu da bize gösteriyor ki iş kazalarını önleme adına yapılan çalışmalarda minimal düzeydeki

tehlikelerin dahi göz ardı edilemeyeceği ve çözüm için YSA tabanlı sistemlere ihtiyacın bir zorunluluk haline geldiği açıkça görülmektedir.

Çalışanların güvenli çalışma davranışını tahmin etmek için geliştirilen başka bir modelde, üç katmanlı bir yapay sinir ağı kullanılmış ve yeterli sayıda veri ile entegre edilerek, çalışanların güvenli çalışma davranışıyla ilgili tatmin edici oranda olumlu geri bildirim alınmıştır [13]. Denetim ortamı, iş baskısı, çalışanların katılımı, riskin kişisel takdiri ve destekleyici ortam gibi güvenlik iklimi yapısı unsurlarının çalışanların güvenli çalışma davranışı ile önemli ölçüde ilişkili olduğu bulunmuştur. İş kazalarını önlemede ya da minimize etmede çalışanın güvenli çalışma davranışının önemli bir etkisi olduğu gerçeğiyle birlikte katmanlı YSA'nın kullanımın doğruluk oranlarını yüksek mertebelere çıkardığı gözlemlenmiştir.

Metal sektöründeki iş kazalarını önleme adına kaza veri seti oluşturulan başka bir çalışmada, çok değişkenli veri analizi yöntemleri kullanılarak çeşitli çıkarımlar elde edilmiş, veri kümesinde değişkenlerin indirgenmesi yapılmıştır. Yapılan çalışmalar neticesinde, en iyi tahmin modelini üreten algoritmanın YSA olduğu tespit edilmiştir [12]. Metal sektöründe faaliyet gösteren işyerlerinin güvenliği için yapılan bu çalışma, diğer sektörlerdeki iş kazalarına da ilham kaynağı olabilir.

Tüm bu çalışmalar ortak olarak analiz edildiğinde iş kazalarını önlemede YSA'nın oldukça etkin bir rolü olduğu görülmüştür. İş kazalarının meydana gelme nedenlerini analiz eden, bu kazaların önlenmesi için alınacak önlemleri belirleyen YSA, iş kazalarının olası sonuçlarını da tahmin ederek, olası riskleri en aza indirmeye yardımcı olabilmektedir. Bu nedenle, yapay sinir ağlarının etkinliği İSG konusunda oldukça önem arz etmektedir. İş kazalarının nedenlerini ve olası sonuçlarını analiz etmek için kullanılabilecek birçok veriye sahip olabilen yapay

sinir ağlarında işçi davranışları, ekipman arızaları, çevresel faktörler gibi verilerin sıklığı ve sayısı arttıkça etkinlik noktasında daha olumlu sonuçlar alınabileceği kanaati oluşmuştur.

4. SONUÇ

Yapay sinir ağı, verilere dayalı olarak öğrenme yeteneğine sahip olan bir modeldir. Bulanık mantık ise belirsiz veya kesin olmayan verileri işlemek için kullanılan bir yöntemdir. Bu iki sistematik yaklaşımın birleşmesiyle oluşan ANFIS bir yapay sinir ağı modeli üzerine bulanık mantık kurallarını entegre etmektedir. Çalışma prensibi, verilerin yapay sinir ağı üzerindeki ağırlıklarının ve bulanık mantık kurallarının adaptif olarak güncellenmesine dayanır. Bu sayede, verilere dayalı olarak en iyi sonuçları elde etmek için sistem sürekli olarak kendini iyileştirir. Literatür incelendiğinde görülmüştür ki ANFIS'in ve YSA'nın birçok uygulama alanı bulunmaktadır. Özellikle tahmin, sınıflandırma ve kontrol gibi problemlerin çözümünde etkili olabilen bu sistemler, ayrıca karmaşık sistemlerin modellenmesi ve analizi için de kullanılabilir. Çalışmaya bütün olarak bakıldığında iş kazalarının önlenmesinde veya minimize edilmesinde ANFIS'in bulanık tabanlı olmasından ötürü iş kazalarının önlenmesinde YSA kadar etkili olmadığı kanaatine varılmıştır.

Geleneksel yöntemlerin yani matematiksel metotlarla yapılan risk analizinin kazaların önlenmesinde etkili olduğu görülmüş olsa da yapay zekanın, özellikle de uzmanlığın gerektirdiği karar vermede öğrenme imkânı sağlayan yapay sinir ağlarının, uzmanın doğru kararlar almasında daha etkili olduğu görülmüştür. Önceden belirlenmiş parametrelerle beslenerek olası riskleri belirleyebilen YSA, iş kazalarının önlenmesinde etkin bir role sahip olduğu görülmüştür. Yapılan çalışmalarda

yapay sinir ağlarının klasik yöntemlere göre daha etkili sonuçlar verdiğini görmem mümkündür. Bunun ana sebebi, yapay sinir ağlarının veri analizi yaparken daha fazla parametreyi dikkate alabilmesi, böylece olası risklerin daha hızlı ve doğru bir şekilde belirlenebilmesinden kaynaklanmaktadır. Yapay sinir ağlarının geleneksel(klasik) yöntemlerle birlikte kullanıldığında daha etkili sonuçlar ürettiği tespit edilmiştir. Kullanılan parametrelerin sayısının artmasıyla birlikte karmaşıklığın da arttığı, geleneksel yöntemlerin bu aşamada yetersiz kaldığı, bu sebeplerden ötürü olası risklerin tespit edilmesinde yapay sinir ağlarının etkin rol aldığı açıkça görülmüştür. Yapay sinir ağlarından maksimum derecede verim alınabilmesi için, daha doğru ve etkin kararlar alma noktasında klasik yöntemlerle entegre edilerek kullanılmasının daha uygun olacağı, böylelikle YSA'nın iş kazalarını minimize etmede daha etkin bir rol alacağı sonucuna varılmıştır.

KAYNAKÇA

1. Ceylan, H. (2011). Türkiye'deki iş kazalarının genel görünümü ve gelişmiş ülkelerle kıyaslanması. *International Journal of Engineering Research and Development*, 3(2), 18-24.
2. Horozoğlu, K. (2017). İş kazalarının iş sağlığı ve güvenliği açısından analizi. *Karabük Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, 8(2), 265-281.
3. Şen, Z. (2004). Yapay sinir ağları. *Su Vakfı*.
4. Elmas, Ç. (2003). Yapay Sinir Ağları (Kuram, Mimari, Eğitim, Uygulama). Seçkin Yayıncılık, Ankara, 22-37.
5. Chen, M. Y. (2013). A hybrid ANFIS model for business failure prediction utilizing particle swarm optimization and subtractive clustering. *Information Sciences*, 220, 180-195.
6. Özkan, G., & İnal, M. (2014). Comparison of neural network application for fuzzy and ANFIS approaches for multi-criteria decision-making problems. *Applied Soft Computing*, 24, 232-238.
7. Yegnanarayana, B. (2009). Artificial neural networks. PHI Learning Pvt. Ltd..Zadeh, L. A. (1965). Fuzzy sets. *Information and control*, 8(3), 338-353.
8. Baykal, N., & Beyan, T. (2004). Bulanık mantık ilke ve temelleri. *Bıçaklar Kitabevi*.
9. Zadeh, L. A. (1965). Fuzzy sets. *Information and control*, 8(3), 338-353.
10. Khan, N., Saleem, M. R., Lee, D., Park, M. W., & Park, C. (2021). Utilizing safety rule correlation for mobile scaffolds monitoring leveraging deep convolution neural networks. *Computers in Industry*, 129, 103448.

11. Tokdemir, O. B., & Ayhan, B. U. (2019). Keskin Bir Cisim ile Temas Sonucu Yaralanma Kazalarının Analitik Hiyerarşi Prosesi ve Yapay Sinir Ağları ile Analizi. Dicle Üniversitesi Mühendislik Fakültesi Mühendislik Dergisi, 10(1), 323-334.
12. Ayanoğlu, C. C., & Kurt, M. (2019). Metal Sektöründe Veri Madenciliği Yöntemleri İle Bir İş Kazası Tahmin Modeli Önerisi. Ergonomi, 2(2), 78-87.
13. Patel, D. A., & Jha, K. N. (2015). Neural network model for the prediction of safe work behavior in construction projects. Journal of construction engineering and management, 141(1), 04014066.
14. Lee, K., & Han, S. (2021). Convolutional neural network modeling strategy for fall-related motion recognition using acceleration features of a scaffolding structure. Automation in Construction, 130, 103857
15. Azadeh, A., Rouzbahman, M., Saberi, M., & Valianpour, F. (2014). An adaptive algorithm for assessment of operators with job security and HSEE indicators. Journal of Loss Prevention in the Process Industries, 31, 26-40.
16. Jeelani, I., Asadi, K., Ramshankar, H., Han, K., & Albert, A. (2021). Real-time vision-based worker localization & hazard detection for construction. Automation in Construction, 121, 103448.
17. Dizdar, E. N., & Koçar, O. (2018). İş sağlığı ve güvenliği yönetim sistemlerinde risklerin yapay sinir ağlarıyla değerlendirilmesi. Academic Platform-Journal of Engineering and Science, 6(3), 73-83.
18. Debnath, J., Biswas, A., Sivan, P., Sen, K. N., & Sahu, S. (2016). Fuzzy inference model for assessing occupational

risks in construction sites. *International Journal of Industrial Ergonomics*, 55, 114-128.

19. Alaeddinoğlu, M. F., Sincar, S., & Naralan, A. (2015). İş Sağlığı ve Güvenliğinde Risk Analizi ve Değerlendirmesi için Geliştirilmiş Bir Karar Destek Sistemi (Yapay Sinir Ağı)-Atatürk Üniversitesi Örneği. *Çukurova Üniversitesi Mühendislik-Mimarlık Fakültesi Dergisi*, 30(2), 275-292.
20. Türker, M., & Kanıt, R. (2020). Yapi Üretim Sürecindeki İş Kazaları Şiddetinin Ön Bilgilendirilmiş Yapay Öğrenme Metodu İle Tahmini. *Konya Journal of Engineering Sciences*, 8(4), 943-956.
21. Savadkoohi, M., Oladunni, T., & Thompson, L. A. (2021). Deep neural networks for human's fall-risk prediction using force-plate time series signal. *Expert systems with applications*, 182, 115220.
- Fausett, L. V. (2006). *Fundamentals of neural networks: architectures, algorithms and applications*. Pearson Education India.

İŞBİRLİKÇİ FİLTRELEME İLE FİLM ÖNERİ SİSTEMİ

Bekir PARLAK¹

Fadimana DİKİCİ²

Esmâ AKSAN³

1. GİRİŞ

Öneri sistemleri (Sharma vd., 2013), kullanıcıların tercihleri ve ihtiyaçlarına göre çeşitli ürün, hizmet veya içerik önerileri sunan algoritmalarlardır. Günümüzde bu sistemler, e-ticaret sitelerinden sosyal medya platformlarına, müzik ve video akış servislerinden haber sitelerine kadar birçok farklı alanda yaygın olarak kullanılmaktadır. Öneri sistemlerinin temel amacı, kullanıcı deneyimini iyileştirerek kullanıcıların aradıkları veya ilgilenebilecekleri içeriklere daha hızlı ve kolay ulaşmalarını sağlamaktır. İşbirlikçi filtreleme (Schafer vd., 2007), kullanıcının geçmişteki davranışları ve diğer kullanıcıların davranışları üzerine odaklanmaktadır. Kullanıcı tabanlı işbirlikçi filtreleme, benzer ilgi alanlarına sahip kullanıcıları belirlemek için kullanıcıların geçmiş verilerini kullanır. Bu benzerliklere dayanarak, belirli bir kullanıcı tarafından beğenilen içerik, aynı ilgi alanlarını paylaşan diğer kullanıcılara önerilir. Öneri sistemleri aldıkları veriye göre farklı kategorilere ayrılmaktadır. Bu kategorilerden en başarılı olanlar içerik tabanlı filtreleme ve

¹ Doç. Dr., Amasya Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü, bekir.parlak@amasya.edu.tr, ORCID: 0000-0001-8919-6481.

² Amasya Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü, fdmndikici@gmail.com, ORCID: 0009-0006-2413-7847.

³ Amasya Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü, aksanesma812@gmail.com, ORCID: 0009-0005-7194-1574.

işbirlikçi filtrelemedir. İçerik tabanlı filtreleme (Deniz vd., 2021), verinin içeriğine göre arama yaparken, kullanıcı tabanlı filtreleme (Cingiz vd., 2021), ise kullanıcılar arasındaki benzerliğe göre değer üretir.

2. LİTERATÜR ÖZETİ

İnternet çağında, öneri sistemleri özellikle e-ticaret siteleri için vazgeçilmez bir teknoloji haline geldi. Kullanıcıların geçmişteki tercihlerini analiz ederek kişiselleştirilmiş öneriler sunan bu sistemler, hem kullanıcı deneyimini zenginleştiriyor hem de şirketlerin karlılığını arttırmaktadır. Film ve müzik platformları, öneri sistemlerinin gücünü ilk keşfedenler arasında yer aldı.

2018 yılında yapılan bir çalışmada (Reddy vd., 2019), içerik tabanlı filtreleme kullanılarak film öneri sistemleri gerçekleştirilmiştir. Film önerileri sunan sistemler, kullanıcıların daha önce keyif aldığı filmlerin özelliklerini analiz ederek, izleme zevklerine uygun yeni filmler önerir. Özellikle geniş kullanıcı kitlesine sahip ve bireylere en doğru önerileri sunarak memnuniyet sağlamayı hedefleyen platformlar için oldukça faydalıdır. Bu sistemler, kullanıcıların beğenebileceği filmleri belirlerken, filmin türü, oyuncu kadrosu ve hatta yönetmen gibi pek çok unsuru değerlendirerek isabetli öneriler sunmayı amaçlamaktadır. İçerik tabanlı filtreleme ile kullanıcının geçmişte beğendiği veya etkileşimde bulunduğu ürünlerin, hizmetlerin veya içeriklerin özelliklerini analiz ederek çalışır. Sistem, kullanıcının tercih ettiği özelliklere sahip benzer öğeleri tanımlar ve bunları gelecekte beğenebileceği varsayımıyla önerir.

2020 yılında yapılan farklı bir film öneri sistemleri çalışmasında (Gupta vd., 2020) ise öneri sistemlerinin performansını artırmak ve özellikle sonuçların isabet oranını yükseltmek amacıyla, içerik tabanlı filtrelemeye alternatif olarak

K-NN algoritmaları ve işbirlikçi filtreleme yöntemleri kullanılmıştır. Bu yaklaşım, temel olarak, kullanıcıların benzer zevklere sahip diğer kullanıcıların tercihlerinden faydalanarak film önerileri sunan işbirlikçi filtreleme mantığına dayanmaktadır. Benzerlik ölçütü olarak ise, Öklid mesafesine kıyasla daha hassas sonuçlar verdiği ve filmler arasındaki mesafeyi daha doğru yansıttığı için kosinüs benzerliği tercih edilmiştir. Bu sayede, içerik tabanlı filtrelemenin bazı dezavantajları da ortadan kaldırılmış olur. 2023 'de yayınlanan bir çalışmada (Özgöbek vd., 2015)ise kişiselleştirilmiş bir haber deneyimi sunmak amacıyla, kullanıcının internet kullanım alışkanlıklarını temel alan bir haber öneri sistemi geliştirilmiştir. Sistem, kullanıcıların ziyaret ettiği web sitelerini, arama motorlarına girdiği kelimeleri ve kaydettiği yer imlerini analiz ederek ilgi alanlarını belirler. Bu amaçla, haber kategorilerini ve içeriklerini içeren geniş bir veri seti kullanılarak bir makine öğrenmesi modeli eğitilmiştir.

2021 yılında yayınlanan bir çalışmada (Thakker vd., 2021) film öneri sistemlerinde makine öğrenmesi algoritmalarının sunduğu potansiyele odaklanmaktadır. Model tabanlı filtreleme yöntemleri kullanılarak, kullanıcıların henüz değerlendiremediği filmlere nasıl puan verilebileceği ve filmlerin kullanıcı tercihlerine göre nasıl sıralanabileceği veya ikiye ayrılabilceği incelenmektedir. Film öneri sistemlerinde makine öğrenmesi algoritmalarının sunduğu potansiyele odaklanmaktadır. Model tabanlı filtreleme yöntemleri kullanılarak, kullanıcıların henüz değerlendiremediği filmlere nasıl puan verilebileceği ve filmlerin kullanıcı tercihlerine göre nasıl sıralanabileceği veya ikiye ayrılabilceği incelenmektedir.

2022 yılında yayınlanan bir çalışmada (Budak vd., 2022) kişiselleştirilmiş ürün önerileri için literatürde ve iş dünyasında sıklıkla kullanılan yöntemlerden biri olan, kullanıcı tabanlı işbirlikçi dağıtılabılır noktalar adına k-means ile kullanıcı tabanlı

işbirlik olanağı sunanlarını birlikte kullanan hibrit bir tedavi yöntemi önerme yöntemi. Kullanıcı tabanlı işbirlikçi oluşturma ve k-means yöntemleri kullanılır.

2013 yılında yapılan bir öneri sistemi adı altında yapılan çalışmada (Odic vd., 2013) hangi bilginin kullanılacağına karar verilmesine yardımcı olacak birkaç yeni teorik kavramın yanı sıra istatistiksel testlere dayalı olarak derecelendirmelerdeki varyansın açıklanmasına hangi bağlamsal bilginin katkıda bulunduğunu tespit etmeye yönelik bir metodoloji öneriyoruz. Deney, 12 farklı bağlamsal bilgi içeren gerçek film veri seti üzerinde gerçekleştirildi. Tespit için güç analizine sahip iki istatistiksel test ve derecelendirmelerin tahmini için biraz farklı gerekçelere sahip üç bağlamsal matris çarpanlara ayırma algoritması kullandık. Sonuçlar, yöntemimiz tarafından alakalı olarak tespit edilen bağlamdaki derecelendirmelerin tahmininde ve ilgisiz olarak tespit edilen bağlamda anlamlı bir farklılık gösterdi; bu durum, güç analizinin önemine ve önerilen yöntemin faydalarına işaret etti.

Bu alanda yapılan farklı bir çalışmada daha ise (Khanal vd., 2020) makine öğrenimi tekniklerinin, algoritmaların, veri kümelerinin, değerlendirmenin, değerlemenin ve çıktının gerekli bileşenler olduğu bulunmuştur. Bu makale, araştırmanın mevcut durumu ve kalan zorluklar hakkında çok ihtiyaç duyulan bir genel bakış sunarak alana önemli bir katkı sağlamaktadır. Bu alanda yapılan farklı çalışmalara bakıldığında öneri sistemlerinde yaygın olarak işbirlikçi filtreleme kullanılmaktadır.

3. İŞBİRLİKÇİ FİLTRELEME YAPISI

İşbirlikçi filtreleme, kullanıcıların geçmiş etkileşim verilerini analiz ederek, benzer ilgi alanlarına sahip kullanıcıları gruplayan ve bu gruplama sonucunda belirli bir kullanıcıya kişiselleştirilmiş öneriler sunan bir yöntemdir. Bu yöntem,

kullanıcının daha önce beğendiği veya etkileşimde bulunduğu içeriklere dayanarak, benzer zevklere sahip diğer kullanıcıların tercihlerini göz önünde bulundurur. İşbirlikçi filtreleme, öneri sistemlerinde oldukça etkili ve yaygın bir yöntem olarak kabul edilir. Bu yöntem, üç temel aşamadan oluşur: Benzerlik hesaplaması, komşuluk seçimi ve tahminleme yöntemi.

3.1.Benzerlik Hesaplanması

Benzerlik hesaplaması yapılırken aktif kullanıcılar arasındaki yakınlığa bakılır. Bu aşamada mevcut kullanıcı ile filmi değerlendirmiş diğer kullanıcılar arasında ortak olarak değerlendirdikleri filmler üzerinden bir benzerlik analizi gerçekleştirilir. Benzerlik hesaplaması yapılırken iki ana yöntem kullanılır: Kosinüs benzerliği (Rahutomo vd., 2012) ve Pearson korelasyonu (Sedgwick vd., 2012). Bu iki benzerlik hesaplama yöntemi arasında Pearson korelasyonu daha çok tercih edilmektedir.

3.1.1.Kosinüs Benzerliği

Kosinüs benzerliği (Rahutomo vd., 2012), iki kullanıcının ortak değerlendirdiği ürünler için oluşturulan değerlendirme vektörleri arasındaki açıyı ölçerek benzerliklerini belirleyen bir yöntemdir. Bu yöntem, vektörlerin uzunluklarını değil, sadece yönlerini dikkate alır. Yani, iki kullanıcının değerlendirme ölçekleri farklı olsa bile (örneğin, biri 1-5 arası, diğeri 1-10 arası değerlendirme kullanıyor olsa da), vektörlerinin yönü benzer ise yüksek bir Kosinüs benzerliği elde edilir. Bu, kosinüs benzerliğinin değerlendirme ölçeklerinden bağımsız olarak sadece kullanıcıların tercihlerinin yönüne odaklandığını göstermektedir.

3.1.2.Pearson Korelasyon Katsayısı

Pearson korelasyonu (Sedgwick vd., 2012) işbirlikçi filtreleme de kullanıcılar arasındaki yakınlığı hesaplayan bir

ölçüdür. Kullanıcıya tahmin üretirmek için izlenilen yolda en önemli aşamalardan birisidir. Korelasyonda, bir değişkenin değeri artarken diğer değişkenin değeri ne yönde ve ne kadar oranda değişiyor ise buna göre değerlendirilir. Bu değerlendirme sonucuna göre elde edilen katsayının +1 ile -1 arasında olması gerekir. Elde edilen katsayının +1 ile -1 aralığında olması durumunda kullanıcılar arasında bir ilişki olduğu anlamına gelir. Aksi takdirde kullanıcılar arasında bir ilişki bulunmamaktadır.

$$r_{xy} = \frac{\sum(x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum(x_i - \bar{x})^2 \sum(y_i - \bar{y})^2}} \quad (1)$$

Denklem (1)'de üst kısmı kullanıcıların ortak derecelendirmeleri arasındaki uyumu, alt kısım ise kullanıcıların derecelendirme dağılımlarının değişkenliğini temsil eder. Elde edilen test sonuçlarına göre Pearson korelasyonunun kosinüs benzerliğine göre daha iyi sonuç verdiği görülmüştür. Geri kalan aşamalarda Pearson korelasyon katsayısından çıkan sonuca göre ilerlenmiştir.

3.2.Komşuluk Seçimi

Benzerlik hesaplaması adımından sonra gelen bir sonraki aşama ise komşuluk seçimidir. Komşuluk seçimi, doğru tahminleme yapabilmek için benzerlik sonuçlarını kullanarak kullanıcılar arasında bir matris oluşturur. Bu komşuluk matrisi kullanıcılar arasındaki en iyi benzerlik sonuçlarına göre oluşturulur. Bu matris oluşturulurken en yaygın yöntem olan k en yakın komşu(KNN) algoritması(Guo vd., 2003)kullanılır.KNN, kullanıcılara benzerliklerine göre en yakın komşularını bulmayı hedefleyen bir makine öğrenmesi algoritmasıdır. Kullanıcının izlediği filmler, verdiği puanlar gibi bilgilerden oluşan bir profil oluşturulur ve diğer kullanıcı profilleri ile karşılaştırılır. Bu karşılaştırma sonucunda, belirtilen bir K değeri kadar en yakın komşu bulunur. KNN algoritması, bu yakın komşuların izlediği

ve beğendiği filmleri kullanıcılara önerir. K değeri, algoritmanın kaç tane komşu kullanılacağını belirler ve yüksek bir K değeri daha genel öneriler sunarken, düşük bir K değeri kişiselleştirilmiş, ancak daha riskli öneriler sunulabilir.

3.3.Tahmin Hesaplama

Öneri sistemlerinde, işbirlikçi filtreleme kullanılarak tahminde bulunma, kullanıcıların henüz izlemediği ya da puanlamadığı filmlere olan ilgilerini tahmin etme sürecidir. Bu süreç, kullanıcıların geçmişte yaptıkları puanlamalar, tercih ettikleri filmler ve diğer etkileşimlerine dayanarak gerçekleştirilir. Özellikle, kullanıcı tabanlı işbirlikçi filtreleme tekniğinde, benzer kullanıcıların tercihleri ve beğenileri göz önünde bulundurularak tahminler yapılır. Bu yöntem, benzerlik ölçütlerine dayanarak kullanıcılar arasında bir bağlantı kurar ve bu benzerlikler üzerinden önerilerde bulunur.

Tahminleme sürecinde birden fazla yöntem kullanılmaktadır. Bilinen en önemli yöntemler basit ortalama ve ağırlıklı ortalamadır (Bulut vd., 2016). Basit ortalama ile kullanıcının her bir filme verdiği puanlar toplanır. Bu toplama, kullanıcının tüm puanlarının toplamını elde ederiz. Ardından, bu toplam değer, kullanıcının puan verdiği toplam film sayısına bölünür. Bu işlem, kullanıcının tüm filmlere verdiği puanların genel bir ortalamasını sağlar. Bu yöntem ile kullanıcının puanlarının değerleri bulunabilir fakat çalışmamızda daha hassas ve doğru sonuçlar elde edebilmek için basit ortalama yerine ağırlıklı ortalama kullanıldı. Ağırlıklı ortalama ile de kullanıcıların filmlere verdiği puanlar farklı ağırlıklarla değerlendirilir. Tahminlerin doğruluğunu artırmak için önemli bir rol oynar. Bu ortalamanın hesaplanması, benzer kullanıcıların puanlarının, kullanıcıların benzerlik derecesiyle çarpılarak yapılır. Bu sayede, benzerliği daha yüksek olan kullanıcıların puanlarını tahmin etme de daha iyi sonuçlar verir. Örneğin, iki

kullanıcı arasındaki benzerlik yüksekse, bu kullanıcının verdiği puan tahmin hesaplamasında daha ağır basar. Bu hesaplama şu şekilde formüle edilebilir:

$$P_{u,i} = \frac{\sum_{v \in N(u)} s_{u,v} \cdot r_{v,i}}{\sum_{v \in N(u)} s_{u,v}} \quad (2)$$

Denklem (2)'de $P_{u,i}$, kullanıcının (u) henüz puanlamadığı ürün veya içeriğe (i) olan tahmin edilen puanını ifade eder. $N(u)$, kullanıcının (u) benzer kullanıcılar kümesini temsil eder. $s_{u,v}$ kullanıcı (u) ile kullanıcı (v) arasındaki benzerlik skorunu, $r_{v,i}$ ise benzer kullanıcı (v)'nin ürün veya içerik (i) için verdiği puanı ifade eder. Benzer kullanıcıların verdikleri puanları ve benzerlik derecelerini dikkate alarak daha doğru sonuçlar elde edilir. Sonuç olarak, tahminlemede kullanılan ağırlıklı ortalama yöntemi, kullanıcıların benzerliklerini ve geçmiş davranışlarını dikkate alarak daha kişiselleştirilmiş öneriler sunmaktadır. Bu yaklaşım, öneri sistemlerinin daha isabetli sonuçlar sunarak kullanıcı memnuniyetini ve etkinliğini artırmasına olanak tanır.

4. KULLANILAN VERİ SETİ VE DEĞERLENDİRME ÖLÇÜTLERİ

Bu başlık altında öneri sistemleri çalışmasında kullanılan veri seti ve yapılan deneysel çalışma sonucunda elde edilen değerlerin hata metrikleri anlatılmıştır.

4.1.Kullanılan Veri Seti

Öneri sistemleri alanında çalışılırken genelde veri seti olarak MovieLens'in 2 milyon veri içeren veri seti kullanılmaktadır. Bu çalışmada ise IMDB'nin 100000 satırlık veri seti kullanılmıştır. Bu veri setinin içerisinde kullanıcı id, film id ve filmlere verilen puanlar bulunmaktadır. Puanlar 1 ile 5 aralığındadır. 610 kullanıcı ve 193609 film(ürün) mevcuttur.

100837 değerlendirmenin 70585'i eğitim, 30251'i test verisi olarak kullanılmıştır. Bu eğitim ve test verileri sonucunda yapılan tahminlerin doğruluk oranı kullanılan değerlendirme ölçütleri ile ölçülmüştür.

4.2.Değerlendirme Ölçütleri

Öneri sistemlerinde elde edilen test verilerinin tahminleme değerlerinin doğruluğu ve güvenilirliği önemlidir. Bunlar sayesinde model doğruluğunun belirlenmesi, sonuca göre model iyileştirme işlemleri gerçekleştirilebilir. Doğru tahminler, kullanıcıların öneri sistemine olan güvenini ve memnuniyetini artırır. Hata ölçütleri, modelin kullanıcı beklentilerini ne kadar karşıladığını değerlendirmek için kritik bir araçtır. Daha düşük hata oranları, kullanıcıların ihtiyaçlarına daha uygun öneriler sunulmasına olanak tanımaktadır. En yaygın kullanılan hata ölçütleri Mean Absolute Error (MAE) (Chai vd., 2014), Root Mean Square Error (RMSE) (Chai vd., 2014), Normalized Mean Absolute Error (NMAE)'dır. Bunlar modelin tahmin ettiği değerlerin gerçek değerlere ne kadar yakın olduğunu gösterir. Düşük hata değerleri, modelin yüksek doğrulukta tahminler yaptığını belirtir.

İstatistik ve tahminleme modellerinin doğruluğunu değerlendirirken kullanılan önemli ölçütlerden birisi olan MAE, gerçek değerler ile tahmini değerler arasındaki farkın mutlak değerini alarak elde ettiği değerler ile bu değerlerin ortalamasını alarak sonuç üretmektedir.

$$MAE = \frac{1}{N} \sum_{i=1}^N |R_{u,i} - R'_{u,i}| \quad (3)$$

Yukarıda gösterilen denklem (3) ile hesaplama işlemi gerçekleştirilmektedir. Burada n değeri test verilerini, $R_{u,i}$ gerçek

test verilerini $R_{u,i}$ ' model eğitimi sonucunda elde edilen tahmini test değerleri vermektedir.

Diğer ölçütlerden birisi olan RMSE ise MSE'nin kareköküdür ve tahmin edilen değerler ile gerçek değerler arasındaki ortalama karesel hatayı ifade eder. RMSE, MSE gibi büyük hatalara duyarlıdır ancak birimi orijinal veri birimiyle aynıdır, bu da yorumlamayı kolaylaştırır. Düşük RMSE, daha doğru tahminler anlamına gelir.

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (R_{u,i} - R'_{u,i})^2} \quad (4)$$

Denklem (4) ile hesaplanmaktadır MAE ve RMSE sonuçlarının değerlendirilmesi, tahminleme modellerinin performansını anlamak için kritik bir adımdır. Bu adımlar sonucunda elde edilen hesaplamalardan yola çıkarak kullanıcının izlemek için hangi filmleri tercih edebileceğini önerebiliriz.

5. DENEYSEL SONUÇLAR

Bu bölümde kullanılan tahminleme algoritmasının test verilerinin değerlendirme ölçütleri ile doğruluk oranlarının hesaplanması sonuçlarını karşılaştırmaktadır.

Tablo 1. Kullanılan Yöntemin MAE Sonucu

Tahminleme yöntemi MAE
Ağırlıklı ortalama 0.5522

Tablo 1 de değerler MAE ile hesaplanmıştır. Bu çalışmada tahminleme yöntemi olarak ağırlıklı ortalama kullanıldığından dolayı ağırlıklı ortalamanın MAE ile hesaplanmış değeri verilmiştir. Burada çıkan değer küçük olması bu çalışmada hatanın az olduğu anlamına gelmektedir. İşbirlikçi filtreleme de girdinin çok olması tahmin hesaplamasının daha iyi sonuçlar

vereceğini gösterir. Aynı zamanda girdi verisinin çok olması dışında komşuluk matrisinin de sadece benzer kullanıcılardan değil aralarında ilişki olmayan kullanıcıların da matrise alınması tahmin üretirken daha iyi sonuçlar verebileceğini gösterir.

Tablo 2. Kullanılan Yöntemin RMSE Sonucu

Tahminleme yöntemi RMSE
Ağırlıklı Ortalama 0.9465

Tablo 2 de değerler RMSE ile hesaplanmıştır. MAE’de yapılan işlemlerin aynısı RMSE’de de yapılmıştır. Bu sonuçlara göre iki değerlendirme ölçütü kullanılarak eğitim sonucunda üretilen tahminlerin doğruluk oranı hesaplanmıştır. İki değerlendirme ölçütü sonuçları tablo olarak verilmiştir. RMSE değeri tahminlerin gerçek değerlerden ne kadar sapma gösterdiğini, bu sapmanın karelerinin ortalama olarak kökünü alarak ölçer. 0.9465 RMSE değeri, ortalama olarak tahminlerin gerçek değerlerden yaklaşık olarak 0.9465 puan sapma gösterdiğini gösterir. Bu değer, öneri sisteminin tahminlerinin genel doğruluğunu ve sapma miktarını detaylandırır.

MAE ve RMSE değerleri birlikte değerlendirildiğinde, 0.5522 MAE ve 0.9465 RMSE değerleri öneri sisteminin genel doğruluğunun ve tahminlerin gerçek değerlere ne kadar yakın olduğunun iyi bir göstergesidir. Ancak, RMSE değeri genellikle daha büyük hataların etkisini daha fazla gösterirken, MAE daha dengeli bir ölçü sunar. Bu metriklerin düşük olması, öneri sisteminin performansının iyileştirilmesine olanak sağlayacak alanları belirlemek için önemlidir.

6. SONUÇLAR VE ÖNERİLER

Yapılan çalışma sonucunda kullanılan veriler ve algoritma tekniği ile elde edilen tahmini puanların belirli hata ölçütlerine bakılarak istenilen doğrulukta sonuçlar elde edilmiştir.

Çalışmamızda sadece bir yöntem kullanılmıştır. Literatürde bulunan uygulamalara bakılarak daha farklı algoritmaların ve yöntemlerin kullanılmasıyla elde edilen sonuçların karşılaştırılması ve modelin daha da iyileştirilmesi önerilir.

İşbirlikçi filtreleme yönteminde genel olarak sadece benzer kullanıcılar üzerinden komşu matrisi oluşturulup tahmin hesaplanmaktadır ama komşu matrisinin benzer kullanıcıların yanında aralarında ilişki bulunmayan kullanıcılarında eklenmesi durumu önerilir.

KAYNAKÇA

- Bozkurt, M., & Acı, Ç. İ. (2021). Öneri Algoritmalarının Film Önerme Problemi Üzerinde Karşılaştırılması: MovieLens Örneği. *Bilgisayar Bilimleri ve Teknolojileri Dergisi*, 2(2), 36-42.
- Bozkurt, M., & Acı, Ç. İ. (2021). Öneri Algoritmalarının Film Önerme Problemi Üzerinde Karşılaştırılması: MovieLens Örneği. *Bilgisayar Bilimleri ve Teknolojileri Dergisi*, 2(2), 36-42.
- Budak, H., & GUMUSTAS, E. (2022). Kişiselleştirilmiş Ürün Öneri Sistemi İçin Kullanıcı Bazlı İşbirlikçi Filtreleme Ve Kümeleme Kullanan Hibrit Bir Yaklaşım. *İstanbul Ticaret Üniversitesi Sosyal Bilimler Dergisi*, 21(43), 253-268.
- Bulut, H., & Milli, M. (2016). İşbirlikçi filtreleme için yeni tahminleme yöntemleri. *Pamukkale Üniversitesi Mühendislik Bilimleri Dergisi*, 22(2), 123-128.
- Chai, T. ve Draxler, RR (2014). Kök ortalama kare hatası (RMSE) veya ortalama mutlak hata (MAE). *Yerbilimsel model geliştirme tartışmaları*, 7 (1), 1525-1534.
- Cingiz, M. Ö., & Marangoz, K. (2021). Kullanıcı Tabanlı ve Öğe Tabanlı İşbirlikçi Filtreleme ile Kümeleme Algoritmalarının Değerlendirilmesi. *Avrupa Bilim ve Teknoloji Dergisi*, (28), 453-458.
- Deniz, E., ÖZ, V. K., Keser, S. B., Okyay, S., & Kartal, Y. (2021). İçerik tabanlı bilimsel yayın öneri sisteminde benzerlik ölçümlerinin incelenmesi. *Dicle Üniversitesi Mühendislik Fakültesi Mühendislik Dergisi*, 12(2), 221-228.
- Guo, G., Wang, H., Bell, D., Bi, Y. ve Greer, K. (2003). Sınıflandırmada KNN modeline dayalı yaklaşım. Anlamli İnternet Sistemlerine Geçiş 2003: CoopIS, DOA ve ODBASE: OTM Konfederasyon Uluslararası Konferansları, CoopIS, DOA ve ODBASE 2003, Catania, Sicilya, İtalya, 3-7 Kasım 2003. Bildiriler (s. 986-996). Springer Berlin Heidelberg.
- Gupta, M., Thakkar, A., Gupta, V., & Rathore, D. P. S. (2020, July). Movie recommender system using collaborative filtering. In *2020 international conference on electronics*

- and sustainable communication systems (ICESC)* (pp. 415-420). IEEE.
- <https://www.kaggle.com/code/colinmorris/movielens-preprocessing>
- <https://www.kaggle.com/code/kanncaa1/recommendation-systems-tutorial/input>
- Khanal, S. S., Prasad, P. W. C., Alsadoon, A., & Maag, A. (2020). A systematic review: machine learning based recommendation systems for e-learning. *Education and Information Technologies*, 25(4), 2635-2664.
- Odić, A., Tkalčič, M., Tasič, JF ve Košir, A. (2013). Bir film öneri sisteminde ilgili bağlamsal bilgilerin tahmin edilmesi ve tespit edilmesi. *Bilgisayarlarla Etkileşim Kurmak*, 25 (1), 74-90.
- Özgöbek, Ö., & Erdur, R. C. (2015). Öneri sistemleri ve bir uygulama alanı olarak haber öneri sistemleri. *Akademik Bilişim Konferansları, Eskişehir*, 31.
- Rahutomo, F., Kitasuka, T. ve Aritsugi, M. (2012, Ekim). Anlamsal kosinüs benzerliği. İleri bilim ve teknoloji ICAST 7. uluslararası öğrenci konferansında (Cilt 4, Sayı 1, s. 1). Güney Kore: Seul Üniversitesi.
- Reddy, SRS, Nalluri, S., Kuniseti, S., Ashok, S. ve Venkatesh, B. (2019). Tür korelasyonunu kullanan içerik tabanlı film öneri sistemi. *Akıllı Akıllı Hesaplama ve Uygulamalar: İkinci Uluslararası SCI 2018 Konferansı Bildirileri*, Cilt 2 (s. 391-397). Springer Singapur.
- Schafer, J. B., Frankowski, D., Herlocker, J., & Sen, S. (2007). Collaborative filtering recommender systems. In *The adaptive web: methods and strategies of web personalization* (pp. 291-324). Berlin, Heidelberg: Springer Berlin Heidelberg.
- Sedgwick, P. (2012). Pearson korelasyon katsayısı. *Bmj*, 345.
- Sharma, L. ve Gera, A. (2013). Öneri sistemi araştırması: Araştırma zorlukları. *Uluslararası Mühendislik Trendleri ve Teknolojisi Dergisi (IJETT)*, 4 (5), 1989-1992.
- Thakker, U., Patel, R., & Shah, M. (2021). A comprehensive analysis on movie recommendation system employing collaborative filtering. *Multimedia Tools and Applications*, 80(19), 28647-28672.

Yalçın, E. (2021). İşbirlikçi Filtreleme Algoritmalarının Çok-Beğenilen Ürünlere Yönelik Yanlılığı. *Bilecik Şeyh Edebali Üniversitesi Fen Bilimleri Dergisi*, 8(1), 279-291.

TÜRKÇE VE İNGİLİZCE HABER VERİLERİNİN SINIFLANDIRILMASINDA ÖZNİTELİK SEÇİMİ VE GLOBALLEŞTİRME TEKNİKLERİNİN ROLÜ

Bekir PARLAK¹

1. GİRİŞ

Web'in ve internet kullanımının hızla artması, internetteki yapılandırılmamış veri miktarının her geçen gün hızlı bir şekilde artmasına neden olmaktadır. Uluslararası Veri Şirketi'ne göre İnternet'teki yapılandırılmamış veri miktarı 2020 yılına kadar 40 zeta baytı aşacak ve bu, 2010'da İnternet'teki yapılandırılmamış veri miktarından 50 kat daha fazla olacağı anlamına gelmektedir[1]. Bu yapılandırılmamış verilerin manuel olarak sınıflandırılması neredeyse imkansızdır ve bu verilerin daha yönetilebilir ve erişilebilir hale getirilmesi için sürekli bir otomatik kategorizasyon sürecine ihtiyaç vardır. Özellikle haber portallarına bakıldığında, bazen teknoloji, spor, siyaset, sağlık gibi kategorilerle ilgili belgeler yanlış kategoride gibi görünüyor veya belgeler, diğerleri olarak adlandırılan genel bir kategoride yer alıyor. Bazen de aynı haber içeriğinden dolayı birden fazla kategoride yer alabilmektedir. Örneğin; bir spor haberi aynı zamanda ekonomi haberi olarak ta servis edilebiliyor. Bu noktada önemli bir araştırma alanı olan Metin Madenciliği[2] alanındaki yaklaşım ve yöntemlere ihtiyaç duyulmaktadır. Literatürde Akıllı Metin Analizi, Metinde Bilgi Keşfi ve Metin Veri Madenciliği olarak da bilinen Metin Madenciliğinin amacı, yapılandırılmamış

¹ Doç. Dr., Amasya Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü, bekir.parlak@amasya.edu.tr, ORCID: 0000-0001-8919-6481.

metin belgelerinden değerli ve önemli bilgi ve bilgiler çıkarmaktır. Metin Madenciliği, makine öğrenimini, hesaplamalı dilbilimi, bilgi erişimini ve istatistikleri karma olarak kullanabilen disiplinler arası bir alandır. Metin Madenciliği çalışmalarında en yaygın olarak kullanılan yöntemlerden biri, makine öğrenmesi alanında denetimli öğrenme kategorisi içerisinde yer alan Metin Kategorizasyonu/Sınıflandırma yöntemidir. Metin Sınıflandırma, önceden tanımlanmış bir veri kümesinden yararlanarak bir model oluşturur ve kategorize edilmemiş verileri doğru bir kategoriye atamayı amaçlar[3]. Başka bir deyişle, kategorize edilmemiş verileri içeriğine göre değerlendirir ve kategorilere ayırır. Metin sınıflandırması birçok araştırmacı tarafından incelenmiş ve spam filtreleme[4, 5], yazar tanımlama[6], web sayfası sınıflandırması[7], tıbbi doküman sınıflandırması[8], haber verilerinin sınıflandırılması[35] ve duygu sınıflandırması[9] gibi çeşitli alanlara uygulanmıştır.

Metin sınıflandırma çalışmaları da öznitelik çıkarma ve sınıflandırma adımı içerir. Ancak bu iki adımın yanında çoğunlukla bir öznitelik seçimi ve/veya bir öznitelik dönüştürme adımı bulunmaktadır. Öznitelik çıkarma adımı genellikle yaygın olarak bilinen kelime çantası (BoW=bag-of-words) yaklaşımı kullanılmaktadır. Daha sonra özniteliklerin sırası dikkate alınmaz ve metin dokümanları bu özniteliklerin ağırlıklı frekansları ile temsil edilir. Küçük koleksiyonlar bile binlerce benzersiz kelimeden oluşabileceğinden, metin sınıflandırması yüksek boyutlu bir sınıflandırma problemi olarak bilinir. Bu nedenle, araştırmacılar performansı artırmak ve çalışma süresini azaltmak için hala boyut küçültmenin etkili yollarını bulmaya çalışıyorlar. Bazı yeni öznitelik seçim yöntemleri[10-12] araştırmacılar tarafından halen önerilmekle birlikte, farklı globalleştirme teknikleriyle[13] lokal öznitelik seçim metotlarını global hale getirmek, mevcut öznitelik seçme yöntemlerini öznitelik dönüştürme yöntemleriyle[14] veya diğer öznitelik

seçme yöntemlerini farklı şekillerde birleştiren iki aşamalı öznitelik seçme yöntemleri[15] de önerilmektedir.

Bu çalışmada kullanılan her iki veri kümesi de dengeli yani sınıflardaki doküman sayısı eşittir. Bu sebeple globalleştirme tekniği olarak TOPLAM ve MAKSİMUM olarak iki farklı teknik kullanıldı. Diğer bir globalleştirme tekniği olan AĞIRLIKLİ TOPLAM, dengeli veri kümelerinde TOPLAM tekniğiyle aynı sonuçları verdiğinden kullanılmamıştır. Ayrıca bu globalleştirme teknikleri dört farklı lokal öznitelik seçim metotlarına uygulanmıştır. Sınıflandırıcı olarak da, metin sınıflandırma alanında başarılı üç farklı örüntü sınıflandırıcı kullanılmıştır.

Bu çalışmada, bu alanda kıyaslama için kullanılan iki kamuya açık Türkçe ve İngilizce haber veri kümesi olan TTC-3600[1] ve 20Newsgroup[16] veri kümelerini sınıflandırdık. Her iki veri kümesinin sonuçlarını karşılaştırarak, sunulan globalleştirme tekniklerinin Türkçe ve İngilizce haber metinlerini sınıflandırmada etkisi incelenmiştir. Böylece hem Türkçe hem de İngilizce haber verileriyle deneysel çalışmalar yaparak bu alanda literatüre katkı sağlanmıştır.

2. İLGİLİ ÇALIŞMALAR

Literatürde Türkçe haber veri kümeleri ile ilgili bir takım çalışmalar bulunmaktadır. Ancak, İngilizce ve diğer birkaç dillerde haber veri kümeleri ile ilgili daha çok çalışma bulunmaktadır. Bunun nedeni Türkçe dilinde ciddi sayıda veri kümelerinin olmamasından kaynaklanmaktadır. Türkçe haber verilerinin sınıflandırılması üzerine bir araştırmada, Kılınç ve arkadaşları[1], Türkçe haber ve makale içeriğinin metin sınıflandırma araştırmasında yaygın olarak kullanılabilecek TTC-3600 adlı yeni bir veri kümesi oluşturdular. TTC-3600'de metin sınıflandırma alanında beş başarılı sınıflandırıcı ve iki öznitelik seçme yöntemi değerlendirilmiştir. Deneysel sonuçlar, Rastgele

Orman sınıflandırıcı ve özellik sıralamasına dayalı öznitelik seçim yönteminin kombinasyonu ile ön işleme ve öznitelik seçimi adımlarından sonra yapılan tüm karşılaştırmalarda en iyi doğruluk değerinin %91.03'e ulaştığını göstermektedir. Başka bir çalışmada[17], TTC-3600 veri kümesi üzerinde Evrişimli Sinir Ağları (CNN=Convolutional Neural Network) ve Word2Vec yöntemi kullanılarak bir metin sınıflandırma çalışması gerçekleştirilmiştir. CNN ve Word2Vec yöntemi, klasik istatistiksel ve makine öğrenimi tabanlı sınıflandırma algoritmalarından daha iyi performans (%93,3 doğruluk) gösterdi. Köksal[18], TTC-4900 veri kümesi ile deneyler yaptı. Bu veri kümesi, TTC-3600 ile karşılaştırılabilir. Sonuç olarak, orijinal veri kümesi yüzde 90'lık bir F1 skoru aldı, ancak verileri kök ayırma kullanmadan düzelttikten sonra F1 skoru yüzde 91.77'ye yükseldi. Yıldırım ve Yıldız[19] Türkçe metin sınıflandırması için geleneksel BoW yaklaşımı ile Sinir Ağı tabanlı yaklaşımları karşılaştırmak için TTC-4900 ve TTC-3600 veri kümelerini kullanmışlardır. Hem TTC-3600 hem de TTC-4900 veri kümeleri üzerinde yaptıkları deneylerin, BoW yaklaşımını ve çok terimli naive Bayes (ÇTNB) kullanarak TTC-3600'de %93.1 F1 skoru ve TTC-4900 veri kümelerinde %90.0 F1 skoruna ulaşarak daha iyi sonuçlar ortaya koyduğunu iddia ettiler. Kılınç[20], topluluk öğrenme modellerinin Türkçe metin sınıflandırmasına etkisini araştırmıştır. Artırmalı KA sınıflandırıcı ile elde edilen en iyi değer %85,52'tir. Bu çalışma, topluluk öğrenme modellerinin, Türkçe metin sınıflandırmasında temel sınıflandırıcıların performansını artırarak genel olarak daha doğru sonuçlar verdiğini ortaya koymuştur.

İngilizce haber verilerinin sınıflandırılmasıyla ilgili bir araştırma, Dadgar ve arkadaşları[2], kullanıcıların herhangi bir zamanda istedikleri ülkedeki en popüler haber grubunu belirleyebilmeleri için haberleri çeşitli gruplara ayırmayı amaçlar. Haryanto ve diğerleri[21], DVM ve Ki-kare öznitelik seçimi

kullanılarak metin sınıflandırmasının, kök çıkarma ve lemmatizasyon dahil olmak üzere ön işleme veri sonuçları üzerinde etkinliğini artırabildiğini gösterdiler. Bir başka çalışmada[22], öznitelik terimlerini oluşturmak için yeni bir iki-gram alfabe yaklaşımı kullanılmıştır. Önerilen yöntem metin sınıflandırma alanına iki önemli katkı sağlamıştır. İlk olarak, geleneksel alfabeye dayalı sabit özniteliklerin doküman sözlüklerine gerek kalmadan kullanılabilirliği gösterilmiştir; bu, büyük veri kümeleri için vektör uzayı boyutlarını önemli ölçüde azaltır. İkincisi ise, doğal dil işleme araçlarına gerek kalmadığını göstermektedir. Mevcut çalışma, Arapça veya İngilizce metin belgelerinin koleksiyonlarını başarıyla sınıflandırmıştır. Arapça veri kümesinde kaydedilen en iyi sonuçlarla karşılaştırıldığında, vektör alanından yaklaşık %80 tasarruf etti ve performansı %2 artırdı. Bir diğer çalışmada[23] ise, araştırmacılar, haber makalelerinde cümle düzeyindeki olumsuzlukları belirlemeye odaklandılar. Destek Vektör Makinesi yüzde 96,46 doğruluğa sahipken, Naive Bayes yüzde 94,16 doğruluğa sahip olduğu görülmüştür.

3. ÖZNİTELİK SEÇİM METOTLARI

Metin sınıflandırma görevlerinde öznitelik seçimi, metin sınıflandırma hatasını optimize etmek için ilgili metin özniteliklerinin minimum boyutunu arayan bir süreç olarak tanımlanır. Öznitelik seçim yöntemleri genellikle filtreler, sarmalayıcılar ve gömülü yöntemler olarak sınıflandırılır[24].

Sarmalayıcı yöntemleri, bir değerlendirme fonksiyonu ve bir arama stratejisi kullanarak alt kümeleri seçer. Arama stratejisi ne olursa olsun, öznitelik sayısı arttıkça olası sonuçların sayısı geometrik olarak artacaktır. Gömülü yöntemler süreci, sınıflandırıcıların eğitiminde gömülüdür[25]. Sarmalayıcılar ve gömülü yöntemler, filtrelerden çok daha fazla hesaplama süresi

gerektirir ve yalnızca belirli bir sınıflandırıcı ile çalışabilir. Bu nedenle filtreler, metni sınıflandırmak için en yaygın öznitelik seçme yöntemidir.

Filtreler global veya lokal şeklinde iki grupta incelenir[13]. Global yöntemler, sınıf sayısına bakılmaksızın her özniteliğe tek bir skor atarken, lokal yöntemler ilgili öznitelik için her sınıfa özgü bir skoru olduğu için sınıf sayısı kadar skor atar[26]. Global yöntemler tipik olarak her öznitelik için skoru hesaplar ve daha sonra N'nin genellikle deneysel olarak belirlendiği öznitelik kümesi olarak ilk N özniteliği seçer. Lokal yöntemler benzerdir ancak ilk N öznitelikleri seçmeden önce bir özniteliğin birden çok skorunu tek bir skora dönüştürmeyi gerektirir. Bunun için globalleştirme teknikleri kullanılır. Yaygın olarak kullanılan bazı global öznitelik seçim yöntemleri, doküman frekansı (DF=Document Frequency), bilgi kazancı (IG=Information Gain), Gini indeksi (GI=Gini Index), kapsamlı öznitelik seçici (EFS=Extensive Feature Selector) ve ayırt edici öznitelik seçici (DFS=Distinguishing Feature Selector) [18]. Lokal yöntemler, Ki-kare (CHI=Chi-Square) , olasılık oranı (OR=Odds Ratio), ayırmacı öznitelik seçici (DFSS = Discriminative Feature Selection) ve ağırlıklı log-olasılık oranı (WLLR=Weighted Log-Likelihood Ratio) algoritmalarını içerir.

Tablo 1. Öznitelik Seçim Yöntemleri İçin Ön-Gösterimler

Gösterim	Anlamı
a	C_j sınıfında t terimini içeren doküman sayısı
b	C_j sınıfı dışında t terimini içeren doküman sayısı
c	C_j sınıfında t terimini içermeyen doküman sayısı
d	C_j sınıfı dışında t terimini içermeyen doküman sayısı
e	C_j sınıfında t teriminin frekansı
f	C_j sınıfı dışında t teriminin frekansı
N	Bütün sınıflardaki toplam doküman sayısı
M	Toplam sınıf sayısı
$f(t, c_j)$	t teriminin C_j sınıfı için skoru
$p(t, C_j)$	t teriminin C_j sınıfında olma olasılığı

$p(t, \bar{C}_j)$	t teriminin C_j sınıfı dışında olma olasılığı
$p(\bar{t}, C_j)$	t teriminin C_j sınıfında olmama olasılığı
$p(\bar{t}, \bar{C}_j)$	t teriminin C_j sınıfı dışında olmama olasılığı
$p(t C_j)$	C_j sınıfı mevcut olduğunda t teriminin olma olasılığı
$p(\bar{t} C_j)$	C_j sınıfı mevcut olduğunda t teriminin olmama olasılığı
$p(t \bar{C}_j)$	C_j sınıfı mevcut olmadığıda t teriminin olma olasılığı
$p(\bar{t} \bar{C}_j)$	C_j sınıfı mevcut olmadığıda t teriminin olmama olasılığı
$p(C_j t)$	t terimi mevcut olduğunda C_j sınıfının olma olasılığı
$p(\bar{C}_j t)$	t terimi mevcut olduğunda C_j sınıfının olmama olasılığı
$p(C_j \bar{t})$	t terimi mevcut olmadığıda C_j sınıfının olma olasılığı
$p(\bar{C}_j \bar{t})$	t terimi mevcut olmadığıda C_j sınıfının olmama olasılığı

Çalışmamızda, dört farklı lokal öznitelik seçim metodunu kullandık. Bu metotlar; CHI2, OR, DFSS ve WLLR. Bu metotları global hale getirmek için ise TOPLAM ve MAKSİMUM olmak üzere iki farklı teknik kullandık. TOPLAM ve MAKSİMUM tekniği şu şekilde hesaplanmaktadır:

$$TOPLAM = \sum_{j=1}^M f(t, c_j) \quad (1)$$

$$MAKSİMUM = \max_{j=1}^M f(t, c_j) \quad (2)$$

3.1.Chi-Squared (CHI2)

CHI2, t teriminin C_i sınıfı ile korelasyonunu hesaplayan denetimli, tek taraflı bir öznitelik seçim yöntemidir[21]. CHI2 şu şekilde hesaplanır:

$$CHI2(t, C_j) = \frac{N \cdot [p(t, C_j) \cdot p(\bar{t}, \bar{C}_j) - p(t, \bar{C}_j) \cdot p(\bar{t}, C_j)]^2}{p(t) \cdot p(\bar{t}) \cdot p(C_j) \cdot p(\bar{C}_j)} \quad (3)$$

3.2.ODDS Ratio (OR)

OR, sırasıyla payı ve paydası ile belirli bir sınıfa üyelik ve üye olmama hesaplanarak elde edilen denetimli ve tek taraflı bir yöntemdir[8]. Şu şekilde hesaplanır:

$$OR(t, C_j) = \left| \log \frac{p(t|C_j) \cdot (1 - p(t|\bar{C}_j))}{p(t|\bar{C}_j) \cdot (1 - p(t|C_j))} \right| \quad (4)$$

3.3.WLLR

WLLR[27], Nigam ve diğerleri tarafından önerilmiştir. Bu denetimli ve tek taraflı yöntem ve şu şekilde hesaplanır:

$$WLLR(t, c_j) = P(t|c_j) \log \frac{P(t|c_j)}{P(t|\bar{c}_j)} \quad (5)$$

3.4.Discriminative Feature Selection (DFSS)

DFSS, Zong ve diğerleri tarafından önerilmiştir[28]. Bu yöntem ise şu şekilde hesaplanmaktadır:

$$DFSS(t, c_j) = \left(\frac{e/a}{f/b}\right) \cdot P(C_j|t) \cdot P(t|C_j) \cdot |P(C_j|t) - P(C_j|\bar{t})| \quad (6)$$

4. MATERYA VE METOTLAR

4.1.Verİ Kümesi

Türkçe haber sınıflandırma çalışmalarında yaygın olarak kullanılmak üzere hazırlanan TTC-3600 veri kümesi Kılınç ve diğerleri tarafından 2015 yılında derlenmiştir[1]. Deneylerde bu veri kümesinin eğitim ve test bölümleri %50-%50 olacak şekilde ayarlandı.

20Newsgroup veri kümesi, Joachims[29] tarafından toplanan metin sınıflandırması için klasik çok sınıflı İngilizce veri kümesidir. Her sınıfta 1000'e yakın örnek vardır. Ancak deneylerde sadece 10 sınıfa ait toplam 10000 doküman üzerinden deneysel işlemler yapıldı. Deneylerde bu veri kümesinin eğitim ve test bölümlerini benzer şekilde %50-50 olacak şekilde ayarlandı.

TTC-3600 ve 20Newsgroup veri kümelerinin kategorilere göre doküman sayıları Tablo 2-3'de gösterilmiştir.

Tablo 2. TTC-3600 Veri Kümesi

Kategoriler	Toplam	Eğitim	Test
Kültür	600	300	300
Ekonomi	600	300	300
Sağlık	600	300	300
Politika	600	300	300
Spor	600	300	300
Teknoloji	600	300	300
Toplam	3600	1800	1800

Tablo 3. 20Newsgroup Veri Kümesi

Kategoriler	Toplam	Eğitim	Test
alt.atheism	1000	500	500
comp.graphics	1000	500	500
comp.os.ms-windows.misc	1000	500	500
comp.sys.ibm.pc.hardware	1000	500	500
comp.sys.mac.hardware	1000	500	500
comp.windows.x	1000	500	500
misc.forsale	1000	500	500
rec.autos	1000	500	500
rec.motorcycles	1000	500	500
rec.sport.baseball	1000	500	500
Toplam	10000	5000	5000

4.2. Veri Ön-İşleme

Veri ön işleme, performansı etkileyen en önemli faktörlerden biridir. Bu nedenle, TTC-3600 ve 20Newsgroup veri kümeleri vektör hale gelmeden önce bazı metin ön işleme adımları uygulandı. Gereksiz ya da durak sözcükleri her sınıfta oldukça fazla bulunur. Bu kelimeler ayırt edici bir özelliği olmadığı ve başarı oranını olumsuz etkilediği için veri kümesinden çıkarılmıştır. Ayrıca tüm kelimeler küçük harfe dönüştürülmüş, harfler dışında sayı, sembol ve noktalama işaretleri gibi tüm karakterler temizlenmiştir. Bir diğer önemli ön işleme adımı olan kök bulma için Türkçe veri kümesinde Zemberek [28] kütüphanesi kullanılmıştır. İngilizce veri kümesinde, kök bulma için Java ile yazılmış Porter kütüphanesi

kullanılmıştır. Bunun için öncelikle orijinal veri kümesindeki terimler kök haline dönüştürüldü. Bu ön işlemler uygulandıktan sonra, öznitelik vektörleri oluşturuldu. Böylece, metin dokümanları vektör formuna dönüştürülerek sınıflandırma işlemine hazır hale geldi.

4.3.Sınıflandırma Algoritmaları

Bu çalışmada üç farklı sınıflandırıcı kullanılmıştır. Bu sınıflandırıcılar Çok Terimli Naive Bayes (ÇTNB), Karar Ağacı (KA) ve Destek Vektör Makinesi (DVM)'dir.

ÇTNB, metin sınıflandırma çalışmalarında en verimli sınıflandırıcılardan biridir. Ayrıca ÇTNB[30], Naive Bayes sınıflandırıcısının bir çeşididir. Geleneksel Naive Bayes bir metin ile ilgili öznitelikleri dahil ederek modellerken, ÇTNB onu öznitelik sayılarını kullanarak açıkça modeller. Çok değişkenli Bernoulli olay modeli doküman frekanslarını kullandığından, ÇTNB terim frekanslarını dikkate alır. Sınıf kümesi C ve öznitelik boyutu ise N ile temsil edilmektedir.

KA, dahili düğümlerin özniteliklere karşılık geldiği ve her yaprak düğümün bir sınıf etiketine karşılık geldiği bir ağaç temsili kullanır. KA, en yorumlanabilir algoritmalarından biridir. En iyi öznitelik alt kümesi seçmek için kullanılabilir.

DVM, metin sınıflandırma çalışmalarında en verimli sınıflandırıcılardan biridir. DVM sınıflandırıcısının doğrusal ve doğrusal olmayan iki versiyonu vardır. DVM sınıflandırıcısının odak noktası marj kavramıdır[31].

Doğrusal çekirdek içeren DVM için LibSVM kütüphanesi kullanılır. Doğrusal çekirdek esas olarak doğrusal sınıflandırma için kullanılır. Birkaç parametresi olduğu için hızlı hesaplar.

4.4.Deneysel Sonuçlar

Bu çalışmada, dört lokal öznitelik seçim yöntemi, Türkçe ve İngilizce dillerinde iki farklı veri kümesi ve üç adet başarılı

sınıflandırıcı kullanılarak kapsamlı bir analiz yapılmıştır. Ayrıca lokal öznitelik seçim yöntemlerini global hale getirmek için iki farklı globalleştirme tekniği kullanıldı. Öznitelik seçimi aşamasında lokal öznitelik seçim yöntemlerine örnek olan CHI2, OR, WLLR ve DFSS öznitelik seçim yöntemleri kullanılmıştır. Ancak bu çalışmada ÇTNB, KA ve DVM olmak üzere üç farklı başarılı sınıflandırıcı kullanılmıştır. Globalleştirme teknikleri[13] olarak TOPLAM ve MAKSİMUM teknikleri kullanılmıştır. Türkçe veri kümesi için Zemberek algoritması[32] uygulanırken, İngilizce veri kümesi için Porter algoritması[33] uygulandı. Şekil 2’de çalışmayı özetleyen akış diyagramı çizilmiştir. Gereksiz kelimeler, belirlenen gereksiz kelime listesine göre kaldırıldı ve öznitelik ağırlıklandırma aşaması için TF-IDF kullanıldı. Ayrıca, ön işleme adımı olarak küçük harf dönüşümü ve alfabetik olmayan karakterlerin kaldırılması kullanıldı[34]. Globalleştirme tekniklerinin performansını üç farklı sınıflandırıcı üzerinde değerlendirmek için F1 skoru kullanılmıştır.

Dört farklı lokal öznitelik seçim yöntemi ile seçilen farklı öznitelik vektörleri ÇTNB, KA ve DVM sınıflandırıcıları ile beslenmiştir. Öznitelik boyutları, 50, 100, 300, 500 ve 1000 olmak üzere beş farklı boyutta oluşturulmuştur. 19605 ve 39912 sırasıyla TTC-3600 ve 20Newsgroup veri kümeleri için toplam öznitelik sayısıdır. F1 skorlarının sonuçları, TTC-3600 ve 20Newsgroup veri kümesi için Tablo 4-9’da gösterilmektedir.

Tablo 4. TTC-3600 Veri Kümesi İçin ÇTNB Sınıflandırıcı Kullanarak F1 Skorları

Öznitelik Seçim Metotları	50	100	300	500	1000
	TOPLAM				
CHI2	0.636	0.717	0.807	0.824	0.849
OR	0.631	0.711	0.783	0.813	0.841
DFSS	0.700	0.803	0.840	0.848	0.851
WLLR	0.362	0.465	0.697	0.767	0.822
	MAKSİMUM				
CHI2	0.623	0.720	0.812	0.830	0.846
OR	0.445	0.580	0.756	0.787	0.822
DFSS	0.700	0.803	0.838	0.846	0.858
WLLR	0.563	0.643	0.748	0.792	0.827

**Tablo 5. TTC-3600 Veri Kümesi İçin KA Sınıflandırıcı
Kullanarak F1 Skorları**

Öznitelik Seçim Metotları	50	100	300	500	1000
	TOPLAM				
CHI2	0.616	0.667	0.700	0.713	0.760
OR	0.625	0.659	0.707	0.713	0.688
DFSS	0.647	0.723	0.712	0.713	0.687
WLLR	0.310	0.370	0.579	0.650	0.704
	MAKSİMUM				
CHI2	0.610	0.683	0.701	0.713	0.717
OR	0.396	0.492	0.639	0.688	0.706
DFSS	0.647	0.723	0.709	0.715	0.691
WLLR	0.495	0.574	0.650	0.707	0.699

**Tablo 6. TTC-3600 Veri Kümesi İçin DVM Sınıflandırıcı
Kullanarak F1 Skorları**

Öznitelik Seçim Metotları	50	100	300	500	1000
	TOPLAM				
CHI2	0.627	0.704	0.739	0.758	0.755
OR	0.632	0.702	0.747	0.729	0.738
DFSS	0.677	0.757	0.746	0.739	0.761
WLLR	0.366	0.456	0.648	0.704	0.732
	MAKSİMUM				
CHI2	0.617	0.702	0.740	0.753	0.748
OR	0.448	0.543	0.706	0.738	0.738
DFSS	0.677	0.757	0.739	0.752	0.760
WLLR	0.541	0.616	0.694	0.719	0.743

TTC-3600 veri kümesi için DFSS öznitelik seçim yöntemi, MAKSİMUM globalizasyon tekniği, 1000 öznitelik kullanan ÇTNB sınıflandırıcı kombinasyonundan elde edilen en yüksek F1 skoru sırasıyla 0.858'dir. ÇTNB sınıflandırıcı çoğu durumda KA ve DVM sınıflandırıcısına göre daha başarılıdır. Ayrıca ÇTNB sınıflandırıcıda skorlar daha yüksek boyutlarda elde edilirken, KA ve DVM sınıflandırıcılarında daha düşük boyutlarda da elde edilmiştir. Öznitelik seçim yöntemlerinde, genel olarak DFSS daha yüksek performans göstermiştir. Globalizasyon tekniği olarak ise, TOPLAM tekniği

MAKSİMUM tekniğine göre daha yüksek performans göstermiştir.

Tablo 7. 20Newsgroup Veri Kümesi İçin ÇTNB Sınıflandırıcı Kullanarak F1 Skorları

Öznitelik Seçim Metotları	50	100	300	500	1000
	TOPLAM				
CHI2	0.940	0.905	0.903	0.898	0.887
OR	0.919	0.885	0.876	0.861	0.848
DFSS	0.894	0.933	0.907	0.909	0.888
WLLR	0.523	0.934	0.877	0.861	0.854
	MAKSİMUM				
CHI2	0.938	0.903	0.909	0.907	0.894
OR	0.951	0.946	0.922	0.917	0.905
DFSS	0.890	0.932	0.909	0.910	0.887
WLLR	0.950	0.936	0.885	0.875	0.866

Tablo 8. 20Newsgroup Veri Kümesi İçin KA Sınıflandırıcı Kullanarak F1 Skorları

Öznitelik Seçim Metotları	50	100	300	500	1000
	TOPLAM				
CHI2	0.978	0.978	0.973	0.972	0.971
OR	0.976	0.979	0.976	0.970	0.971
DFSS	0.976	0.979	0.977	0.976	0.972
WLLR	0.642	0.945	0.944	0.940	0.942
	MAKSİMUM				
CHI2	0.979	0.977	0.977	0.976	0.974
OR	0.977	0.978	0.974	0.974	0.973
DFSS	0.972	0.979	0.977	0.980	0.973
WLLR	0.973	0.945	0.946	0.946	0.969

Tablo 9. 20Newsgroup Veri Kümesi İçin DVM Sınıflandırıcı Kullanarak F1 Skorları

Öznitelik Seçim Metotları	50	100	300	500	1000
	TOPLAM				
CHI2	0.969	0.967	0.967	0.963	0.961
OR	0.969	0.969	0.966	0.957	0.956
DFSS	0.966	0.968	0.968	0.966	0.962
WLLR	0.694	0.971	0.955	0.954	0.956
	MAKSİMUM				
CHI2	0.965	0.967	0.966	0.970	0.960
OR	0.969	0.969	0.973	0.971	0.968
DFSS	0.963	0.966	0.969	0.964	0.961
WLLR	0.969	0.969	0.955	0.960	0.959

20Newsgroup veri kümesi için DFSS öznitelik seçim yöntemi, MAKSİMUM globalizasyon tekniği, 500 öznitelik kullanan KA sınıflandırıcı kombinasyonundan elde edilen en yüksek F1 skoru sırasıyla 0.980'dir. KA sınıflandırıcı çoğu durumda ÇTNB ve DVM sınıflandırıcısına göre daha başarılıdır. Ayrıca tüm sınıflandırıcılarda en yüksek performans daha düşük boyutlarda elde edilmiştir. Öznitelik seçim yöntemlerinde, genel olarak OR ve DFSS daha yüksek performans göstermiştir. Globalizasyon tekniği olarak ise, MAKSİMUM tekniği TOPLAM tekniğine göre daha yüksek performans göstermiştir.

5. TARTIŞMA

Çalışmada, Türkçe ve İngilizce veri kümelerinin otomatik sınıflandırılması için dört farklı lokal öznitelik seçme yönteminin performansa etkisi analiz edilmiştir. Deneysel çalışmalara baktığımızda herhangi bir öznitelik seçim yönteminin her boyutta en yüksek performansı göstermediğini görmekteyiz. Çünkü boyut ve globalleştirme tekniği gibi parametreler de performansa etki etmektedir. Ancak en yüksek performansı her iki veri kümesi için DFSS metodu farklı sınıflandırıcı ve farklı boyutlarda elde etmiştir. Benzer şekilde, sınıflandırıcılar için de aynı durum

geçerlidir. Çünkü, TTC-3600 veri kümesi için en yüksek skor ÇTNB sınıflandırıcı ile elde edilirken, 20Newsgroup veri kümesi için KA sınıflandırıcı en yüksek skoru elde etmektedir. Ek olarak, öznelik boyutlarının performansa etkisi veri kümesine göre değişmektedir. Örneğin; TTC-3600 veri kümesi için ÇTNB sınıflandırıcıda skorlar daha yüksek boyutlarda elde edilirken, KA ve DVM sınıflandırıcılarında daha düşük boyutlarda da elde edilmiştir. Ancak, 20Newsgroup veri kümesinde ise tüm sınıflandırıcılarda en yüksek performans daha düşük boyutlarda elde edilmiştir. Son olarak, MAKSİMUM tekniği her iki veri kümesi için en yüksek skoru gösterse de boyutlar bazında bu durum değişmektedir.

Sonuç olarak, en yüksek performansı elde etmek için belirli bir öznelik seçim metodu, globalleştirme tekniği ya da belirli bir sınıflandırıcı yoktur. En yüksek skoru gösteren şema, veri kümesine göre değişiklik göstermektedir.

6. SONUÇLAR

Çalışmada, Türkçe ve İngilizce veri kümelerinin otomatik sınıflandırılması için dört farklı lokal öznelik seçim yönteminin etkinliği analiz edilmiş ve en başarılı sınıflandırma şemaları belirlenmiştir. Öznelik seçme yöntemlerini ve bu yöntemleri global hale getirirken kullanılan iki farklı globalleştirme tekniklerinin performanslarını karşılaştırmak için üç farklı sınıflandırma algoritması kullanılmıştır. Deneysel sonuçlar, ÇTNB sınıflandırıcısının TTC-3600 veri kümesinde en yüksek performansı gösterirken, KA sınıflandırıcısı ise 20Newsgroup veri kümesinde en yüksek performansı göstermiştir. Genel olarak, TOPLAM globalleştirme tekniği TTC-3600 veri kümesi için daha yüksek performans gösterirken, MAKSİMUM globalleştirme tekniği 20Newsgroup veri kümesi için daha yüksek performans

sergilemiştir. Ancak en yüksek skorlar her iki veri kümesi için de MAKSİMUM tekniğiyle elde edilmiştir.

Gelecekteki bir çalışma olarak, metin sınıflandırma alanı için yeni bir globalleştirme tekniği oluşturmak ve bu tekniği farklı lokal öznitelik seçim metotları üzerine uygulamak ve performans analizi yapmak olacaktır.

KAYNAKÇA

- [1] Kılınç, D., Özçift, A., Bozyigit, F., Yıldırım, P., Yücalar, F., and Borandag, E. TTC-3600: A new benchmark dataset for Turkish text categorization. *Journal of Information Science*, 2017. 43(2): p. 174-185.
- [2] Dadgar, S.M.H., M.S. Araghi, and M.M. Farahani. A novel text mining approach based on TF-IDF and Support Vector Machine for news classification. in *2016 IEEE International Conference on Engineering and Technology (ICETECH)*. 2016. IEEE.
- [3] Tang, B., He, H., Baggenstoss, P. M., and Kay, S. A Bayesian classification approach using class-specific features for text categorization. *IEEE Transactions on Knowledge and Data Engineering*, 2016. 28(6): p. 1602-1606.
- [4] Subasi, A., Alzahrani, S., Aljuhani, A., and Aljedani, M. Comparison of decision tree algorithms for spam e-mail filtering. in *2018 1st International Conference on Computer Applications & Information Security (ICCAIS)*. 2018. IEEE.
- [5] Kawade, K.O. and K.S. Oza, Content-based SMS spam filtering using machine learning technique. *International Journal of Computer Engineering and Applications*, 2018. 7: p. 4.
- [6] Gupta, S.T., J.K. Sahoo, and R.K. Roul. Authorship Identification using Recurrent Neural Networks. in *Proceedings of the 2019 3rd International Conference on Information System and Data Mining*. 2019.
- [7] Hashemi, M., Web page classification: a survey of perspectives, gaps, and future directions. *Multimedia Tools and Applications*, 2020: p. 1-25.

- [8] Parlak, B. and A.K. Uysal, On classification of abstracts obtained from medical journals. *Journal of Information Science*, 2020. 46(5): p. 648-663.
- [9] Khan, J., A. Alam, and Y. Lee, Intelligent Hybrid Feature Selection for Textual Sentiment Classification. *IEEE Access*, 2021.
- [10] Parlak, B., Class-index corpus-index measure: A novel feature selection method for imbalanced text data. *Concurrency and Computation: Practice and Experience*, 2022: p. e7140.
- [11] Parlak, B. and A.K. Uysal, A novel filter feature selection method for text classification: Extensive Feature Selector. *Journal of Information Science*, 2021: p. 0165551521991037.
- [12] Rehman, A., et al., Selection of the most relevant terms based on a max-min ratio metric for text classification. *Expert Systems with Applications*, 2018. 114: p. 78-96.
- [13] Parlak, B. and A.K. Uysal, The effects of globalisation techniques on feature selection for text classification. *Journal of Information Science*, 2020: p. 0165551520930897.
- [14] Günel, S., Hybrid feature selection for text classification. *Turkish Journal of Electrical Engineering and Computer Science*, 2012. 20(Sup. 2): p. 1296-1311.
- [15] Uysal, A.K., On two-stage feature selection methods for text classification. *IEEE Access*, 2018. 6: p. 43233-43251.
- [16] Asuncion, A. and D. Newman, UCI machine learning repository. 2007.
- [17] Çiğdem, A. and A. Çırak, Türkçe haber metinlerinin konvolüsyonel sinir ağları ve Word2Vec kullanılarak

- sınıflandırılması. Bilişim Teknolojileri Dergisi, 2019. 12(3): p. 219-228.
- [18] Köksal, Ö. Tuning the Turkish Text Classification Process Using Supervised Machine Learning-based Algorithms. in 2020 International Conference on INnovations in Intelligent SysTems and Applications (INISTA). 2020. IEEE.
- [19] Yıldırım, S. and T. Yıldız, Türkçe için karşılaştırmalı metin sınıflandırma analizi. Pamukkale Üniversitesi Mühendislik Bilimleri Dergisi, 2018. 24(5): p. 879-886.
- [20] Kılınç, D., The effect of ensemble learning models on Turkish text classification. Celal Bayar University Journal of Science, 2016. 12(2).
- [21] Haryanto, A.W. and E.K. Mawardi. Influence of word normalization and chi-squared feature selection on support vector machine (svm) text classification. in 2018 International Seminar on Application for Technology of Information and Communication. 2018. IEEE.
- [22] Elghannam, F., Text representation and classification based on bi-gram alphabet. Journal of King Saud University-Computer and Information Sciences, 2021. 33(2): p. 235-242.
- [23] Shirsat, V.S., R.S. Jagdale, and S.N. Deshmukh, Sentence level sentiment identification and calculation from news articles using machine learning techniques, in Computing, Communication and Signal Processing. 2019, Springer. p. 371-376.
- [24] Sebastiani, F., Machine learning in automated text categorization. ACM computing surveys (CSUR), 2002. 34(1): p. 1-47.

- [25] Yang, J., Liu, Y., Zhu, X., Liu, Z., and Zhang, X. A new feature selection based on comprehensive measurement both in inter-category and intra-category for text categorization. *Information Processing & Management*, 2012. 48(4): p. 741-754.
- [26] Taşcı, Ş. and T. Güngör, Comparison of text feature selection policies and using an adaptive framework. *Expert Systems with Applications*, 2013. 40(12): p. 4871-4886.
- [27] Nigam, K., McCallum, A. K., Thrun, S., and Mitchell, T. Text classification from labeled and unlabeled documents using EM. *Machine learning*, 2000. 39(2): p. 103-134.
- [28] Zong, W., Wu, F., Chu, L. K., and Sculli, D. A discriminative and semantic feature selection method for text categorization. *International Journal of Production Economics*, 2015. 165: p. 215-222.
- [29] Joachims, T., A Probabilistic Analysis of the Rocchio Algorithm with TFIDF for Text Categorization. 1996, Carnegie-mellon univ pittsburgh pa dept of computer science.
- [30] Zhao, L., Huang, M., Yao, Z., Su, R., Jiang, Y., and Zhu, X. Semi-supervised Multinomial Naive Bayes for text classification by leveraging word-level statistical constraint. in *Proceedings of the Thirtieth AAAI Conference on Artificial Intelligence*. 2016. AAAI Press.
- [31] Gabrilovich, E. and S. Markovitch. Text categorization with many redundant features: using aggressive feature selection to make SVMs competitive with C4. 5. in *Proceedings of the twenty-first international conference on Machine learning*. 2004. ACM.

- [32] Akın, A.A. and M.D. Akın, Zemberek, an open source NLP framework for Turkic languages. Structure, 2007. 10: p. 1-5.
- [33] Porter, M.F., An algorithm for suffix stripping. Program, 1980. 14(3): p. 130-137.
- [34] Uysal, A.K. and S. Gunal, The impact of preprocessing on text classification. Information Processing & Management, 2014. 50(1): p. 104-112.
- [35] Toğaçar, M., Eşidir, K. A., & Ergen, B., Yapay Zekâ Tabanlı Doğal Dil İşleme Yaklaşımını Kullanarak İnternet Ortamında Yayınlanmış Sahte Haberlerin Tespiti. Journal of Intelligent Systems: Theory and Applications, 2022. 5(1), 1-8.

NEW OPPORTUNITIES AND CHALLENGES IN QUANTUM NATURAL LANGUAGE PROCESSING

Ramazan KATIRCI¹

Taha OĞUZ²

1. INTRODUCTION

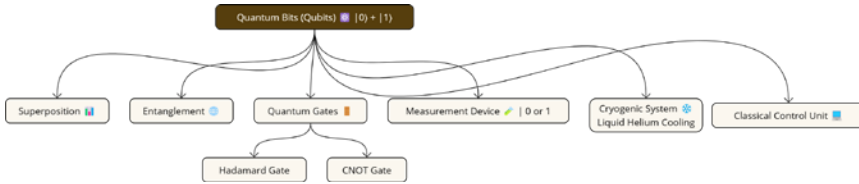
The advent of quantum computing has provided new opportunities and opened different pathways in various computational fields, particularly within the domain of natural language processing (NLP) (Grover, 1996; Sim, Johnson, & Aspuru-Guzik, 2019). The fundamental components and working principles of quantum computers are crucial for understanding their applications in complex fields such as NLP. These components enable efficient data processing and problem-solving. These components are summarized in Figure 1, which provides a visual overview of these foundational elements, including quantum gates, qubits, and their interactions. With the introduction of noisy intermediate-scale quantum (NISQ) devices (Ding et al., 2024), quantum machine learning (QML) has begun to attract significant attention for its potential role in enhancing and complementing traditional machine learning techniques (Aasim, Katirci, Acar, & Ali, 2024a; Coecke, Sadrzadeh, & Clark, 2010; Romito, 2023). NLP, a specialized subfield of artificial

¹ Prof. Dr., Sivas University of Science and Technology, Faculty of Engineering and Natural Sciences, Department of Computer Engineering, ramazankatirci@sivas.edu.tr, ORCID: 0000-0003-2448-011X.

² Res. Asst., Sivas University of Science and Technology, Faculty of Engineering and Natural Sciences, Department of Metallurgical and Materials Engineering, taha.oguz@sivas.edu.tr., ORCID: 0000-0003-4447-645X.

intelligence (AI) that aims to provide machines with the capacity to understand and generate human language, has experienced noteworthy developments through the application of quantum technologies. However, classical NLP approaches face limitations in terms of scalability and computational efficiency due to the inherently high dimensionality of linguistic data (Chang T., 2023; Shor, 1997). At this point, quantum natural language processing (QNLP) seeks to overcome these limitations by leveraging unique properties of quantum computing, such as superposition and entanglement.

Figure 1. Fundamental Components and Working Principles of Quantum Computers



The purpose of this review is to highlight the potential of QNLP and evaluate its role in overcoming the limitations of classical NLP methods. In this context, the particular advantages of quantum computing in handling large-scale linguistic data will be examined. Furthermore, this chapter will address gaps in the existing literature and explore under-researched aspects of QNLP, particularly in complex tasks such as sentence-level processing and machine translation. By presenting how QNLP can offer more efficient and scalable solutions compared to classical approaches, this review aims to provide a broader perspective on the potential of quantum computing to advance NLP.

2. CONTENT OF THE STUDY

This study aims to deeply examine the potential of quantum computing in the field of NLP and comprehensively

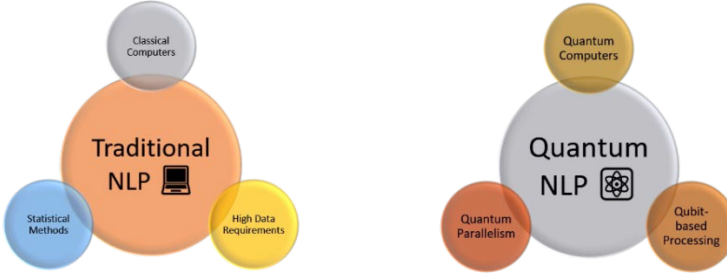
evaluate the numerous studies conducted in this area. It will investigate how QNLP approaches can be utilized to overcome the limitations of traditional NLP methods, such as scalability and computational efficiency. In the first section, the fundamental properties of quantum computers and their impact on NLP will be addressed, followed by an assessment of QNLP's position in the existing literature and an evaluation of previous studies from a broad perspective. In this context, the content of the study will be presented comprehensively, focusing on the theoretical foundations of QNLP, its application areas, and the challenges faced, along with a discussion of significant works in the existing literature.

2.1. Quantum Computing and NLP

The rise of quantum computing has unlocked new possibilities and opened up diverse avenues across various computational domains, particularly in NLP (Grover, 1996; Sim et al., 2019). With the introduction of NISQ devices (Ding et al., 2024), QML has begun to attract significant attention for its potential role in enhancing and complementing traditional machine learning techniques (Aasim et al., 2024a; Coecke et al., 2010; Romito, 2023). NLP, which is a specialized subfield of AI aimed at providing machines with the capacity to understand and generate human language, has experienced noteworthy developments through the application of quantum technologies. Classical NLP approaches have been quite effective in areas such as machine translation, sentiment analysis, and text classification. However, these methods encounter limitations in terms of computational efficiency and scalability due to the inherently high dimensionality of linguistic data, which often requires significant computational power and extensive storage capacity (Chang T., 2023; Shor, 1997). The fundamental differences between classical NLP and quantum NLP are significant in terms

of data processing methods and scalability, as illustrated in Figure 2.

Figure 2. The Main Differences Between Traditional NLP and QNLP



2.2. Quantum Natural Language Processing (QNLP)

QNLP seeks to overcome the limitations inherent in classical NLP methods by taking advantage of the unique properties of quantum computing, such as superposition and entanglement. These quantum characteristics enable a more efficient encoding and processing of complex linguistic information, which can significantly reduce computational demands when dealing with large and intricate datasets. QNLP has shown promise in achieving a quadratic speedup compared to classical approaches, offering substantial computational advantages, especially in managing large-scale linguistic data and intricate sentence structures that are challenging for traditional methods (Lukac, Abdiyeva, & Kameyama, 2018; Zheng, Gao, & Miao, 2023). Recent studies have concentrated on developing quantum neural networks (QNNs) as well as variational quantum circuits (VQCs) to address a variety of NLP tasks, including but not limited to text classification and sentence generation, which are crucial components in building intelligent language systems (Lukac et al., 2018; Metawei, Taher, ElDeeb, & Nassar, 2023; Romito, 2023)

One of the fundamental challenges in NLP lies in effectively representing the meaning derived from text, a task that becomes significantly more complicated as the dimensionality of the linguistic data increases. Classical approaches generally make use of techniques such as word embeddings or recurrent neural networks (RNNs) to capture and process linguistic information. However, as the size and complexity of the data grow, these methods can become computationally intensive and require considerable processing power, which leads to scalability issues (Romito, 2023; Sim et al., 2019). Quantum computing offers an alternative solution by enabling the encoding of text into quantum states, thereby allowing linguistic information to be processed in parallel within exponentially large quantum spaces, which greatly enhances efficiency (Bauer, Bravyi, Motta, & Kin-Lic Chan, 2020; Schuld, Bocharov, Svore, & Wiebe, 2020). The mathematical principles necessary for representing grammatical structures and word meanings through tensor products in a compositional way were initially proposed by Lambek (Lambek, 1958) and later extended by Coecke and colleagues (Coecke et al., 2010), laying a foundational framework for the quantum application of compositional semantics. This quantum approach has demonstrated potential in decreasing computational costs while simultaneously maintaining, or in some cases even improving, the accuracy of NLP models (Chang T., 2023; Coecke, de Felice, Meichanetzidis, & Toumi, 2020).

2.3. Advances in QNLP Models

The quantum self-attention neural network (QSAN) model, which has been proposed in recent research, marks a notable advancement in the utilization of quantum computing within the realm of NLP. This model integrates quantum circuits with self-attention mechanisms, which are essential components in state-of-the-art deep learning frameworks, including Transformer architectures (Romito, 2023). By harnessing the

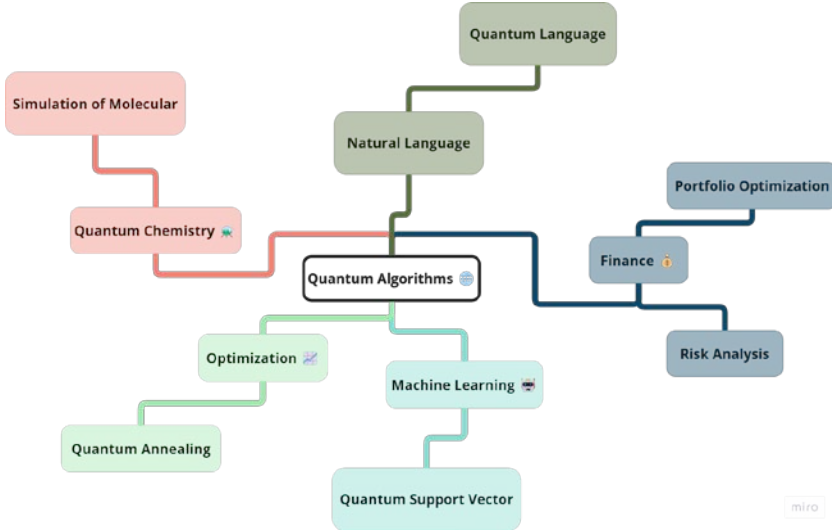
computational capabilities of quantum circuits, QSAN is able to exhibit enhanced performance in tasks such as text classification, demonstrating both high levels of accuracy and impressive scalability, while utilizing a relatively small number of parameters compared to traditional models (Widdows, Alexander, Zhu, Zimmerman, & Majumder, 2024; Zheng et al., 2023). Additionally, this quantum-based approach helps address one of the primary challenges faced by current QNLP models, which often depend on rigid syntactical structures that can restrict their flexibility in handling diverse linguistic inputs. By moving beyond such syntactical dependencies, QSAN offers a more adaptable and efficient framework for dealing with the complexity of natural language (Bauer et al., 2020; Zheng et al., 2023).

As advancements in quantum hardware continue, characterized by increasing numbers of available qubits and steadily decreasing error rates (Gokhale, Pande, & Pramod, 2020; Shenoy, Sheth, Behera, & Panigrahi, 2020), quantum natural language processing (QNLP) is becoming well-positioned to evolve into a mainstream approach for natural language understanding. Recent experimental developments involving the Quantum Instrumentation Control Kit (QICK) for superconducting quantum hardware highlight the steady progress being made in quantum systems (Ding et al., 2024). The hybrid quantum-classical models that are currently prevalent in this field facilitate incremental improvements, effectively leveraging both classical optimization methods and the quantum advantages in data encoding and processing (Lukac et al., 2018; Metawei et al., 2023). Furthermore, adaptive estimation of quantum observables is enhancing the efficiency of quantum algorithms, which has promising implications for QNLP applications (Shlosberg et al., 2023).

While the majority of current research efforts in QNLP have primarily focused on text classification tasks, there is still a notable lack of studies that address more complex areas such as sentence-level processing, especially in the realm of machine translation (Lambek, 1958; Peral-García, Cruz-Benito, & García-Peñalvo, 2024a, 2024b). This gap highlights the need for novel methodologies that can fully exploit the potential of quantum computing to manage the complexities associated with both the syntactic and semantic structures that are fundamental to effective language translation.

Quantum computing has emerged as a promising paradigm for advancing a wide range of fields, including NLP (Grover, 1996; Sim et al., 2019). The applications of quantum algorithms in various fields, such as agriculture, finance, image processing, and high-energy physics, are shown in Figure 3. The capability of quantum computers to efficiently process and analyze linguistic data presents new opportunities for developing models that may surpass classical models in both efficiency and overall capability. Quantum algorithms have already been successfully applied across various domains such as quantum chemistry and materials science (Yabaş, Biçer, & Katırcı, 2021), agriculture (Aasim, Katırcı, Acar, & Ali, 2024b), finance (Emmanoulopoulos & Dimoska, 2022), image processing (Gokhale et al., 2020), and high-energy physics (Wu et al., 2021), which illustrates the broad applicability and adaptability of quantum computing. In recent years, there has been an increasing number of studies investigating the confluence of quantum computing and NLP, with several frameworks and methodologies being proposed to utilize quantum principles for enhancing language understanding and generation (Kartsaklis et al., 2021; Widdows et al., 2024).

Figure 3. Applications of Quantum Algorithms in Various Fields



The study by Peral-García et al. (Peral-García et al., 2024a) presents a comparative analysis of classical NLP techniques versus QNLP approaches, specifically focusing on text classification tasks. This work investigates the potential of QNLP in executing NLP tasks by utilizing quantum computing, incorporating tools such as DisCoPy (Toumi, de Felice, & Yeung, 2022) and lambeq (Ganguly, n.d.; Kartsaklis et al., 2021) to effectively transform sentences into quantum circuits. The findings demonstrate that QNLP provides promising outcomes, particularly regarding enhanced processing speed and reduction in computational complexity. However, the practical applicability of these techniques remains constrained by the limitations of current quantum hardware. The researchers emphasize the need for advancements in quantum technologies to enable the handling of more complex and larger datasets effectively. They conclude that although quantum models have shown potential for simpler tasks, their current levels of accuracy and computational efficiency still fall short compared to classical models when

dealing with more complex sentence structures and extensive datasets.

One of the foundational contributions in this area is the development of the Categorical Compositional Distributional (DisCoCat) model (Coecke et al., 2010; Schuld et al., 2020). This model integrates category theory with distributional semantics to represent the meaning of sentences in a compositional way, effectively capturing the interplay between grammar and semantics. The DisCoCat framework offers a robust mathematical basis for representing both grammatical structures and word meanings through the use of tensor products, thereby laying the groundwork for quantum implementations of compositional semantics.

Coecke and Meichanetzidis expanded upon this foundational work by introducing the concept of QNLP, where they proposed the utilization of quantum circuits to represent both grammatical structures and the meanings of words. In their approach, individual words are encoded into quantum states, while the grammatical relationships between these words are modeled through quantum entanglement. This framework enables the effective use of quantum parallelism and entanglement, allowing linguistic information to be processed in a more efficient manner by leveraging quantum computational advantages (Coecke et al., 2020; K. Meichanetzidis et al., 2021).

Another important contribution in this field comes from by Meichanetzidis et al. (Konstantinos Meichanetzidis, Toumi, Felice, & Coecke, 2021), who implemented a grammar-aware question-answering system on quantum computers. In their work, they used the DisCoCat model to encode both the questions and the corresponding answers into quantum states, which allowed the quantum system to effectively retrieve the correct answer by leveraging the grammatical structure and semantic content of the

inputs. This research demonstrated the practical viability of quantum models for specific NLP tasks, illustrating how quantum mechanics can be used to enhance linguistic comprehension and processing. Their work set an important precedent for future studies in quantum-assisted language processing, inspiring continued exploration in applying quantum capabilities to more complex NLP challenges.

In the field of quantum machine learning applied to NLP, an overview was provided regarding the use of quantum approaches in deep learning models. The study explored how quantum algorithms could improve various NLP tasks, including text classification, sentiment analysis, and language modeling. It particularly highlighted the potential computational speedups that quantum computing could offer during both the training and inference stages, suggesting that quantum-enhanced models may be capable of handling larger datasets and capturing more complex language structures more efficiently than traditional classical models (Ganguly, n.d.; Kaisar, Masum, Maurya, & Murthy, n.d.; Peral-García et al., 2024a).

Lorenz et al. investigated the application of quantum deep learning techniques within the context of NLP, specifically focusing on how quantum neural networks could be trained to perform various language-related tasks. They proposed novel architectures where quantum circuits were utilized to model the inherently hierarchical and sequential nature of language. By leveraging the properties of quantum superposition and entanglement, these architectures aimed to effectively capture the complex linguistic patterns that emerge in natural language. Their research highlighted the potential for achieving a quantum advantage in NLP applications through the use of well-crafted quantum algorithms (Lorenz, Pearson, Meichanetzidis, Kartsaklis, & Coecke, 2021).

2.4. Challenges and Limitations in QNLP

Despite these advancements, there are still several limitations that continue to challenge existing research in quantum NLP. One of the most significant shortcomings is the scalability of quantum models. Many of the frameworks that have been proposed, including those based on the DisCoCat model, demand a substantial number of qubits to represent even moderately sized sentences. This requirement makes these approaches impractical given the current limitations of quantum hardware, which typically have a limited number of available qubits and are still subject to relatively high error rates (Coecke et al., 2010; K. Meichanetzidis et al., 2021).

Moreover, much of the current research predominantly focuses on theoretical models and small-scale demonstrations, frequently relying on simplified or artificially constructed datasets. For example, the grammar-aware question-answering system developed by Meichanetzidis et al. was evaluated using a limited set of sentences, which raises questions regarding the performance and scalability of these models when applied to real-world, large-scale language processing tasks (Konstantinos Meichanetzidis et al., 2021).

Another significant challenge lies in the complexity involved in encoding classical linguistic data into quantum states. The methods proposed for this purpose often require complex mappings and transformations, which can be computationally demanding and, as a result, may diminish the potential speed advantages that quantum processing aims to offer (Rath & Date, 2019). In addition to these difficulties, training quantum neural networks presents another major obstacle. Challenges such as vanishing gradients, the inherent noise present in quantum circuits, and the difficulties associated with implementing efficient optimization algorithms on quantum hardware all

contribute to making the training process notably challenging (Ghosh, n.d.).

3. RESULTS

This review has evaluated developments in QNLP and assessed its potential compared to classical NLP methods. QNLP has been found to hold significant promise, particularly in leveraging quantum properties such as superposition and entanglement to process high-dimensional linguistic data more efficiently. This advantage is crucial in overcoming the scalability and computational power limitations faced by classical methods. Innovative approaches, such as the QSAN model, have shown promising results in terms of accuracy and scalability for NLP tasks.

However, the current limitations of quantum hardware, including restricted qubit capacity and high error rates, continue to limit the practical application of QNLP in large-scale scenarios. Moreover, most existing studies have been conducted on small-scale, artificially constructed datasets, indicating a need for further research to meet the complex demands of real-world language processing. The complexity involved in converting linguistic data into quantum states also remains a significant challenge.

QNLP has the potential to surpass classical NLP methods, especially when supported by more advanced quantum hardware and hybrid quantum-classical approaches. Therefore, future research should focus on improving both quantum technology and QNLP models to realize this potential.

REFERENCES

- Aasim, M., Katırcı, R., Acar, A. Ş., & Ali, S. A. (2024a). A comparative and practical approach using quantum machine learning (QML) and support vector classifier (SVC) for Light emitting diodes mediated in vitro micropropagation of black mulberry (*Morus nigra* L.). *Industrial Crops and Products*, 213, 118397. Retrieved from <https://doi.org/https://doi.org/10.1016/j.indcrop.2024.118397>
- Aasim, M., Katırcı, R., Acar, A. Ş., & Ali, S. A. (2024b). A comparative and practical approach using quantum machine learning (QML) and support vector classifier (SVC) for Light emitting diodes mediated in vitro micropropagation of black mulberry (*Morus nigra* L.). *Industrial Crops and Products*, 213(January). Retrieved from <https://doi.org/10.1016/j.indcrop.2024.118397>
- Bauer, B., Bravyi, S., Motta, M., & Kin-Lic Chan, G. (2020). Quantum Algorithms for Quantum Chemistry and Quantum Materials Science. *Chemical Reviews*, 120(22), 12685–12717. Retrieved from <https://doi.org/10.1021/acs.chemrev.9b00829>
- Chang T., D. (2023). Variational Quantum Classifiers for Natural -Language Text.
- Coecke, B., de Felice, G., Meichanetzidis, K., & Toumi, A. (2020). Foundations for Near-Term Quantum Natural Language Processing. *ArXiv Preprint ArXiv:2012.03755*. Retrieved from <http://arxiv.org/abs/2012.03755>
- Coecke, B., Sadrzadeh, M., & Clark, S. (2010). Mathematical Foundations for a Compositional Distributional Model of Meaning. Retrieved from <http://arxiv.org/abs/1003.4394>

- Ding, C., Di Federico, M., Hatridge, M., Houck, A., Leger, S., Martinez, J., ... Cancelo, G. (2024). Experimental advances with the QICK (Quantum Instrumentation Control Kit) for superconducting quantum hardware. *Physical Review Research*. Retrieved from <https://doi.org/10.1103/PhysRevResearch.6.013305>
- Emmanoulopoulos, D., & Dimoska, S. (2022). Quantum Machine Learning in Finance: Time Series Forecasting. *ArXiv Preprint ArXiv:2202.00599*. Retrieved from <http://arxiv.org/abs/2202.00599>
- Ganguly, S. (n.d.). Quantum Natural Language Processing based Sentiment Analysis using lambeq Toolkit, 5–10.
- Ghosh, K. J. B. (n.d.). Encoding classical data into a quantum computer, 1–13.
- Gokhale, A., Pande, M. B., & Pramod, D. (2020). Implementation of a quantum transfer learning approach to image splicing detection. *International Journal of Quantum Information*, 18(5), 2050024. Retrieved from <https://doi.org/10.1142/S0219749920500240>
- Grover, L. K. (1996). A fast quantum mechanical algorithm for database search. In *Proceedings of the Annual ACM Symposium on Theory of Computing* (Vol. Part F1294, pp. 212–219). Retrieved from <https://doi.org/10.1145/237814.237866>
- Kaisar, A., Masum, M., Maurya, A., & Murthy, D. S. (n.d.). Hybrid Quantum-Classical Machine Learning for Sentiment Analysis.
- Kartsaklis, D., Fan, I., Yeung, R., Pearson, A., Lorenz, R., Toumi, A., ... Coecke, B. (2021). lambeq: An Efficient High-Level Python Library for Quantum NLP. *ArXiv*

Preprint ArXiv:2110.04236. Retrieved from
<http://arxiv.org/abs/2110.04236>

Lambek, J. (1958). The Mathematics of Sentence Structure. *The American Mathematical Monthly*, 65(3), 154–170. Retrieved from
<https://doi.org/10.1080/00029890.1958.11989160>

Lorenz, R., Pearson, A., Meichanetzidis, K., Kartsaklis, D., & Coecke, B. (2021). QNLP in Practice: Running Compositional Models of Meaning on a Quantum Computer. Retrieved from
<https://doi.org/10.1613/jair.1.14329>

Lukac, M., Abdiyeva, K., & Kameyama, M. (2018). CNOT-measure quantum neural networks. In *Proceedings of The International Symposium on Multiple-Valued Logic* (Vol. 2018-May, pp. 186–191). Retrieved from
<https://doi.org/10.1109/ISMVL.2018.00040>

Meichanetzidis, K., de Felice, G., Toumi, A., Coecke, B., Gogioso, S., & Chiappori, N. (2021). Quantum natural language processing on near-term quantum computers. *Electronic Proceedings in Theoretical Computer Science, EPTCS*, 340, 213–229. Retrieved from
<https://doi.org/10.4204/EPTCS.340.11>

Meichanetzidis, Konstantinos, Toumi, A., Felice, G. De, & Coecke, B. (2021). Grammar-aware sentence classification on quantum computers, 1–19.

Metawei, M. A., Taher, M., ElDeeb, H., & Nassar, S. M. (2023). A topic-aware classifier based on a hybrid quantum-classical model. *Neural Computing and Applications*, 35(25), 18803–18812. Retrieved from
<https://doi.org/10.1007/s00521-023-08706-7>

- Peral-García, D., Cruz-Benito, J., & García-Peñalvo, F. J. (2024a). Comparing Natural Language Processing and Quantum Natural Processing approaches in text classification tasks. *Expert Systems with Applications*, 254. Retrieved from <https://doi.org/10.1016/j.eswa.2024.124427>
- Peral-García, D., Cruz-Benito, J., & García-Peñalvo, F. J. (2024b). Systematic literature review: Quantum machine learning and its applications. *Computer Science Review*, 51. Retrieved from <https://doi.org/10.1016/j.cosrev.2024.100619>
- Rath, M., & Date, H. (2019). Quantum Data Encoding : A Comparative Analysis of Classical-to-Quantum Mapping Techniques and Their Impact on Machine Learning Accuracy.
- Romito, A. (2023). Quantum hardware measures up to the challenge. *Nature Physics*, 19(9), 1234–1235. Retrieved from <https://doi.org/10.1038/s41567-023-02090-8>
- Schuld, M., Bocharov, A., Svore, K. M., & Wiebe, N. (2020). Circuit-centric quantum classifiers. *Physical Review A*, 101(3), 32308. Retrieved from <https://doi.org/10.1103/PhysRevA.101.032308>
- Shenoy, K. S., Sheth, D. Y., Behera, B. K., & Panigrahi, P. K. (2020). Demonstration of a measurement-based adaptation protocol with quantum reinforcement learning on the IBM Q experience platform. *Quantum Information Processing*, 19(5), 146. Retrieved from <https://doi.org/10.1007/s11128-020-02657-x>
- Shlosberg, A., Jena, A. J., Mukhopadhyay, P., Haase, J. F., Leditzky, F., & Dellantonio, L. (2023). Adaptive estimation of quantum observables. *Quantum*, 7, 906.

Retrieved from <https://doi.org/10.22331/q-2023-01-26-906>

- Shor, P. W. (1997). Polynomial-time algorithms for prime factorization and discrete logarithms on a quantum computer. *SIAM Journal on Computing*, 26(5), 1484–1509. Retrieved from <https://doi.org/10.1137/S0097539795293172>
- Sim, S., Johnson, P. D., & Aspuru-Guzik, A. (2019). Expressibility and Entangling Capability of Parameterized Quantum Circuits for Hybrid Quantum-Classical Algorithms. *Advanced Quantum Technologies*, 2(12), 1900070. Retrieved from <https://doi.org/10.1002/qute.201900070>
- Toumi, A., de Felice, G., & Yeung, R. (2022). DisCoPy for the quantum computer scientist. *ArXiv Preprint ArXiv:2205.05190*. Retrieved from <http://arxiv.org/abs/2205.05190>
- Widdows, D., Alexander, A., Zhu, D., Zimmerman, C., & Majumder, A. (2024). Near-term advances in quantum natural language processing. *Annals of Mathematics and Artificial Intelligence*. Retrieved from <https://doi.org/10.1007/s10472-024-09940-y>
- Wu, S. L., Sun, S., Guan, W., Zhou, C., Chan, J., Cheng, C. L., ... Wei, T. C. (2021). Application of quantum machine learning using the quantum kernel algorithm on high energy physics analysis at the LHC. *Physical Review Research*, 3(3). Retrieved from <https://doi.org/10.1103/PhysRevResearch.3.033221>
- Yabaş, E., Biçer, E., & Katırcı, R. (2021). Experimental and In Silico studies on optical properties of new thiadiazole tetrasubstituted metal-free and zinc phthalocyanine

compounds. *Optical Materials*, 122(November), 111808.
Retrieved from
<https://doi.org/10.1016/j.optmat.2021.111808>

Zheng, J., Gao, Q., & Miao, Z. (2023). Design of a Quantum Self-Attention Neural Network on Quantum Circuits. In *Conference Proceedings - IEEE International Conference on Systems, Man and Cybernetics* (pp. 1058–1063). Institute of Electrical and Electronics Engineers Inc. Retrieved from
<https://doi.org/10.1109/SMC53992.2023.10393989>

TRANSFORMER MİMARİSİ

Ramazan KATIRCI¹

Hilal ÇELİK²

1. GİRİŞ

Transformer mimarisi, bilgi işleme sürecini tamamen dikkat (attention) mekanizmaları üzerine inşa eden bir yapay zeka modelidir. Bu mekanizma, RNN (Tekrarlayan Sinir Ağı) ve CNN (Evrişimsel Sinir Ağı) gibi geleneksel yöntemlerden farklı olarak, bilgiyi sıralamadan bağımsız bir biçimde işleyebilir. Dikkat mekanizmaları, her bir girdiyi diğer girdilerle olan ilişkileri üzerinden değerlendirerek paralel bilgi işleme yeteneği sağlar. Ancak, doğal dil işleme çalışmalarında sıralama bilgisi önemli olduğundan, Transformer modeli bu bilgi eksikliğini gidermek için konumsal kodlama (positional encoding) kullanır. Konumsal kodlama, belirteçlerin sırasını modele dahil eder. Böylece model, işlem sırasında belirteçlerin sıralama bilgisini bilir. Dikkat mekanizmaları ve Konumsal kodlamanın yanı sıra, Transformer mimarisi, geleneksel makine öğrenimi yöntemlerine kıyasla daha fazla hiperparametre içerir ve katmanlı bir yapıya sahiptir. Bu bileşenler, modelin etkili veri işleme yeteneklerini artırarak, karmaşık görevleri daha başarılı bir şekilde yerine getirmesini sağlar.

¹ Prof. Dr, Sivas Bilim ve Teknoloji Üniversitesi, Mühendislik ve Doğa Bilimleri Fakültesi, Bilgisayar Mühendisliği Bölümü, ramazankatirci@sivas.edu.tr, ORCID: 0000-0003-2448-011X.

² Arş. Gör. Sivas Bilim ve Teknoloji Üniversitesi, Mühendislik ve Doğa Bilimleri Fakültesi, Bilgisayar Mühendisliği Bölümü, hilalcelik@sivas.edu.tr, ORCID: 0000-0001-5428-3411.

Makine öğrenmesi algoritmaları, farklı disiplinlerde önemli yenilikler sağlayarak çeşitli alanlarda dikkate değer ilerlemelere imkân tanımaktadır. Sağlık alanında, bu algoritmalar hastalıkların sınıflandırılması (Kaur et al. 2022) ve görüntüleme (Tuncer et al. 2024) süreçlerinde başarıyla uygulanmaktadır. Tarımda ise, bitki yetiştirme süreçlerinin iyileştirilmesi (Aasim et al. 2024) ve verimliliğin artırılması (Aasim et al. 2022) amacıyla kullanılmıştır. E-ticaret sektöründe, doğru müşteri kitlesine ulaşım (ÇELİK and ÇINAR 2021) ve müşteri kaybını tahmin etme (Nagaraj et al. 2023) gibi uygulamalarda yer bulmuştur. Finans alanında kaynak yatırımlarının tahmin edilmesinde (Karkan, Power, and Systems 2022) ve eğitimde kişiselleştirilmiş öğrenme deneyimlerinin geliştirilmesinde (Essa, Celik, and Human-Hendricks 2023) katkı sağlamaktadır. Bu çerçevede, makine öğrenmesi, birçok alanda süreçlerin etkinliğini artıran yenilikçi çözümler sunmaktadır.

Makine öğreniminin ilk dönemlerinde, kural tabanlı sistemler yaygın olarak kullanılıyordu. Bu sistemler, belirli kurallara dayalı algoritmalarla bilgi işleme süreçlerini yönlendiriyor ve belirli hedeflere yönelik çalışıyordu. Kural tabanlı sistemler, yazım düzeltmede (Kukich 1992) kelimeler arasındaki benzerlikleri saptayıp hataları düzeltmek, makine çevirisinde (Hutchins 1995) dil kurallarını kullanarak diller arası çeviri yapmak ve sohbet robotlarında (Adamopoulou and Moussiades 2020) kullanıcı sorularını belirli kurallar çerçevesinde yanıtlamak gibi alanlarda kullanılıyordu. Bu yaklaşımlar, belirli bir kurallar dizisi kullanarak bilgi işleme görevlerini gerçekleştirirken, daha karmaşık ve bağlama dayalı anlamlandırma gerektiren durumlarda sınırlı kalıyordu. Bu sınırlamaların birçoğu, özellikle bağlamsal anlamı daha iyi yakalayabilen ve daha büyük veri setlerinde daha etkili performans sergileyen erişim tabanlı (retrieval based) yaklaşımların geliştirilmesi ile aşılmıştır (Lewis et al.

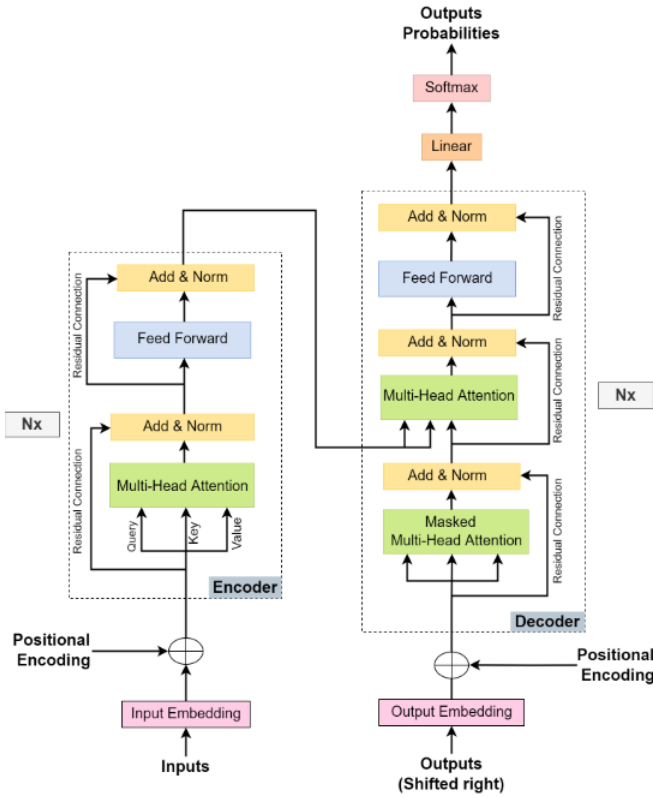
2021)(Voorhees and Tice 2000). Erişim tabanlı yaklaşım, Google ekibi tarafından geliştirilen konuşma modellemesi ile doğal dil işleme (NLP) alanında bir dönüm noktası oluşturmıştır (Vinyals and Le 2015). Daha ileri bir versiyon olarak, üretim tabanlı yaklaşım, erişim tabanlı yöntemlerin eksikliklerini geride bırakmayı başarmıştır. Erişim tabanlı sistemler, belirlenmiş bir veri setinden ve bağlamdan yanıtlar elde ederken, üretim tabanlı yöntemler kullanıcı sorgularından ve önceki kullanıcı yanıtlarından yeni yanıtlar üretmektedir. Erişim tabanlı sistemler, yeni içerik oluşturmaz, mevcut bilgi havuzundan en uygun yanıtları seçerken buna karşılık, üretim tabanlı yaklaşımlar, dil modelleri kullanarak yanıtlar üretir (Pandey and Sharma 2023). Üretim tabanlı yöntemler; özetleme (Ioannidis et al. 2023), metin-görüntü üretimi (Liu and Chilton 2022), sohbet robotları (Adamopoulou and Moussiades 2020) (Le, Boureau, and Nickel 2019) , otomatik konuşma tanıma (Gulati et al. 2020), makine çevirisi (Raffel et al. 2023), metin tamamlama (Radford et al. 2019), metin özetleme (Sun and Platoš 2024) ve bilgisayarlı görü (Yu et al. 2022), gibi alanlarda başarılı sonuçlar ortaya koymaktadır. Günümüzde kullanılan üretim tabanlı modellerin en önemlilerinden biri, Transformer tabanlı Generative Pre-Training (GPT) modelidir (Radford et al. 2018) (Radford et al. 2019)(Brown et al. 2020). Bu model, büyük veri setlerinden öğrenme yeteneği ve bağlama dayalı insan benzeri metinler üretmesiyle Transformers'ın potansiyelini etkileyici bir biçimde kullanmaktadır.

2. TRANSFORMER MİMARİSİ VE BİLEŞENLERİ

Transformer, tamamen dikkat mekanizmalarına temel alan geleneksel makine öğrenimi yöntemlerinden daha fazla bileşen içeren yenilikçi bir mimaridir. Bu bileşenler arasında kodlayıcı (encoder) ve kod çözücü (decoder) yapıları yer alır;

encoder, giriş dizisini işleyerek gizli bir temsile dönüştürürken, kod çözücü bu temsili alıp çıkış dizisini üretir. Gömme ve konumsal kodlama bölümü, her kelimenin vektörel bir temsilini oluşturarak modelin kelimelerin bağlamını anlamasını sağlar (ÇELİK, KATIRCI, and İŞLEK 2024). Kendi dikkat (Self-Attention) mekanizması ve

Şekil 1. Transformer Mimarisi



Kaynak: (Islam et al. 2024)

Çoklu Dikkat Başlıkları (Multi-Head Attention), birden fazla kendine dikkat mekanizmasının paralel olarak çalışmasından oluşur. ADD & Katman Normalizasyon (Layer Normalization) bölümü, modelin performansını artırmak için normalizasyon ve artık bağlantı işlemleriyle verimliliği yükseltir.

Artık bağlantı, önceki katmanların bilgilerini koruyarak daha derin ağların eğitimini kolaylaştırır(He et al. 2015). Son olarak, FNN (Feedforward Neural Network), verilerin giriş katmanından başlayarak gizli katmanlar üzerinden ileri doğru aktığı ve sonuçların çıkış katmanında üretildiği, ileri beslemeli bir yapay sinir ağıdır(GÜLTEPE 2019). Transformer mimarisinin paralelleştirme yeteneği, özellikle çok başlı dikkat mekanizması sayesinde, verinin aynı anda farklı özelliklerine odaklanarak işlenmesini mümkün kılar. Bu özellik, giriş dizisinin farklı bölümlerine aynı anda farklı başlıklarla (head) odaklanarak, her başlığın bağımsız olarak dikkat hesaplaması yapmasını sağlar. Böylece model, veri içerisindeki farklı ilişkileri ve temsilleri paralel olarak öğrenir. Bu süreç, her kelimenin diğer tüm kelimelerle olan ilişkisini aynı anda hesapladığı için, sıraya bağlı hesaplamaların (örneğin, RNN'lerde olduğu gibi) aksine, daha hızlı işlem yapar.(Islam et al. 2024).

3. TRANSFORMER MİMARİSİ VE BİLEŞENLERİ

Transformer, tamamen dikkat mekanizmalarına temel alan geleneksel makine öğrenimi yöntemlerinden daha fazla bileşen içeren yenilikçi bir mimaridir. Bu bileşenler arasında kodlayıcı (encoder) ve kod çözücü (decoder) yapıları yer alır; encoder, giriş dizisini işleyerek gizli bir temsile dönüştürürken, kod çözücü bu temsili alıp çıkış dizisini üretir. Gömme ve konumsal kodlama bölümü, her kelimenin vektörel bir temsilini oluşturarak modelin kelimelerin bağlamını anlamasını sağlar (ÇELİK et al. 2024). Kendi dikkat (Self-Attention) mekanizması ve Çoklu Dikkat Başlıkları (Multi-Head Attention), birden fazla kendine dikkat mekanizmasının paralel olarak çalışmasından oluşur. ADD & Katman Normalizasyon bölümü, modelin performansını artırmak için normalizasyon ve artık bağlantı işlemleriyle verimliliği yükseltir. Artık bağlantı, önceki

katmanların bilgilerini koruyarak daha derin ağların eğitimini kolaylaştırır(He et al. 2015). Son olarak, FNN (Feedforward Neural Network), verilerin giriş katmanından başlayarak gizli katmanlar üzerinden ileri doğru aktığı ve sonuçların çıkış katmanında üretildiği, ileri beslemeli bir yapay sinir ağıdır(GÜLTEPE 2019). Transformer mimarisinin paralelleştirme yeteneği, özellikle çok başlı dikkat mekanizması sayesinde, verinin aynı anda farklı özelliklerine odaklanarak işlenmesini mümkün kılar. Bu özellik, giriş dizisinin farklı bölümlerine aynı anda farklı başlıklarla (head) odaklanarak, her başlığın bağımsız olarak dikkat hesaplaması yapmasını sağlar. Böylece model, veri içerisindeki farklı ilişkileri ve temsilleri paralel olarak öğrenir. Bu süreç, her kelimenin diğer tüm kelimelerle olan ilişkisini aynı anda hesapladığı için, sıraya bağlı hesaplamaların (örneğin, RNN'lerde olduğu gibi) aksine, daha hızlı işlem yapar.(Islam et al. 2024).

3.1.Kodlayıcı-kod çözücü (Encoder-Decoder) Yapısı

Seq2seq(Sequence to Sequence), farklı uzunluklardaki girdi ve çıktı dizilerini eşleştirmek üzere tasarlanmış bir makine öğrenimi modelleme yöntemidir. Bu yaklaşım, girdileri çıktı ile ilişkilendirmeyi hedefler ve iki temel bileşenden oluşur: kodlayıcı ve kod çözücü (Choudhary and Chauhan 2023). Kodlayıcı, bir girdi dizisini işleyerek, bu dizinin anlamını özetleyen ve genellikle "bağlam" olarak adlandırılan bağlamsallaştırılmış bir temsil üretir. Kod çözücü ise, kodlayıcı tarafından oluşturulan bu bağlamı kullanarak hedef çıktıyı üretir(Jurafsky and Martin 2021) (Cho et al. 2014).

Transformer mimarisinde kodlayıcı ve kod çözücü, alt katmanlardan oluşan bir yığınla yapılandırılmıştır. Kodlayıcı, iki temel alt katman içerir: ilki çoklu başlı dikkat mekanizması (multi-head self-attention), ikincisi ise konumuna dayalı tam bağlantılı ileri beslemeli bir ağıdır. Bu alt katmanların etrafında

bir kalıntı bağlantısı (residual connection) uygulanmakta ve ardından katman normalizasyonu yapılmaktadır. Kod çözücü da benzer şekilde alt katmandan oluşmakla birlikte bir üçüncü alt katman olan kodlayıcının çıktısı üzerinde çok başlı dikkat uygular ve böylece kod çözücü, kodlayıcı tarafından sağlanan bilgiyi kullanarak daha anlamlı ve bağlama dayalı çıktılar üretmesine olanak tanır. Aynı zamanda, kod çözücünün kendine dikkat alt katmanı, gelecekteki konumlarına dikkat edilmesini engelleyen bir maskeleyme mekanizması uygular. Maskeleyme tekniği, Transformer modellerinde, i konumundaki çıktıyı tahmin ederken, modelin yalnızca i konumundan önceki bilgilere dayanmasını sağlar. Bu, dil modelleme gibi görevlerde, bir kelimenin tahmininin yalnızca kendisinden önceki kelimelere bağlı olması gerektiğini garanti eder (Vaswani et al. 2017).

3.2.Gömme ve Konumsal Kodlama

Gömme kavramı, özellikle kelime gömme (word embedding) olarak bilinir ve belirteçleri sabit uzunlukta sayısal vektörler şeklinde temsil eder. Bu temsiller, belirteçlerin anlamlarını, bağlamlarını ve aralarındaki ilişkileri matematiksel olarak ifade etmeyi sağlar (Almeida and Xex 2015). Ancak Transformer, sıralı veri yapısına sahip olmadığından sadece kelime gömme yeterli olmaz. Bu nedenle modelde iki bileşen yer alır: kelime gömme, kelimeleri sabit boyutlu vektörlere dönüştürerek anlamlarını yakalarken, konumsal kodlama ise kelimenin cümledeki sırasını belirtmek için eklenen konumsal gömmedir. Böylece model, kelimenin hem anlamını hem de cümledeki konumunu öğrenir. Sinüs ve kosinüs fonksiyonlarıyla hesaplanan bu sinyaller, kelime gömmeye eklenir. Böylece model, kelimenin anlamını sadece bağlamına göre değil, aynı zamanda cümledeki sırasına göre de öğrenebilir (Vaswani et al. 2017).

Gömme boyutu, modelin her belirteci temsil etmek için kullandığı vektörün uzunluğudur. Büyük boyutlar daha fazla bilgiyi yakalayabilir, ancak aynı zamanda hesaplama maliyetini artırabilir ve aşırı öğrenmeye neden olabilir. Küçük boyutlar ise daha az bilgi taşır, ancak daha hızlı hesaplama ve daha az bellek kullanımı sağlar(Cheema and Khan 2020). Dolayısıyla, bu değer problem alanına ve kullanılacak verinin büyüklüğüne göre optimize edilmelidir.

Transformer mimarisinde işlenecek her belirteç kendisine atanan benzersiz numaralarla temsil edilir. Seçilen gömme boyutuna göre her kelime, bir gömme vektörü ile temsil edilir. Gömme vektörlerinin başlangıç değerleri genellikle 0 ile 1 arasında rastgele olarak belirlenir ve daha sonra model, kelimeler arasındaki anlamları öğrenmeye başladıkça bu değerler güncellenir(Pennington, Socher, and Christopher 2014) (Choudhary and Chauhan 2023).

Konumsal kodlama (positional encoding), Transformer mimarisinde her belirtecin dizideki konumunu belirten bilgiyi sağlamaktadır. Konumsal kodlama hesaplama formülü denklem 1 ve denklem 2 de verilmiştir. Denklem 1, belirli bir konum için çift indeksli bileşenlerin değerini sinüs fonksiyonu kullanarak hesaplar. Denklem 2 ise aynı konum için tek indeksli bileşenlerin değerini kosinüs fonksiyonu kullanarak hesaplar(Vaswani et al. 2017).

Her bir pozisyon ve gömme boyutunun her bir boyutu i için konumsal kodlama PE (pos, i) şöyle tanımlanır:

$$PE(pos, 2i) = \sin\left(\frac{konum}{10000^{\frac{2i}{d_{model}}}}\right) \quad (1)$$

$$PE(pos, 2i + 1) = \cos\left(\frac{konum}{10000^{\frac{2i}{d_{model}}}}\right) \quad (2)$$

pos: Dizideki belirtecin konumu (0 tabanlı indeks).

i: Boyut indeksi (0, 1, 2, ..., $d_{\text{model}}-1$).

d_{model} : Gömme boyutunun (embedding dimension) büyüklüğü.

Bu kodlama, modelin farklı pozisyonlardaki belirteçleri ayırt etmesine olanak tanır. Sinüs ve kosinüs fonksiyonları, her bir boyutun farklı bir frekansa karşılık gelmesini sağlar, bu da belirteçlerin göreceli pozisyonlarını etkili bir şekilde yakalamaya yardımcı olur. Sonraki adım, gömme vektörüne konumsal kodlama eklemektir. Seçilen gömme boyutu kadar konumsal kodlama, kelimelerin konumunu belirlemek için kullanılır. Son olarak, konumsal gömülme hesaplandıktan sonra, kelime gömümleri ve konumsal gömümler bir araya getirilir. Bu adım, modelin hem kelimelerin anlamını hem de sırasını anlamasını sağlamak için kritik öneme sahiptir(Vaswani et al. 2017).

3.3.Kendine Dikkat Mekanizması ve Çoklu Dikkat Başlıkları

Dikkat mekanizması, geleneksel ardışık yapıların yerine, cümlelerin parçaları arasında ilişkilendirmeler yaparak daha dinamik bir anlamlandırma süreci sunmaktadır. Bu yaklaşım, dil modellerinin uzun bağımlılıkları daha etkin bir şekilde öğrenmesine olanak tanımıştır(Bahdanau, Cho, and Bengio 2016). 2017 yılında Vaswani ve arkadaşları tarafından geliştirilen Transformer mimarisi, dikkat mekanizmasını merkezine alarak dil modellerinin paralel işlemlerle daha hızlı ve verimli çalışmasını sağlamıştır (Shaw, Uszkoreit, and Vaswani 2018). Bu mekanizma sayesinde, metin içindeki belirteçlerin birbirleriyle olan ilişkilerini daha iyi kavramak mümkün hale gelmiş. Transformer mimarisinin sunduğu bu yenilik, dil modellerinin daha karmaşık yapılarla başa çıkmasına ve daha iyi performans göstermesine katkıda bulunmuştur(Zhang et al. 2020).

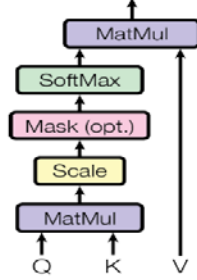
Dikkat fonksiyonu, bir sorgu ile bir dizi anahtar-değer çiftini bir çıktıya eşlemek olarak tanımlanabilir. Bu süreçte sorgu, anahtarlar, değerler ve çıktı hepsi vektörler olarak temsil edilir. Öz-dikkat (self-attention) alt katmanları, Transformer mimarisinde, genellikle h kadar dikkat başlığı kullanarak çalışır. Her dikkat başlığı, giriş verisinin farklı yönlerine dikkat eder ve bu sayede model, bilgi temsillerini zenginleştirir. Çok başlı dikkat (multi-head attention) mekanizması, birçok tek başlı dikkat mekanizmasının birleşimiyle ortaya çıkar. Her bir başlıktan elde edilen sonuçlar, alt katmanın çıktısını oluşturmak için birleştirilir. Bu yapı, modelin farklı bakış açılarıyla dikkat etmesine olanak tanırken, daha zengin bir bilgi temsili sağlar. Böylece model, farklı kelimelerin bağlam içindeki anlamlarını daha iyi anlayabilir ve bu da uzun bağımlılıkların etkin bir şekilde öğrenilmesini sağlar (Shaw et al. 2018). Çok başlı dikkat mekanizması, sinir ağlarının sıralı veri içerisindeki farklı özellikleri yakalamasına olanak tanır. Bu mekanizma, dikkat başlıklarının belirli aralıklarda (kısa veya uzun) bilgi toplamasını sağlayarak, giriş bağlamlarının temsillerini zenginleştirir. Her dikkat başlığı, sorgu (query) ve anahtar (key) temsilleri arasındaki ilişkilere dayalı olarak farklı dikkat ağırlıkları atar, böylece her başlık, verinin farklı yönlerine odaklanarak modelin daha kapsamlı bir öğrenme süreci gerçekleştirmesine katkıda bulunur. Bu yaklaşım, model performansını genel olarak iyileştirir (Islam et al. 2024) (Lu et al. 2024).

Sorgu (Query): Bir kelimenin diğer kelimelerle olan ilişkisini bulmak için kullanılan bir vektördür. Bu vektör, modelin bir kelimeye ne kadar dikkat etmesi gerektiğini belirlemek için kullanılır.

Anahtar (Key): Diğer kelimelerle ilişkili bilgileri içeren vektördür. Model, sorgu(query) ile anahtar(key) vektörünü karşılaştırarak, hangi kelimelerin dikkate alınması gerektiğine karar verir.

Değer (Value): Modelin dikkatini yönlendirdiği kelimelere ait bilgiyi taşıyan vektördür. Dikkat dağılımı sonucunda, en alakalı kelimelere karşılık gelen değer(value) vektörleri seçilir ve sonuç olarak hesaplanan çıkış (output) elde edilir.

Şekil 2. Ölçeklendirilmiş Nokta-Çarpım Dikkat



Kaynak: (Vaswani et al. 2017)

Sorgu hesabı $\implies$ Query = (WE + PE) X W_q

Anahtar hesabı $\implies$ Key = (WE + PE) X W_k

Değer hesabı $\implies$ Value = (WE + PE) X W_v

WE: Word Embedding (Kelime Gömmesi)

PE: Positional Embedding (Konum Gömmesi)

W_q, W_k, W_v: Sırasıyla sorgu, anahtar ve değer için ağırlıklar

Ölçeklenmiş nokta-çarpımı(scaled dot-product)

$$scale = \frac{Query \times (Key)^T}{\sqrt{d_{model}}} \quad (3)$$

Ölçekleme sonucuna softmax gibi bir aktivasyon fonksiyonu uygulanır.

$$\text{softmax} \left(\frac{Query \times (Key)^T}{\sqrt{d_{model}}} \right) \quad (4)$$

Dikkati bulmak için ölçekleme sonucu ile değer çarpılır.

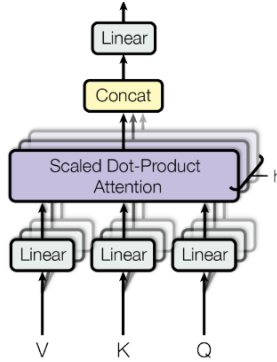
$$\text{Attention}(\text{Query, Key, Value}) = \text{softmax}\left(\frac{\text{Query} \times (\text{Key})^T}{\sqrt{d_{\text{model}}}}\right) \times \text{Value} \quad (5)$$

Çok başlı dikkat (şekil 3) modülü içinde ölçeklenmiş nokta-çarpımı dikkat fonksiyonunun paralel uygulanması, giriş dizisindeki farklı bölümler arasındaki maksimum bağımlılıkların çıkarılması için kritik bir öneme sahiptir. Her bir dikkat başlığı (h ile gösterilen), kendi öğrenilebilir ağırlıkları olan W^{Q_i} , W^{K_i} , W^{V_i} kullanarak dikkat mekanizmasını uygular. Her başlık tarafından hesaplanan dikkat çıktıları, sonrasında birleştirilir ve sonuçlar tek bir matris haline getirilmek üzere doğrusal bir dönüşümle işlenir(Vaswani et al. 2017)

$$\text{head}_i = \text{Attention}(QW^{Q_i}, KW^{K_i}, VW^{V_i})$$

$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \text{head}_2, \dots, \text{head}_H)W^0$$

Şekil 3. Çok Kafalı Dikkat



Kaynak: (Vaswani et al. 2017)

3.4.ADD & Katman Normalisyon

Add katmanı, artık bağlantı (residual connection) kullanarak, önceden hesaplanan bir katmanın çıktısını sonraki katmanın çıktısıyla toplar. Bu, modelin derin katmanlarda bilgi kaybını önlemesine ve derin öğrenme ağlarında öğrenme sürecini

kolaylaştırmaya yardımcı olur. Makine öğrenmesi modellerinde artık bağlantılar, gradyanların daha etkili bir şekilde yayılmasına olanak tanır ve bu sayede eğitim sürecinde daha derin ağların kullanılmasını mümkün kılar (Islam et al. 2024).

Transformer mimarisinde, her iki alt katman olan kendine dikkat mekanizması ve ileri besleme katmanı etrafında birer adet "Add & Norm" katmanı bulunur. Her alt katmanda bir "Add & Norm" işlemi gerçekleştirilir. Kendine dikkat katmanı, dikkat skorlarını üretir ve bu çıktıya, katmana giren orijinal giriş artık bağlantı ile eklenir. Ardından, bu toplam sonuç katman normalizasyon yapılır. Feed-forward katmanı ise çıktıyı lineer dönüşüm ve aktivasyon işlevlerinden geçirir; ardından orijinal giriş yine artık bağlantı ile eklenir ve normalizasyon işlemi yapılır. Bu nedenle, her iki alt katmanda birer "Add & Norm" işlemi uygulanır (Vaswani et al. 2017). Dikkat ve feed-forward katmanlarından sonra uygulanan normalizasyon Transformer mimarisinin eğitim sürecini hızlandırır, modelin stabilitesini artırır ve daha iyi performans elde edilmesine katkıda bulunur (Raffel et al. 2023) (Ba, Kiros, and Hinton 2016).

Add katmandan sonra matrisi normalleştirmek için, her pozisyon için ortalama ve standart sapmayı pozisyon bazında hesaplamak gerekir (Islam et al. 2024). Ortalama (denklem 6), bir veri setindeki tüm değerlerin toplanarak, veri sayısına bölünmesi ile hesaplanır ve standart sapma(denklem 7) ise veri noktalarının ortalamadan ne kadar saptığını ölçer(Jurafsky and Martin 2023).

$$\text{Ortalama: } x = \frac{1}{N} \sum_{i=1}^N X_i \quad (6)$$

$$\text{Standart Sapma: } \sigma = \sqrt{\frac{1}{N} \sum_{k=1}^N (x_i - \mu)^2} \quad (7)$$

μ = Ortalama (Mean)

N = Veri setindeki toplam değer sayısı

x_i = i'inci veri noktası

σ = Standart sapma (Standard Deviation)

3.5.Artık Bağlantılar (Residual Connections)

Artık bağlantı, bir katmandan elde edilen çıktının, birkaç katman ileriye taşınarak, daha sonraki katmanlardaki çıktıya eklenmesi prensibine dayanır. Yani, bir katmandan gelen girdiyi, aynı katmanın çıktısına ekler. Bu sayede model, her katmanda öğrenmesi gereken farkları (residuals) öğrenir ve bilgiyi daha verimli bir şekilde işler. Bu yapı, giriş verisinin direkt olarak sonraki katmanlara iletilmesine olanak tanır ve modelin derinliğinin artması durumunda bile öğrenmenin etkin bir şekilde devam etmesini sağlar(He et al. 2015).

Transformer mimarisinde kaybolan gradyanlar sorununu ele almak için, artık bağlantı hem dikkat hem de ileri besleme katmanlarının her çıktısına uygulanır (Islam et al. 2024). Katman normalizasyonu, artık bağlantıdan sonra uygulanarak her katmandaki çıktıyı normalize eder ve modelin daha kararlı bir şekilde öğrenmesini sağlar (He et al. 2021)(Vaswani et al. 2017)(Sonkar and Baraniuk 2023).

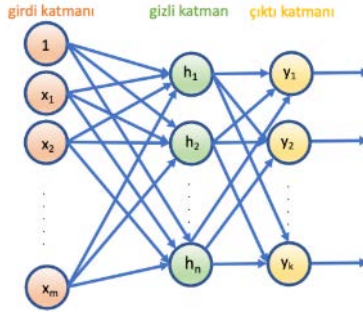
3.6.Feed-Forward Ağlar

FNN (Feedforward Neural Network) (şekil 4), Türkçede "İleri Beslemeli Sinir Ağı" olarak adlandırılır ve yapay sinir ağlarının en temel türlerinden biridir. Bu ağlarda veri, giriş katmanından başlayarak gizli katmanlar üzerinden ilerleyip çıkış katmanına ulaşır. Bu süreçte veri yalnızca ileri doğru ilerler, geriye doğru bir akış olmadığı için bu yapıya "ileri beslemeli" denir. FNN'ler, genellikle giriş, gizli ve çıkış katmanlarından oluşur. Giriş katmanı ham veriyi alırken, gizli katmanlar veriyi işler ve çıkış katmanı sonuçlar üretir. Her bir katmanda bulunan nöronlar, aktivasyon fonksiyonları yardımıyla girdi verilerini işler(GÜLTEPE 2019)(Guanabara et al. n.d.).

$$FFN(x) = \text{aktivasyon_fonk}(0, xW1 + b1)W2 + b2$$

Bu denklemde, W_1 ve W_2 , öğrenilen ağırlık matrislerini; b_1 ve b_2 , bias terimlerini temsil eder. Bu katman, her pozisyonda aynı hesaplamayı yapar, ancak her katmanda farklı parametreler kullanılır.

Şekil 4. İleri Beslemeli Sinir Ağı



Kaynak: (Çakır and Angın 2021)

Transformer mimarisinde, her kodlayıcı ve kod çözücü katmanı, dikkat alt katmanlarının yanı sıra tam bağlantılı (fully connected) bir ileri beslemeli ağ (FFN) içerir. İleri beslemeli ağ, iki doğrusal dönüşümden (linear transformation) ve bunların arasında ReLU gibi bir aktivasyon fonksiyonundan oluşur (Vaswani et al. 2017).

3.7. Transformer Mimarisinde Parallelleştirme

Parallelleştirme, modelin giriş verisinin farklı bölümlerini eşzamanlı olarak ele almasını sağlayarak eğitim ve çıkarım süreçlerinde verimliliği artırır (Kora and Mohammed n.d.). Büyük boyutlu modellerin geniş eğitim verileriyle çalışabilmesi için, çoklu GPU üzerinde dağıtılarak paralel eğitim yapılması mümkündür. Bu yöntem, büyük modellerin daha kısa sürede eğitilmesine olanak tanır (Shoeybi et al. 2020). Transformer modelinde paralel eğitim, özellikle dikkat mekanizmaları sayesinde mümkündür. Bu mimaride paralelleştirme, giriş dizisindeki tüm girişleri arasında paralel hesaplamalar yaparak, geleneksel RNN veya LSTM (Uzun Kısa Süreli Bellek) gibi

ardışık yöntemlerden farklı olarak her belirtecin diğerleriyle olan ilişkisini aynı anda değerlendirmeye olanak tanır (Vaswani et al. 2017)(Xiong et al. 2020). Ayrıca Transformer mimarisindeki her kodlayıcı ve kod çözücü katmanındaki ileri beslemeli ağlar bağımsız olarak çalışarak hesaplamaların paralel olarak gerçekleştirilmesine olanak tanır(Islam et al. 2024)

Transformer mimarisinde, çoklu dikkat başlıkları modelin giriş verisindeki çeşitli özellikleri eşzamanlı olarak öğrenmesini sağlar. Her bir başlık, verinin farklı bir boyutunu analiz eder ve bu başlıklardan elde edilen sonuçlar bir araya getirildiğinde, daha zengin ve kapsamlı bir temsil oluşturulur. Ölçeklendirilmiş nokta çarpımı dikkati fonksiyonu, çok başlıklı dikkat modülü içerisinde paralel olarak uygulanarak, giriş dizisinin farklı bölümleri arasındaki maksimum bağımlılıkların çıkarılmasına önemli katkı sağlar. Her dikkat başlığı, kendi öğrenilebilir ağırlıkları (W_kQ , W_kK , W_kV) ile bağımsız olarak dikkat hesaplamalarını gerçekleştirir ve bu sayede model, verinin farklı temsillerini paralel şekilde öğrenme kapasitesine sahip olur (Vaswani et al. 2017).

KAYNAKÇA

- Adamopoulou, Eleni, and Lefteris Moussiades. 2020. "Chatbots: History, Technology, and Applications." *Machine Learning with Applications* 2:100006. doi: 10.1016/J.MLWA.2020.100006.
- Almeida, Felipe, and Geraldo Xex. 2015. "Word Embeddings: A Survey." (1991).
- Ba, Jimmy Lei, Jamie Ryan Kiros, and Geoffrey E. Hinton. 2016. "Layer Normalization."
- Bahdanau, Dzmitry, Kyunghyun Cho, and Yoshua Bengio. 2016. "Neural Machine Translation by Jointly Learning to Align and Translate."
- Brown, Tom B., Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. "Language Models Are Few-Shot Learners." *Advances in Neural Information Processing Systems* 2020-Decem.
- ÇAKIR, Berna, and Pelin ANGIN. 2021. "Zamansal Evrişimli Ağlarla Saldırı Tespiti: Karşılaştırmalı Bir Analiz." *European Journal of Science and Technology* (22):204–11. doi: 10.31590/ejosat.848784.
- ÇELİK, Hilal, Ramazan KATIRCI, and Berrin İŞLEK. 2024. "Effect Of Parameters On Performance In Question-Answer Model With Simple Rnn Deep Learning

Method.” *INTERNATIONAL CONFERENCE ON SCIENTIFIC AND INNOVATION RESEARCH-III* 161–69.

Cheema, Sunbal, and Gul N. Khan. 2020. “Power and Performance Analysis of Deep Neural Networks for Energy-Aware Heterogeneous Systems.” *Conference Proceedings - IEEE International Conference on Systems, Man and Cybernetics* 2020-Octob:2184–89. doi: 10.1109/SMC42975.2020.9283092.

Cho, Kyunghyun, Dzmitry Bahdanau, Fethi Bougares, Holger Schwenk, and Yoshua Bengio. 2014. “Learning Phrase Representations Using RNN Encoder–Decoder for Statistical Machine Translation.” *Journal of Clinical Microbiology* 28(4):828–29. doi: 10.1128/jcm.28.4.828-829.1990.

Choudhary, Prakash, and Sumit Chauhan. 2023. “An Intelligent Chatbot Design and Implementation Model Using Long Short-Term Memory with Recurrent Neural Networks and Attention Mechanism.” *Decision Analytics Journal* 9(May):100359. doi: 10.1016/j.dajour.2023.100359.

Guanabara, Editora, Koogan Ltda, Editora Guanabara, and Koogan Ltda. n.d. “Evrişimli Ve Yinelemeli Sinir Ağları İle Görüntülere Başlık Atama.”

Gulati, Anmol, James Qin, Chung Cheng Chiu, Niki Parmar, Yu Zhang, Jiahui Yu, Wei Han, Shibo Wang, Zhengdong Zhang, Yonghui Wu, and Ruoming Pang. 2020. “Conformer: Convolution-Augmented Transformer for Speech Recognition.” *Proceedings of the Annual Conference of the International Speech Communication Association, INTERSPEECH* 2020-Octob:5036–40. doi: 10.21437/Interspeech.2020-3015.

- GÜLTEPE, Yasemin. 2019. “Makine Öğrenmesi Algoritmaları Ile Hava Kirliliği Tahmini Üzerine Karşılaştırmalı Bir Değerlendirme.” *European Journal of Science and Technology* (16):8–15. doi: 10.31590/ejosat.530347.
- He, Kaiming, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2015. “Deep Residual Learning for Image Recognition.” *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition* 2016-Decem:770–78. doi: 10.1109/CVPR.2016.90.
- He, Ruining, Anirudh Ravula, Bhargav Kanagal, and Joshua Ainslie. 2021. “RealFormer: Transformer Likes Residual Attention.” *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021* 929–43. doi: 10.18653/v1/2021.findings-acl.81.
- Ioannidis, Jules, Joshua Harper, Ming Sheng Quah, and Dan Hunter. 2023. “Gracenote . Ai : Legal Generative AI for Regulatory Compliance.”
- Islam, Saidul, Hanae Elmekki, Ahmed Elsebai, Jamal Bentahar, Nagat Drawel, Gaith Rjoub, and Witold Pedrycz. 2024. “A Comprehensive Survey on Applications of Transformers for Deep Learning Tasks.” *Expert Systems with Applications* 241. doi: 10.1016/j.eswa.2023.122666.
- Jurafsky, Daniel, and James Martin. 2021. “Machine Translation and Encoder-Decoder Models.”
- Jurafsky, Daniel, and James H. Martin. 2023. *Speech and Language Processing An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition*.
- Kora, Rania, and Ammar Mohammed. n.d. “A Comprehensive Review on Transformers Models For Text Classification.” *2023 International Mobile, Intelligent, and Ubiquitous*

- Computing Conference (MIUCC)* 1–7. doi: 10.1109/MIUCC58832.2023.10278387.
- Le, Matthew, Y. Lan Boureau, and Maximilian Nickel. 2019. “Revisiting the Evaluation of Theory of Mind through Question Answering.” *EMNLP-IJCNLP 2019 - 2019 Conference on Empirical Methods in Natural Language Processing and 9th International Joint Conference on Natural Language Processing, Proceedings of the Conference* 5872–77. doi: 10.18653/v1/d19-1598.
- Liu, Vivian, and Lydia B. Chilton. 2022. *Design Guidelines for Prompt Engineering Text-to-Image Generative Models*. Vol. 1. Association for Computing Machinery.
- Lu, Yi, Xin Zhou, Wei He, Jun Zhao, Tao Ji, Tao Gui, Qi Zhang, and Xuanjing Huang. 2024. “LongHeads: Multi-Head Attention Is Secretly a Long Context Processor.”
- Pandey, Sumit, and Srishti Sharma. 2023. “A Comparative Study of Retrieval-Based and Generative-Based Chatbots Using Deep Learning and Machine Learning.” *Healthcare Analytics* 3(April). doi: 10.1016/j.health.2023.100198.
- Pennington, Jeffrey, Richard Socher, and D. Manning Christopher. 2014. “GloVe: Global Vectors for Word Representation.” *British Journal of Neurosurgery* 31(6):682–87. doi: 10.1080/02688697.2017.1354122.
- Radford, Alec, Karthik Narasimhan, Tim Salimans, and Ilya Sutskever. 2018. *Improving Language Understanding by Generative Pre-Training*.
- Radford, Alec, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. “Language Models Are Unsupervised Multitask Learners.”
- Raffel, Colin, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li,

- and Peter J. Liu. 2023. "Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer." *Journal of Machine Learning Research* 21:1–67.
- Shaw, Peter, Jakob Uszkoreit, and Ashish Vaswani. 2018. "Self-Attention with Relative Position Representations." *NAACL HLT 2018 - 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies - Proceedings of the Conference* 2:464–68. doi: 10.18653/v1/n18-2074.
- Shoeybi, Mohammad, Mostofa Patwary, Raul Puri, Patrick LeGresley, Jared Casper, and Bryan Catanzaro. 2020. "Megatron-LM: Training Multi-Billion Parameter Language Models Using Model Parallelism."
- Sonkar, Shashank, and Richard G. Baraniuk. 2023. "Investigating the Role of Feed-Forward Networks in Transformers Using Parallel Attention and Feed-Forward Net Design." (2021):1–6.
- Sun, Yujia, and Jan Platoš. 2024. "Abstractive Text Summarization Model Combining a Hierarchical Attention Mechanism and Multiobjective Reinforcement Learning." *Expert Systems with Applications* 248(August 2023). doi: 10.1016/j.eswa.2024.123356.
- Vaswani, Ashish, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. "Attention Is All You Need." *Advances in Neural Information Processing Systems* 2017-Decem(Nips):5999–6009.
- Vinyals, Oriol, and Quoc Le. 2015. "A Neural Conversational Model." 37.

- Xiong, Bo, Yimin Huang, Hanrong Ye, Steffen Staab, and Zhenguo Li. 2020. “MOFA: Modular Factorial Design for Hyperparameter Optimization.” 1–13.
- Yu, Jiahui, Zirui Wang, Vijay Vasudevan, Legg Yeung, Mojtaba Seyedhosseini, and Yonghui Wu. 2022. “CoCa: Contrastive Captioners Are Image-Text Foundation Models.” (1):1–19.
- Zhang, Yu, Mingxiang Tuo, Qingyu Yin, Le Qi, Xuxiang Wang, and Ting Liu. 2020. “Keywords Extraction with Deep Neural Network Model.” *Neurocomputing* 383:113–21. doi: 10.1016/j.neucom.2019.11.083.

POWERLIFTING (GÜÇ KALDIRMA) SPORCULARININ YARIŞMA KATEGORİLERİNİN GEÇMİŞ SPORCU PERFORMANSLARINA GÖRE BELİRLENMESİ

Ege Kaan SUCU¹

Ümit YILMAZ²

Atınç YILMAZ³

1. GİRİŞ

Powerlifting, göreceli ve mutlak maksimum güç üzerine odaklanan bir spordur. Göreceli güç, bir sporcunun vücut ağırlığına kıyasla ne kadar güçlü olduğunu ifade ederken mutlak güç ise bir kişinin vücut ağırlığından bağımsız olarak kaldırabileceği toplam ağırlık miktarını ifade etmektedir. Sporcular kilo kategorilerine, yaşlarına ve deneyim seviyelerine göre sınıflandırılmaktadır.

Powerlifting, müsabakalarda squat, bench press ve deadlift olmak üzere üç çok eklemlili kaldırışın gerçekleştirildiği dinamik bir güç sporudur. Bir powerlifting müsabakası sırasında, yarışmacılara her bir kaldırışı tam bir hareket aralığında gerçekleştirmeleri için 3 deneme hakkı verilir (Beausejour vd.,

¹ Yüksek Lisans Öğrencisi, İstanbul Beykent Üniversitesi, Lisansüstü Eğitim Enstitüsü, Bilgisayar Mühendisliği Yüksek Lisans Programı, egekaansucu@gmail.com, ORCID: 0000-0000-0000-0000.

² Öğr. Gör. Dr., Balıkesir Üniversitesi, Bigadiç Meslek Yüksekokulu, Yönetim ve Organizasyon Bölümü, umityilmaz@balikesir.edu.tr, ORCID: 0000-0003-4268-8598.

³ Doç. Dr., İstanbul Beykent Üniversitesi, Mühendislik-Mimarlık Fakültesi, Bilgisayar Mühendisliği (Türkçe) Bölümü, atincyilmaz@beykent.edu.tr, ORCID: 0000-0003-0038-7519.

2024: e244). Her bir kaldırma, toplam kaldırma performansının sırasıyla yaklaşık %36, %24 ve %40'ını temsil etmektedir. Belirli değerlendirme kriterlerini kabul eden her sporcunun amacı, yarışma sırasında mümkün olan en ağır yükü kaldırmaktır. Ulusal ve uluslararası alanda raw (ham) ve donanımlı yarışmalar düzenlemekte olup ham yarışmalarda atletin yanı sıra dizlik, kemer ve bilekliklerin giyilmesine izin verilmektedir (Tromaras, Zaras, Stasinaki, Mpampoulis, & Terzis, 2024: 1). Bu çalışma ham powerlifting yarışması hakkındadır. Powerlifting, basit kuralları ve düşük maliyeti nedeniyle uluslararası alanda oldukça popülerdir (Nichols, Pavlidis, & Nowak, 2024: 256).

Modern powerlifting 1960'larda halterden sonra ortaya çıkmıştır. Günümüzde ise powerlifting için ulusal ve uluslararası yarışmalar koordine eden ve düzenleyen birçok federasyon ve organizasyon vardır. Bu nedenle, yarışma kuralları federasyon üyeliğine ve yarışma türüne (örn. ham/klasik ve donanımlı etkinlikler) bağlı olarak biraz farklılık göstermektedir. Bununla birlikte, sporcu performansı sporcunun üst ve alt vücut gücünün maksimum ifadesine dayanmaktadır (van den Hoek vd., 2024: 1-2).

Çalışmada, donanımsız powerlifting yarışmalarından birinden elde edilen verileri kullanarak farklı yaşlardaki ve ağırlık gruplarındaki powerlifting sporcu popülasyonlarında yaş, vücut ağırlığı, squat, bench press, deadlift ve toplam ağırlıktan faydalanarak göreceli kümelerinin oluşturulması ve müsabakalara yeni katılan sporcuların mevcut bilgilerinden faydalanarak belirlenen bu kümelere sınıflandırılmasına ilişkin önerilerde bulunmaktadır. Yapılan analizler neticesinde, powerlifting sporcularının kategorilerinin daha doğru bir tanımının sunulması ve sporcuların performanslarının çağdaş bir şekilde kıyaslanması amaçlanmaktadır.

2. MATERYALLER VE YÖNTEMLER

Bu çalışmada, 2023 Avrupa Powerlifting Federasyonu Avrupa Açık Klasik Powerlifting Şampiyonası'na katılan ve yarışmayı başarıyla tamamlayan 96 erkek sporcunun mevcut yaşları ve vücut ağırlıkları ile yarışmada elde ettikleri en iyi squat (çömelme), en iyi bench press (göğüsten itiş) ağırlığı ve en iyi deadlift (barı yerden kesme) ağırlıkları ile bu üç hareketteki en iyi kaldırışların ağırlıklarının toplamı (en iyi squat ağırlığı + en iyi bench press ağırlığı + en iyi deadlift ağırlığı) verilerinden faydalanılmıştır (OpenPowerlifting, 2023). En iyi squat, bench press ve deadlift ağırlıkları, yarışmacıların her bir harekette gerçekleştirdikleri üç denemeden elde ettikleri en yüksek kaldırışları temsil etmektedir. Powerlifting yarışmalarında genellikle ağırlık, cinsiyet, yaş ve ekipman kategorileri bulunmaktadır. Daha ağır sporcuların daha fazla kas kütesine sahip olması ve daha hafif sporculara kıyasla genellikle daha güçlü olmasının beklenmesi sebebiyle powerlifting sporcuları belirli vücut ağırlığı kategorilerinde yarışmaktadır (Tromaras vd., 2024: 1). Powerlifting yarışmalarında cinsiyet kategorilerinin olması, kadınlar ve erkekler arasındaki fizyolojik farklılıklar nedeniyle gereklidir. Araştırmalar, erkeklerin kadınlara kıyasla genellikle daha fazla kas kütesine ve kalınlığına sahip olduğunu göstermektedir. Bu da erkeklerin daha yüksek güç performansı sergilemelerine neden olmaktadır (Bartolomei, Grillone, Di Michele, & Cortesi, 2021: 1). Powerlifting yarışmalarında, yaşa bağlı fizyolojik farklılıklar nedeniyle yaş kategorilerine ayrılma gerekliliği görülmüştür. Genç sporcuların kas gelişimi ve güçleri, yaşla birlikte değiştiğinden, powerlifting sporcularının yaş kategorilerine ayrılması adil bir değerlendirme için gerekli görülmüştür. Ayrıca, yaş gruplarına ayırmanın, sporda uzun vadeli katılımı teşvik ettiği ve farklı yaş aralıklarından sporcuların rekabet etmesini sağladığı görüşü hakimdir. Bu doğrultuda, sporcuların performanslarının kendi yaş grubunda daha iyi

değerlendirilebileceği ileri sürülmektedir (Latella vd., 2024: 754). Ayrıca yarışma, sadece raw (ham) kaldırışları içermektedir. Yani sporcular şampiyonada squat takımı, bench press gömleği, deadlift takımı ve diz sarma gibi özel destekleyici kıyafet ve ekipmanlar kullanmadan, dizlik, kemer ve bileklik gibi temel ekipmanlarla sadece fiziksel güç ve tekniklerini kullanarak yarışmıştır. Bu sayede yarışmacıların kümelenmesinde, ekipmanların yarışmacıların performansları üzerindeki etkisinin belirsizliği ortadan kaldırılmıştır. Bahsi geçen şampiyonada ağırlık ve cinsiyet kategorilerinde ayrıma gidilmiştir. Şampiyonada erkek sporcular -59 kg, -66 kg, -74 kg, -83 kg, -93 kg, -105 kg, -120 kg ve 120+ kg ağırlık kategorilerine ayrılarak yarışmışlardır.

2.1. Materyaller

Veriler, 18 yaş ve üzeri tüm sporcular analizlere dahil edilerek, yalnızca uyuşturucu testi yapılmış, donanımsız yarışmalardan elde edilen tüm yarışma sonuçlarını içerecek şekilde filtrelenmiştir. Filtreleme sonrası gerçekleştirilen müsabakadan squat, bench press ve deadlift için maksimum başarılı kaldırışlar ile bu başarılı kaldırışların toplamı ve sporcu yaşı ve vücut ağırlığı elde edilmiş ve çalışmanın veri kümesini oluşturmuştur.

2.2. Yöntemler

Çalışmada müsabakaya katılan yarışmacıların kümelenmesi için K-Means kümeleme algoritmasından, kümeleri belli olan yarışmacıların mevcut girdiler yardımıyla doğru bir şekilde sınıflanıp sınıflanmadığını tespit etmek amacıyla da KNN ve Yapay Sinir Ağları (YSA) sınıflandırma algoritmalarından faydalanılmıştır. Geliştirilen KNN modeli yardımıyla modelin eğitiminde ve testinde verileri kullanılmayan 4 yarışma için de ek sınıflandırma işlemi yapılmıştır.

2.2.1.K-Means Kümeleme Algoritması

Kümeleme, veri madenciliği ve makine öğrenmesinde önemli bir konudur ve kümelemenin birincil amacı, verilerin anlamlı yapısını öğrenmek ve verilen verileri farklı gruplara dağıtmaktır (Yu, Wei, Wu, Feng, & Zhang, 2024: 1). K-means kümeleme algoritması, veri kümesindeki farklı kategorileri ayıran özellikleri bulmak için etiketsiz bir eğitim veri kümesi kullanan popüler denetimsiz öğrenme tekniklerinden biridir (Wang vd., 2024: 5). Bu algoritma, veri kümesindeki örnek noktaları arasındaki ana özellikleri etrafında kümelemek için kullanılmaktadır (Tao, Zhang, Chen, Qi, & Liu, 2024: 6). Yinelemeli süreci, her bir veri noktasının en yakın küme merkezine atanmasıyla başlar ve ardından her bir kümeye bağlı tüm noktaların ortalamasıyla belirlenen merkezlerin yeniden hesaplanmasıyla devam eder (Khan vd., 2024: 6).

2.2.2.KNN Sınıflandırma Algoritması

Sınıflandırma, sınıflandırıcının eğitim sırasında topladığı bilgilere dayanarak belirli bir test örneğine bir sınıf atama işlemidir. Görevi, vektör tabanlı bir giriş modelini önceden tanımlanmış birkaç sınıftan birine atamaktır (Deepa vd., 2022: 1823). Bu araştırmada powerlifting sporcularını, kümeleme aşaması sonrasında oluşturulan kümelere sınıflandırmak için K-en yakın komşu (KNN) sınıflandırma algoritması kullanılmıştır. KNN, mesafe metriklerine dayalı olarak hem sınıflandırma hem de regresyon problemlerinde yaygın olarak kullanılan denetimli bir makine öğrenmesi algoritmasıdır (Tsalera, Papadakis, & Samarakou, 2020: 226). KNN, her bir etiketsiz veri noktası ile veri kümesindeki diğer tüm noktalar arasındaki mesafeyi hesaplayarak etiketsiz verileri sınıflandırır. Ardından, veri kümesindeki örüntüleri bularak her bir etiketsiz veri noktasını en benzer şekilde etiketlenmiş verinin sınıfına atar (Almomany, Ayyad, & Jarrah, 2022).

2.2.3. YSA Sınıflandırma Algoritması

YSA'lar, beyindeki biyolojik nöronların davranışını taklit eden yapay nöron katmanlarından oluşan bir hesaplama modelidir. YSA'ların benzersiz doğrusal olmayan uyarlanabilir bilgi işleme yeteneği, doğrusal olmayan ilişkileri en iyi şekilde modelleyerek son derece doğru tahminler sağlar (Yılmaz & Orbak, 2023).

YSA'da ağ, belirli görevleri yerine getirmek için öğrenme kuralları kullanılarak eğitilmektedir. Öğrenme kuralı, ağırlıkların değiştirilmesi süreci olarak tanımlanmaktadır. Bu kural, algoritmaların eğitimi için bir işlemdir. Öğrenme süreci, sınıflandırma için hayati bir rol oynamaktadır ve ağ mimarisinin, bağlantıların ve ağırlıkların güncellenmesini içermektedir (Murawwat, Asif, Ijaz, Imran Malik, & Raahemifar, 2022: 2817).

YSA'lar girdi ve çıktı uzayı arasında karmaşık eşlemeler oluşturabilir ve böylece keyfi olarak karmaşık doğrusal olmayan karar sınırları oluşturabilirler. YSA'lar gerçek dünyadaki karmaşık sınıflandırma problemlerini çözmek için günümüzde en sık kullanılan yöntemlerden biridir (Naik, Nayak, Behera, & Abraham, 2016: 70).

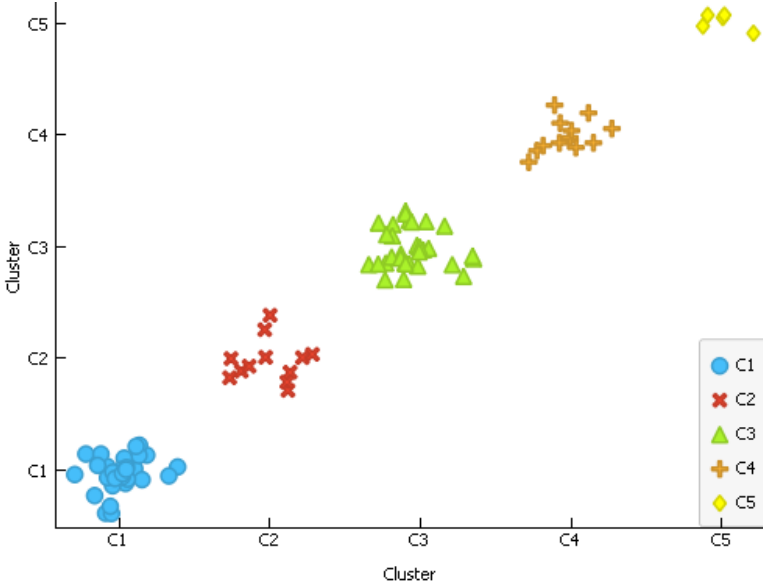
3. BULGULAR VE ANALİZLER

Tüm veriler Orange veri madenciliği programı kullanılarak analiz edilmiştir. Açık kaynaklı olan Orange yazılımı, iş akışlarını görsel olarak oluşturarak makine öğrenmesi, veri madenciliği, veri analizi ve veri görselleştirme işlemlerinin gerçekleştirilmesine olanak sağlamaktadır (Dobesova, 2024: 3-4).

Çalışmada öncelikle müsabakaya katılan yarışmacıların kümelenmesi için K-Means kümeleme algoritmasından yararlanılmıştır. Araştırmada kullanılan veri setindeki sütunlar

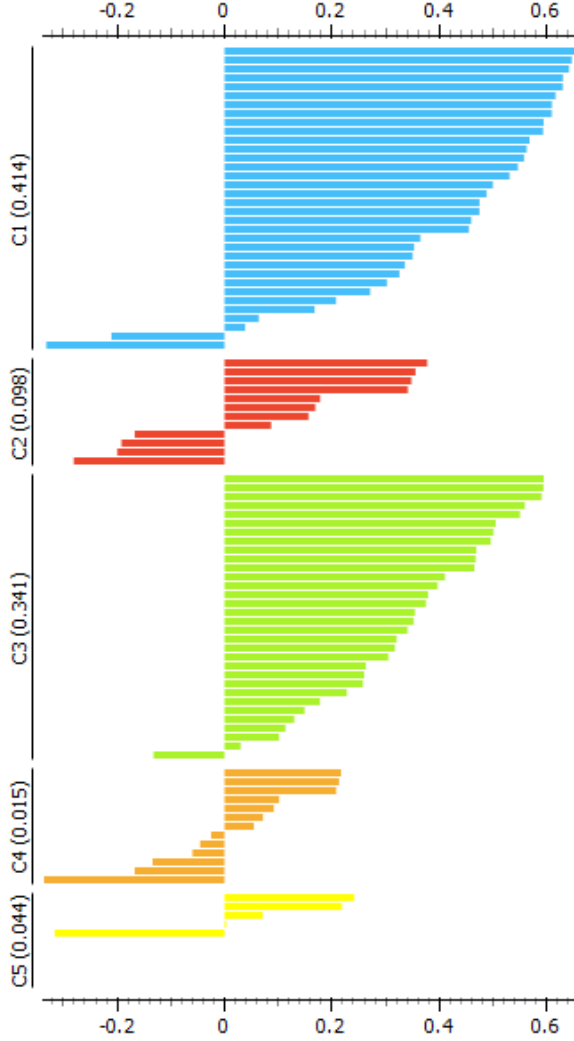
normalize edildikten sonra kümeleme işlemi gerçekleştirilmiştir. Küme sayısı 5 olarak belirlenmiştir. Küme merkezleri ve kümelerle ilişkin bilgiler Şekil 1’de gösterilmiştir.

Şekil 1. Kümelere İlişkin Bilgiler



Kümeleme işlemi sonucunda kümelemenin kalitesini ölçmek için Silhouette grafiğinden faydalanılmıştır. Her nokta için siluet değeri, o noktanın diğer kümelerdeki noktalara kıyasla kendi kümesindeki noktalara ne kadar benzer olduğunun bir ölçüsüdür. Silhouette değeri, komşu kümelerden çok uzak noktaları belirten +1’den, bir kümede veya diğerinde belirgin bir şekilde bulunmayan noktaları gösteren 0’a ve muhtemelen yanlış kümeye atanan noktaları gösteren -1 arasında değişmektedir (Lletí, Ortiz, Sarabia, & Sánchez, 2004). Öklid mesafesi veya Manhattan mesafesi gibi herhangi bir mesafe ölçüsüyle hesaplanabilir. Silhouette değeri, Öklid mesafesi, Manhattan mesafesi ve Cosine mesafesi gibi herhangi bir mesafe ölçüsüyle hesaplanabilmektedir. Şekil 2’de Öklit mesafesi ile hesaplanan bir Silhouette grafiği görülmektedir.

Şekil 2. Silhouette Grafiği



Geliştirilen KNN ve YSA modellerinin çalıştırılması sonucu pozitif kategoriye doğru bir şekilde tahmin eden örnek sayısını DP (doğru pozitif), negatif kategoriye doğru bir şekilde tahmin eden örnek sayısını DN (doğru negatif), pozitif kategoriye yanlış bir şekilde tahmin eden örnek sayısını YP (yanlış pozitif)

ve son olarak negatif kategoriye yanlış bir şekilde tahmin eden örnek sayısını YN (yanlış negatif) temsil etmektedir (Chang, Ganatra, Hall, Golightly, & Xu, 2022).

Yapılan sınıflandırma tahminlerine ilişkin doğruluk, duyarlılık, özgüllük, kesinlik ve F1 skoru gibi bulgulardan faydalanılarak sınıflandırma durumu tahlil edilmiştir. Doğruluk, toplam tahminlerin doğru olan kısmıdır. Kesinlik, gerçekten pozitif sınıfa ait olan pozitif tahminlerin tüm pozitif tahminlere olan oranıdır. Duyarlılık, veri kümesindeki pozitif örneklerin pozitif olarak tahmin edilen kısmıdır. Özgüllük, negatif sınıfın ne kadar iyi tahmin edildiğini özetleyen doğru negatif oranıdır. F1 skoru, kesinlik ve duyarlılığın harmonik ortalamasıdır (Eldin Rashed, Elmorsy, & Mansour Atwa, 2023; Jangam & Annavarapu, 2021). Bu metriklerin hesaplanması için kullanılan formüller sırasıyla Denklem 1-Denklem 5 aralığında sıralanmıştır (Navazi, Yuan, & Archer, 2023; Qian vd., 2024).

$$\text{Doğruluk} = \frac{(DP + DN)}{(DP + DN + YP + YN)} \quad (1)$$

$$\text{Kesinlik} = \frac{DP}{(DP + YP)} \quad (2)$$

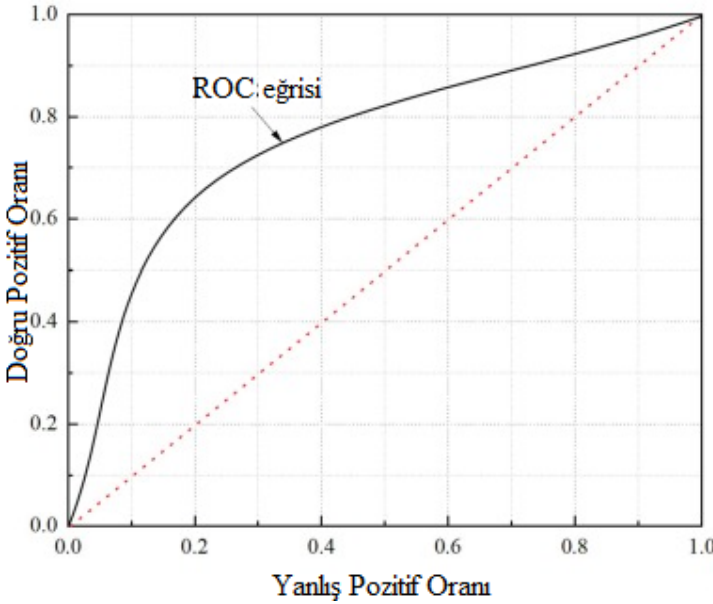
$$\text{Duyarlılık} = \frac{DP}{(DP + YN)} \quad (3)$$

$$\text{Özgüllük} = \frac{DN}{(DN + YP)} \quad (4)$$

$$\text{F1 Skoru} = \frac{(2 * \text{Kesinlik} * \text{Duyarlılık})}{(\text{Kesinlik} + \text{Duyarlılık})} \quad (5)$$

Bu performans ölçüt değerlerinin yanı sıra ROC eğrileri de sınıflandırmanın durumunu göz önüne sürmede etkin bir şekilde kullanılmaktadır. ROC (Receiver Operating Characteristic – Alıcı Çalışma Karakteristiği) eğrisi, çok sınıflı bir sınıflandırıcı modelinin değişen eşik değerlerindeki performansını gösteren bir grafikdir. Şekil 3’te gösterildiği gibi, x ekseninde yanlış pozitif oranı ile y ekseninde doğru pozitif oranına karşı çizilir. Bir sınıflandırıcının sınıflandırma yapabilmesi için 0 ile 1 arasında bir eşik aralığı tanımlanmıştır. Her nokta için yanlış pozitif oranı ve doğru pozitif oranı birbirlerine karşı çizilir. ROC eğrisinin sol üst köşesi iyi sınıflandırmayı, sağ alt köşesi ise kötü sınıflandırmayı temsil etmektedir. Çizimdeki köşegen rastgele tahmini temsil etmektedir. Herhangi bir sınıflandırıcının ROC eğrisi köşegenin altındaysa, o sınıflandırıcı rastgele tahminden daha kötü performans gösteriyor demektir ki bu da amacı tamamen boşa çıkarır (Li vd., 2023: 4-5).

Şekil 3. ROC Eğrisinin Temel Gösterimi



ROC eğrisi, makine öğreniminde karışıklık matrisi adı verilen temel bir matrise dayalı olarak oluşturulmaktadır (Haixiang, Yijing, Yanan, Xiao, & Jinling, 2016: 179). Tablo 1, veri kümelerinin sınıf etiketi tahmininde seçilen sınıflandırıcının ne kadar iyi olduğunu tespit etmede kullanılan faydalı araçlardan biri olan karışıklık matrisini temsil etmektedir (Banerjee & Sarathee Bhowmik, 2025: 8). İkili sınıflandırmanın yanı sıra çok sınıflı sınıflandırma problemleri için de uygulanabilmektedir.

Tablo 1. İkili Sınıflandırma Probleminin Karışıklık Matrisi

	Öngörülen Pozitif Sınıf	Öngörülen Negatif Sınıf
Pozitif Sınıf	DP	YN
Negatif Sınıf	YP	TN

KNN ve YSA sınıflandırma algoritmaları yardımıyla yapılan sınıflandırma neticesinde her bir sınıfa ait doğruluk, kesinlik, duyarlılık, özgülük ve F1 skoru performans ölçüt değerleri Tablo 2’de gösterilmiştir.

Tablo 2. Sınıfların Doğruluk, Kesinlik, Duyarlılık, Özgülük ve F1 Skoru Değerleri

Sınıflar	Doğruluk		Kesinlik		Duyarlılık		Özgülük		F1 Skoru	
Yöntemler	KNN	YSA	KNN	YSA	KNN	YSA	KNN	YSA	KNN	YSA
C1	0.990	0.969	1.000	0.943	0.971	0.971	1.000	0.968	0.985	0.957
C2	0.969	0.969	1.000	0.909	0.750	0.833	1.000	0.988	0.857	0.870
C3	0.958	0.938	0.889	0.906	1.000	0.906	0.938	0.953	0.941	0.906
C4	0.938	0.990	0.733	0.929	0.846	1.000	0.952	0.988	0.786	0.963
C5	0.938	0.990	0.333	1.000	0.200	0.800	0.978	1.000	0.250	0.889
Sınıflar Üzerinden Ortalama	0.896	0.927	0.892	0.927	0.896	0.927	0.971	0.970	0.889	0.926

Veri setinin KNN yardımıyla sınıflandırılması sonucu elde edilen örnek sayısına göre, öngörülen orana göre, gerçekleşen orana göre ve son olarak olasılıkların toplamına göre karışıklık matrisleri sırasıyla Tablo 3, Tablo 4, Tablo 5 ve Tablo 6’da gösterilmiştir. Tablo 4’te mor renkli değerler doğru tanımlanan veri kümelerinin yüzdesini, pembe renkliler ise yanlış tanımlanan veri kümelerinin yüzdesini göstermektedir. Karışıklık

matrisinden, verilerin %89,6'sının her bir sınıfta gerçekten sınıflandırılabilirdiğini ve aynı değerin C1 ve C2 sınıflarında %100 olduğu gözlemlenmiştir.

Tablo 3. Örnek Sayısına Göre Karışıklık Matrisi Sonuçları

		Öngörülen					
		C1	C2	C3	C4	C5	Σ
Gerçekleşen	C1	33	0	1	0	0	34
	C2	0	9	3	0	0	12
	C3	0	0	32	0	0	32
	C4	0	0	0	11	2	13
	C5	0	0	0	4	1	5
Σ		33	9	36	15	3	96

Tablo 4. Öngörülen Orana Göre Karışıklık Matrisi Sonuçları

		Öngörülen					
		C1	C2	C3	C4	C5	Σ
Gerçekleşen	C1	100,0%	0,0%	2,8%	0,0%	0,0%	34
	C2	0,0%	100,0%	8,3%	0,0%	0,0%	12
	C3	0,0%	0,0%	88,9%	0,0%	0,0%	32
	C4	0,0%	0,0%	0,0%	73,3%	66,7%	13
	C5	0,0%	0,0%	0,0%	26,7%	33,3%	5
	Σ	33	9	36	15	3	96

Tablo 5. Gerçekleşen Orana Göre Karışıklık Matrisi Sonuçları

		Öngörülen					Σ
		C1	C2	C3	C4	C5	
Gerçekleşen	C1	97,1%	0,0%	2,9%	0,0%	0,0%	34
	C2	0,0%	75,0%	25,0%	0,0%	0,0%	12
	C3	0,0%	0,0%	100,0%	0,0%	0,0%	32
	C4	0,0%	0,0%	0,0%	84,6%	15,4%	13
	C5	0,0%	0,0%	0,0%	80,0%	20,0%	5
	Σ	33	9	36	15	3	96

Tablo 6. Olasılıkların Toplamına Göre Karışıklık Matrisi

		Öngörülen					
		C1	C2	C3	C4	C5	Σ
Gerçekleşen	C1	32,7	0,0	1,3	0,0	0,0	34
	C2	0,0	8,4	3,6	0,0	0,0	12
	C3	1,6	0,7	29,7	0,0	0,0	32
	C4	0,0	0,0	0,0	9,3	3,7	13
	C5	0,0	0,0	0,0	4,0	1,0	5
Σ		34	9	35	13	5	96

Veri setinin YSA yardımıyla sınıflandırılması sonucu elde edilen örnek sayısına göre, öngörülen orana göre, gerçekleşen orana göre ve son olarak olasılıkların toplamına göre karışıklık matrisleri sırasıyla Tablo 7, Tablo 8, Tablo 9 ve Tablo 10’da gösterilmiştir. Tablo 8’de mor renkli değerler doğru tanımlanan veri kümelerinin yüzdesini, pembe renkliler ise yanlış tanımlanan veri kümelerinin yüzdesini göstermektedir. Karışıklık matrisinden, verilerin %92,7’sinin her bir sınıfta gerçekten sınıflandırılabilildiğini ve aynı değer C5 sınıfında %100 olduğu gözlemlenmiştir.

Tablo 7. Örnek Sayısına Göre Karışıklık Matrisi Sonuçları

		Öngörülen					
		C1	C2	C3	C4	C5	Σ
Gerçekleşen	C1	33	0	1	0	0	34
	C2	0	10	2	0	0	12
	C3	2	1	29	0	0	32
	C4	0	0	0	13	0	13
	C5	0	0	0	1	4	5
	Σ	35	11	32	14	4	96

Tablo 8. Öngörülen Orana Göre Karışıklık Matrisi Sonuçları

		Öngörülen					
		C1	C2	C3	C4	C5	Σ
Gerçekleşen	C1	94,3%	0,0%	3,1%	0,0%	0,0%	34
	C2	0,0%	90,9%	6,2%	0,0%	0,0%	12
	C3	5,7%	9,1%	90,6%	0,0%	0,0%	32
	C4	0,0%	0,0%	0,0%	92,9%	0,0%	13
	C5	0,0%	0,0%	0,0%	7,1%	100,0%	5
	Σ	35	11	32	14	4	96

Tablo 9. Gerçekleşen Orana Göre Karışıklık Matrisi Sonuçları

		Öngörülen					
		C1	C2	C3	C4	C5	Σ
Gerçekleşen	C1	97,1%	0,0%	2,9%	0,0%	0,0%	34
	C2	0,0%	83,3%	16,7%	0,0%	0,0%	12
	C3	6,2%	3,1%	90,6%	0,0%	0,0%	32
	C4	0,0%	0,0%	0,0%	100,0%	0,0%	13
	C5	0,0%	0,0%	0,0%	20,0%	80,0%	5
	Σ	35	11	32	14	4	96

Tablo 10. Olasılıkların Toplamına Göre Karışıklık Matrisi

		Öngörülen					
		C1	C2	C3	C4	C5	Σ
Gerçekleşen	C1	26,2	0,4	3,5	2,8	1,2	34
	C2	0	9,2	2,8	0	0	12
	C3	3,3	3,1	25,2	0,1	0,3	32
	C4	1,9	0	0	10,4	0,7	13
	C5	0,2	0	0	1,1	3,6	5
	Σ	32	13	31	14	6	96

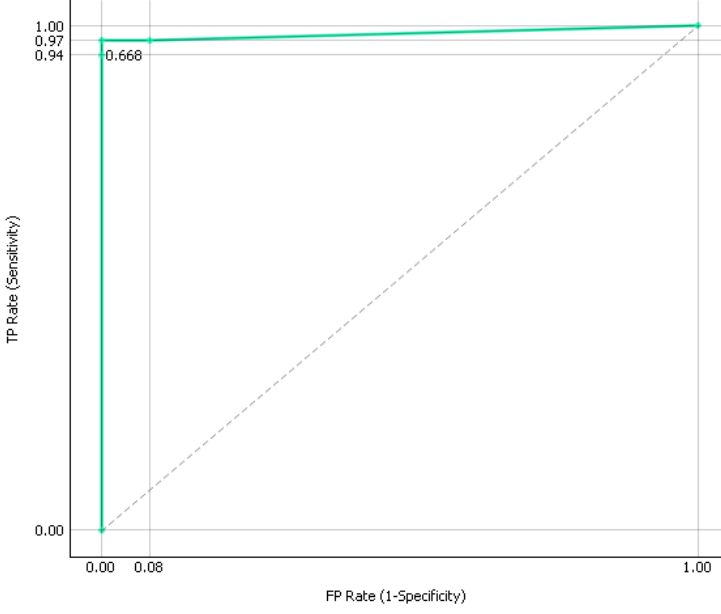
Gerçekleştirilen sınıflandırma işlemleri neticesinde tüm sınıflara ait AUC (Area Under Curve – Eğri Altında Kalan Alan) değerleri Tablo 11’de sıralanmıştır.

Tablo 11. Sınıfların AUC Değerleri

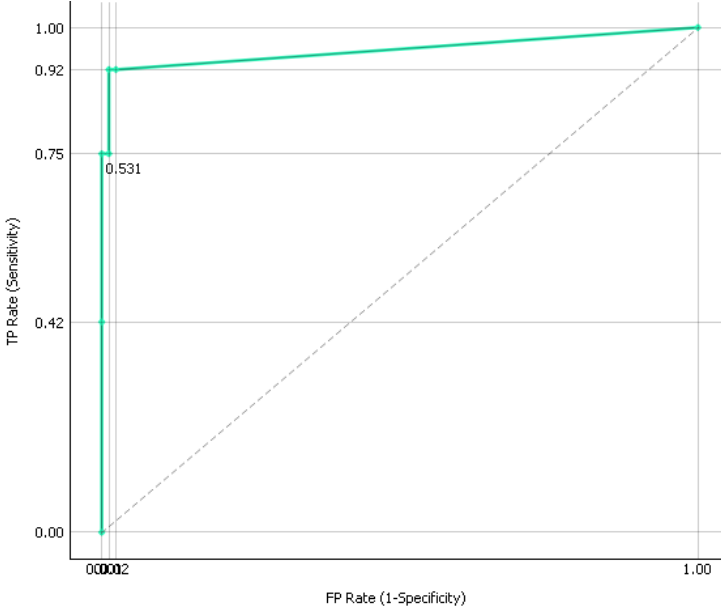
Sınıflar	AUC	
	KNN	YSA
C1	0.986	1.000
C2	0.967	0.988
C3	0.978	0.980
C4	0.962	0.994
C5	0.656	1.000
Sınıflar Üzerinden Ortalama	0.952	0.990

Beş sınıfın her birinin ROC eğrileri KNN sınıflandırma işlemi neticesinde elde edilmiş ve sırasıyla Şekil 4, Şekil 5, Şekil 6, Şekil 7 ve Şekil 8’de gösterilmiştir.

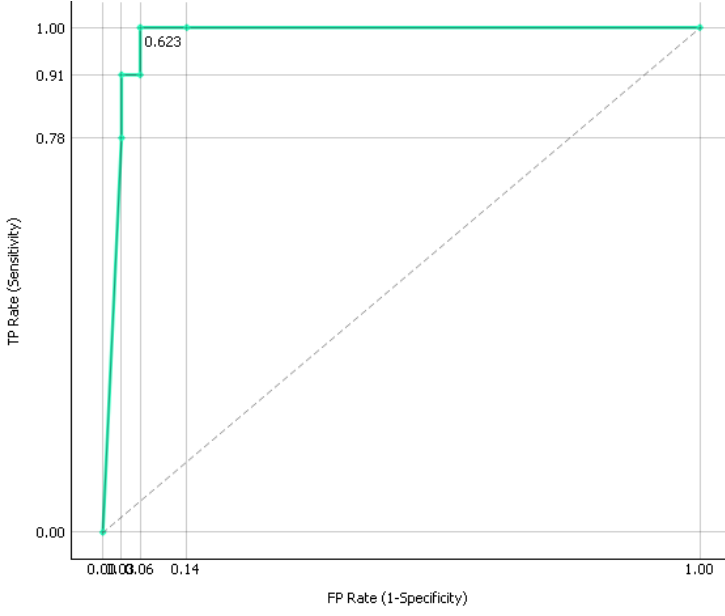
Şekil 4. C1 Kümesi için ROC Grafiği (KNN)



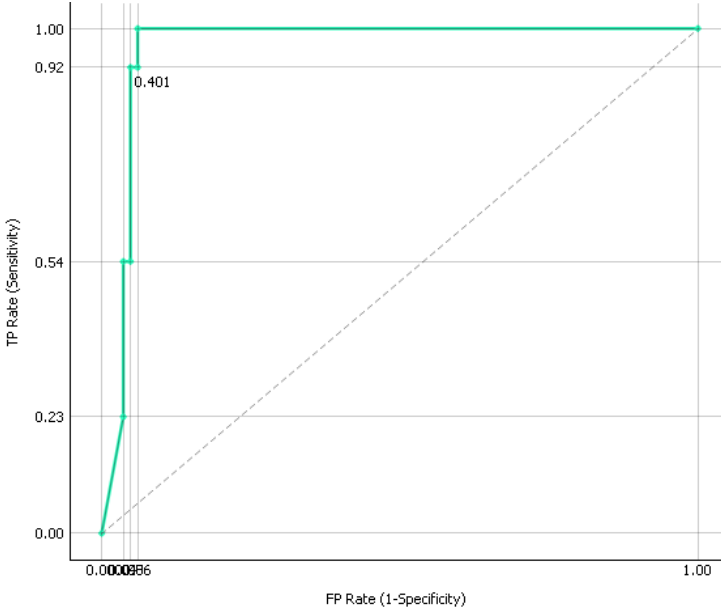
Şekil 5. C2 Kümesi için ROC Grafiği (KNN)



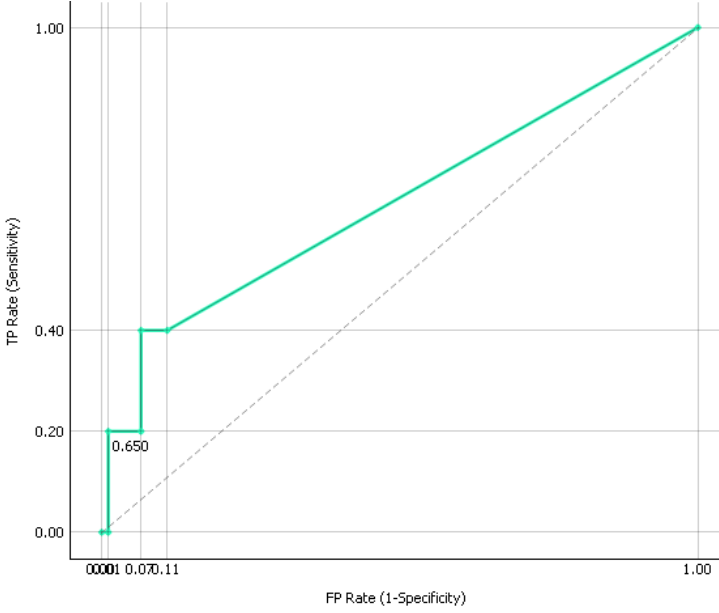
Şekil 6. C3 Kümesi için ROC Grafiği (KNN)



Şekil 7. C4 Kümesi için ROC Grafiği (KNN)



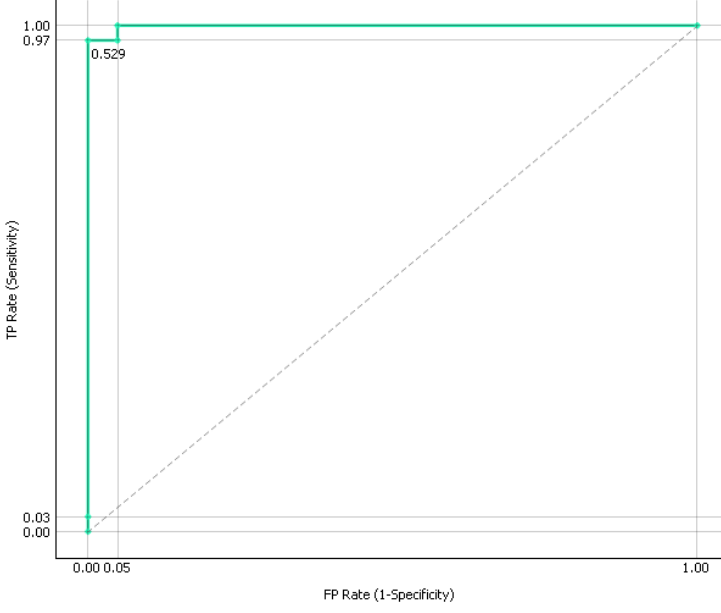
Şekil 8. C5 Kümesi için ROC Grafiği (KNN)



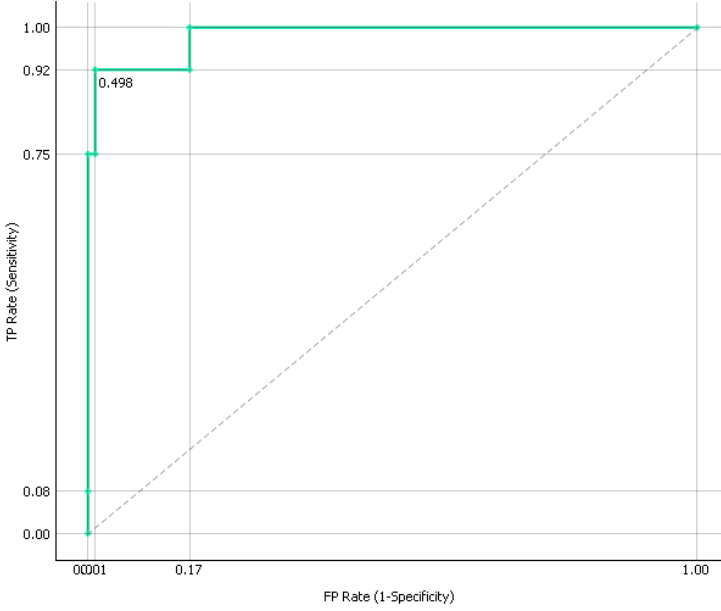
Beş sınıftan dördünün eğri alanı şekillerde ve tabloda görüldüğü gibi %96-%99 aralığında gerçekleşmiştir. Eğri alanı yalnızca C5 sınıfında %66 bulunmuştur. Sınıflar üzerinden ortalama AUC değeri ise %95 elde edilmiştir. Buradan hareketle KNN sınıflandırıcısının ele alınan probleme ait sınıfları sınıflandırmak için oldukça iyi bir performans sergilediğini göstermektedir.

Beş sınıfın her birinin ROC eğrileri YSA sınıflandırma işlemi neticesinde elde edilmiş ve sırasıyla Şekil 9, Şekil 10, Şekil 11, Şekil 12 ve Şekil 13'te gösterilmiştir.

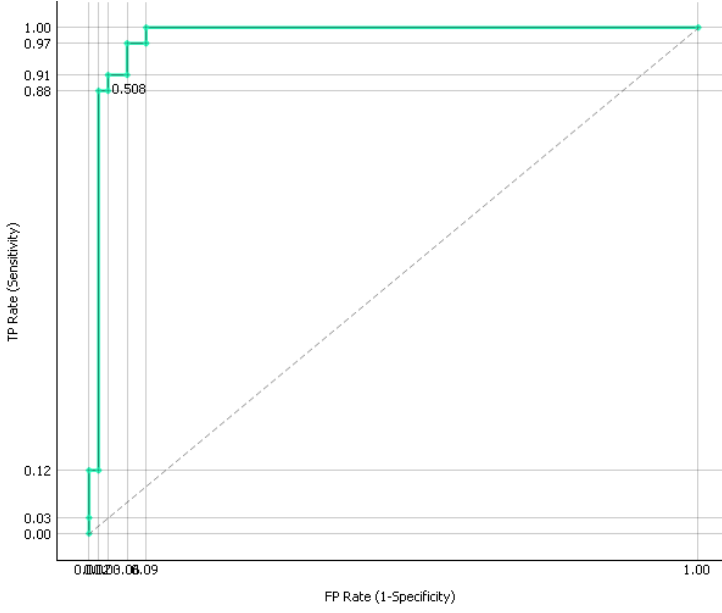
Şekil 9. C1 Kümesi için ROC Grafiği (YSA)



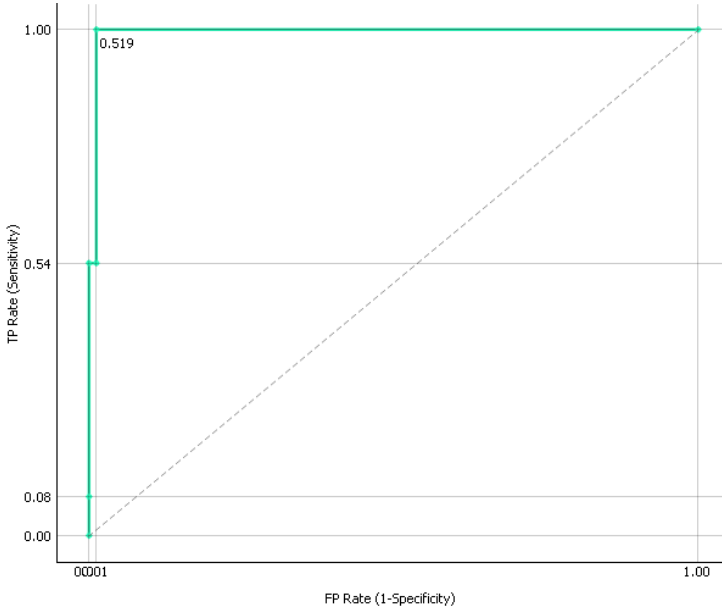
Şekil 10. C2 Kümesi için ROC Grafiği (YSA)



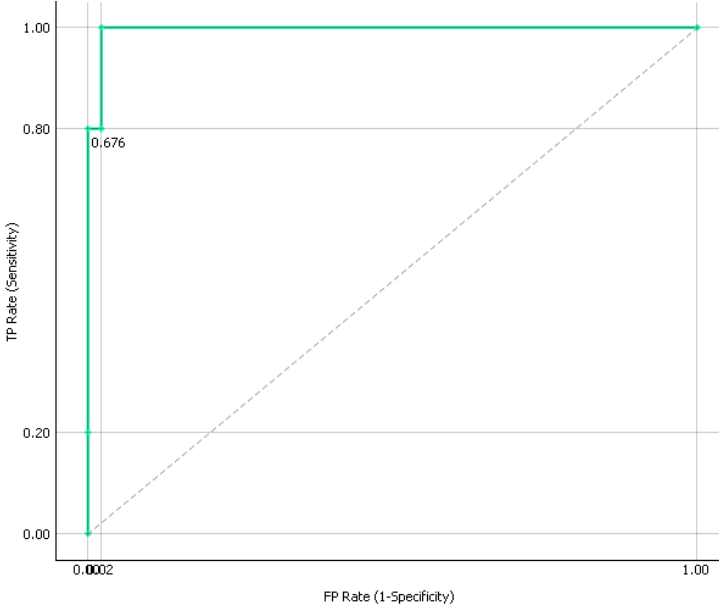
Şekil 11. C3 Kümesi için ROC Grafiği (YSA)



Şekil 12. C4 Kümesi için ROC Grafiği (YSA)



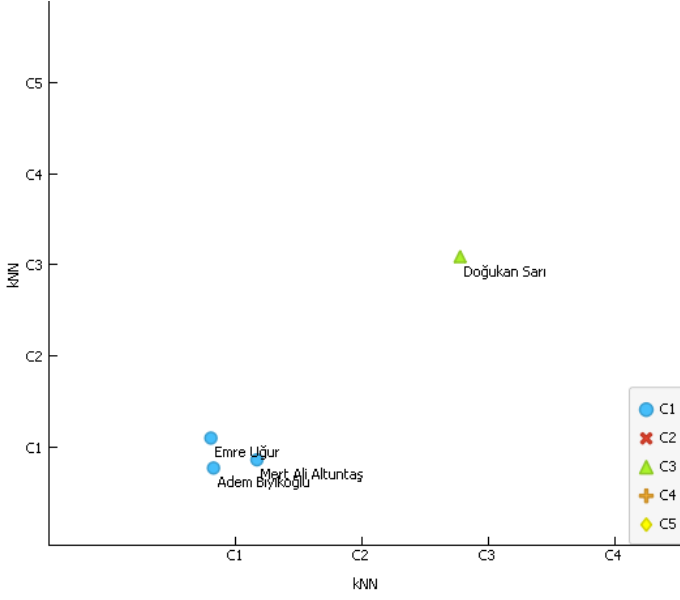
Şekil 13. C5 Kümesi için ROC Grafiği (YSA)



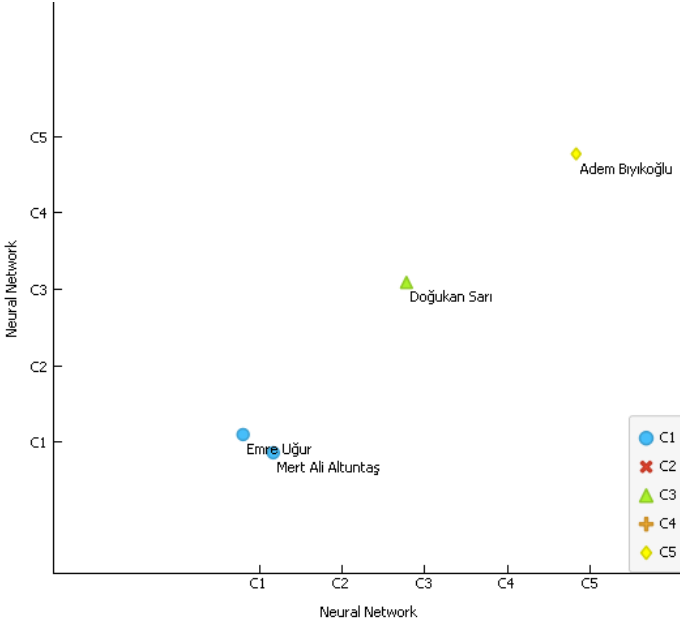
Beş sınıftan dördünün eğri alanı şekillerde ve tabloda görüldüğü gibi %98-%100 aralığında gerçekleşmiştir. Sınıflar üzerinden ortalama AUC değeri ise %99 elde edilmiştir. Buradan hareketle YSA sınıflandırıcısının ele alınan probleme ait sınıfları sınıflandırmak için çok iyi bir performans sergilediğini göstermektedir.

Geliştirilen KNN ve YSA modellerinin etkinliğinin gözler önüne serilmesi amacıyla dikkate alınan müsabakaya katılmayan 4 adet sporcunun kategorizasyonu yapılmıştır. Kategorizasyon sonuçları sırasıyla Şekil 14 ve Şekil 15'te gösterilmiştir.

Şekil 14. Yeni Sporcuların Atandığı Sınıflar (KNN)



Şekil 15. Yeni Sporcuların Atandığı Sınıflar (YSA)



4. SONUÇLAR

Powerlifting sporcularının kategorizasyonu geleneksel normlara göre belirlenmiştir. Bu normal bir amaca hizmet etse de powerlifting sporcularını doğru bir şekilde kategorize etmediği inanılmaktadır. Çalışmada, powerlifting sporcularının doğru bir şekilde kümelenmesi ve bu kümelere doğru şekilde sınıflandırılması, bu spor dalının kategorilerinin oluşturulmasında bir norm önerisinde bulunulmuştur. Oluşturulan bu normlar powerlifting sporcularında gücün doğru bir şekilde kıyaslanması ve değerlendirilmesine olanak sağlayacaktır. Özetle çalışma, powerlifting sporcularının kuvvetlerinin benzer bir popülasyonla kıyaslanabilmesi için güncellenmiş referans kategoriler sağlamaktadır.

Çalışmada, 96 adet sporcu K-Means ile kümelenmek için seçilmiştir. Gerçekleştirilen işlem sonrasında her bir sporcunun kategorisi belirlenmiş ve KNN ve YSA ile sınıflandırmaya tabi tutulmuştur. İşlem sonucunda sporcuların kümelendikleri kategoriye KNN ile %89,6 doğruluk oranıyla, YSA ile %92,7 doğruluk oranıyla sınıflandırıldığı tespit edilmiştir. Bu oran, yarışmaya katılmaya hazırlanan sporcuların hangi kategoride yarışacaklarının sağlıklı bir şekilde belirlenmesine olanak sağlayacağı ve yarışmacı performanslarının kıyaslanmasında etkili olabileceği inancına dayanak oluşturmaktadır. Son olarak da geliştirilen KNN ve YSA modelleri ile müsabakaya katılmayan 4 adet sporcunun kategorizasyonu yapıp modellerin çalıştığı gözler önüne serilmiştir.

Bu araştırmanın birçok güçlü yönü olmasına rağmen, verilerin geçmişe dönük olarak toplanmış olması ve yapılan analizlerin ve sonuçların bazı sınırlamalara sahip olabileceği göz önünde bulundurulmalıdır. Araştırma, yalnızca 2023 Avrupa Powerlifting Federasyonu Avrupa Açık Klasik Powerlifting Şampiyonası'na katılan ve yarışmayı başarıyla tamamlayan 96

erkek sporcuya ilişkin verileri incelemektedir ve yarışmacıların katılmış olabileceği diğer ulusal ve/veya uluslararası uluslararası yarışmaları hesaba katmamaktadır. Gelecekteki çalışmalarda aktif powerlifting sporcularının yarıştıkları tüm müsabakalardaki performanslarını içerecek bir veri tabanı hazırlanarak oluşturulan normların daha da geliştirilmesi mümkün olacaktır.

KAYNAKÇA

- Almomany, A., Ayyad, W. R., & Jarrah, A. (2022). Optimized implementation of an improved KNN classification algorithm using Intel FPGA platform: Covid-19 case study. *Journal of King Saud University - Computer and Information Sciences*, 34(6, Part B), 3815-3827. doi:<https://doi.org/10.1016/j.jksuci.2022.04.006>
- Banerjee, S., & Sarathée Bhowmik, P. (2025). Machine learning based classifiers for dynamic and transient disturbance classification in smart microgrid system. *Measurement*, 240, 115576. doi:<https://doi.org/10.1016/j.measurement.2024.115576>
- Bartolomei, S., Grillone, G., Di Michele, R., & Cortesi, M. (2021). A Comparison between Male and Female Athletes in Relative Strength and Power Performances. *Journal of Functional Morphology and Kinesiology*, 6(1), 17. Retrieved from <https://www.mdpi.com/2411-5142/6/1/17>
- Beausejour, J. P., Guinto, G., Artrip, C., Corvalan, A., Furtado Mesa, M., Lebron, M. A., & Stock, M. S. (2024). Successful Powerlifting in a Unilateral, Transtibial Amputee: A Descriptive Case Series. *The Journal of Strength & Conditioning Research*, 38(5), e243-e252. doi:10.1519/jsc.0000000000004733
- Chang, V., Ganatra, M. A., Hall, K., Golightly, L., & Xu, Q. A. (2022). An assessment of machine learning models and algorithms for early prediction and diagnosis of diabetes using health indicators. *Healthcare Analytics*, 2, 100118. doi:<https://doi.org/10.1016/j.health.2022.100118>
- Deepa, G., Mary, G. L. R., Karthikeyan, A., Rajalakshmi, P., Hemavathi, K., & Dharanisri, M. (2022). Detection of brain tumor using modified particle swarm optimization

- (MPSO) segmentation via haralick features extraction and subsequent classification by KNN algorithm. *Materials Today: Proceedings*, 56, 1820-1826. doi:<https://doi.org/10.1016/j.matpr.2021.10.475>
- Dobesova, Z. (2024). Evaluation of Orange data mining software and examples for lecturing machine learning tasks in geoinformatics. *Computer Applications in Engineering Education*, 32(4), e22735. doi:<https://doi.org/10.1002/cae.22735>
- Eldin Rashed, A. E., Elmorsy, A. M., & Mansour Atwa, A. E. (2023). Comparative evaluation of automated machine learning techniques for breast cancer diagnosis. *Biomedical Signal Processing and Control*, 86, 105016. doi:<https://doi.org/10.1016/j.bspc.2023.105016>
- Haixiang, G., Yijing, L., Yanan, L., Xiao, L., & Jinling, L. (2016). BPSO-Adaboost-KNN ensemble learning algorithm for multi-class imbalanced data classification. *Engineering Applications of Artificial Intelligence*, 49, 176-193. doi:<https://doi.org/10.1016/j.engappai.2015.09.011>
- Jangam, E., & Annavarapu, C. S. R. (2021). A stacked ensemble for the detection of COVID-19 with high recall and accuracy. *Computers in Biology and Medicine*, 135, 104608. doi:<https://doi.org/10.1016/j.combiomed.2021.104608>
- Khan, F., Khan, O., Parvez, M., Ahmad, S., Yahya, Z., Alhodaib, A., . . . Ağbulut, Ü. (2024). K-means clustering optimization of various quantum dots and nanoparticles-added biofuels for engine performance, emission, vibration, and noise characteristics. *Thermal Science and Engineering Progress*, 54, 102815. doi:<https://doi.org/10.1016/j.tsep.2024.102815>

- Latella, C., van den Hoek, D., Wolf, M., Androulakis-Korakakis, P., Fisher, J. P., & Steele, J. (2024). Using Powerlifting Athletes to Determine Strength Adaptations Across Ages in Males and Females: A Longitudinal Growth Modelling Approach. *Sports Medicine*, 54(3), 753-774. doi:10.1007/s40279-023-01962-6
- Li, Y., Zou, Z., Zhang, J., He, Y., Huang, G., & Li, J. (2023). Refined evaluation methods for preventive maintenance of project-level asphalt pavement based on confusion-regression model. *Construction and Building Materials*, 403, 133105. doi:https://doi.org/10.1016/j.conbuildmat.2023.133105
- Lletí, R., Ortiz, M. C., Sarabia, L. A., & Sánchez, M. S. (2004). Selecting variables for k-means cluster analysis by using a genetic algorithm that optimises the silhouettes. *Analytica Chimica Acta*, 515(1), 87-100. doi:https://doi.org/10.1016/j.aca.2003.12.020
- Murawwat, S., Asif, H. M., Ijaz, S., Imran Malik, M., & Raahemifar, K. (2022). Denoising and classification of Arrhythmia using MEMD and ANN. *Alexandria Engineering Journal*, 61(4), 2807-2823. doi:https://doi.org/10.1016/j.aej.2021.08.014
- Naik, B., Nayak, J., Behera, H. S., & Abraham, A. (2016). A self adaptive harmony search based functional link higher order ANN for non-linear data classification. *Neurocomputing*, 179, 69-87. doi:https://doi.org/10.1016/j.neucom.2015.11.051
- Navazi, F., Yuan, Y., & Archer, N. (2023). An examination of the hybrid meta-heuristic machine learning algorithms for early diagnosis of type II diabetes using big data feature selection. *Healthcare Analytics*, 4, 100227. doi:https://doi.org/10.1016/j.health.2023.100227

- Nichols, E., Pavlidis, A., & Nowak, R. (2024). "It's like Lifting the Power": Powerlifting, Digital Gendered Subjectivities, and the Politics of Multiplicity. *Leisure Sciences*, 46(3), 254-273. doi:10.1080/01490400.2021.1945982
- OpenPowerlifting. (2023). *2023 EPF European Open Classic Powerlifting Championships*. Retrieved from: <https://www.openpowerlifting.org/m/epf/2310>
- Qian, L., Bai, J., Huang, Y., Zeebaree, D. Q., Saffari, A., & Zebari, D. A. (2024). Breast cancer diagnosis using evolving deep convolutional neural network based on hybrid extreme learning machine technique and improved chimp optimization algorithm. *Biomedical Signal Processing and Control*, 87, 105492. doi:<https://doi.org/10.1016/j.bspc.2023.105492>
- Tao, Y., Zhang, Q., Chen, Q., Qi, C., & Liu, Y. (2024). An intelligent approach using micro-seismic monitoring signal clustering and an optimized K-means model to guide the selection of support patterns in underground mines. *Tunnelling and Underground Space Technology*, 154, 106095. doi:<https://doi.org/10.1016/j.tust.2024.106095>
- Tromaras, K., Zaras, N., Stasinaki, A.-N., Mpampoulis, T., & Terzis, G. (2024). Lean Body Mass, Muscle Architecture and Powerlifting Performance during Preseason and in Competition. *Journal of Functional Morphology and Kinesiology*, 9(2), 89. Retrieved from <https://www.mdpi.com/2411-5142/9/2/89>
- Tsalera, E., Papadakis, A., & Samarakou, M. (2020). Monitoring, profiling and classification of urban environmental noise using sound characteristics and the KNN algorithm.

Energy Reports, 6, 223-230.
doi:<https://doi.org/10.1016/j.egy.2020.08.045>

- van den Hoek, D. J., Mallard, A., Garrett, J. M., Beaumont, P. L., Howells, R. J., Spathis, J. G., . . . Latella, C. (2024). Powerlifting participation and engagement across all ages: A retrospective, longitudinal, population analysis with comparison to community strength norms. *International Journal of Sports Science & Coaching*, 0(0), 1-11. doi:10.1177/17479541241244481
- Wang, X., Li, K., Jiang, J., Cui, X., Pan, Y., & Xiong, K. (2024). Detecting natural gas storage microleakage based on K-means clustering under constraint of Jeffries-Matusita distance criterion using mobile LiDAR data. *Journal of Environmental Management*, 370, 122539. doi:<https://doi.org/10.1016/j.jenvman.2024.122539>
- Yılmaz, Ü., & Orbak, Â. Y. (2023). Prediction of Turkish mutual funds' net asset value using the fund portfolio distribution. *Neural Computing and Applications*, 35(26), 18873-18890. doi:10.1007/s00521-023-08716-5
- Yu, Y.-F., Wei, P., Wu, X., Feng, Q., & Zhang, C. (2024). Discriminative fuzzy K-means clustering with local structure preservation for high-dimensional data. *Knowledge-Based Systems*, 304, 112537. doi:<https://doi.org/10.1016/j.knosys.2024.112537>

BİLGİSAYAR BİLİMLERİ VE MÜHENDİSLİĞİ

yaz
yayınları

YAZ Yayınları
M.İhtisas OSB Mah. 4A Cad. No:3/3
İscehisar / AFYONKARAHİSAR
Tel : (0 531) 880 92 99
yazyayinlari@gmail.com • www.yazyayinlari.com