



**YAPAY ZEKÂ VE TOPLUMSAL DÖNÜŞÜM:
ETİK İHLALLER**

Rabia KILIÇ

Yüksek Lisans Tezi

Radyo Televizyon Sinema Ana Bilim Dalı

Doç. Dr. Zeynep DEMİRCİOĞLU BİRİCİK

2025

(Her Hakkı Saklıdır.)

T.C
ATATÜRK ÜNİVERSİTESİ
SOSYAL BİLİMLER ENSTİTÜSÜ
RADYO TELEVİZYON SİNEMA ANA BİLİM DALI

Rabia KILIÇ

YAPAY ZEKÂ VE TOPLUMSAL DÖNÜŞÜM:
ETİK İHLALLER

YÜKSEK LİSANS TEZİ

TEZ DANIŞMANI
Doç. Dr. Zeynep DEMİRCİOĞLU BİRİCİK

ERZURUM – 2025



TEZ BEYAN FORMU

SOSYAL BİLİMLER ENSTİTÜSÜ MÜDÜRLÜĞÜNE

BİLDİRİM

Atatürk Üniversitesi Lisansüstü Eğitim ve Öğretim Uygulama Esaslarının ilgili maddelerine göre hazırlamış olduğum "YAPAY ZEKÂ VE TOPLUMSAL DÖNÜŞÜM: ETİK İHLALLER" adlı tezin/raporun tamamen kendi çalışmam olduğunu ve her alıntıya kaynak gösterdiğimi taahhüt eder, tezimin/raporumun kâğıt ve elektronik kopyalarının aşağıda belirttiğim koşullarda saklanmasına izin verdiğimi onaylarım.

Gereğini bilgilerinize arz ederim *.

☐ Tezimin/Raporumun tamamı her yerden erişime açılabilir.

☒ Tezimin/Raporumun makale için **altı ay**, patent için **iki yıl** süreyle erişiminin ertelenmesini istiyorum.

19/06/2025

Aslı ıslak imzalıdır

Rabia KILIÇ

*** LİSANSÜSTÜ TEZLERİN ELEKTRONİK ORTAMDA TOPLANMASI, DÜZENLENMESİ VE ERİŞİME AÇILMASINA İLİŞKİN YÖNERGE**

ÜÇÜNCÜ BÖLÜM

Çeşitli ve Son Hükümler

Lisansüstü tezlerin erişime açılmasının ertelenmesi **MADDE 6– (1)** Lisansüstü teze ilgili patent başvurusu yapılması veya patent alma sürecinin devam etmesi durumunda, tez danışmanının önerisi ve enstitü anabilim dalının uygun görüşü üzerine enstitü veya fakülte yönetim kurulu iki yıl süre ile tezin erişime açılmasının ertelenmesine karar verebilir.

(2) Yeni teknik, materyal ve metotların kullanıldığı, henüz makaleye dönüşmemiş veya patent gibi yöntemlerle korunmamış ve internetten paylaşılması durumunda 3. şahıslara veya kurumlara haksız kazanç imkanı oluşturabilecek bilgi ve bulguları içeren tezler hakkında tez danışmanının önerisi ve enstitü anabilim dalının uygun görüşü üzerine enstitü veya fakülte yönetim kurulunun gerekçeli kararı ile altı ayı aşmamak üzere tezin erişime açılması engellenebilir.

Gizlilik dereceli tezler MADDE 7– (1) Ulusal çıkarları veya güvenliği ilgilendiren, emniyet, istihbarat, savunma ve güvenlik, sağlık vb. konulara ilişkin lisansüstü tezlerle ilgili gizlilik kararı, tezin yapıldığı kurum tarafından verilir. Kurum ve kuruluşlarla yapılan işbirliği protokolü çerçevesinde hazırlanan lisansüstü tezlere ilişkin gizlilik kararı ise, ilgili kurum ve kuruluşun önerisi ile enstitü veya fakültenin uygun görüşü üzerine üniversite yönetim kurulu tarafından verilir. Gizlilik kararı verilen tezler Yükseköğretim Kuruluna bildirilir.

(2) Gizlilik kararı verilen tezler gizlilik süresince enstitü veya fakülte tarafından gizlilik kuralları çerçevesinde muhafaza edilir, gizlilik kararının kaldırılması halinde Tez Otomasyon Sistemine yüklenir.



SOSYAL BİLİMLER ENSTİTÜSÜ

Graduate School of Social Sciences

TEZ KABUL TUTANAĞI

SOSYAL BİLİMLER ENSTİTÜSÜ MÜDÜRLÜĞÜNE

Doç. Dr. Zeynep DEMİRCİOĞLU BİRİCİK danışmanlığında, Rabia KILIÇ tarafından hazırlanan bu çalışma 20 / 05 / 2025 tarihinde aşağıda isimleri yazılı jüri tarafından, Radyo TV ve Sinema Ana Bilim Dalı'nda Yüksek Lisans Tezi olarak kabul edilmiştir.

Başkan: Doç. Dr. Zeynep DEMİRCİOĞLU BİRİCİK	İmza:	Aslı ıslak imzalıdır
Jüri Üyesi : Doç. Dr. İrfan HİDİROĞLU	İmza:	Aslı ıslak imzalıdır
Jüri Üyesi : Dr. Öğr. Üyesi Zühal ERGÜN	İmza	Aslı ıslak imzalıdır

Prof. Dr. Hanifi ŞAHİN
Enstitü Müdürü

İÇİNDEKİLER

İÇİNDEKİLER	IV
ÖZET.....	VII
ABSTRACT	VIII
ÖN SÖZ.....	IX
ŞEKİLLER DİZİNİ	X
TABLolar DİZİNİ	XI

GİRİŞ

BİRİNCİ BÖLÜM

1.1. YAPAY ZEKÂ	3
1.2. YAPAY ZEKÂ'NIN TARİHİ GELİŞİMİ	4
1.3. TEMEL YAPAY ZEKÂ TEKNİKLERİ	8
1.3.1. Doğal Dil İşleme (Natural Language Processing - NLP).....	8
1.3.2. Görüntü İşleme (Image Processing).....	8
1.3.3. Doğal Dil Üretimi (Natural Language Generation - NLG).....	8
1.3.4. Veri Madenciliği (Data Mining)	9
1.3.5. Uzaktan Algılama (Remote Sensing)	9
1.3.6. Uzman Sistemler	9
1.3.7. Bulanık Mantık.....	10
1.3.8. Genetik Algoritmalar.....	12
1.3.9. Yapay Sinir Ağları	13
1.3.9.1. Yapay Sinir Ağlarının Özellikleri	16
1.3.9.1.1. Öğrenme Yeteneği	17
1.3.9.1.2. Uyarlanabilirlik.....	18
1.3.9.1.3. Hata Toleransı.....	18

1.3.9.1.4. Donanım ve Hız.....	18
1.3.9.1.5. Genelleme Yeteneği.....	18
1.3.9.1.6. Doğrusal Olmama	19
1.3.9.2. Yapay Sinir Ağlarının Avantajları ve Dezavantajları	19
1.3.10. Makine Öğrenmesi	20
1.3.10.1. Makine Öğrenmesi Algoritma Türleri.....	21
1.3.10.1.5.1. Denetimli Öğrenme	22
1.3.10.1.5.2. Denetimsiz Öğrenme	24
1.3.10.1.5.3. Pekiştirmeli (Takviyeli) Öğrenme	25
1.3.10.2. İşlevleri Açısından Makine Öğrenmesi Alanında Kullanılan Algoritmalar	26
1.3.11. Derin Öğrenme	28

İKİNCİ BÖLÜM

YAPAY ZEKÂ, TOPLUM VE ETİK

2.1. YAPAY ZEKÂNIN TOPLUM ÜZERİNDEKİ ETKİLERİ.....	32
2.2. ETİK NEDİR?	32
2.3. YAPAY ZEKÂ ETİĞİ.....	37

ÜÇÜNCÜ BÖLÜM

ALGORİTMİK ÖNYARGI BAĞLAMINDA ETİK İHLALLER

3.1 ARAŞTIRMANIN SORUNSALI	47
3.2. ARAŞTIRMANIN SORULARI.....	47
3.3. ARAŞTIRMANIN AMACI.....	48
3.4. ARAŞTIRMANIN ÖNEMİ.....	48
3.5. ARAŞTIRMANIN KAPSAM VE SINIRLILIKLARI.....	49
3.6. ARAŞTIRMANIN YÖNTEMİ	50

3.7. BULGULAR.....	52
3.7.1. Vaka I: Cinsiyet Temelli Ayrımcılık: Amazon'un İşe Alım Algoritması.....	52
3.7.2. Vaka II: Irk Temelli Ayrımcılık: COMPAS Vakası	55
3.7.3. Vaka III: MIT ve Stanford Araştırması : "Yüz Tanıma Sistemlerinde Cinsiyet ve Irk Temelli Algoritmik Önyargı"	57
3.7.4. Vaka IV: Google Fotoğraflarda "Goril" Etiketleme Skandalı.....	60
3.7.5. Vaka V: Sağlık Kararlarını Yönlendiren Algoritmada Irksal Önyargı	62
3.7.6. Vaka VI: Sahne Tanıma Algoritmalarında Sosyoekonomik Önyargı.....	65
3.8. BULGULARIN DEĞERLENDİRİLMESİ	68
SONUÇ.....	71
KAYNAKÇA	75
İNTERNET KAYNAKLARI	78

ÖZET

YÜKSEK LİSANS TEZİ

YAPAY ZEKÂ VE TOPLUMSAL DÖNÜŞÜM:
ETİK İHLALLER

Rabia KILIÇ

Danışman: Doç. Dr. Zeynep DEMİRCİOĞLU BİRİCİK

2025, 92 sayfa

Bu çalışma, yapay zekânın toplum üzerindeki etkilerini ve bu teknolojilerin beraberinde getirdiği etik sorunları incelemektedir. Yapay zekâ, insan benzeri düşünme yetenekleri ve öğrenme yetenekleri sağlayan bir dizi algoritma ve model içermektedir. Yapay zekâ teknolojisi, endüstri, sağlık, eğitim ve birçok diğer alan üzerinde derin etkiler yaratmıştır.

Ancak, bu etkilerle birlikte önemli etik sorunlar da gündeme gelmektedir. Özellikle, kişisel verilerin korunması, algoritmik önyargı, yapay zekâ tarafından üretilen içeriklerin etik kullanımı ve sosyal medyanın insanlar üzerindeki psikolojik etkileri gibi konular etik tartışmaların odağında yer almaktadır.

Bu çalışmada, yapay zekâ teknolojilerinin yol açtığı etik ihlallerden biri olan algoritmik önyargı, seçilen vakalar üzerinden nitel içerik analizi yöntemiyle incelenmiştir. Vakalar, medya haber siteleri ve akademik literatür gibi çeşitli kaynaklardan toplanmıştır. Analizlerde, yapay zekâ algoritmalarının hangi tür önyargılarda bulunduğu, ortaya çıkan etik sorunlar ve bu sorunların çözümüne yönelik öneriler tartışılmaktadır.

Bu çalışma, teknolojik ilerlemenin insan hayatına olan etkilerini anlamak ve yönlendirmek isteyen herkes için önemli bir kaynak olacaktır.

Anahtar Kelimeler: Yapay zekâ, toplum, etik, algoritmik önyargı, nitel içerik analizi

ABSTRACT

**MASTER’S THESIS
ARTIFICIAL INTELLIGENCE AND SOCIAL TRANSFORMATION:
ETHICAL VIOLATIONS**

Rabia KILIÇ

Advisor: Assoc. Prof. Dr. Zeynep DEMİRCİOĞLU BİRİCİK

2025, 92 pages

This study examines the impact of artificial intelligence on society and the ethical issues raised by these technologies. Artificial intelligence includes a set of algorithms and models that provide human-like thinking abilities and learning capabilities. AI technology has had profound impacts on industry, health, education and many other fields.

However, with these impacts come significant ethical issues. In particular, issues such as the protection of personal data, algorithmic bias, the ethical use of AI-generated content, and the psychological effects of social media on people are at the center of ethical debates.

In this thesis, algorithmic bias, one of the ethical violations caused by artificial intelligence technologies, is analyzed through qualitative content analysis method on selected cases. The cases were collected from various sources such as media news sites and academic literature. The analysis discusses the types of biases that artificial intelligence algorithms engage in, the ethical problems that arise, and suggestions for solving these problems.

This study will be an important resource for anyone who wants to understand and guide the effects of technological progress on human life.

Keywords: Artificial intelligence, society, ethics, algorithmic bias, qualitative content analysis

ÖNSÖZ

Çalışmam, yalnızca akademik bir zorunluluğun değil, içimde uzun süredir büyüyen bir sorgulamanın, merakın ve farkındalığın ürünü oldu. Yapay zekânın hayatımıza bu kadar nüfuz ettiği bir dönemde, sadece “ne yapabiliyor?” sorusunu değil, “neye mal oluyor?” sorusunu da sormak gerektiğini hissettim. İşte bu çalışma da tam olarak bu hissin peşinden gitmemin bir sonucu.

Araştırma süreci boyunca çok şey öğrendim; sadece algoritmaların nasıl çalıştığını değil, aynı zamanda insanların sistemler karşısında nasıl görünmez olabildiğini, verilerin ne kadar "tarafsız" görünüp aslında ne kadar yüklü olabileceğini ve teknolojinin etikle olan karmaşık ilişkisini... Bazen yoruldum, bazen takıldım, bazen de yazdıklarımı tekrar tekrar sorguladım. Ama tüm bunlar, bu süreci gerçekten benim için anlamlı kıldı.

Bu yolculukta en büyük teşekkürüm, her zaman yapıcı eleştirileriyle, sakinliğiyle ve yol göstericiliğiyle yanımda olan tez danışmanım Doç. Dr. Zeynep Demircioğlu Biricik’e. Gerek akademik bilgeliği, gerekse yönlendirmeleriyle bana yalnızca bir rehber değil, aynı zamanda ilham kaynağı oldu.

Ailem, bu süreçte sabrıyla ve anlayışıyla en büyük dayanağım oldu. Kimi zaman sessizlikleriyle, kimi zaman bir fincan çayla, kimi zaman da sadece orada olarak bana güç verdiler. Destekleri hep içimi ısıttı. Özellikle sevgili teyzem Ayşe Özdemir’e ayrı bir parantez açmak isterim; hem akademik bilgisiyle bana rehberlik ettiği, hem de her anlamda yanımda olduğu için. Bazen bir cümleyle zihnimi berraklaştırdı, bazen sadece dinleyerek yükümü hafifletti. Varlığı bu süreci benim için daha anlamlı ve daha katlanılır kıldı. Ayrıca yanımda olan tüm dostlarıma da teşekkür ederim; kafam karıştığında beni dinledikleri, moral verdikleri, “olur bu” dedikleri için...

Çalışmamın, yapay zekâya yalnızca bir teknoloji değil, aynı zamanda bir toplumsal aktör olarak bakmaya vesile olmasını umuyorum. Belki büyük değişimlere yol açmaz; ama doğru bir bakış açısıyla atılmış küçük bir adımın bile anlamlı olacağına inanıyorum.

Rabia KILIÇ

ŞEKİLLER DİZİNİ

Şekil 1.1. Beyin Antrenmanı Uygulamaları Gerçekten İşe Yarıyor mu?	3
Şekil 1.2. Turing Testi Tasarımı	5
Şekil 1.3. Uzman Sistemin Genel Yapısı.....	10
Şekil 1.4. Bulanık Mantık Programının Şematik Yapısı	12
Şekil 1.5. Genetik Algoritmanın Genel Akış Şeması	13
Şekil 1.6. Yapay Sinir Hücresinin Yapısı	14
Şekil 1.7. Biyolojik Sinir Hücresi ve Yapay Sinir Ağı.	16
Şekil 1.8. Makine Öğrenmesi Algoritmaları Hiyerarşisi	21
Şekil 1.9. Denetimli Makine Öğrenmesi.....	22
Şekil 1.10. Denetimsiz Makine Öğrenimi	24
Şekil 1.11. Makine Öğrenmesi ile Derin Öğrenmenin Farkı	30
Şekil 1.12. Yapay Zekâ, Makine Öğrenmesi ve Derin Öğrenme Karşılaştırması	30
Şekil:2.1 Otonom Silahlar	42

TABLOLAR DİZİNİ

Tablo 1.1. İşlevlerine Göre Algoritmalar	26
Tablo 3.1. Tematik Kodlama ve Bulgular	53
Tablo 3.2. Etik İlkelere Göre Değerlendirme	54
Tablo 3.3. COMPAS Algoritmasının Yanılma Oranları	55
Tablo 3.4. Tematik Kodlama ve Bulgular	56
Tablo 3.5. Etik İlkelere Göre Değerlendirme	57
Tablo 3.6. Tematik Kodlama ve Bulgular	58
Tablo 3.7. Etik İlkelere Göre Değerlendirme	59
Tablo 3.8. Tematik Kodlama ve Bulgular	61
Tablo 3.9. Etik İlkelere Göre Değerlendirme	63
Tablo 3.10. Tematik Kodlama ve Bulgular	63
Tablo 3.11. Etik İlkelere Göre Değerlendirme	66
Tablo 3.12. Tematik Kodlama ve Bulgular	67
Tablo 3.13. Etik İlkelere Göre Değerlendirme	66

GİRİŞ

Yapay zekâ, son yılların en etkileyici ve tartışmalı teknolojilerinden biri olarak karşımıza çıkmaktadır. Giderek hızlanan dijitalleşme sürecinin merkezine yerleşen bu teknoloji, sadece teknik bir yenilik değil, aynı zamanda toplumsal yapıları dönüştüren güçlü bir aktör hâline gelmiştir. Endüstri 4.0 kavramıyla birlikte otomasyon sistemlerinde, üretim hatlarında, sağlık hizmetlerinden finansal analiz süreçlerine, eğitimden iletişim teknolojilerine kadar pek çok alanda yapay zekâ destekli sistemlerin kullanımı olağan hâle gelmiştir. Bu dönüşüm, iş gücü piyasasında önemli kırılmalar yaratmakta, bireylerin gündelik yaşam pratiklerini yeniden şekillendirmekte ve toplumsal kurumların işleyiş biçimlerini kökten değiştirmektedir.

Bu yeni dönemde en dikkat çekici meselelerden biri, yapay zekâ teknolojilerinin yalnızca işlevsel fayda ve verimlilik artışı ile sınırlı kalmayıp aynı zamanda etik sorumluluklar, mahremiyet, adalet ve eşitlik gibi temel toplumsal değerleri de doğrudan etkiliyor olmasıdır. Gelişmiş algoritmalar, görünüşte nesnel veri işleme süreçleri sunarken, gerçekte veri setlerinin geçmiş önyargıları taşıma potansiyeli nedeniyle yeni türden ayrımcılık biçimlerini pekiştirebilmektedir. Bu durum, özellikle işe alım, kredi verme, sağlık hizmeti sunumu ve adli yargı gibi hassas karar alanlarında ciddi sorunlara yol açmaktadır. Yapay zekâ tarafından alınan kararların insanlar üzerinde kalıcı etkiler yarattığı örnekler, teknolojik gelişmenin toplumsal adaletle buluşmadığı takdirde nasıl zararlı sonuçlara neden olabileceğini gözler önüne sermektedir.

Yapay zekâ teknolojilerinin kültürel ve iletişimsel boyutları da en az teknik yönleri kadar önemli hâle gelmiştir. Özellikle medya alanında kullanılan içerik öneri sistemleri, kullanıcı davranışlarını yönlendirmekte ve kamusal alandaki bilgi dolaşımını şekillendirmektedir. Sosyal medya algoritmalarının bireylerin görüşlerini kutuplaştırdığı, dezenformasyonun yayılmasını kolaylaştırdığı ve toplumun ortak gerçeklik algısını zayıflattığı yönündeki eleştiriler giderek daha görünür hâle gelmektedir. Bu çalışma, yapay zekânın toplum üzerindeki etkilerini çok boyutlu bir yaklaşımla ele alarak, özellikle etik sorunlar bağlamında ortaya çıkan dönüşümleri anlamaya çalışmaktadır. Yapay zekâ, sadece bir araç değil, aynı zamanda toplumsal ilişkilerin yeniden üretildiği bir sistem olarak değerlendirilmektedir. Bu bakış açısıyla çalışma, teknolojinin tarafsız ve apolitik bir yapı olmadığı; aksine içinde bulunduğu toplumsal bağlamdan etkilenen ve

onu yeniden şekillendiren bir güç olduğunu ortaya koymaktadır. Bu bağlamda çalışmanın temel amacı, yapay zekâ teknolojilerinin toplumdaki etik etkilerini, özellikle algoritmik önyargılar çerçevesinde sorgulamak ve bu alandaki boşluklara eleştirel bir bakışla katkı sunmaktır.

Araştırma, yapay zekâyâ ilişkin temel kavramlar ve teknik alt yapıya dair genel bir çerçeve çizdikten sonra, bu teknolojilerin toplumsal, kültürel ve etik boyutlarını ele alacaktır. Çalışmanın merkezinde, algoritmik önyargıların nasıl ortaya çıktığı, hangi koşullarda daha belirgin hâle geldiği ve bunların bireyler ile topluluklar üzerindeki etkileri yer almaktadır. Bu amaç doğrultusunda, nitel içerik analizi yöntemi kullanılarak seçilmiş somut vakalar incelenecektir. Amazon'un işe alım sisteminde kadın adaylara yönelik ayrımcılık yapması, COMPAS algoritmasının Afro-Amerikan mahkûmlar için daha yüksek risk skorları üretmesi ya da Google'ın görüntü tanıma sisteminde ırksal etiketleme hataları yapması gibi örnekler üzerinden, algoritmaların nasıl toplumsal eşitsizlikleri yeniden ürettiği tartışılacaktır. Bu vakalar, teknolojinin teknik olduğu kadar politik bir boyutunun da olduğunu açıkça ortaya koymaktadır.

Bu bağlamda çalışma, yalnızca mevcut sorunları teşhis etmekle kalmayıp aynı zamanda yapay zekâ sistemlerinin daha adil, şeffaf ve hesap verebilir hâle gelmesi için çeşitli çözüm önerilerini de değerlendirmeyi amaçlamaktadır. Gelişmekte olan etik denetim mekanizmaları, algoritmik şeffaflık ilkeleri ve insan-merkezli tasarım yaklaşımları bu çözüm yolları arasında yer almaktadır. Böylece çalışma, yalnızca akademik değil, aynı zamanda pratik fayda sağlayacak öneriler sunmayı hedeflemektedir.

Yapay zekâ alanında yaşanan hızlı ilerlemeler, beraberinde pek çok yeniliği ve sorunu getirmektedir. Çalışma, söz konusu teknolojik gelişmeleri salt bir ilerleme hikâyesi olarak değil; aynı zamanda eleştirel bir toplumsal sorgulama süreci olarak değerlendirmektedir. Bu çerçevede sunulan analizler, hem bireylerin hem de kurumların yapay zekâ ile kurdukları ilişkilere dair daha bilinçli ve sorumlu bir yaklaşım geliştirmelerine katkı sağlamayı amaçlamaktadır.

BİRİNCİ BÖLÜM

1.1. YAPAY ZEKÂ

Yapay zekâ, bilgisayar biliminin bir disiplindir ve zekâ gerektiren işlevlerin otomasyonu ile ilgilidir. Bu tür işlevler, sonuç çıkarma, bilgi kazanma sürecinde öğrenme, öz-gelişim, gözlem ve algılama yoluyla nesneleri tanıma gibi genel hedefleri içerir. Tüm bu genel amaçlı çalışmalar, bilgisayar biliminde "Yapay Zekâ" olarak bilinir.

Yapay zekâ, genellikle insan zekâsının gerektirdiği işleri, konuşma, tanıma, görsel algılama, karar verme ve çeşitli diller arasında tercüme etme gibi alanlarda gerçekleştirebilen bilgisayar sistemlerinin teorik çalışması ve geliştirilmesini içerir (Rouhiainen, 2020).



Şekil 1.1. Beyin Antrenmanı Uygulamaları Gerçekten İşe Yarıyor mu?

(<https://fykmobile.com/articles/brain-training.html>)

Temel amaç, insan zekâsının belirgin özelliklerini bilgisayarlara aktaracak algoritmaları geliştirmek ve insanlar gibi akıllıca davranabilen sistemler oluşturmaktır. Bazen bu hedef, insanlar üzerinde egemenlik kuran sistemlerin geliştirilme hayallerini içerebilir. Ancak bilimsel ve gerçekçi bir bakış açısıyla bakıldığında, yapay zekâ, deneyimlerden öğrenme yeteneği, mantıklı analiz yeteneği, görüntülerin, şekillerin ve desenlerin tanıma yeteneği, karmaşık problemleri çözebilme yeteneği, dil anlayışı, kelime işleme yeteneği ve bilgisayar bilimine farklı bir bakış açısı getiren bir disiplindir.

Başka bir deyişle yapay zekânın, makinelerden bir şeyler öğrenmek için algoritmalar kullanma ve insanların yaptığı gibi karar verirken

öğrendiklerinden faydalanma becerisi olduğunu söyleyebiliriz. Ancak insanların aksine yapay zekâ ile çalışan makineler, mola vermeye veya dinlenmeye ihtiyaç duymadığı gibi muazzam miktarda veriyi analiz de edebilirler. Yine aynı işleri yapan makinelerin hata oranı, insan emsallerine göre hatırı sayılır boyutta daha azdır (Rouhiainen, 2020).

Bütün bu tanımlamalara göre yapay zekâ, iki temel yaklaşımla ilgilidir. İlk yaklaşım, zekânın doğası hakkında anlayış kazanmak için insan düşünme süreçlerinin incelenmesi, ikinci yaklaşım ise bu insan düşünme süreçlerini bilgisayarlar, robotlar ve benzeri teknolojiler aracılığıyla uygulayarak somutlaştırılmasıdır.

1.2. YAPAY ZEKÂ'NIN TARİHİ GELİŞİMİ

"Yapay zekâ çalışmalarının başlangıcını tespit etmek oldukça karışık bir görevdir. Yapay zekânın doğuşunu Alan Turing'in programları saklayabilen bilgisayarın keşfine dayandırmak mümkündür. İlk bilgisayarlar, belirli problemleri çözmek amacıyla özel olarak tasarlanmış ve ayrılmış makinelerdi. Turing'in yaklaşımı, programların bilgisayar belleğinde veri gibi saklanmasına dayanmaktaydı. Bu yaklaşım, daha sonra tüm modern bilgisayarların temel ilkesi haline geldi. Programlar bellekte saklandığında, bilgisayar kolayca çalışan programı değiştirip yeni bir programı çalıştırabilirdi. Bu, bilgisayarın işlevini değiştirerek farklı görevleri yerine getirebilme yeteneği anlamına geliyordu. İşte bu durum, bilgisayar dünyasında öğrenme veya düşünme sürecinin başlangıcı olarak kabul edilebilir.

Yapay zekânın günümüzdeki popülerliği ve çok çeşitli disiplinlerde araştırma konusu olması, felsefi temellerine ve ilk somut adımlarına 17. yüzyılda rastlayan bir geçmişe sahiptir. 17. yüzyılın başlarına dayanan bir dönemde, toplumun farklı kesimlerinde, özellikle yönetici ve aristokrat sınıfı arasında, hayvan ve insan davranışlarını taklit eden otomatlar oluşturulması o dönemin felsefi düşünce akımlarına yansıyan bir yarışın başladığını gösterir.

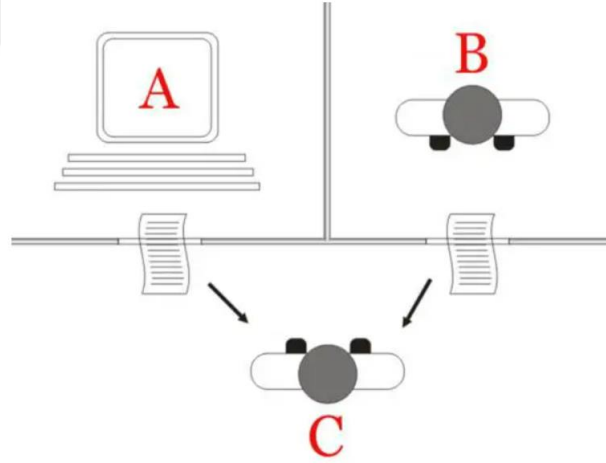
Dönemin önde gelen filozoflarından Descartes (1596-1650), insanı saat benzeri bir düzenekle çalışan makinelerle benzetmiştir. İnsan davranışlarının birçok yönü bu dönemde taklit edilmeye başlandıktan sonra, İngiliz matematikçi Charles Babbage (1792-1871), insana ait fiziksel özellikler yerine zihinsel yeteneklerin taklit edilmesini hedeflemiş ve "Fark Motoru" olarak adlandırdığı ilk hesap makinesini geliştirmiştir. Babbage'in bu hesap makinesi, temel matematiksel işlemleri gerçekleştirmenin ötesinde, ara işlem sonuçlarını saklayabilen bir hafızaya sahipti ve ayrıca satranç ve dama

oyunları gibi etkileşimli görevleri yerine getirebilme yeteneğine sahipti (Coşkun, Gülleroğlu, 2021).

Charles Babbage'in tasarladığı bu hesap makinesi, o dönemin standartlarının ötesinde bir işlem kapasitesine sahip olup, insan zihni özelliklerini taklit etme amacını taşıyarak, o zamanın yapay zekâ çalışmaları için önemli bir ilerleme adımı temsil eder.

Yapay zekâ çalışmalarının kökeni, Cezeri'nin (1136-1206) robot tasarımlarına kadar uzanmasına rağmen, çağdaş anlamda yapay zekâ çalışmalarının özellikle İkinci Dünya Savaşı sırasında ve sonrasında önem kazanmaya başlamıştır.

"Alan Turing, İkinci Dünya Savaşı esnasında "Bombe" ismiyle adlandırdığı ve o dönemde oldukça kritik bir rol oynayan ilk otomatik kod çözme makinesini icat ederek savaşın seyrini değiştiren önemli bir figür haline gelmiştir. Savaşın ardından, başta Alan Turing olmak üzere diğer birçok araştırmacı bağımsız bir şekilde yapay zekâ konusunda çalışmalar yapmışlardır. Turing, 1947 yılında yapay zekâ ile bilgisayar programlarının birleştirilerek akıllı makinelerin geliştirilebileceği konusunda ilk konferansını vermiş ve bu kavramı açıklamıştır (McCarthy, 2007).



Şekil 1.2. Turing Testi Tasarımı

(<https://www.matematiksel.org/turing-testi-insani-robottan-nasil-ayirir/>)

Turing'in yapay zekâ araştırmalarına önemli bir katkısı, "Turing Testi" olarak bilinen deneydir (Bakınız Şekil 1.2). Alanında yayınlanan ilk çalışmalardan biri olan Turing'in çalışması, özellikle dijital bilgisayarlar açısından mekanik zekâ konusunda yapılan tartışmalara önemli bir katkı sunmuştur. Turing, "taklit oyunu" olarak adlandırdığı deneyde, bir bilgisayarı ve bir insanı sorgucu adını verdiği bir kişi tarafından

görülemeyen bölgelere yerleştirmiştir. Sorgucu, bu iki varlıkla doğrudan iletişim kuramamakta ve onlarla bir terminal aracılığıyla etkileşimde bulunmaktadır. Eğer bir sorgu sistemi, gelen cevapları kaynağın bilgisayar mı yoksa insan mı olduğunu ayırt edemiyorsa, Turing, bu durumun bilgisayarın zekice davranabildiği sonucuna varılması gerektiğini ileri sürmüştür. (Yavuz,1995).

1956'da ABD'de, ABD'li bilgisayar bilimcisi John McCarthy tarafından düzenlenen Dartmouth Konferansı, yapay zekâ teriminin ilk kez tanıtıldığı ve yapay zekânın potansiyelini keşfetmek amacıyla araştırma merkezlerinin oluşturulmasının önünü açan bir etkinlikti. Aynı dönemde, araştırmacılar Allen Newell ve Herbert Simon, yapay zekâyı dünya çapında dönüştürücü bir bilgisayar bilimi olarak tanıtmada etkili oldular.

1951 yılında, Ferranti Mark 1 adlı bir makine, satranç ve dama oyunlarında amatör oyunculara karşı başarılı bir algoritma kullanarak dikkat çekti. Ancak, o tarihten günümüze gelirken, Google'ın AlphaGo adlı yapay zekâ sistemi, özellikle Doğu Asya kökenli "Go" oyununda kazandığı zaferlerle yapay zekânın oyunlarda üstünlüğünü gösterdi.

Ferranti Mark 1' den sonra, Simon ve Newell matematiksel problemleri çözmek amacıyla bir 'Genel Problem Çözücü' algoritması oluşturdular. Aynı dönemde, John McCarthy, makine öğrenmesi alanında ciddi bir rol alan LISP programlama dilini yarattı. 1960'larda, araştırmacılar geometrik teoremleri ve matematiksel problemleri çözmek amacıyla algoritmalar geliştirme konusuna odaklandılar. 1960'lı yılların sonlarında bilgisayar bilimcileri, robotlarda makine öğrenimini geliştirmeyi, Machine Vision ve Learning gibi alanlarla geliştirmeye çalıştılar. Japonya'da 1972'de inşa edilen ilk insansı akıllı robot WABOT-1, bu çabaların bir ürünü olarak ortaya çıktı. Ancak, uzun yıllar süren büyük miktarda finansman ve çabaya rağmen bilgisayar bilimcileri, makinelerde zekâ oluşturma son derece zor bir görev olduğunu fark ettiler (Sargın, Göçen, 2020).

Başarıya ulaşmak için yapay zekâ uygulamaları, büyük miktarlardaki verilerin işlenmesini gerektiriyordu; ancak o dönemde mevcut bilgisayarlar, bu büyük veri işleme kapasitesine sahip değillerdi. Bundan dolayı, şirketlerin ve devletlerin, yapay zekâyı olan güvenleri sarsılmaya başladı. Bu durum, bilgisayar bilimcilerinin 1970 yılının ortalarından ve 1990 yılının ortalarına kadar yeterli finansmanı bulmakta zorlandıkları bir dönemi beraberinde getirdi. Bu döneme "Yapay Zekâ Kışı" (AI Winters) denir.

1981 yılında Japonların gerçekleştirdiği beşinci kuşak bilgisayar projesi, PROLOG programlama dili ile önemli bir adım atmış ve PROLOG'un popülaritesini artırmıştır. Ayrıca, bu projenin bir yeniliği de yapay zekâ çalışmaları ile paralel bilgi işleme arasında bir bağ oluşturmuştur. Bu çalışmayı destekleyen Electrotechnical Laboratories (ETL) kurumu, bu projenin bir parçası olarak geliştirdiği EM-4 isimli paralel işlemcili, veri akışı temelli makineyi 1990'ların başında hayata geçirmiştir. Bu makine, paralel işlemi desteklemek için C⁺⁺ programlama dili kullanmıştır (Yavuz,1995).

1990'ların sonlarına doğru, Amerikan şirketleri, "Yapay Zekâ Kışı" adı verilen dönemin ardından yeniden yapay zekâ konusuna ilgi göstermeye başladı. Aynı dönemde, Japon hükümeti, makine öğrenimini geliştirmek amacıyla beşinci nesil bilgisayar geliştirme planını duyurdu. Yapay zekâ uzmanları, bilgisayarların yakın gelecekte konuşma yeteneği, dilleri çevirme yeteneği, görselleri yorumlama yeteneği ve insanlar gibi düşünme yeteneği kazanabileceğine inanıyorlardı. Gerçekten de, 1997 yılında IBM tarafından geliştirilen Deep Blue, dünya satranç şampiyonu Garry Kasparov'a karşı galip gelen ilk bilgisayar oldu ve bu başarı, bilgisayarların bu yeteneklere sahip olabileceğini kanıtladı.

Bilgisayar depolama kapasitesinde ve işlem gücünde yaşanan büyük gelişmeler, şirketlere büyük miktarda veriyi depolama olanağı sunmuş ve son 15 yılda Amazon, Google, Baidu gibi şirketlerin liderliğinde birçok makine öğrenimi uygulamasına olanak tanımıştır. Bu şirketler, sadece tüketici davranışlarını anlamakla kalmamış, aynı zamanda doğal dil işleme, görüntü analizi ve diğer yapay zekâ uygulamaları üzerinde çalışarak bu gelişmeleri daha da ileri taşımışlardır. Sonuç olarak, makine öğrenimi artık kullandığımız birçok çevrimiçi hizmetin temel bir parçası haline gelmiştir. Bu nedenle, günümüzdeki elektronik ve bilgisayar altyapısı, hayallerdeki projeleri gerçekleştirmek için mümkün kılan bir temel oluşturmuştur.

Yapay zekâ hala hızla gelişen bir alandır ve sürekli yeni ilerlemeler yaşanmaktadır. Bu ilerlemeler, daha iyi algoritmaların geliştirilmesi, daha fazla verinin kullanılabilir hale gelmesi, daha hızlı ve güçlü bilgisayarların ortaya çıkması ve daha iyi öğrenme tekniklerinin keşfedilmesi ile sağlanmaktadır. Ayrıca, yapay zekâ sistemlerinin gerçek dünya uygulamalarına entegre edilmesi, yapay zekânın hızla büyümesine katkıda bulunmaktadır.

Yapay zekâ, gelecekte daha da önemli bir rol oynayacak gibi görünmektedir. Otomasyon, sağlık hizmetleri, taşımacılık ve birçok diğer sektörde yapay zekâ kullanımı artmaya devam edecektir. Ancak bu gelişmeler, etik, güvenlik ve sosyal etkiler gibi sorunları da beraberinde getirecektir. Bu nedenle, bu alanlarda çalışmalar ve düzenlemeler devam etmektedir.

1.3. TEMEL YAPAY ZEKÂ TEKNİKLERİ

Temel yapay zekâ teknikleri, yapay zekâ uygulamalarını geliştirmek için kullanılan temel yaklaşımları ve yöntemleri ifade etmektedir. Başlıca doğal dil işleme, görüntü işleme, doğal dil üretimi, veri madenciliği, uzaktan algılama, makine öğrenmesi, bulanık mantık, uzman sistemler, yapay sinir ağları, genetik algoritmalar ve derin öğrenme gibi teknikler yapay zekânın temel tekniklerini oluşturmaktadır.

1.3.1. Doğal Dil İşleme (Natural Language Processing - NLP)

Doğal Dil İşleme, bilgisayarların insanların konuştuğu dilin anlamını anlamalarını, işlemelerini, yorumlamalarını ve hatta cümleler üretebilmelerini sağlama sürecidir. Bu, bir makinenin verilen örneklerden öğrenerek işlem yapma yeteneği açısından, temel bir makine öğrenimi uygulaması olarak kabul edilebilir (Cengiz,2023).

1.3.2. Görüntü İşleme (Image Processing)

Görüntü işleme, görsel verileri analiz etmek ve işlemek için kullanılır. Yüz tanıma, nesne tanıma, görüntü sınıflandırma ve görüntü denetimi gibi uygulamalarda yaygın olarak kullanılır.

1.3.3. Doğal Dil Üretimi (Natural Language Generation - NLG)

Doğal dil üretimi (NLG), bir veri kümesinden yazılı veya sözlü anlatılar üretmek için yapay zekâ programlmasının kullanımıdır. NLG hakkındaki araştırmalar genellikle veri noktalarına bağlam sağlayan bilgisayar programları oluşturma üzerine odaklanır. İleri düzey NLG yazılımı büyük miktarda sayısal veri madenciliği yapabilir, desenleri tanımlayabilir ve bu bilgiyi insanların anlayabileceği şekilde paylaşabilir. NLG

yazılımının hızı, özellikle internet üzerinde haber ve diğer zaman duyarlı hikâyeler üretmek için çok kullanışlıdır. En iyi durumda, NLG çıktısı web içeriği olarak kelimesi kelimesine yayınlanabilir (Wigmore,2013).

1.3.4. Veri Madenciliği (Data Mining)

Veri madenciliği, büyük veri kümelerinden gizli desenleri ile bilgileri çıkarmak için istatistiksel ve matematiksel yöntemleri kullanır. Pazar analizi, müşteri segmentasyonu ve hızlı karar destek sistemleri oluşturma gibi alanlarda kullanılır.

1.3.5. Uzaktan Algılama (Remote Sensing)

Uzaktan algılama, uzaydan veya diğer uzaktan gözlem platformlarından elde edilen verileri analiz ederek yeryüzündeki değişiklikleri izlemek ve haritalandırmak için kullanılır. Bu, tarım, çevre izleme ve coğrafi bilgi sistemleri (GIS) uygulamalarında önemlidir.

1.3.6. Uzman Sistemler

Yapay zekânın en önemli tekniklerinden biri uzman sistemlerdir.

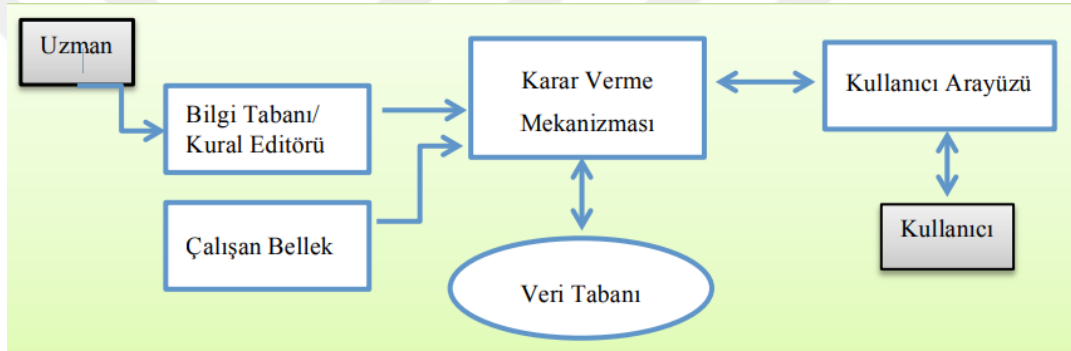
Uzman sistemler, çözümü bir uzmanın bilgi ve yeteneğini gerektiren problemleri, bilgi ve mantıksal çıkarım kullanarak o uzman gibi çözebilen sistemlerdir. Yani problemi çözmede uzman kişi veya kişilerin bilgi ve mantıksal çıkarım mekanizmasının modellenmesi amaçlanmaktadır (Atalay ve Çelik, 2017).

Uzman sistemler, belirli ön koşulların eşzamanlı olarak mevcut olduğu bir konuyla ilgili bir uzmanın hangi kararı alacağını belirleyen tüm kuralları içeren bir programı kullanarak gelen problemlere uygulamak temeline dayanır. Bu yaklaşımın avantajı, her bir kararın hangi kuralların uygulandığının kolayca anlaşılabilmesidir. Bu, özellikle birçok kurala dayalı bürokratik kararlar için organizasyonlar için uygulamalar geliştirmeyi kolaylaştırır.

Uzman sistemleri oluşturan temel ögeler (Bkz. Şekil 1.3) şunlardır: Bilgi tabanı (kural tabanı), çalışan bellek (yardımcı yorumlama modülü), veri tabanı, karar verme mekanizması (çıkartım motoru) ve kullanıcı arayüzü. Bilgi tabanı bilgilerin saklandığı,

yeni bilgilerin üretimine olanak tanıyan birim olduğundan bir uzman sistemin merkezidir. Çalışan hafızada, problem ile alakalı soruların yanıtları ve tanısal testlerin bulguları gibi güncel problem verileri muhafaza edilmektedir. Karar verme mekanizmasında akılcı sonuçlar edinilmesine destek olur. Veri tabanı, bilgi tabanıyla daima etkileşim halinde olmalıdır. Çıkarım motorunun (karar verme mekanizması) temel görevi, bilgi tabanını incelemek ve mevcut problem verilerden sonuçlar çıkarmaktır (Atalay ve Çelik, 2017).

Kullanıcı arayüzü, bilgi toplama, bilgi tabanını sorgulama ve hataları ayıklama, deneme ve test senaryolarını yürütme, özet sonuçlar oluşturma, sonuçları etkileyen faktörleri açıklama ve sistem performansını değerlendirme gibi görevleri yerine getiren bir bileşendir.



Şekil 1.3. Uzman Sistemin Genel Yapısı (Atalay ve Çelik, 2017)

Stanford Üniversitesi'nde 1970 yılında bir grup doktor tarafından geliştirilen MYCIN, ilk uzman sistem olarak kabul edilmiştir. Bu sistem bakteriyolojik ve menenjit hastalıklarının teşhisi ve tedavisi için tasarlanmıştır. Ancak günümüzde bu tür uzman sistemler daha karmaşık ve daha büyüklükte daha fazla veriyi ele alabilmek için daha gelişmiş yapay zekâ teknikleri kullanılmaktadır. MYCIN, tıbbi teşhis ve tedavi alanında uzman sistemlerin gelişimine önemli katkılarda bulunmuş bir öncüdür (Pirim,2006).

1.3.7. Bulanık Mantık

Bulanık mantık, yapay zekâ alanında sıkça kullanılan bir tekniktir. Bu teknik, belirsiz veya bulanık verilerle çalışmayı mümkün kılar. Yapay zekâ uygulamalarında bulanık mantık, karar verme süreçlerini iyileştirmek ve daha gerçekçi sonuçlar elde etmek için kullanılır.

Bulanık mantıkta, herhangi bir şeyin kesin bir doğru veya yanlış olmadığını kabul eder ve bu nedenle belirli bir değere sahip olma olasılığını ifade ederiz. Bu, özellikle çeviri, görüntü işleme ve robotik gibi karmaşık yapay zekâ uygulamalarında kullanışlıdır. Bulanık mantık, belirsizliği ve insan benzeri düşüncüyü modellemek için önemli bir araçtır.

Bu tür bir üyelik fonksiyonu, "kesinlikle üye" ve "kesinlikle üye değil" arasındaki hassasiyeti ayarlayabilme yeteneği sunar. Bulanık mantık, isminden de anlaşılacağı gibi mantık kurallarının daha bulanık ve esnek bir yaklaşımla gerçekleştirilmesini ifade eder. Klasik mantık, önermeleri sadece "doğru" veya "yanlış" olarak ele alır ve genellikle "1" ve "0" ile temsil eder. Ancak bulanık mantık, bu iki keskin kategori yerine, bir önermenin ne kadar doğru veya yanlış olduğunu ifade eden bir olasılık dağılımını kabul eder. Bu, gerçek dünya problemlerini daha iyi yansıtır (Türkiye Mühendislik Haberleri, 2003).

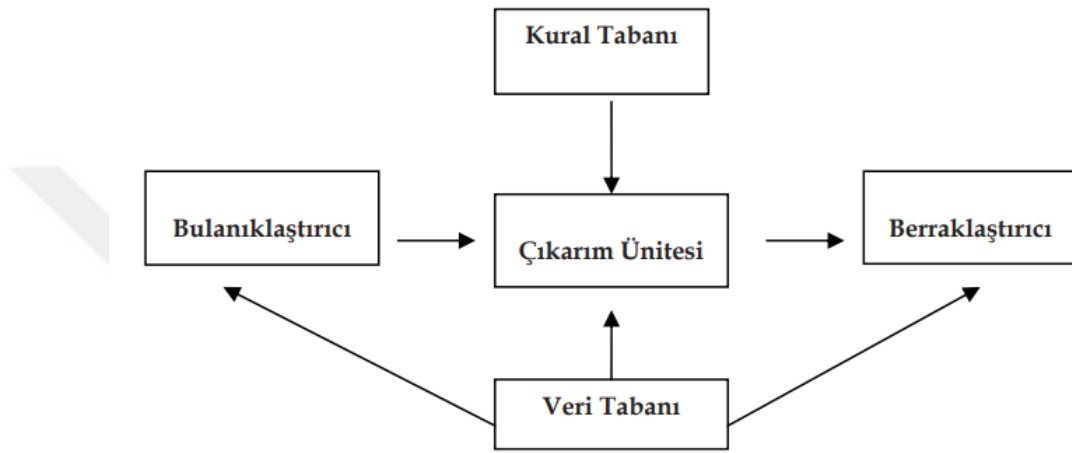
Gerçek hayatta önermeler genellikle kesin doğru veya yanlış değildir; genellikle kısmen doğru veya belirli bir olasılıkla doğru olarak ifade edilir. Bu nedenle, bulanık mantık, klasik mantığın yetersiz kaldığı durumlar için geliştirilmiştir. Bulanık mantık, belirsizlik, karmaşıklık ve gerçek dünya problemleriyle başa çıkmak için önemli bir araç olarak kullanılır.

Bulanık mantık, birçok uygulama alanında geniş bir kullanım bulmaktadır. Örnekler arasında hava tahmini, otomasyon sistemleri, yapay zekâ ve makine öğrenimi, tıp gibi alanlar bulunmaktadır.

Hava tahmini, hava koşullarının belirsiz ve değişken doğası nedeniyle, bulanık mantık modelleri sayesinde daha doğru tahminler yapabilir. Otomasyon sistemlerinde, fabrika otomasyonu veya trafik ışıkları gibi uygulamalarda, bulanık mantık, değişken koşullar ve belirsizlikleri ele almak için önemli bir rol oynar. Ayrıca, yapay zekâ ve makine öğrenme alanlarında bulanık mantık, veri madenciliği ve sınıflandırma problemlerinde kullanılır, bu sayede karmaşık veriler daha iyi işlenebilir. Tıp alanında ise, tanı koyma ve tedavi planlamasında hastalık belirtileri üzerinde çalışırken bulanık mantık, belirsizlikleri ve çeşitli parametreleri hesaba katarak daha iyi sonuçlar elde edilmesine yardımcı olur.

Bu uygulama alanlarında bulanık mantık, gerçek dünyanın karmaşıklığını ve belirsizliğini daha iyi ele almak için önemli bir araç olarak hizmet verir.

Bilgisayar tabanlı uygulamalarda, veri tabanları, bulanıklaştırıcı yazılımları, kural tabanlı sistemler, çıkarım motorları ve berraklaştırıcı yazılımlar gibi bileşenler kullanılarak işlem uygulanır. Genel olarak, bir bulanık mantık sisteminin yapısı Şekil 2'de gösterildiği gibidir.



Şekil 1.4. Bulanık Mantık Programının Şematik Yapısı (Serhatlıoğlu, Ozan ve Hardalaç, 2005)

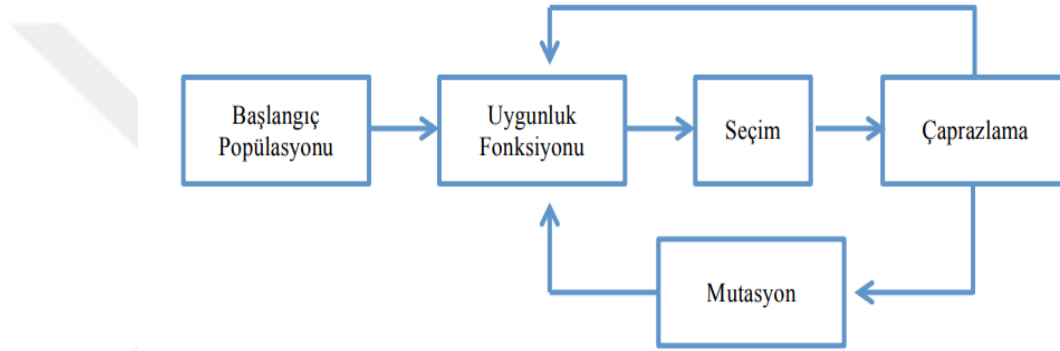
Bulanık mantık, insan düşüncesinin esnek ve değişken yapısını dikkate alan bir algoritmadır. Bu yaklaşım, veriler arasındaki neden-sonuç ilişkilerini tanımlayarak, mantıklı ve doğru bir çıktı üretebilir. Bu işlem, verilerin belirlenmesiyle başlar. Bu veriler, belirli sınırlar içinde gruplandırılarak bulanık kümeler haline getirilir. Tüm olası senaryolar göz önünde bulundurularak bir kural tabanı oluşturulur. Bu kurallar, bir kontrol algoritması tarafından değerlendirilir ve sonuç bilgisi üretilir (Serhatlıoğlu, Ozan ve Hardalaç, 2005).

1.3.8. Genetik Algoritmalar

Genetik algoritmalar, canlıların genetik süreçlerinin bilgisayar ortamında taklit edilmesi işlemidir. Bu algoritmalar, doğada "en uygun olanın hayatta kalması" prensibi üzerine kurulmuştur. Genetik algoritmalar, bir problemde birden fazla çözüm seçeneği içinden en iyi çözümü bulmak amacıyla kullanılır. Rastgele arama yöntemleri kullanarak mevcut çözümleri geliştirerek, en iyi çözüme ulaşmayı hedeflerler.

Genetik algoritma (GA) süreci doğal evrime benzetilir. Bu nedenle Üreme, Çaprazlama, Mutasyon gibi doğal evrimde kullanılan operatörleri içerir. Üreme, uygunluk değerlerine bakılarak stokastik yöntemlerle seçilen bireylerden yeni bir popülasyon oluşturma işlemidir. Bu işlem, ilerleyen generasyonlarda daha yüksek uygunluk değerlerine sahip bireylerin oluşmasına neden olur. Çaprazlama, çoğunlukla rastgele olarak seçilen iki bireyin kromozomları çaprazlanarak gerçekleşir. Bu işlemde, bireylerin kromozomunu oluşturan dizilerin değişik kısımlar yer değiştirilerek yeni döl üretimi sağlanır. Bu döl popülasyonunda daha az uygunluk değerine sahip “zayıf” bireylerin yerine konabilir. Çaprazlama, genetik algortmada en önemli operatördür ve generasyonda yeni çözümlerinin üretiminden sorumludur (Çalışkan, Acar, 2006).

Genetik algoritmaların şematik yapısı Şekil-3'te görüldüğü gibidir.



Şekil 1.5. Genetik Algoritmanın Genel Akış Şeması (Atalay ve Çelik, 2017)

Genetik algoritmalar, genellikle bilinen çözüm yöntemleriyle ele alınması zor veya büyük problemler için kullanılan bir yöntemdir. Bu algoritmalar, genellikle NP (Polinomdışı) problemler gibi zorlu matematiksel problemleri ele alabilir ve genellikle kesin çözüme çok yakın sonuçlar üretebilir. Bu nedenle, gezgin satıcı problemi, yerleşim problemleri, üretim planlama, karesel atama, mekanik öğrenme, atölye çizelgeleme, elektronik tasarım, hücreli üretim ve finansal analiz gibi bir dizi alanda uygulanabilir. Genetik algoritmalar, bu tür karmaşık problemleri ele alırken etkili bir yaklaşım sunarlar ve pratikte işe yarar sonuçlar elde etmek için kullanılırlar.

1.3.9. Yapay Sinir Ağları

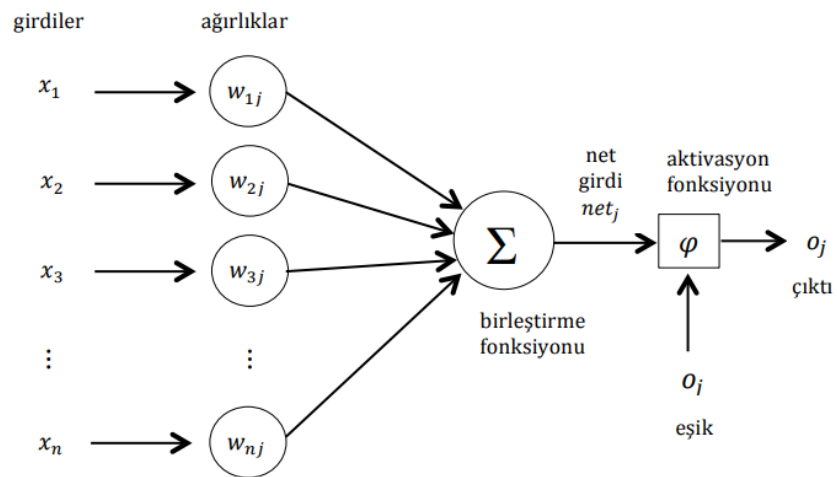
Yapay sinir ağları, yapay zekâ alanında önemli bir rol oynayan matematiksel modellerdir. Yapay sinir ağları, insan beyninin sinir hücrelerinin çalışma prensiplerinden esinlenerek tasarlanmıştır. Bu ağlar insan beyninin herhangi bir görevini

gerçekleştirmek amacı taşımaktadır. Yapay sinir ağları, büyük miktarda veri işleme, öğrenme ve karmaşık görevleri gerçekleştirmek için kullanılmaktadır.

Yapay sinir ağları, örnekler üzerinde gözlem yaparak bilgi edinir, bu bilgileri genellemelere dönüştürür ve daha sonra hiç karşılaşmadığı yeni örnekler üzerinde kararlar almak için öğrendiklerini kullanır. Yapay sinir ağlarının kritik bir özelliği, yapay sinir ağlarının, deneyimlerden öğrenebilme kabiliyetine sahip olmalarıdır, bu yetenek, insan beyninin öğrenme sürecine benzer bir mekanizmayı kullanır. Bununla birlikte, sadece öğrenme yeteneği değil, aynı zamanda bilgiler arasında karmaşık ilişkiler kurabilme yeteneği de yapay sinir ağlarının dikkate değer bir özelliğidir. Bu özellikler sayesinde, yapay sinir ağları karmaşık verileri işleyebilir ve yeni verileri anlamlandırarak kararlar alabilirler.

İlk yapay sinir ağı modeli, 1943 yılında sinirbilimci Warren McCulloch ve matematikçi Walter Pitts tarafından "Sinir Aktivitesinde Düşüncelere Ait Mantıksal Bir Hesap (A Logical Calculus of Ideas Immanent in Nervous Activity)" başlıklı makale ile tanıtılmıştır (Vizyoner Genç, 2019).

Yapay sinir ağlarında, biyolojik sinir ağlarındaki sinir hücrelerine benzer şekilde yapay sinir hücreleri bulunur. Her bir yapay sinir hücresinin temel bileşenleri şunlardır: girdiler, ağırlıklar, birleştirme (toplama) fonksiyonu, aktivasyon (transfer) fonksiyonu ve hücrenin çıktısı.



Şekil 1.6. Yapay Sinir Hücresinin Yapısı (Atalay ve Çelik, 2017)

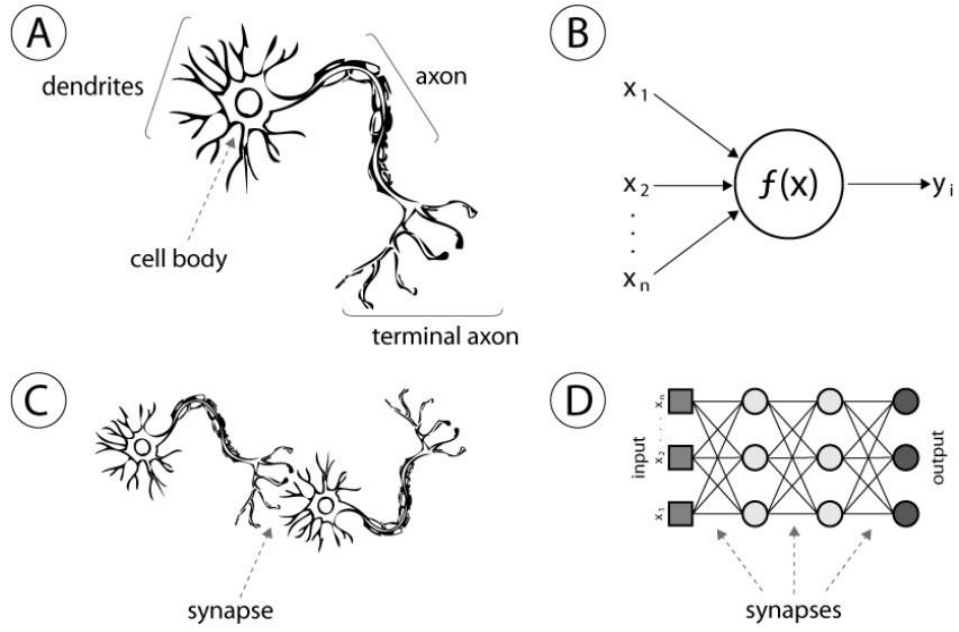
Girdiler, dışarıdan ya da başka hücreden aldığı bilgiyi sinire getirir. Bir sinir çoğunlukla çok sayıda yerden girdileri alır. Bu girdiler birleştirilmek için nöron çekirdeğine transfer edilir.

Ağırlıklar, yapay sinire gelen girişlerin ağırlıklarını saptayan katsayılardır. Her girişin kendine özgü bir katsayısı bulunur. Bu katsayılar, yapay hücreye ulaşan bilginin önemini ve hücreye etkisini ifade eder. Ağırlıkların değeri pozitif olabildiği gibi negatif ya da sıfır da olabilir. Bir girişin katsayısı yüksekse bu giriş yapay sinirle daha güçlü bir bağlantı kurmuştur. Katsayı düşükse bu bağlantının zayıf olduğu anlamına gelir.

Birleştirme fonksiyonu aynı zamanda bir araya getirme işlevi veya toplama işlevi olarak bilinir ve sinir sistemlerindeki her ağırlığın ilgili girişlerle çarpıldıktan sonra toplandığı ve sonrasında bir eşik değeri ile karşılaştırılarak, söz konusu hücreye ait net girişin hesaplandığı bir fonksiyonu ifade eder. Aktivasyon fonksiyonu, birleştirme fonksiyonunun hücreye gönderdiği net bilgiyi işler ve hücrenin bu girdiye karşı üretmesi gereken çıktıyı belirler.

Çıktı, aktivasyon fonksiyonu sonucunun diğer sinirlere ya da dış dünyaya gönderilmesidir. Bir sinirin yalnızca bir çıktısı vardır. Bir sinire ait olan çıktı, başka bir sinire giriş olabilir.

Yapay sinir ağı, katmanlar halinde birleşen yapay sinir hücrelerinden oluşur. Bir yapay sinir ağı genellikle üç katmandan oluşur: Girdi katmanı, gizli (ara) katmanlar ve çıktı katmanı. (Bkz.Şekil-5) Girdi ve çıktı katmanındaki nöron sayısı, bağımsız ve bağımlı değişkenlerin sayısına bağlıdır, ancak gizli katmanların sayısı ve her bir gizli katmandaki nöron sayısı, kullanıcının belirlediği en iyi performansı elde etmek amacıyla ayarlanır.



Şekil 1.7. Biyolojik Sinir Hücresi ve Yapay Sinir Ağı (Matarollo, 2013).

Girdi katmanı, yapay sinir ağına gelen bilgileri tutmaktadır. Veri seti baz alınarak ilgili özelliklerin her biri giriş katmanında bir düğüm olarak gösterilmektedir. İlgili girdilerin hepsinin belirli ağırlık değerleri vardır. Bu girdiler ara katmandaki düğümlere bu ağırlıklar yoluyla bağlanmıştır. Ağırlıklar ilgili özelliğin önem derecesi vermektedir. Gizli (ara) katmanlarda, girdi katmanından gelen veriler işlenip bir çıktıya dönüşmektedir. Bu dönüşüm ağırlık değerleri kullanılarak yapılmaktadır. Problemin zorluk derecesi ara katman sayısında belirleyici etken olmaktadır. Çıktı katmanı, çıktı değeri ya da değerlerinin tutulduğu yerdir. Ara katmandan gelen sonuç verisi çıktı değeri olarak ifade edilmektedir (Erhandı, 2020).

1.3.9.1. Yapay Sinir Ağlarının Özellikleri

Yapay sinir ağları modelleri, biyolojik sinir ağlarının çalışma prensiplerinden ilham alınarak oluşturulmuştur. Bu yapay sinir ağları, biyolojik olmayan yapı taşların düzenli bir tasarımla birbirlerine yoğun bir şekilde bağlandığı ağlar olarak inşa edilir. Sinir sisteminin bu modellemesi, yapay sinir ağlarının biyolojik sinir sisteminin avantajlarına da sahip olduğunu gösterir. Bu avantajları aşağıdaki şekilde özetleyebiliriz: İlk olarak, yapay sinir ağlarının bir üstünlüğü, paralellik özelliğine sahip olmalarıdır.

Yapay sinir ağları modelinde her bir eleman kendi kendinin işlemcisidir ve aynı katmanlar arasında zaman bağımlılığı olmadan tamamen eşzamanlı olarak çalışabilirler.

Bu özellik, yapay sinir ağlarının hız açısından önemli bir avantaj sağladığını gösterir. İkinci bir üstünlüğü ise, yapay sinir ağlarının öğrenebilme yeteneğine sahip olmalarıdır. İnsan sinir sistemi gibi, yapay sinir ağlarının da problemleri çözmek için öğrenme yeteneği vardır. Üçüncü bir üstünlüğü, paralel olarak çalışan yapay sinir ağlarının karmaşık işlevleri yerine getirmesi gerekmeksizin daha basit işlemleri gerçekleştirebilmesidir. Bir diğer avantajı ise, yapay sinir ağlarının ayrı ayrı elemanlarda meydana gelen hasarın, performansta ciddi bir düşüğe neden olmamasıdır. Öte yandan, bir bilgisayarın herhangi bir işlemcisini çıkarmak, o makinenin işe yaramaz hale gelmesine yol açabilir.

Yapay sinir ağlarının en önemli özellikleri arasında paralel çalışma yeteneği, öğrenme sırasında kendi kendini geliştiren eğitim yöntemleri, donanım açısından kolayca uygulanabilir olmaları, genelleme yeteneği ve sistem cevabının hücre kaybına karşı direnci sayılabilir. Ancak, öğrenme algoritmalarının yerel en uygunlara sıkça takılabileceği ve yapının gereksiz yere karmaşılaşarak genelleme yeteneğinin azaldığı gözlenmektedir. Bu nedenle, yapay sinir ağlarının ihtiyaca göre büyüme yeteneğini kazanması ve yerel en uygunlardan sıyrılabilmesi için genetik algoritmalar kullanılarak eğitilir.

Genetik algoritmalar, en uygun çözümü arayarak ağın her bir iterasyonda içerdiği hücre sayısını artırır. Bu sayede, en düşük hücre sayısı ile optimum çözüme ulaşma amacı güdüdür (Bozüyük vd, 2005).

Aşağıda yapay sinir ağlarının özellikleri açıklamalarıyla sıralanmıştır.

1.3.9.1.1. Öğrenme Yeteneği

Yapay Sinir Ağı'nın (YSA) istenen bir davranışı sergileyebilmesi için, doğru bağlantılar kurulmalı ve bu bağlantıların uygun ağırlıklarla ayarlanmalıdır. Ancak, Yapay Sinir Ağı'nın karmaşık yapısı nedeniyle bu bağlantılar ve ağırlıklar önceden belirlenemez veya özel olarak tasarlanamaz. Bu yüzden Yapay Sinir Ağı, belirli bir problemi çözebilmek için verilen eğitim örneklerini kullanarak istenen davranışı öğrenmelidir.

1.3.9.1.2. Uyarlanabilirlik

Yapay Sinir Ağları (YSA), hedeflenen problemin çözümü için ağırlıklarını problemdeki değişikliklere göre ayarlayabilir. Başka bir deyişle, bir problemi çözmek amacıyla eğitilen YSA, problemdeki değişikliklere uyum sağlamak için yeniden eğitilebilir ve eğer bu değişiklikler sürekli ise gerçek zamanlı olarak sürekli eğitim alabilir. Bu özellik, YSA'nın adaptif sinyal işleme, örnek tanıma, sistem teşhisi ve kontrol gibi alanlarda etkili bir şekilde kullanılmasını sağlar.

1.3.9.1.3. Hata Toleransı

Yapay Sinir Ağları (YSA), çok sayıda hücrenin farklı şekillerde birbirine bağlandığı bir yapıya sahiptir ve ağına içerdiği bilgi, bu bağlantılara dağılmış bir şekilde bulunur. Bu nedenle, eğitilmiş bir YSA'nın bazı bağlantıları hatta bazı hücreleri işlevsiz hale gelirse bile, ağına doğru bilgiyi üretme yeteneği önemli ölçüde etkilenmez. Bu, geleneksel yöntemlere göre hatalara karşı daha fazla tolerans sağlar.

1.3.9.1.4. Donanım ve Hız

Yapay Sinir Ağları (YSA), büyük ölçekli entegre devre (VLSI) teknolojisi ile uygulanabilme yeteneğine sahiptir, bu da YSA'nın hızlı veri işleme kapasitesini artırır ve gerçek zamanlı uygulamalarda önemli bir avantaj sağlar.

1.3.9.1.5. Genelleme Yeteneği

Yapay Sinir Ağları (YSA), ilgilendiği problemi öğrendikten sonra, eğitim sürecinde karşılaşmadığı test örnekleri için de istenilen tepkileri üretebilme yeteneğine sahiptir. Örnek verilecek olursa, karakter tanıma için eğitilen bir YSA, bozuk karakter girişlerini doğru bir şekilde tanıyabilir veya bir sistemin eğitilmiş YSA modeli, eğitim sürecinde verilmeyen giriş sinyalleri için de sistemin aynı davranışını sergileyebilir. Bu, YSA'nın öğrendiği desenleri genellenme yeteneği sayesinde mümkün olur. Yani, YSA, öğrendiği bilgiyi yeni ve farklı durumlarla ilişkilendirebilir ve bu nedenle öğrenme sonrası uygulamalarda başarılı olabilir (Kabalcı,2015).

1.3.9.1.6. Doğrusal Olmama

Yapay sinir ağlarının en belirgin özelliklerinden biri, gerçek dünyadaki karmaşık ve doğrusal olmayan ilişkileri modelleyebilme yeteneğidir. Yapay sinir ağlarının ana yapı taşı olan sinir hücreleri, doğrusal olmayan işlemleri gerçekleştirebilir. Bu nedenle, yapay sinir ağları, doğrusal olmayan ilişkileri modellenen güçlü yapıya sahiptir. Bu doğrusal olmayan özellik, tüm yapay sinir ağına yayılmıştır. Bu nedenle, yapay sinir ağları, doğrusal olmayan ve karmaşık problemlerin çözümünde önemli bir araç olarak kabul edilir. Bu özellikleri sayesinde yapay sinir ağları, birçok uygulama alanında etkili bir şekilde kullanılabilir, özellikle de veriler arasındaki doğrusal olmayan ilişkileri modelleme gerekliliği bulunduğu durumlarda (Bayır,2006).

Bu özellikler, yapay sinir ağlarının çok çeşitli uygulamalarda kullanılmasını sağlar. Yapay zekâ, görüntü işleme, doğal dil işleme, tahmin analizi, oyun geliştirme ve daha birçok alanda yapay sinir ağlarının kullanılmasını içerir.

1.3.9.2. Yapay Sinir Ağlarının Avantajları ve Dezavantajları

Yapay sinir ağları, insan beynini taklit etmedeki ustalığıyla bilinen karmaşık bilgi işleme sistemleridir. İnsan beyninin çalışma prensiplerinden ilham alarak tasarlanan bu yapay sinir ağları, karmaşıklığı ve veri analizi kapasitesini büyük ölçüde artırarak bilim dünyasına pek çok avantaj sunmuştur.

Avantajları

Yapay sinir ağları, verilerin aynı anda birçok görevi gerçekleştirme kabiliyetine sahiptir, çünkü verileri paralel olarak işleyebilirler. Yapay sinir ağları direnç gösterir. Bu, bir veya daha fazla hücrenin veya sinir ağının kaybının yapay sinir ağlarının performansını etkilediği anlamına gelir. Yapay sinir ağları, ağ içinde bilgi saklamak için kullanılır; böylece bir veri çifti mevcut olmasa bile bu, ağın sonuç üretmediği anlamına gelmez. Yapay sinir ağları yavaşça bozulur, yani aniden çalışmaları durmaz; bunun yerine bu ağlar zaman içinde yavaşça performanslarını kaybedebilirler. Yapay Sinir Ağları, bu ağların geçmiş olaylardan öğrenmesi ve karar vermesi konusunda eğitilebilirler (GeeksforGeeks, 2021).

Dezavantajları

Yapay sinir ağları, paralel işleme ve donanım bağımlılık gibi faktörler nedeniyle güvenilirlik konusunda bazı zorluklarla karşılaşabilir. Ayrıca, doğru ağ yapısının belirlenmesi ve problem ifadesinin anlaşılması gibi konularda da belirsizlikler olabilir. Bu nedenle, bir yapay sinir ağının nasıl çözüm üreteceğini tam olarak bilemeyiz. Bu durumda, yapay sinir ağının güvenilirliği sorgulanabilir (GeeksforGeeks, 2021).

1.3.10. Makine Öğrenmesi

Makine öğrenmesi, bilgisayarların verileri analiz edip öğrenmelerini ve deneyime dayalı olarak gelecekteki görevleri daha iyi gerçekleştirmelerini sağlayan bir yapay zekâ tekniğidir. Temelde, makine öğrenimi, bilgisayarların belirli bir görevi uygulamak için verileri kullanarak kendilerini özelleştirmelerini ifade eder.

Makine öğrenmesi, insan karar alma süreçlerini taklit ederek, geçmiş verilere dayalı örnekler ve gözlemler aracılığıyla anlamlı ilişkileri ve desenleri otomatik olarak öğrenmeye çalışan bir tekniktir. Makine öğrenmesi temelli yöntemler, özünde örüntü tanıma, sınıflandırma, kümeleme ve tahmin yaklaşımlarını içerir. Bu teknikler, pazarlama, mühendislik, finans ve tıp gibi çok çeşitli alanlarda başarıyla uygulanmaktadır (Gülpınar Demirci, 2022).

Makine öğrenimi ilk olarak 1980'lerde adını duyurmuş ve popülerliği veri madenciliği ile birlikte hızla yükselmiştir. Bu teknoloji, sistem tarafından önerilen veriler ve parametrelerle gerçekleştirilen benzetimler yoluyla, insandan daha iyi saptama yapabilen, öğrenme yeteneğine sahip, programlananları ortaya çıkarabilen, kendini geliştirebilen ve eğitebilen yapısıyla tanınmıştır.

Makine öğrenmesinin gerçekleştirilebilmesi için, makinenin verilerle birlikte mantıklı ve rasyonel sonuçlar üretebilecek bir algoritmaya sahip olması gerekliliği unutulmamalıdır. Mesela bir markette müşterilerin alışveriş verileri algoritmalar yoluyla sisteme işlenmektedir. Bu algoritmalara bakıldığında içecek alan müşterilerin aynı anda çikolata aldıkları da görülmektedir. Böyle bir durumda içecek dolapları ile çikolataların bulunduğu raflar yakın olacak şekilde ayarlanmaktadır. Dolayısıyla satışlarda ciddi bir yükseliş görülmektedir.

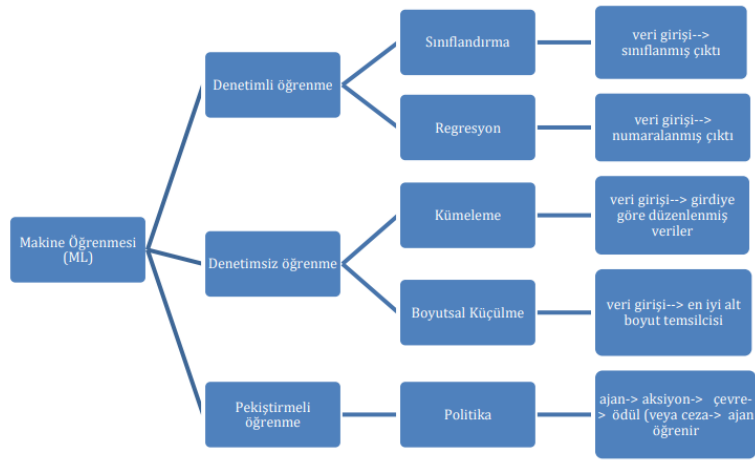
Makine öğrenmesinin temel amacı, verilerden öğrenilen bilgileri kullanarak tahminlerde bulunmak, kararlar almak veya görevleri otomatize etmektir. Makine

öğrenimi, çeşitli uygulamalara sahiptir, bunlar arasında doğal dil işleme, görüntü işleme, otonom araçlar, sağlık sektörü, finans ve daha fazlası bulunmaktadır. Makine öğrenimi modelleri, daha fazla veri ve daha fazla eğitimle genellikle daha iyi hale gelmektedir. Makine öğrenimi, iş süreçlerini otomatikleştirmek, verilerden değerli bilgiler çıkarmak ve karar alma süreçlerini geliştirmek gibi birçok faydalı kullanım alanına sahip olmaktadır.

1.3.10.1. Makine Öğrenmesi Algoritma Türleri

Makine öğrenimi algoritmaları, verilerdeki örüntüleri tanımlama ve öğrenme görevlerini gerçekleştirir. Makine öğrenimi alanında birçok farklı algoritma bulunmaktadır ve bu algoritmalar farklı görevleri çözmek veya veri türlerini işlemek için tasarlanmıştır.

Makine öğrenmesi şekline göre algoritmalara ait hiyerarşi Şekil 1.8’de gösterilmektedir.

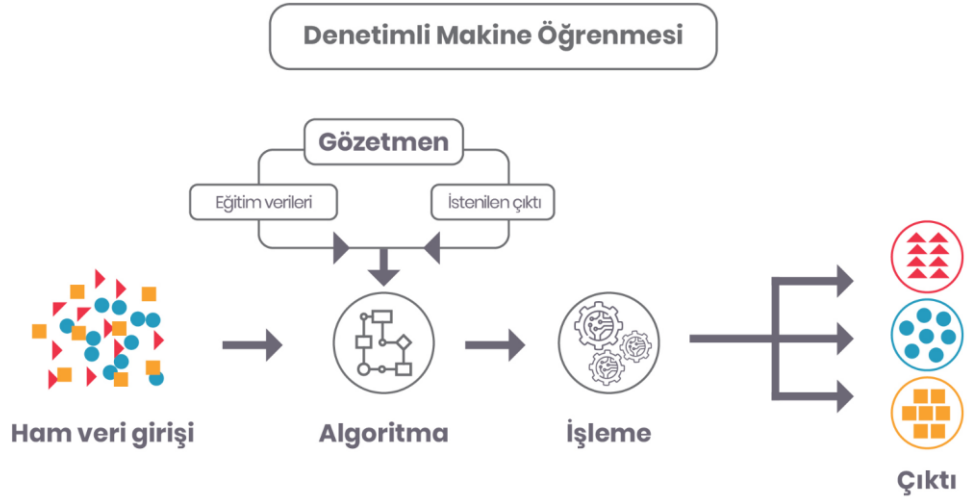


Şekil 1.8. Makine Öğrenmesi Algoritmaları Hiyerarşisi (Tosunoğlu, Yılmaz, Özeren ve Sağlam, 2021)

Makine öğrenmesinde birçok farklı modeller vardır, bunlar denetimli öğrenme, denetimsiz öğrenme ve pekiştirmeli (takviyeli) öğrenme olarak makinelere 3 taktik ile öğretilmektedir.

1.3.10.1.5.1. Denetimli Öğrenme

Birçok veri örneği kullanılarak bir modelin eğitildiği ve bu modelin daha sonra yeni verilere tahminler yapmak için kullanıldığı bir öğrenme türüdür. Denetimli öğrenme, verilerin belirli girdi-çıkı ilişkilerini öğrenmeye çalışır. *Denetimli öğrenmede makineye hem girdi hem çıktı değeri verilir (Tosunoğlu, Yılmaz, Özeren ve Sağlam, 2021).* Denetimli öğrenme, sınıflandırma (bir girdinin belirli bir kategoriye ait olup olmadığını tahmin etme) ve regresyon (bir girdinin sürekli bir değerini tahmin etme) gibi birçok uygulama alanında kullanılır. Bu yöntem, özellikle veriler arasındaki karmaşık ilişkilerin öğrenilmesi gereken durumlarda etkilidir.



Şekil 1.9. Denetimli Makine Öğrenmesi

(<https://www.turhost.com/blog/makine-ogrenmesi-machine-learning-nedir/>)

Denetimli öğrenme algoritmaları aşağıdaki şekilde açıklanabilir.

K En Yakın Komşuluk Algoritması (KNN), belirli bir veri örneğinin sınıfını tahmin etmek için özellikler arasındaki benzerliği değerlendiren, denetimli makine öğrenme alanında temel bir algoritmadır.

Bir örneği komşularına olan mesafesini hesaplayarak komşularına göre tanımlarken algoritmasında, K parametresi modelin performansını etkiler. K değeri çok küçük ise model aşırı uyuma duyarlı olabilir. K değeri çok büyük ise örneğinin yanlış sınıflandırılmasına neden olabilir. Ayrıca bu algoritma sınıflandırmanın yanı sıra regresyon için de kullanılmaktadır. KNN, genellikle Öklid mesafesini kullanan denetimli öğrenme tekniğine dayalı en basit makine öğrenimi algoritmalarından biridir. KNN, istatistiksel olarak

parametrik olmayan bir algoritmadır, yani temel veriler üzerinde herhangi bir varsayımda bulunmaz. KNN'nin kullandığı Öklid uzaklığı, bilinmeyen düğümlere ortalama bir değer belirler. Örneğin, verilerde bulunan herhangi bir düğüm eksik ise, bu eksik düğümü en yakın komşunun değerinden tahmin edilir. Bu değer kesin değildir, ancak olası bir eksik düğümü ortaya çıkarır ise tanımlamaya yardımcı olur (Tuğrul, Ahmed,2022).

Naive Bayes, istatistiksel sınıflandırma problemlerini çözmek için kullanılan bir olasılık temelli bir sınıflandırma algoritmasıdır. Naive Bayes, özellikle metin sınıflandırma (örneğin spam filtresi) gibi uygulamalarda yaygın olarak kullanılır. Temel olarak Bayes Teoremi' ne dayanır ve "Naive" (saf) sıfatı, temel varsayımlarının basitliğine atıfta bulunur.

Naive Bayes sınıflandırıcısı, bir tahmin edicinin (x) değerinin belirli bir sınıf (c) üzerindeki etkisinin diğer tahmin edicilerin değerlerinden bağımsız olduğunu varsayar. Bu varsayıma sınıf koşullu bağımsızlık denir. Bu varsayıma rağmen, NB güçlü bir sınıflandırıcıdır ve özellikle metin madenciliğinde yaygın olarak kullanılmaktadır (Gülpınar Demirci, 2022).

Naive Bayes, özellikle sınıflandırma problemlerinde iyi çalışan ve hızlı bir algoritma olabilir. Ancak, bağımsızlık varsayımı gerçek dünyadaki durumların karmaşıklığına uymayabilir ve bu nedenle bazı durumlarda diğer algoritmalar daha iyi sonuçlar verebilir.

Karar ağaçları, sınıflandırma ve tahmin problemlerinde başarıyla kullanılan bir tekniktir. Bu yöntem, değişkenler arasındaki değerlerin görülme sıklığı ve dağılımına dayanarak bir karar şeması oluşturur. Her bir düğüm, bir soruyla ilişkilidir ve tüm olası cevaplar, ağacın farklı seviyelerinde yer alan yapraklar aracılığıyla temsil edilir. Kökten yaprağa uzanan her yol, bir karar kuralını ifade eder. Bu karar kuralları takip edilerek sınıflandırma işlemi gerçekleştirilir.

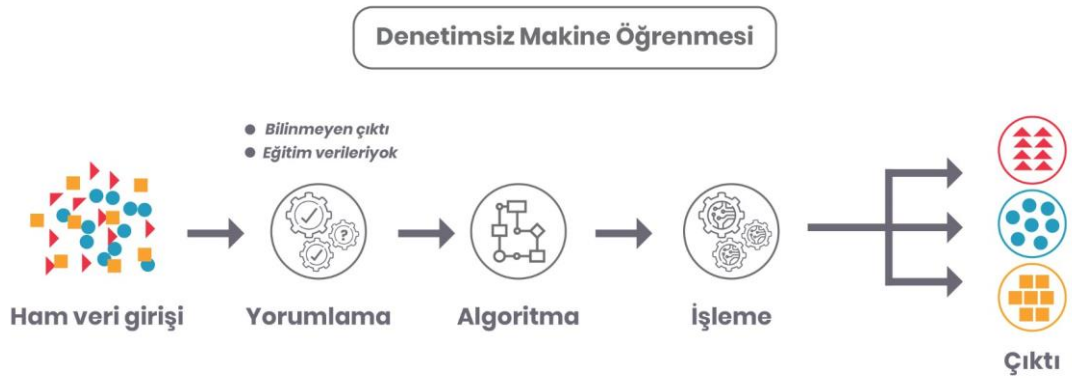
Rastgele Orman, makine öğrenimi ve veri madenciliği alanlarında sıkça kullanılan ve birden çok öğrenme modelini bir araya getirerek daha iyi sonuçlar elde etmeyi amaçlayan bir algoritmadır. Rastgele Orman, karar ağaçlarını bir araya getirerek daha güçlü ve kararlı bir tahmin yapma modeli oluşturur. Rastgele Orman'ın ana özelliği rastgelelik kullanmasıdır. Her ağaç eğitilirken, rastgele bir alt kümesi veriler kullanılır ve ağaçlar farklı şekillerde eğitilir. Ayrıca, her bir karar düğümünde rastgele özelliklerin seçilmesiyle ağaçlar çeşitlendirilir. Bu, modelin daha genelleştirilebilir ve daha dayanıklı olmasına yardımcı olur. Rastgele Orman, sınıflandırma, regresyon, özellik seçimi ve aykırı değer tespiti gibi birçok veri madenciliği ve makine öğrenimi uygulamasında

yaygın olarak kullanılır. Veri bilimi projelerinde model oluştururken, veri keşfi ve model performansı iyileştirmesi yaparken Rastgele Orman'ı göz önünde bulundurmak oldukça faydalı olabilir.

1.3.10.1.5.2. Denetimsiz Öğrenme

Denetimsiz öğrenme, giriş verilerinin önceden belirlenmiş etiketler veya hedefler olmadan öğrenildiği bir makine öğrenimi türüdür. Bu, verileri incelemek, kalıpları tanımak ve bu kalıplar aracılığıyla anlamlı bilgiler elde etmek amacıyla kullanılır.

Denetimsiz öğrenme, yalnızca giriş verilerinin bulunduğu ve bu verilere karşılık gelen çıktıların önceden bilinmediği bir öğrenme yaklaşımıdır. Bu yaklaşımın temel özelliği, öğrenme sürecini desteklemek için etiketlenmiş (yani belirli bir sonucu gösteren) bir veri kümesinin eksik olmasıdır. Bu nedenle, veri setleri etiketlenmez; bunun yerine, veriler benzerlikleri veya yapısal özellikleri temel alınarak kümeler ayrılır ve veri içindeki kalıplar, ilişkiler ve yapılar tespit edilmeye çalışılır.



Şekil 1.10. Denetimsiz Makine Öğrenimi

(<https://www.turhost.com/blog/makine-ogrenmesi-machine-learning-nedir/>)

Denetimsiz öğrenme algoritmaları aşağıdaki şekilde açıklanabilir.

K-ortalamlar algoritması kümeleme problemlerini çözmek için yaygın olarak kullanılan bir makine öğrenimi algoritmasıdır. K-ortalama, veri noktalarını K adet küme içine bölmek amacıyla kullanılır. Her bir küme, benzer veri noktalarını içermelidir ve bu benzerlik, genellikle iki nokta arasındaki uzaklığı hesaplamak için kullanılan öklidyen uzaklık ölçüsünden yararlanılarak hesaplanır. K-ortalama, genellikle başlangıç

merkezlerin rastgele seçildiği birçok tekrar sonucunda farklı sonuçlara yol açabilir. Bu nedenle, algoritmanın sonuçlarını iyileştirmek için başlangıç merkezlerinin daha iyi bir şekilde seçilmesi veya algoritmanın birden fazla kez çalıştırılması ve en iyi sonucun seçilmesi gibi yöntemler kullanılabilir. K-ortalama algoritması, veri madenciliği, görüntü işleme, biyoinformatik, pazarlama analizi ve diğer birçok uygulama alanında kullanılır.

Hiyerarşik kümeleme algoritması, veri madenciliği ve veri analitiği alanlarında sıklıkla kullanılan bir kümeleme yöntemidir. Bu algoritma, benzer özelliklere sahip veri noktalarını gruplandırarak veri setini bir ağaç yapısı (dendrogram) kullanarak görsel bir şekilde temsil eder.

Hiyerarşik kümeleme, veri setindeki yapıları ve ilişkileri anlamak, benzer grupları tanımlamak ve daha sonra analiz veya karar verme süreçlerine rehberlik etmek için kullanılır.

Konu modelleme, metin belgeleri içindeki temel temaları ve konuları otomatik olarak tanımlamak için kullanılan bir makine öğrenimi ve doğal dil işleme tekniğidir. Konu modelleme, büyük metin verilerinin analizi, metin madenciliği ve içerik anlama konularında yaygın olarak kullanılır. Konu modelleme, büyük metin verilerinin daha iyi anlaşılmasına ve bilgi çıkarılmasına yardımcı olur. Bu, hem akademik araştırmalarda hem de endüstriyel uygulamalarda büyük öneme sahiptir.

1.3.10.1.5.3. Pekiştirmeli (Takviyeli) Öğrenme

Pekiştirmeli Öğrenme, dinamik çevrelerde belirsizlikle karşılaşılan karar verme problemlerini çözmek için kullanılan bir makine öğrenimi dalıdır.

Klasik planlama ve kontrol yöntemleri, çözülmek istenilen problemi matematiksel olarak modeller ve ardından bu model üzerinde çeşitli optimizasyon işlemleri yaparak, değişik durumlar için uzun vadede en çok getirisi olacak kararları Pekiştirmeli Öğrenme algoritmaları ise genellikle çözülecek problem üzerinde hiçbir ön kabul yapmaz, problemin bulunduğu ortam ile etkileşime geçer, birçok değişik kararı dener, ortamdan veri toplar ve ardından bu veriyi analiz ederek optimal kararları bulmaya çalışır (Savunma Sanayii Yapay Zekâ Platformu).

Derin Q ağları, derin öğrenme tekniklerinin takviyeli öğrenme problemlerine uygulanmasının örneklerinden biridir. Daha karmaşık ve büyük veri setleriyle çalışabilen ve öğrenmeyi iyileştiren derin öğrenme ağları kullanarak ajanların karmaşık görevleri

daha iyi gerçekleştirmelerine olanak tanır. Bu, özellikle oyun oynama ve çevrimiçi reklam verme gibi alanlarda büyük ilgi çekmiştir.

1.3.10.2. İşlevleri Açısından Makine Öğrenmesi Alanında Kullanılan Algoritmalar

Tablo 1.1’de işlevleri açısından makine öğreniminde kullanılan algoritmalar verilmiştir.

Tablo 1.1. İşlevlerine Göre Algoritmalar (Tosunoğlu, Yılmaz, Özeren ve Sağlam, 2021, s.181-183)

Algoritma Kategorisi	Alt algoritmalar	İşlevleri
Regresyon algoritmaları	• Sıradan en küçük kareler regresyonu (OLSR) • Doğrusal regresyon • Lojistik regresyon • Adımsal regresyon • Çok değişkenli uyarlanabilir regresyon (MARS) • Yerel olarak tahmini dağılım grafiği düzenleme (LOESS)	Bağımsız değişkenden yola çıkarak bağımlı değişkene ilişkin kestirimde bulunmak amacıyla kullanılırlar.
Örnek tabanlı algoritmalar	• k-En Yakın Komşu (kNN) • Vektör Niceme (LVQ) Öğrenme • Kendi Kendini Düzenleyen Harita (SOM) • Yerel Ağırlıklı Öğrenme (LWL) • Destek Vektör Makineleri (SVM)	Örnek bir veri tabanı oluşturarak, en iyi eşleşmeyi bulmak ve bir tahminde bulunmak için verileri veri tabanı ile karşılaştırır.
Düzenleştirme algoritmaları	• Sırt Regresyonu • En Küçük Mutlak Büzülme ve Seçim Operatörü (LASSO) • Elastik ağ • En Küçük Açılı Regresyon (LARS)	Karmaşık modeller üzerinde daha iyi olan regresyon algoritmaları.
Karar ağacı algoritmaları	• Sınıflandırma ve Regresyon Ağacı (CART) • Yinelemeli dikotomatör 3 (ID3) • C4.5 ve C5.0 (güçlü bir yaklaşımın farklı sürümleri) • Ki-kare Otomatik Etkileşim Algılama (CHAID) • Güdük karar	Verilerdeki özniteliklerin gerçek değerlerine dayalı olarak bir karar modeli oluşturur. Sınıflandırma ve regresyon problemlerine yönelik veriler

Tablo 1.1. (Devamı)

	• M5 • Koşullu Karar Ağaçları	üzerinde eğitilir. Hızlı ve doğru karar veren bir makine öğrenimi modelidir.
Bayes algoritmaları	• Naive Bayes • Gauss Saf Bayes • Çok terimli Naif Bayes • Ortalamalı Tek Bağımlılık Tahmincileri (AODE) • Bayes İnanç Ağı (BBN) • Bayes Ağı (BN)	Sınıflandırma ve regresyon problemleri için Bayes teoremini uygulayan algoritmalarıdır.
Kümeleme algoritmaları	• k-Means • k-Medians • Beklenti Maksimizasyonu (EM) • Hiyerarşik kümeleme	Grupların ortak özellikleri maksimum olacak şekilde düzenlenmesini sağlar.
Birliktelik kuralı öğrenme algoritmaları	• Apriori algoritması • Eclat algoritması	Verilerdeki değişkenler arasında gözlemlenen ilişkileri en iyi açıklayan kuralları çıkarır. Büyük veri kümeleri üzerindeki önemli ve yararlı ilişkilerin keşfedilmesini sağlar.
Yapay sinir ağı algoritmaları	• Algılayıcı • Çok Katmanlı Algılayıcılar (MLP) • Geri Yayılım • Stokastik Gradyan İniş • Hopfield Ağı • Radyal Temel Fonksiyon Ağı (RBFN)	Biyolojik sinir ağı yapısından ilham alan modeldir. Regresyon ve sınıflandırma problemleri başta olmak üzere her türlü problemde kullanılabilen algoritmalarıdır.
Derin öğrenme algoritmaları	• Evrişimli Sinir Ağı (CNN) • Tekrarlayan Sinir Ağları (RNN'ler) • Uzun Kısa Vadeli Bellek Ağları (LSTM'ler) • Yığılmış Otomatik Kodlayıcılar • Derin Boltzmann Makinesi (DBM) • Derin İnanç Ağları (DBN)	Çok büyük ve karmaşık sinir ağları oluştururlar. Ses, görüntü, metin gibi büyük veri kümeleriyle çalışırlar.
Boyut azaltma algoritmaları	• Temel Bileşen Regresyonu (PCR) • Kısmi En Küçük Kareler Regresyonu (PLSR) • Sammon Haritalama • Çok Boyutlu Ölçekleme (MDS)	Denetimsiz öğrenme ile verilerdeki doğal yapıları arar ve kullanır, daha az bilgi kullanarak

Tablo 1.1. (Devamı)

	• Lineer Diskriminant Analizi (LDA) • Karışım Diskriminant Analizi (MDA) • Kuadratik Diskriminant Analizi (QDA) • Esnek Diskriminant Analizi (FDA)	verileri özetleyebilir.
Topluluk algoritmaları	• Artırma • Önyüklemeli toplama (Bagging) • AdaBoost • Ağırlıklı Ortalama (Karıştırma) • Yığılmış Genelleme (Yığınlama) • Gradient Güçlendirme Makineleri (GBM) • Gradyan Artırılmış Regresyon Ağaçları (GBRT) • Rastgele Orman	Bağımsız olarak eğitilmiş birden fazla zayıf modelin genel tahmini yapması için bir araya getirilmiş modelidir.
Diğer algoritmalar	• Özellik seçim algoritmaları • Algoritma doğruluk değerlendirmesi • Performans ölçüleri • Optimizasyon algoritmaları • Hesaplamalı zeka (evrimsel algoritmalar, vb.) • Bilgisayarla Görme (CV) • Doğal Dil İşleme (NLP) • Öneri Sistemleri • Pekiştirmeli Öğrenme • Grafik Modeller	Özel görevler için geliştirilmiş algoritmalar, sayısı artırılabilir.

En yaygın makine öğrenmesi algoritmaları Bayes sınıflandırıcısı, k-en yakın komşu, karar ağaçları, destek vektör makineleri, lojistik regresyon ve yapay sinir ağlarıdır.

1.3.11. Derin Öğrenme

Derin öğrenme (Deep Learning), yapay zekâ teknikleri arasında en önemlilerinden biridir ve makine öğrenmesinin bir alt dalıdır. Derin öğrenme verileri anlamak, tanımak ve karmaşık görevleri gerçekleştirmek için kullanılan bir tekniktir. Temel olarak, derin öğrenme, büyük veri setleri üzerinde karmaşık görevleri gerçekleştirmek için çok katmanlı yapay sinir ağlarını kullanma yaklaşımıdır. Derin öğrenme, girdi verilerini işlemek, anlamak ve temsil etmek için bu çok katmanlı yapay sinir ağlarını kullanarak karmaşık desenleri ve ilişkileri öğrenir.

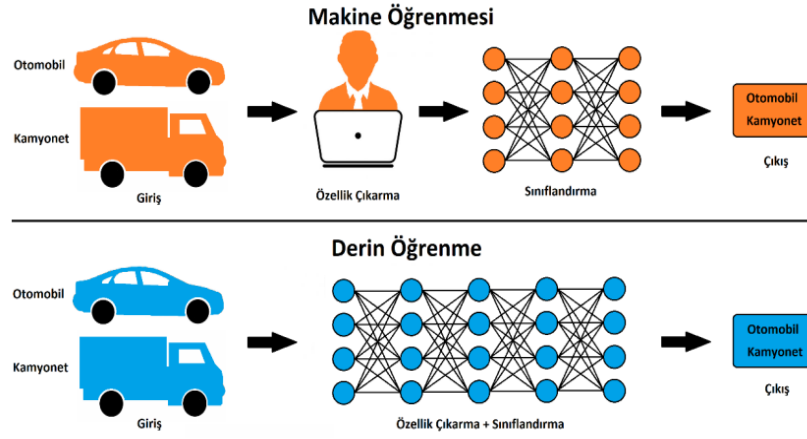
Derin öğrenme için tek katmanlı olarak çalışan makine öğrenmesinin çok katmanlı bir yapıda çalışan hali olduğu söylenebilir. Yani makine öğrenmesi algoritmalarıyla, çözümü için üzerinde çalışılan algoritmada sadece bir durum elde edilirken; derin öğrenme ile birden çok durumu ayrı ayrı öğrenen makine öğrenmesi algoritmaları, birbirleriyle çok katmanlı bir yapıda çalışarak sonuç üretir (Karaca, 2021).

Derin öğrenmenin temeli yapay sinir ağlarıdır. Bu ağlar, beynin sinir hücrelerinin çalışma biçiminden esinlenir. Derin öğrenme beyinde bulunan nöronlar gibi çalışır, çok katmanlı (derin) sinir ağlarının kullanılmasıyla tanınır. Bu katmanlar, girdi verilerini işlemek ve temsil etmek için kullanılır.

Derin öğrenme, öğrenme sürecine dayalı bir tekniktir. Bu, bir yapay sinir ağının, verilerle eğitilirken hata hesaplayıp bu hatayı azaltmaya çalıştığı bir süreci içerir. Ağ, eğitim verileri üzerinde tekrar tekrar eğitilir ve bu süreç ağı daha iyi sonuçlar elde etmek için ayarlar.

Makine öğrenmesi ve derin öğrenmeye baktığımızda birbirine benzer niteliktedirler. Fakat yüksek veri kapasitesine sahip olan yapısı ile derin öğrenme, birden fazla katmanı aynı anda işleyebilen bir yapıya sahiptir. Yani, makine öğrenmesi tek katman üzerinde işlem yaparken, derin öğrenme eşzamanlı olarak birçok katmanda işlem yapabilir.

Makine öğrenmesinde iki nesneyi birbirinden ayırt etmek gerekirse ilgili özellikleri makineye tanıtmak gerekir. Mesela bir kuş ile balık resmini ayırmamız gerekirse makine öğrenmesinde ona bir tanıtmak yapmamız gerekir. Kanadı varsa kuş, yoksa balık gibi. Fakat derin öğrenmede buna gerek kalmaz çünkü kendisi öğrenir. Derin öğrenmede resimleri sisteme göstermek yeterlidir. Sistem ayırıcı özellikleri, kuralları ve farkları kendisi bulmaktadır. Yine kamyon ve araba örneğine bakarsak makine öğrenmesinde kasası yoksa araba, kasası varsa kamyon gibi tanıtmak yapmamız gerekir. Fakat yine derin öğrenmede resimleri göstermek yeterlidir çünkü bu ayrımı kendisi yapabilir. Şekil 7’de derin öğrenme ile makine öğrenmesi arasındaki farklılık gösterilmiştir.



Şekil 1.11. Makine Öğrenmesi ile Derin Öğrenmenin Farkı (Ağgöl, 2021)

Verinin çoğalması, yapay zekânın özelliklerini iyi bir şekilde açığa çıkartmasını sağlamaktadır. İşler karmaşıktıkça; yapay zekâdan makine öğrenmesine, daha da karmaşıktıkça da; derin öğrenmeye geçişler başlayacaktır (Erhandı, 2020).

Bu durum Şekil 1.12’de gösterilmektedir.



Şekil 1.12. Yapay Zekâ, Makine Öğrenmesi ve Derin Öğrenme Karşılaştırması (Erhandı, 2020)

Derin öğrenme modelleri, genellikle büyük veri setleri gerektirir ve yüksek hesaplama gücüyle eğitilir. Bu nedenle, bu teknik özellikle büyük şirketler ve araştırma kurumları tarafından kullanılır. Derin öğrenme, birçok alanda başarıyla kullanılmaktadır. Örnek uygulama alanları arasında görüntü ve nesne tanıma, ses tanıma, otomotivde otonom sürüş, sağlık alanındaki teşhisler, doğal dil işleme, öneri sistemleri ve daha birçok şey bulunmaktadır.

İKİNCİ BÖLÜM

YAPAY ZEKÂ, TOPLUM VE ETİK

Yapay zekâ, son yıllarda teknolojik gelişmelerin merkezinde yer alarak toplumsal yaşamda köklü dönüşümlere neden olmaktadır. İnsan zekâsına özgü öğrenme, problem çözme ve dil işleme gibi becerileri taklit edebilen bu sistemler; sağlık, eğitim, finans, hukuk ve ulaşım gibi birçok alanda verimliliği artırmakta, karar alma süreçlerini hızlandırmakta ve yeni olanaklar sunmaktadır. Ancak bu teknolojik ilerlemenin sunduğu fırsatların yanı sıra, toplumsal yapı üzerinde ciddi etkileri ve etik sorunları da beraberinde getirdiği açıktır.

Yapay zekânın doğrudan etkilerinden biri işgücü piyasasında görülmektedir. Otomasyonun artmasıyla birlikte bazı meslek grupları risk altına girerken, düşük vasıflı işçilerin iş güvencesi zayıflamakta ve gelir eşitsizliği derinleşmektedir. Bunun yanında, YZ sistemlerinin bireylerin kişisel verilerini toplaması ve analiz etmesi, mahremiyetin ihlali ve gözetim toplumu algısının güçlenmesi gibi ciddi endişelere yol açmaktadır.

Bu bağlamda, yapay zekâ teknolojileriyle ilişkili en kritik sorunlardan biri de algoritmik önyargıdır. Yapay zekâ sistemleri, beslendikleri veri setleri doğrultusunda karar alır. Bu veriler geçmişteki toplumsal önyargıları ve eşitsizlikleri barındırıyorsa, sistemler bu kalıpları tekrar üretme eğiliminde olur. Örneğin, işe alım süreçlerinde kullanılan algoritmalar geçmişte belirli grupların dışlandığı verilerle eğitildiğinde, benzer ayrımcı sonuçları tekrar edebilir. Aynı şekilde, suç tahmin algoritmalarının bazı etnik grupları daha fazla hedef alması ya da kredi değerlendirme sistemlerinin kadınları dezavantajlı konuma getirmesi, algoritmik önyargının doğrudan sonuçlarıdır.

Bu tür önyargılar sadece bireyleri değil, tüm toplumları etkileyebilecek yapısal adaletsizliklere neden olabilir. Ayrıca, algoritmaların genellikle "kara kutu" gibi çalışması—yani nasıl karar verdiklerinin dışarıdan anlaşılamaması—bu sorunların tespitini ve çözümünü daha da zorlaştırmaktadır. Yüz tanıma sistemlerinde görülen ırksal tanıma hataları, hem ayrımcılık hem de mahremiyet ihlali açısından çarpıcı örnekler sunmaktadır.

Bu nedenle, yapay zekânın toplumsal etkilerini hafifletmek ve adil kullanımını sağlamak adına bazı temel ilkelerin benimsenmesi gerekmektedir. Bunlar arasında şeffaflık, hesap verebilirlik, denetlenebilirlik, etik ilkeler doğrultusunda sistem tasarımı ve çeşitliliği yansıtan veri setlerinin kullanımı öne çıkmaktadır. İnsan haklarına saygılı, kapsayıcı ve adalet temelli bir yapay zekâ anlayışı, teknolojinin sunduğu potansiyelin toplumun tüm kesimleri için faydalı olmasını sağlayacaktır.

Yapay zekânın sunduğu imkânlar ne kadar büyük olursa olsun, bu teknolojinin etik ve toplumsal etkileri göz ardı edildiğinde mevcut eşitsizlikler daha da derinleşebilir. Bu nedenle "insan odaklı yapay zekâ" anlayışı, sadece teknolojik gelişmenin değil, aynı zamanda adil bir geleceğin temeli olmalıdır.

2.1. YAPAY ZEKÂNIN TOPLUM ÜZERİNDEKİ ETKİLERİ

Teknolojik gelişmelerin paralelinde, toplumsal dönüşüm süreçlerinin yaşanması, yapay zekânın sosyal yaşam üzerindeki etkilerinin belirginleşmesine yol açmaktadır. Bu durum, yapay zekânın bireyler ve topluluklar arasındaki etkileşimleri, iletişim biçimlerini ve genel yaşam standartlarını yeniden şekillendirdiği anlamına gelmektedir. Yapay zekânın sosyal dinamiklere entegre edilmesi, bireylerin günlük yaşamlarında ve toplum genelinde köklü değişimlere sebep olabilecek potansiyellere sahiptir.

Özellikle bulut bilişim, nesnelerin interneti, zeki sensörler, büyük veri sistemleri, 3 boyutlu baskı ve blok zinciri uygulamaları akıllı şehir uygulama mallarında olduğu gibi bazı yeni teknolojilerle birlikte kullanıldığında daha etkili olmaktadır. Ulaşımdan sağlığa, tarımdan imalata, kamu yönetiminden ticarete, enerjiden eğitime yaşamın her alanında daha etkili ve daha işlevsel olan yeni sistemler üretilmektedir. Bu teknolojileri yaşamın bu alanlarında kullanarak insan vücuduna implante edilen Artificial Intelligence (AI) yani yapay zekâ ürünleri, toplantılarda karar kurulu üyelerinin rolünü yerine getirebilecek AI robotları, giyilebilir internet, insan için hayatı kolaylaştıran akıllı şehir uygulamaları, gerçek zamanlı karar vermek için artırılmış gerçeklik kullanan görme sistemleri, 3D baskılı imalat sistemleri, yapay organlar ve sürücüsüz araçlar vb. kendisini gösterecektir. Bu konuda kapsamlı bir değerlendirme Dünya ekonomik Forum tarafından hazırlanan bir raporda bulunabilir (Öztemel, 2020).

Bu gelişmeler incelendiğinde, yapay zekânın bilim dünyasının ilgisini öğrenme, hayal gücü, yaratıcılık, bağımsız düşünme ve icat unsurlarının yanı sıra etkili inovasyon rekabetine yönlendirerek bilgi temelli bir topluma doğru ilerlemelere yön verdiği açıkça görülmektedir.

Yapay zekâ teknolojisinde son yıllarda yaşanan hızlı ilerlemeler, beraberinde hem büyük fırsatlar hem de önemli tartışmaları getirmiştir. Bu tartışmaların merkezinde, yapay zekânın insanlık için taşıdığı potansiyel tehlikeler yer almaktadır. Yapay zekâ, özünde insan eliyle geliştirilen bir sistemdir ve bu nedenle doğası gereği, geliştiricisinin niyeti ve ahlaki değerleri doğrultusunda şekillenir. Bir teknolojik sistemin nasıl davranacağı, onu programlayan birey ya da grubun yaklaşımıyla doğrudan ilişkilidir. Dolayısıyla, iyi niyetle geliştirilen yapay zekâ uygulamaları topluma büyük faydalar sağlayabilirken; kötü niyetle tasarlanan sistemler zarar verici sonuçlar doğurabilir.

Bu bağlamda temel sorun, yapay zekânın kendisinden ziyade, insanın doğasında yatar. Yani tehlike, teknolojiden çok, onu geliştiren insanın zihniyetinde ve etik anlayışındadır. İnsan, kendi zekâ kapasitesi çerçevesinde bir yapay zekâ modeli geliştirebilir. Hiçbir yapay zekâ, geliştiricisinin ötesine geçemez; çünkü o sistem, insan aklının ve bilgisinin bir ürünüdür. Bu durumda, yapay zekâdan kaynaklanabilecek riskleri ortadan kaldırmak ya da en aza indirmek, yine insanın elindedir. Kötü niyetli bir robotu etkisiz hale getirmek, kötü niyetli bir insanı durduraktan daha kolaydır. Çünkü robot, belirli algoritmalar çerçevesinde çalışır ve bu algoritmalar aşılabilir ya da kontrol altına alınabilir. İnsanın geliştirdiği sistemler, yine insan zekâsıyla kontrol edilebilir.

Bu nedenle, teknolojinin güvenli ve etik bir şekilde gelişimi için öncelikli olarak insanın eğitimi önem taşır. Eğitim yoluyla bireylerin etik değerleri benimsemeleri, toplumda sorumluluk bilinciyle hareket etmeleri sağlanabilir. Aynı zamanda, hem insanlardan hem de yapay zekâ sistemlerinden gelebilecek tehditleri gözlemleyip önlem alabilecek sosyal ve teknik denetim mekanizmalarına yatırım yapılmalıdır. Bu noktada, en büyük sorumluluk bilim insanlarına düşmektedir. Bilim insanlarının yalnızca yapay zekâyı geliştirmekle kalmayıp, bu teknolojiyi doğru anlatmaları, toplumda oluşabilecek korkulara karşı bilinç oluşturup, yapay zekânın faydalarını ön plana çıkarmaları da büyük önem taşır.

Robotların insanlardan üstün olacağına dair inançlar, bu teknolojiyi geliştirme motivasyonunu zayıflatabilir. Oysa yapay zekâ, insan zekâsının bir ürünü olduğu sürece, insanın denetiminden çıkamaz. Bu bilinçle hareket eden bireyler, yapay zekâyı daha verimli ve insanlığa hizmet eden sistemler üretmek için bir araç olarak kullanabilir. Yine de, her teknolojiye olduğu gibi, yapay zekânın da olumlu yönlerinin yanında bazı olumsuz etkileri olabileceği gerçeği unutulmamalıdır. Bu etkilerle baş edebilmek için insan merkezli bir yaklaşım benimsemek ve yapay zekânın gelişim sürecini etik, hukuki ve sosyal sorumluluklar çerçevesinde yürütmek gereklidir.

Bu teknolojinin şu alanlarda gerekli önlemler alınmadığı takdirde olumsuz sonuçlar doğurabileceği öne sürülebilir:

- *Robotlar ile ilişkisi güçlü olan toplum bireylerinin toplumsal olaylara karşı duyarlılıklarında önemli oranda düşüş gözlemlenebilecektir. Birlikte olmak, sosyal paylaşımlar gerçekleştirmek, gelenekleri yaşatmak gibi durumlar, teknolojik gelişmeler ile zaten kısıtlanmaktadır. Yapay zekâ eğer kontrol edilmezse buna destek verebilecek en önemli araç olabilecektir.*

- *Mahremiyetin ortadan kalkması ya da azalmasına yol açacak sistemler geliştirilmesinde ve bu yönde sistemlerin maharetleri ile performanslarının artırılmasında yapay zekânın önemli oranda desteği olabilecektir.*

- *Siber saldırıların yaygınlaşması ve bilgi güvenliğinin azalması amacı ile bu teknolojinin kullanılması, kötü niyetlilerin başarısına katkı üretebilecektir.*

- *Bilgi hırsızlığının (özellikle kimlik bilgilerinin ele geçirilmesi) artmasında istenmeyen sonuçların doğmasını kolaylaştıracaktır.*

- *Takip edilme oranlarının artması ve kişisel özgürlüklerin kısıtlanmasına destek verecek yöntemleri yapay zekâ ile daha kolay uygulamak mümkün olabilecektir.*

- *Yapay zekâ, tüm sunulan hizmetlerde 7/24 hizmet sunulması beklentisini önemli oranda artıracak ve bu durum da hizmet üretenlerin üzerinde önemli oranda bir yükün doğmasına neden olabilecektir.*

- *Robotları kullanarak haksız kazanım elde etmek isteyenlere önemli oranda cesaret verebilecektir.*

• *Yakın gelecekte topluma yayılan ve istenmeyen bilgilerden kurtulamama önemli bir sorun alanı olarak görünmektedir. Özellikle toplumların yanlış yönlendirilmesini önlemek için bilgiye erişimin kısıtlanması veya yanlış bilgi yayılımı, artan manipülasyon girişimleri belki de yukarıdakiler arasında yapay zekânın topluma verebileceği en önemli zararlardan olabilecektir. Özellikle sosyal medya üzerinden yukarıda belirtildiği üzere “belirli bir amaca yönelik” olarak dağıtılacak olan bir mesajın direkt muhataplarına ulaşması, hata her muhatabına onun en fazla etkilenebileceği bir üslup ve tarz ile ulaşabileceği bir sistemin geliştirilmesi düşünüldüğünde konunun vahameti daha iyi anlaşılabilir. Üstelik bu mesajın ulaşmasını engellemek, sosyal medya sistemini kapatmak ile de mümkün olamayabilecektir (Öztemel, 2020).*

Yapay zekânın toplum üzerinde olumlu mu yoksa olumsuz mu etkisi olduğu sorusunun kesin bir cevabı olmasa da yapay zekânın yadsınamaz bir etkisi olduğu aşîkâr.

2.2. ETİK NEDİR?

Etik, bireylerin ve toplumların davranışlarını yönlendiren “doğru” ve “yanlış” kavramlarını inceleyen, bu kavramların temel ilkelerini ve genel hükümlerini felsefi bir temelde ele alan bir disiplindir. Bu alan, yalnızca bireylerin davranışlarını değil, aynı zamanda bu davranışların ardındaki niyetleri, değerleri ve sorumlulukları sorgular. Etik, bireyin hem kendine hem de başkalarına karşı olan yükümlülüklerini anlamlandırmasını sağlayarak bireysel ve toplumsal yaşamda bir denge kurulmasına katkıda bulunur.

"Etik" terimi, Yunanca kökenli olup "kişinin karakteri" manasına gelen "ethikos" kelimesinden türetilmiştir. "Ethikos" kelimesi ise, "ahlaki doğa " veya "karakter" anlamına gelen "ethos" teriminden gelmektedir. Etik, doğru, yanlış, iyi, kötü, adalet, suç, ahlaki ve ahlaksız gibi kavramları tanımlamak suretiyle ahlaka dair sorulara yanıt bulmayı amaçlar (Koçak,2024). Bu etimolojik yapı, etik kavramının yalnızca dışsal eylemleri değil, aynı zamanda bireyin içsel dünyasını, karakterini ve ahlaki duruşunu da kapsadığını ortaya koymaktadır. Bu yönüyle etik, bireyin kimliği ile davranışları arasındaki ilişkiyi değerlendiren kapsamlı bir düşünce alanıdır.

Koçak (2024), etiği; doğru, yanlış, iyi, kötü, adalet, suç, ahlaki ve ahlaksız gibi temel kavramları tanımlamak ve bu kavramlara ilişkin ahlaki sorulara tutarlı ve sistematik yanıtlar üretmek amacıyla var olan bir disiplin olarak tanımlar. Bu tanıma göre etik,

bireyin davranışlarını yalnızca kendi vicdanına göre değil, toplumsal değerlerle de uyumlu olacak şekilde değerlendirmesine olanak tanır. Etik, bu yönüyle birey ile toplum arasındaki bağları güçlendirirken, insan davranışlarının evrensel ilkeler ışığında değerlendirilmesine de katkı sunar.

Günümüzde yaşanan hızlı teknolojik ilerlemeler etik kavramına yeni bir boyut kazandırmıştır. Yapay zekâ, robotik sistemler, biyoteknoloji, sosyal medya gibi alanlarda yaşanan gelişmeler, klasik etik yaklaşımların ötesine geçilmesini gerekli kılmaktadır. Etik artık yalnızca bireysel davranışları değil, sistemsel süreçleri ve teknolojik gelişmelerin insanlık üzerindeki etkilerini de kapsayacak şekilde genişlemektedir. Bu bağlamda etik, hem bilimsel araştırmalarda hem de teknolojik uygulamalarda bir yön gösterici olarak önemli bir yere sahiptir.

Etik, bireyler ve toplumlar için günlük yaşamda karşılaşılan sorunları, çatışmaları ve seçimleri inceleyen bir disiplindir. Etik alanının esas amacı, bireylere davranışlarıyla ilgili rehberlik sunmak ve toplumsal normlar, adalet, dürüstlük ile sorumluluk gibi kavramların üzerine düşündürmektir (Darı, Koçyiğit, 2024). Etik bu yönüyle, adalet, dürüstlük, toplumsal sorumluluk ve güven gibi temel değerleri esas alarak bireylere yol gösterir. Bu tanım, etik düşüncenin yalnızca bireysel düzeyde değil, toplumsal yaşamın sürdürülebilirliği açısından da ne denli hayati bir rol üstlendiğini gözler önüne sermektedir.

Ayrıca etik; kişisel değer yargıları, toplumsal normlar, dinî inançlar, kültürel kodlar ve felsefî bakış açıları gibi birçok farklı unsurun etkileşimiyle şekillenen karmaşık bir yapıdadır. Bu durum, insan davranışlarının yalnızca bireyin özgür iradesiyle değil, aynı zamanda bulunduğu sosyal, kültürel ve tarihsel bağlamla da şekillendiğini göstermektedir. Etik, bu karmaşık yapıyı analiz ederek bireyin davranışlarını etkileyen tüm unsurların farkına varılmasına yardımcı olur.

Sonuç olarak etik, bireylere yalnızca “ne yapılmalı?” sorusunu sormaz; aynı zamanda “neden yapılmalı?”, “bu eylemin sonuçları kimleri etkiler?” gibi daha derinlemesine sorularla bireyin karar alma sürecine yön verir. Etik, bireyin sadece kendi vicdanıyla değil, aynı zamanda içinde yaşadığı topluma, doğaya ve geleceğe karşı da sorumlu olduğunu hatırlatır. Özellikle çağımızda bilimsel ve teknolojik ilerlemeler hızla

devam ederken, etik değerlerin bu gelişmelerle birlikte düşünülmesi bir zorunluluk haline gelmiştir. Etik, bilinçli, sorumlu ve adil bir toplumsal yapı inşa etmede vazgeçilmez bir rehberdir.

2.3. YAPAY ZEKÂ ETİĞİ

Yapay zekâ etiği, yapay zekânın geliştirilmesine, sorumlu kullanımına ve sonuçlarına tavsiyelerde bulunan bir dizi ahlaki ilke ve yönergeyi temsil eder. Yapay zekâ etiği, gelişen teknolojik sistemlerin yalnızca işlevsellik açısından değil, insan onuru, toplumsal adalet ve güven temelinde de değerlendirilmesini gerekli kılmaktadır. Yapay zekânın tasarımı, kullanımı ve etkileri bağlamında ortaya çıkan etik sorunlar, hem bireylerin temel haklarını koruma hem de toplumun bütününe ilgilendiren kararların adil şekilde alınmasını sağlama açısından derinlemesine tartışılması gereken bir alan oluşturmaktadır. Bu noktada yapay zekâyâ yönelik etik ilkeler, yalnızca teorik bir yönelim değil, pratik uygulamalarda yol gösterici bir çerçeve sunmaktadır.

IBM'e göre yapay zekâ etiğinin üç temel ilkesi, bu teknolojinin insan merkezli ve sorumlu bir biçimde geliştirilmesine olanak tanımaktadır. İlki olan kişilere saygı ilkesi, her bireyin özerkliğine duyulan saygıyı esas alırken; özerkliği sınırlı olan bireylerin de özel olarak korunmasını zorunlu kılmaktadır. İkinci ilke olan yararlılık, yapay zekânın bireylerin refahına katkı sağlayacak biçimde kullanılmasını, zarar verme riskinin en aza indirilmesini ve etik değerlere uygun biçimde hareket edilmesini vurgulamaktadır. Üçüncü temel ilke olan adalet ise kararların eşitlik ve kapsayıcılık temelinde alınmasını, bireyler arasında ayrımcılık yapılmaksızın hakların ve yükümlülüklerin adil biçimde dağıtılmasını öngörmektedir (<https://pecb.com/article/artificial-intelligence-ethics-and-social-responsibility>).

Bu ilkeler rehberliğinde değerlendirildiğinde, yapay zekâ alanında karşılaşılan en kritik etik sorunlardan birinin algoritmik önyargı olduğu görülmektedir. Algoritmaların beslendiği veri setlerinde var olan tarihsel veya kültürel önyargılar, yapay zekâ sistemleri tarafından yeniden üretilmekte ve bu durum, bireylerin eşitlik ilkesine aykırı şekilde muamele görmesine neden olabilmektedir. Karar destek sistemlerinde, işe alım süreçlerinde, kredi değerlendirmelerinde ya da sağlık hizmetlerinde, algoritmik

önyargılar sistematik ayrımcılık yaratmakta ve teknolojinin yararlılık ve adalet ilkeleriyle çelişmesine yol açmaktadır.

Bu nedenle, yapay zekâ etiğinin uygulanabilirliğini sağlamak için algoritmik önyargının hem tanımlanması hem de giderilmesi öncelikli bir gereklilik haline gelmiştir.

Yapay zekânın en önemli etik zorluklarından bazıları şunlardır:

Yapay zekâ ve yalnızlık:

Yapay zekâ ile ilgili etik problemlerden biri, yalnızlık konusuyla ilgilidir. Yapay zekâ, insanlar arasındaki kişisel ilişkileri etkileyebilir ve toplumsal yalnızlığı artırabilir. Yapay zekânın etik kullanımı ve insanların bu teknolojiyle nasıl etkileşimde bulunacakları konularında daha fazla düşünülmesi gerekmektedir.

Günümüzde yapay zekâ araçlarına olan bağımlılığın artışı, insanların toplumsal bağlarından giderek uzaklaşmasına neden olarak, "dijital insan" olarak adlandırabileceğimiz yeni bir insan türünün ortaya çıkmasına yol açmaktadır (Taşçı, Çelebi, 2020).

Yapay zekânın yaygınlaşması, teknolojik bağımlılık ve anonimlik kavramlarını gündeme getirerek bireylerin yalnızlık hissini artırma potansiyeline sahiptir. Yapay zekâ destekli uygulama ve platformlar, kullanıcıların dikkatini çekmek amacıyla kişiselleştirilmiş içerikler sunarak onların bu platformlarda daha uzun süre vakit geçirmelerini sağlamakta ve böylece dijital içerik bağımlılığının artmasına zemin hazırlamaktadır. Teknolojik bağımlılık, bireylerin doğrudan sosyal etkileşimlerini kısıtlayarak yalnızlık duygusunu artırmakta, anonimlik de bireylerin gerçek hayattaki sosyal bağlantılarını zayıflatarak sosyal izolasyon yaratabilmektedir. Bu durum, bireylerin sanal etkileşimlere daha fazla yönelmesine ve fiziksel dünyadaki sosyal ilişkilerini ihmal etmesine yol açmaktadır (Çağırkan, 2024).

Sonuç olarak, yapay zekâ ve dijital platformların kişilerin sosyal bağlantıları üzerindeki etkileri, toplumda daha derin bir yalnızlık ve sosyal dışlanma riski yaratma potansiyeline sahiptir.

Yapay Zekâ ve Gizlilik

Dijital çağda kişisel veriler inanılmaz derecede değerli bir meta haline geldi. Her gün çevrimiçi olarak üretilen ve paylaşılan büyük miktarda veri, işletmelerin, hükümetlerin ve kuruluşların yeni iç görüler elde etmesini ve daha iyi kararlar almasını sağlamıştır. Ancak bu veriler aynı zamanda bireylerin paylaşmak istemeyebileceği veya kuruluşların kendi rızaları olmadan kullanabileceği hassas bilgiler de içeriyor. İşte bu noktada gizlilik devreye girmektedir.

Gizlilik, kişisel bilgilerin gizli tutulması ve yetkisiz erişime açık olmaması hakkıdır. Bireylerin kişisel verileri ve bunların nasıl kullanılacağı üzerinde kontrol sahibi olmalarını sağlayan temel bir insan hakkıdır.

Günümüzde, toplanan ve analiz edilen kişisel veri miktarı artmaya devam ettiği için gizlilik çeşitli nedenlerden dolayı çok önemlidir. Birincisi, bireyleri kimlik hırsızlığı veya dolandırıcılık gibi zararlardan korur. Ayrıca, kişisel özerkliğin ve kişisel bilgiler üzerindeki kontrolün korunmasına yardımcı olur ki bu da kişisel haysiyet ve saygı için gereklidir. Ayrıca mahremiyet, bireylerin kişisel ve profesyonel ilişkilerini gözetlenme veya müdahale korkusu olmadan sürdürmelerine olanak tanır. Kişisel özerklik, koruma ve adalet için gerekli olan temel bir insan hakkıdır. Yapay zekâ hayatımızda daha yaygın hale gelmeye devam ettikçe, teknolojinin etik ve sorumlu bir şekilde kullanılmasını sağlamak için gizliliğimizi koruma konusunda uyanık olmalıyız (Van Rijmenam, 2023).

Yapay Zekâ ve Eşitsizlik

Yapay Zekâ, modern dijital çağda güçlü ve dönüştürücü bir etki yaratarak, endüstrilerde köklü değişikliklere yol açmış, verimliliği artırmış ve çeşitli alanlarda yeni fırsatlar ortaya çıkarmıştır. Ancak yapay zekâ teknolojilerinin gelişmeye ve yayılmaya devam etmesiyle birlikte, eşitsizliklerin daha da derinleşmesi konusunda bazı endişeler gündeme gelmiştir. Yapay zekâ ve eşitsizlik arasındaki ilişki, toplumdaki erişim, fırsatlar ve güç dengeleriyle ilgili önemli soruları ön plana çıkarmaktadır.

Yapay zekâ teknolojisinin kullanımı, bazı topluluklar veya ülkeler arasında erişim eşitsizliğine neden olabilir. Bu, teknolojik gelişmeye eşit katılım olmayabilir ve bu da toplumsal eşitsizlikleri artırabilir. Aynı zamanda yapay zekâ kullanma noktasında

cinsiyetler arasındaki var olan fırsat eşitsizliği iş dünyasında öne çıkan başlıca sorunlardan biridir.

Dünyanın lider yetenek şirketi Randstad, “Yetenek Kıtlığını Anlamak: Yapay Zekâ ve Eşitlik raporunu yayınladı. Raporda, işverenleri tarafından kadınların yalnızca yüzde 35'ine yapay zekâyâ erişim olanağı sunulduğu, erkeklerde ise bu oranın yüzde 41 olduğu belirtiliyor. Rapor ayrıca, erkeklerin %71'i YZ becerilerine sahip olduğunu belirtirken, kadınlarda bu oran %29 ile sınırlı kaldığını ortaya koyuyor. Erkeklerin %47'si, iş yerinde YZ çözümlerini problem çözme amacıyla kullandığını ifade ederken, kadınlarda bu oran %37'de kalıyor. Üstelik kadınlar aldıkları eğitimin onları kariyerlerinde teknolojiyi kullanmaya yeterince hazırladığı konusunda erkeklere (%35) kıyasla (%30) daha az güvende hissediyorlar.

Dünya Ekonomik Forumu'nun Küresel Cinsiyet Eşitliği Raporu 2024'e göre, ekonomik katılım ve fırsat eşitliği açığı yalnızca %60,5 oranında kapanmış durumda. Kadınların yapay zekâ eğitimine erişiminin artırılmasının, küresel ekonomiye 172 trilyon dolar değerinde katkı sağlayabileceği öngörülüyor (HR Dergi, 2024).

Yapay zekâ, aynı zamanda bazı sektörlerde iş gücü dinamiklerini değiştirebilir. Belirli mesleklerin otomatizasyonu, bazı iş gücü gruplarını olumsuz etkileyebilir ve bu da ekonomik eşitsizliği artırabilir.

Bu sorunları ele almak için, yapay zekâ geliştiricileri ve karar alıcılar, algoritmik önyargıyı azaltmaya yönelik çözümler geliştirmeye ve etik standartlara uymaya odaklanmalıdır. Ayrıca, yapay zekâ teknolojisinin kullanımıyla ilgili politika ve düzenlemelerin, toplumsal eşitsizlikleri azaltacak şekilde tasarlanması da önemlidir (Sheikhi, 2022).

Yapay Zekâ ve Siyasi Propaganda

Yapay zekâ teknolojisinin siyasi propaganda amacıyla kullanımı konusu, teknolojinin güçlenmesiyle birlikte daha fazla önem kazanmıştır. Eskiden önemli bir seçim öncesinde bilginin dağıtımı broşürler ve afişler vasıtasıyla yapılmaktayken günümüzde bu tür bilgilerin yayılımı ağırlıklı olarak dijital platformlarda gerçekleşmekte ve seçmenlerin düşüncelerini daha şahsi bir düzeyde etkilemek için ileri düzey teknolojilerden faydalanılmaktadır.

Yapay zekâ, kullanıcıların çevrimiçi davranışlarından elde ettiği verilerle kişiselleştirilmiş içerik üretebilir. Bu, belirli siyasi görüşleri destekleyen veya karşı çıkan içerikleri belirli hedef kitlelere yönlendirmek için kullanılabilir. Bu, kullanıcıların dünya görüşlerini şekillendirebilir ve güçlendirebilir.

Yapay zekâ algoritmaları, sosyal medya platformlarında belirli siyasi mesajları yayma ve kullanıcılara yönlendirme konusunda etkili olabilir. Bu, yanıltıcı bilgilerin ve propagandanın daha etkili bir şekilde yayılmasına olanak tanır. Facebook gibi platformların sağladığı büyük veri setleri, potansiyel seçmenlerin ilgi alanları, etnik kökenleri ve belirli durumlarda hangi duygusal hallere sahip oldukları gibi ayrıntılı profillerin oluşturulmasına olanak tanımaktadır. Bu tür bilgiler, siyasi kampanyalarda bireysel odaklı mesajların ve hedeflenmiş reklamların tasarlanmasında kullanılabilmektedir (Rouhiainen, 2020).

Yeni dijital teknolojiler, gerçek ve sahte medyayı ayırt etmeyi giderek daha zor hale getirmektedir. Bu soruna katkıda bulunan en son gelişmelerden biri, yapay zekâ (AI) kullanarak bir kişinin gerçekte söylemediği veya yapmadığı şeyleri tasvir eden hiper-gerçekçi videolar, yani deepfake'lerin ortaya çıkışıdır. Deepfake teknolojisi, sosyal medyanın geniş erişim gücü ve hızıyla birleştiğinde, ikna edici sahte içeriklerin hızla milyonlarca insana ulaşmasına olanak tanımaktadır. Bu durum, toplum üzerinde potansiyel olarak ciddi ve olumsuz etkiler doğurabilir (Öztürk, 2024). Yapay zekâ, deepfake teknolojisi aracılığıyla ses ve görüntü manipülasyonu yapabilir. Bu, siyasi figürlerin sözlerini çarpıtarak veya tamamen sahte videolar oluşturarak kamuoyunu yanıltabilir.

Ayrıca yapay zekâyâ entegre edilmiş chatbotlar, dezenformasyonun yayılmasında kritik bir rol oynamaktadır. Bu chatbotlar, içerik üretimi amacıyla sosyal medya platformlarını kullanan yeni nesil otomatik sistemlerdir. Kullanımı kolay sosyal medya platformları, bu botların yaygınlaşmasında büyük etkiye sahiptir. Paylaşılan içerikler, sosyal ağları toplumsal algı yönetiminin önemli bir unsuru haline getirmektedir. Bu çerçevede, chatbotlar toplumsal algı yönetiminin önemli bir bileşeni olarak karşımıza çıkmaktadır. Yapay zekâ destekli chatbotlar siyasi konular üzerine etkileşimde bulunabilmekte; bu durum, siyasi görüşlerin şekillenmesine veya yanıltıcı bilgi içeren etkileşimlerin oluşturulmasına yönelik olarak kullanılabilmektedir. (Köçeri, 2023)

Bu tür durumlar, demokratik süreçlere zarar verebilir, toplumsal bölünmeyi artırabilir ve siyasi kararlar üzerinde istenmeyen etkiler yaratabilir. Bu nedenle, yapay zekânın siyasi propaganda amacıyla kullanımı konusunda dikkatli olunması ve uygun önlemlerin alınması önemlidir. Çeşitli ülkeler ve topluluklar, yapay zekâ etik kuralları ve düzenlemeleri oluşturarak bu tür kötüye kullanımlara karşı korunma sağlamaya çalışmalıdır.

Yapay zekânın silah olarak kullanımı

Yapay zekâ, özellikle askeri alanlarda çeşitli amaçlar için kullanılabilmekte ve bu kullanımlar, tartışmalı etik ve güvenlik sorunlarına yol açabilmektedir.



Şekil:2.1 Otonom Silahlar

<https://www.yenicaggazetesi.com.tr/yapay-zekâ-destekli-robotlar-guvenlik-tehdidi-olusturuyor-866735h.htm>

Yapay zekâ, otonom silah sistemlerinde kullanılabilir. Otonom silahlar, insan müdahalesi olmadan hedef belirleme ve saldırı yapma yeteneğine sahip silah sistemleridir. Bu durum, etik ve hukuki sorunlara yol açabilir ve kontrolsüz otonom silahların yanlış hedeflere yönelmesi gibi tehlikelere neden olabilir. Otonom sistemler, savaş alanlarında hızlı kararlar alarak insan kaynaklı hataları en aza indirmeyi amaçlasa da, insan hakları ihlalleri riskini önemli ölçüde artırma potansiyeline sahiptir. Otonom silahların insan unsurlarını devre dışı bırakması, savaş hukuku ve çatışma kurallarını zayıflatma potansiyeline sahiptir; bu durum, masum sivillerin yaşamlarını tehlikeye sokabilir ve savaşın insani boyutunu ihmal etmeye neden olabilir.

Otonom silah sistemleri, kendilerine özgü etik sorunlar barındırmaktadır. Bu sistemler, düşman ile sivil ayrımını yaparken yanlış değerlendirme yapabilir veya bu ayrımı yerinde gerçekleştiremeyebilir. Özellikle savaş alanlarında kullanılan yapay zekâ destekli silahlar, hatalı kararlar alarak insan hakları ihlallerine neden olma potansiyeline sahiptir. Eğer bir otonom silahın aldığı bir karar neticesinde masum bir sivilin yaşamını yitirmesi durumu söz konusu olursa, bu olayda sorumluluğun kime ait olacağı önemli bir tartışma konusu haline gelmektedir. Otonom sistemlerin kullanım sonucu insan hakları ihlalleri meydana geldiğinde, sorumluluğun yapay zekâ geliştiricilerine, bu teknolojileri kullanan devletlere ya da uluslararası hukuk bağlamında başka aktörlere mi ait olduğu konusu belirsizliğini sürdürmektedir (Turan,2024).

Yapay zekâ ve robotik gelişirken, askeri kuruluşlar da muhtemelen bunları kendi çıkarları için kullanma yolları keşfedecek. Ağustos 2017’de Elon Musk ve yirmi altı ülkeden yüz on altı başka CEO ve Yapay zekâ araştırmacısı bir araya gelerek Birleşmiş Milletler ’den Yapay zekâ silahların kullanımını yasaklamasını isteyen açık bir mektup imzaladılar. Bu açık mektuptan önemli bir paragrafta şöyle deniyor: “ [Ölümcül otonom silahlar] geliştirildiklerinde, silahlı çatışmanın hiç olmadığı kadar büyük bir ölçekte ve insanların kavrayabileceğinden daha hızlı zaman ölçeklerinde gerçekleşmesine neden olacaklar. Bu Pandora’nın Kutusu bir kere açılınca, kapatmak zor olacak. Dolayısıyla Yüksek Akit Taraflardan bizi bu tehlikelerden korumanın bir yolunu bulmalarını rica ediyoruz” (Rouhiainen, 2020).

Yapay zekâ, görüntü analizi, ses analizi ve diğer veri işleme yetenekleri sayesinde hedef tanıma ve analizde kullanılabilir. Bu, hedeflere daha hassas ve etkili saldırılar yapılmasına olanak tanır, ancak aynı zamanda yanlış anlamalara ve sivil hedeflerin zarar görmesine neden olabilir.

Askeri siber harekâtını ilerletmek için Yapay zekâ anahtar teknoloji olarak düşünülmekte, siber uzay çağında sadece insan zekâsına güvenmenin yanlış olacağı değerlendirilmektedir. Klasik siber güvenlik araçları bilinen kötücül kodların uyumuna bakmakta, dolayısıyla korsanlar kodlarının küçük bir kısmını değiştirerek bilgisayardaki savunmayı aşmaya çalışmaktadırlar. Yapay zekâ destekli araçlar diğer taraftan, ağ aktivitesinin daha geniş modelleri üzerinden anomalileri tespit için eğitilmekte ve dolayısıyla ataklara karşı daha etkin olabilmektedirler (Türe, Topuz, 2019). Ancak, aynı

teknoloji siber saldırganlar tarafından da kullanılabilir, bu da siber savaşların daha sofistike ve karmaşık hale gelmesine neden olabilir.

Yapay zekânın silah olarak kullanılması, uluslararası toplumda etik ve hukuki standartlar konusunda bir dizi soru ortaya çıkarmaktadır. Birçok ülke ve uzman, yapay zekâ tabanlı silah sistemlerinin geliştirilmesi ve kullanılması konusunda düzenlemeler yapılması ve uluslararası anlaşmaların oluşturulması gerektiğini savunmaktadır.

Yapay zekânın algoritmik önyargısı ve taraflılığı:

Yapay zekâ sistemlerinin temelini oluşturan algoritmalar, bu teknolojilerin nasıl çalıştığını, nasıl öğrendiğini ve nasıl karar verdiğini belirleyen yapısal mekanizmalardır.

Yapay zekâ, insan zekâsını taklit etmeyi amaçlayan sistemler bütünüdür ifade ederken; algoritmalar, bu sistemlerin bilgi işleme süreçlerini yöneten kural ve işlemler dizisidir. Yani algoritmalar, yapay zekânın düşünmesini ve öğrenmesini sağlayan mantıksal iskeletlerdir. Yapay zekâ sistemlerinin başarısı, büyük ölçüde bu algoritmaların doğruluğuna, verimliliğine ve uyarlanabilirliğine bağlıdır. Algoritmalar, belirli bir problemi çözmek veya bir hedefe ulaşmak için tasarlanmış sistematik adımlardır. Algoritmaların çoğu, örüntü tanıma, tahmin, sınıflandırma gibi işlemler için eğitilmiştir ve genellikle büyük veri kümeleriyle çalışırlar.

Yapay zekâ sistemlerinde algoritmaların işlevselliği kadar etik sonuçları da son yıllarda önemli bir tartışma konusu haline gelmiştir. Çünkü algoritmalar yalnızca matematiksel değil, aynı zamanda toplumsal etkileri olan kararlar verebilmektedir. Algoritmaların beslendiği verilerde bulunan önyargılar, bu sistemlerin çıktılarına doğrudan yansiyabilir. Örneğin, geçmiş verilerde cinsiyet, ırk veya sosyoekonomik durum temelli ayrımcılıklar yer alıyorsa, algoritmalar bu önyargıları yeniden üretebilir. Bu durum, algoritmik önyargı olarak adlandırılmakta ve yapay zekâ etiği çerçevesinde ciddi bir sorun olarak değerlendirilmektedir. Bu bağlamda, yapay zekâ teknolojilerinin geliştiricileri için yalnızca teknik yeterlilik değil, aynı zamanda etik farkındalık da büyük önem taşımaktadır. Algoritmaların ne şekilde oluşturulduğu, hangi verilerle eğitildiği ve nasıl değerlendirildiği soruları, sadece bilimsel değil aynı zamanda etik ve toplumsal sorumluluk çerçevesinde de ele alınmalıdır (Hayata Yönelik Bilim-Teknik, 2024). Bu nedenle yapay zekâ ve algoritmalar üzerine yapılan çalışmalar, disiplinler arası bir

yaklaşım gerektirmekte; mühendislik, bilişim, hukuk, sosyoloji ve felsefe gibi alanların ortak katkılarıyla daha güvenli ve adil sistemler geliştirilmesi mümkündür.

Algoritmik önyargı, bir algoritmanın çıktılarında sistematik olarak belirli bireyler ya da gruplar aleyhine sonuçlar üretmesi olarak tanımlanabilir (Barocas, Selbst, 2016). Bu önyargılar, genellikle eğitim verilerinde yer alan tarihsel eşitsizliklerden, modelleme sürecindeki varsayımlardan ya da yetersiz denetim mekanizmalarından kaynaklanmaktadır. Yapay zekâ algoritmaları, önyargı barındırma potansiyeliyle etik tartışmaların odağında yer almaktadır. Algoritmaların kullanım amacı, daha adil ve tarafsız karar süreçleri oluşturmak olmalıdır; zira bu sistemler teorik olarak tarafsızdır ve önyargı içermemesi beklenir. Ancak algoritmaların çıktıları, bireylerin kimliğini ifşa edecek şekilde verileri işlediğinde ya da cinsiyet, cinsel yönelim, etnik köken, inanç sistemi, politik görüşler ya da sağlık bilgileri gibi hassas veriler üzerinden ayrımcılığa neden olduğunda, ciddi etik sorunlar ortaya çıkmaktadır.

Yapay zekâ ilk ortaya çıktığında, insan kaynaklı yanlılığı ortadan kaldırarak objektif bir şekilde işlev gören bir teknoloji olacağı düşüncesi, onun en güçlü yönlerinden biri olarak kabul ediliyordu. Ancak, yapay zekâ sistemlerinin matematiksel veya değerler açısından tam anlamıyla tarafsız olmadığı ortaya çıkmıştır. Çünkü bu ürünler ve uygulamalar, potansiyel olarak yanlılık, önyargı ve kötü niyet içeren insanlar tarafından geliştirilmektedir. Zaten algoritmalar başlangıçta yanlılıkla oluşturulduğunda, yapay zekânın çıktılarının güvenilirliği de azalmaktadır.

Algoritmalar oluşturmak için kullanılan veriler, algoritmalara ve ürettikleri klinik önerilere yansıyan bir yanlılık içerebilir. Yapay zekâ algoritmaları, onları kimin geliştirdiğine, programcılarının, şirketlerin veya sağlık sistemlerini kuranların/yönetenlerin motivasyonlarına/yararlarına bağlı olarak sonuçları çarpıtmak için de tasarlanabilir. Yapay zekâ tasarımlarında bir yanlılık kaynağı olarak örneğin Volkswagen'in, testler sırasında azot oksit emisyonlarını azalmış göstererek, araçlarının emisyon testlerini geçmelerine izin veren algoritması gibi, belirli sonuçların elde edilmesini amaçlayan algoritmaların üretilmesi olanaklıdır (Güvercin, 2020). Yanlılıkların erken aşamalarda oluşması durumunda, diğer aşamaları da etkilemesi kaçınılmaz olmaktadır.

Yapay zekâ sistemleri, eğitim aldıkları verilerdeki ırk, cinsiyet ya da diğer demografik faktörlere dair önyargıları öğrenip yeniden üretebilirler. Bu da, yapay zekânın toplumdaki mevcut önyargıları devam ettirmesine yol açar. Sonuç olarak, bu durum, farklılıkların ve azınlık gruplarının daha az temsil edildiği bir dünya düzeninin pekişmesine neden olabilir. İnsanlar, toplumsal normlar ve kabul görmüş doğrular çerçevesinden bakmaya eğilimli hale gelirler, bu da büyük veri uçurumunun ortaya çıkmasına zemin hazırlar (Kayıhan, 2024).

Yapay zekâda algoritmik önyargıya bilinen bir örnek, yüz tanıma sistemlerinde cinsiyet veya ırk temelli ön yargılar olabilir. Çoğu yüz tanıma sistemi, eğitim verilerindeki dengesizlikler veya yanlılık nedeniyle bazı grupları daha iyi tanıma eğiliminde olabilir. Bazı yüz tanıma sistemleri, belirli etnik grupları daha iyi tanırken, diğer etnik grupları daha kötü tanır veya hatalı sonuçlar üretebilir. Bu, eğitim verilerinin temsil eksikliği veya çeşitliliğin yetersiz olması nedeniyle ortaya çıkabilir. Bu tür ön yargılar, insanların farklı cinsiyetler, ırklar veya etnik kökenlere sahip olduğu bir dünyada yapay zekânın tarafsız ve adil bir şekilde çalışmasını engelleyebilir.

ÜÇÜNCÜ BÖLÜM

ALGORİTMİK ÖNYARGI BAĞLAMINDA ETİK İHLALLER

Yapay zekâ sistemleri toplumsal yaşamın çeşitli alanlarında giderek daha fazla kullanılmakta ve bu teknolojilerin tarafsız olduğu yönünde yaygın bir algı bulunmaktadır. Ancak, algoritmaların eğitildiği veri setlerinin ve bu sistemlerin geliştirildiği bağlamların, var olan toplumsal önyargıları yeniden üretebildiği görülmektedir. Bu durum, özellikle ceza adaleti, işe alım süreçleri ve kamu güvenliği gibi alanlarda ciddi etik sorunlara yol açmaktadır (Barocas, Selbst, 2016). Bu bölümde, algoritmik önyargının farklı boyutlarını incelemek amacıyla altı örnek vaka, nitel içerik analizi yöntemiyle tematik kodlama kullanılarak derinlemesine değerlendirilmiştir.

3.1 ARAŞTIRMANIN SORUNSALI

Yapay zekâ teknolojilerinin hızla gelişmesi ve toplumsal yaşamın çeşitli alanlarına entegre olması, beraberinde önemli etik sorunları da gündeme getirmektedir. Bu teknolojilerin tarafsız ve objektif olduğu yönündeki yaygın algıya rağmen, yapay zekâ sistemlerinin beslendiği veri setlerindeki tarihsel ve toplumsal önyargılar, algoritmalar aracılığıyla yeniden üretilmekte ve sistematik ayrımcılığa yol açabilmektedir. Özellikle algoritmik önyargı, yapay zekânın adalet, eşitlik ve şeffaflık ilkelerini ihlal etmesine neden olarak, bireylerin temel haklarını ve toplumsal düzeni tehdit eden bir sorun haline gelmiştir.

Bu çalışmanın temel sorunsalı, yapay zekâ sistemlerinde ortaya çıkan algoritmik önyargıların etik ihlaller bağlamında nasıl şekillendiğini ve bu ihlallerin toplumsal dönüşüm sürecinde ne tür sonuçlar doğurduğunu incelemektir. Araştırma, özellikle cinsiyet, ırk ve sosyoekonomik statü temelli ayrımcılık örneklerini merkeze alarak, algoritmaların karar verme süreçlerindeki etik ihlalleri analiz etmeyi amaçlamaktadır.

3.2. ARAŞTIRMANIN SORULARI

1. Yapay zekâ işe alımlarda güvenilir bir iş kaynağı sorumlusu olabilir mi?

2. Sağlık alanında yapay zekâ destekli teşhis sistemlerinde cinsiyet, ırk veya coğrafi kökene dayalı önyargılar hasta sonuçlarını nasıl etkilemektedir?

3. Algoritmik önyargılar, yapay zekâ sistemlerinde hangi etik ilkeleri (adalet, şeffaflık, hesap verebilirlik, eşitlik vb.) ihlal etmektedir ve bu ihlallerin toplumsal sonuçları nelerdir?

4. Yapay zekâ teknolojilerinin beslendiği veri setleri ve tasarım süreçleri, önyargıların yeniden üretilmesine nasıl katkıda bulunmaktadır?

5. Yapay zekâ kaynaklı etik ihlallerin (işe alımda ayrımcılık, yüz tanımda ırksal önyargı) önlenmesi için hangi düzenleyici ve teknik önlemler geliştirilebilir?

3.3. ARAŞTIRMANIN AMACI

Bu çalışmanın amacı, yapay zekâ teknolojilerinin toplumsal dönüşüm üzerindeki etkilerini, özellikle etik ihlaller bağlamında incelemek ve bu ihlallerin temelinde yatan algoritmik önyargı sorununu nitel içerik analizi yöntemiyle değerlendirmektir. Çalışma, yapay zekâ sistemlerinin cinsiyet, ırk ve sosyoekonomik statü gibi faktörlere dayalı ayrımcılık örneklerini analiz ederek, bu teknolojilerin adalet, şeffaflık, hesap verebilirlik ve eşitlik gibi etik ilkelerle uyumunu sorgulamaktadır.

Araştırma, seçilen vakalar üzerinden algoritmik önyargının nasıl ortaya çıktığını, hangi toplumsal eşitsizlikleri pekiştirdiğini ve bu sorunların çözümüne yönelik önerileri tartışmayı hedeflemektedir. Bu kapsamda, yapay zekânın insan hakları, mahremiyet ve sosyal adalet üzerindeki potansiyel risklerini disiplinlerarası bir yaklaşımla ele alarak, teknolojik gelişmelerin etik ve toplumsal sorumluluk çerçevesinde yeniden değerlendirilmesine katkı sağlamaktadır.

Çalışmanın nihai hedefi, yapay zekâ uygulamalarının daha adil, kapsayıcı ve insan odaklı bir şekilde tasarlanabilmesi için politika yapıcılara, geliştiricilere ve akademik çevrelere yol gösterici bir perspektif sunmaktır.

3.4. ARAŞTIRMANIN ÖNEMİ

Bu çalışma, yapay zekânın toplumsal dönüşüm üzerindeki etkilerini etik ihlaller bağlamında ele alarak, disiplinlerarası bir yaklaşım sunmakta ve bu yönüyle literatürdeki

mevcut boşluklardan birini doldurmayı amaçlamaktadır. Yapay zekâ teknolojilerinin yalnızca teknik değil, aynı zamanda etik, sosyolojik ve kültürel boyutlarını da içeren bir değerlendirme çerçevesi oluşturulmuştur. Bu kapsamda çalışma, yapay zekâ temelli uygulamaların birey hak ve özgürlükleri üzerindeki etkilerini sorgulamakta; gözetim, ayrımcılık, mahremiyet ihlalleri ve algoritmik önyargı gibi güncel ve kritik etik sorunları incelemektedir.

Literatürde genellikle teknolojik gelişmelerin avantajlarına odaklanan çalışmalardan farklı olarak, bu çalışma, yapay zekânın toplumun yapısal dinamiklerini nasıl dönüştürdüğünü ve bu dönüşüm sürecinde ortaya çıkan etik sorunları derinlemesine analiz etmektedir. Böylelikle çalışma, teknoloji ve etik ekseninde bir denge kurmayı hedefleyerek, hem akademik çevreler hem de politika yapıcılar için yol gösterici olma potansiyeline sahiptir.

Ayrıca çalışma, mevcut etik değerler ile güncel yapay zekâ uygulamaları arasında bir köprü kurmakta; bu değerlerin ihlallerinin pratik hayatta nasıl karşılık bulduğunu örnek olaylar üzerinden tartışmaktadır. Bu bağlamda çalışma, sadece teorik bir çerçeve sunmakla kalmamakta, aynı zamanda çözüm önerileriyle literatüre özgün bir katkı sağlamaktadır.

3.5. ARAŞTIRMANIN KAPSAM VE SINIRLILIKLARI

Bu çalışma, yapay zekâ teknolojilerinin toplumsal dönüşüm üzerindeki etkilerini, özellikle etik ihlaller bağlamında ele almayı amaçlamaktadır. Çalışma, algoritmik önyargı, dijital gözetim, ayrımcılık ve mahremiyetin ihlali gibi dijital toplumlarda sıkça karşılaşılan etik sorunlar üzerine odaklanmaktadır. Bu doğrultuda, birey-toplum-devlet üçgeninde yapay zekânın etkileri değerlendirilmiş; teknolojik gelişmelerin sosyal yapılar üzerindeki yansımaları disiplinlerarası bir yaklaşımla irdelenmiştir.

Araştırmanın kapsamı, 2014-2024 yılları arasında öne çıkan ve toplumun geniş kesimlerini etkileyen yapay zekâ uygulamalarıyla sınırlıdır. Bu bağlamda, söz konusu uygulamalarla ilişkili etik sorunlar incelenmiş; özellikle algoritmik önyargının yol açtığı etik ihlaller üzerine odaklanılmıştır. Çalışmada yapay zekânın teknik boyutlarına yer verilmekle birlikte, çalışma esas olarak teknolojinin etik ve toplumsal etkilerini

irdelemektedir. Ayrıca, inceleme alanı büyük oranda teorik ve kavramsal çerçevede kalmış, deneysel veri ya da saha çalışmasına yer verilmemiştir.

Coğrafi olarak ise çalışma, küresel düzeydeki gelişmeleri takip etse de, ağırlıklı olarak Batı merkezli uygulama örneklerine ve literatüre dayanmaktadır.

Son olarak, çalışmada yer alan etik değerlendirmeler, yapay zekâ uygulamalarında ortaya çıkan algoritmik önyargılara odaklanmakta ve bu değerlendirmeler belirli örnek olaylar üzerinden yapılmaktadır. Dolayısıyla, ulaşılan bulgular tüm yapay zekâ uygulamalarına genellenmemektedir. Ancak, bu sınırlılıklarına rağmen çalışma yapay zekânın toplumsal boyutlarına ilişkin etik tartışmalar açısından önemli bir teorik zemin sunmaktadır.

3.6. ARAŞTIRMANIN YÖNTEMİ

Araştırmada betimsel analiz yöntemi, veri toplama tekniği olarak da nitel içerik analizi kullanılmıştır.

İçerik analizi, nitel araştırma yöntemleri arasında yer alan ve yazılı, görsel veya işitsel materyallerin içeriğini sistematik bir şekilde analiz etmeyi amaçlayan bir araştırma yöntemidir. Bu yöntem, araştırmacıların belirli bir konu veya sorun hakkında bilgi edinmelerine ve bu bilgilerden anlamlı sonuçlar çıkarmalarına olanak tanır. İçerik analizi, esnek yapısı nedeniyle iletişim bilimleri, siyaset bilimi, psikoloji, sosyoloji, pazarlama gibi çeşitli sosyal bilim alanlarında yaygın olarak kullanılmaktadır. İçerik analizi, araştırma probleminin belirlenmesi, veri toplama, verilerin kodlanması ve analiz edilmesi gibi adımları içeren bir süreçtir. Bu yöntem, özellikle medya araştırmalarında metinlerin ardındaki gizli anlamları, ideolojik alt yapıları veya toplumsal mesajları ortaya çıkarmak için etkili bir araç olarak kabul edilir (Alanka, 2024).

Bu çalışmada, algoritmik önyargının farklı bağlamlardaki etkilerini anlamak amacıyla nitel içerik analizi yöntemi tercih edilmiştir. Özellikle çok yönlü toplumsal etkilerin değerlendirilmesinde sayısal verilerden ziyade anlamın, bağlamın ve yapısal örüntülerin yorumlanmasına odaklanan bu yöntem, çalışmanın bilgi edinme süreciyle örtüşmektedir.

Çalışmada analiz edilen altı vaka, kamuoyunda geniş yankı uyandırmış, akademik literatürde sıklıkla referans verilmiş ve yapay zekâ sistemlerinin farklı bağlamlardaki önyargılı etkilerini gösteren örnekler olarak seçilmiştir. Veriler, medya haberleri ve akademik yayınlardan derlenmiştir. Her bir vakanın kendine özgü bağlamı korunarak, genel eğilimler ve yapısal sorunlar ortaya çıkarılarak çözüm önerileri sunulmuştur.

Nitel içerik analizinde çözümlemelerin yapılması için temalar belirlenmiştir. Belirlenen temalar şunlardır:

Cinsiyet Temelli Ayrımcılık

Algoritmaların eğitildiği veri setleri, geçmişteki toplumsal cinsiyet eşitsizliklerini yansıtıyorsa, yapay zekâ sistemleri bu önyargıları sürdürme eğiliminde olur. Örneğin, bir iş başvuru algoritması geçmişte erkek adayların daha fazla tercih edildiği bir veri setiyle eğitildiyse, kadınları sistematik olarak düşük puanlama eğiliminde olabilir. Bu durum, kadınların istihdamda dezavantajlı hale gelmesine yol açar.

İrk Temelli Ayrımcılık

Veri setlerinde yer alan ırksal önyargılar, algoritmaların belirli etnik veya ırksal grupları ayrımcılığa uğratmasına neden olabilir. Örneğin, bazı ülkelerde polis tahmin sistemlerinde siyahi bireylerin daha yüksek riskli olarak işaretlenmesi, geçmiş tutuklama verilerinin taraflı olmasından kaynaklanabilir. Bu, yapay zekânın ırksal adaletsizlikleri yeniden üretmesine neden olur.

Sosyoekonomik Temelli Ayrımcılık

Algoritmalar, düşük gelirli bireylerin yeterince temsil edilmediği veri setleriyle eğitildiğinde, bu grupları dezavantajlı konuma düşürebilir. Örneğin, sahne tanıma algoritmaları genellikle üst gelir grubuna ait çevre görüntüleriyle eğitildiğinden, düşük gelirli bölgeleri “harabe” ya da “gecekondu” olarak sınıflandırabilir. Bu durum, sosyoekonomik önyargının teknoloji aracılığıyla pekişmesine yol açar.

Veri Seti Önyargısı

Algoritmaların çıktısı doğrudan eğitildiği veriye bağlıdır. Eğer veri seti temsili değilse (örneğin belirli bir grubun örnekleri eksik ya da dengesizse), sonuçlar da önyargılı olacaktır. Bu, sistemin belirli grupları yeterince tanımamasına veya yanlış

genelleştirmeler yapmasına yol açar. Veri temsiliyetinin eksikliği, algoritmik önyargının temel kaynaklarından biridir.

Şeffaflık Eksikliği

Bir algoritmanın nasıl çalıştığını, hangi kriterleri dikkate aldığını ya da hangi verilerle eğitildiğini bilmek çoğu zaman mümkün değildir. Bu "kara kutu" durumu, algoritmik kararların sorgulanamamasına ve toplumda güvensizliğe neden olur. Şeffaflık eksikliği, özellikle yapay zekâ sistemlerinin etik denetimini zorlaştırır.

Hesap Verebilirlik

Algoritmik sistemlerin neden belirli bir kararı verdiğini açıklayamaması, sorumluluğun kime ait olduğunun da belirsizleşmesine yol açar. Bu durum, yanlış kararlar karşısında bireylerin hak arama mekanizmalarını zayıflatır. Hesap verebilirlik eksikliği, algoritmik önyargının tespit edilmesini ve düzeltilmesini engelleyebilir.

Adalet

Algoritmik sistemlerin karar alma süreçlerinde eşit muamele sağlanması, adalet ilkesinin gereğidir. Ancak önyargılı verilerle eğitilen veya taraflı algoritmalarla çalışan sistemler, bu ilkeyi ihlal edebilir. Bu, özellikle kamu hizmetlerinde (eğitim, sağlık, hukuk) kullanılan algoritmalar için ciddi etik ve hukuki sorunlar doğurur.

3.7. BULGULAR

Seçilen vakalar cinsiyet temelli ayrımcılık, ırk temelli ayrımcılık, sosyoekonomik temelli ayrımcılık, veri seti önyargısı, şeffaflık eksikliği, hesap verebilirlik ve adalet temaları çerçevesinde nitel içerik çözümlemesine tabi tutulmuştur.

3.7.1. Vaka I: Cinsiyet Temelli Ayrımcılık: Amazon'un İşe Alım Algoritması

Amazon, 2014 yılında, işe alım sürecini otomatikleştirecek bir yapay zekâ sistemi geliştirmeye başladı. Şirketin Edinburgh'daki mühendislik merkezi, 2015 yılına kadar aday özgeçmişlerini analiz etmek ve iş başvurusu yapan kişileri sıralamak için bir algoritma geliştirdi. Bu sistem, başvuruların cinsiyet, ırk ve diğer demografik özelliklerine dayalı önyargılara yol açabilecek şekilde çalıştı. Örneğin, algoritma, kadınlarla ilişkilendirilen terimleri (örneğin "kadın satranç kulübü kaptanı") negatif

puanlayarak onları daha düşük puanlarla değerlendirdi. Ayrıca, kadınların mezun olduğu okullardan gelen başvuruları olumsuz şekilde sıraladı. Bu durumun temel nedeni, algoritmanın veri setinin büyük çoğunluğunun erkek başvurularından oluşmasıydı. Sonuç olarak, algoritma erkek adayları daha fazla tercih etti. Bu durum, Amazon'un, algoritmanın şeffaflık ve adalet ilkesine aykırı olduğunu fark etmesiyle sonuçlandı. Proje sonlandırıldı (Reuters, 2018).

Amazon'un işe alım algoritması vakasından elde edilen veriler, Reuters haberleri, çeşitli medya haber siteleri ve akademik literatür gibi kaynaklardan derlenmiştir. Veriler tematik kodlama ile gruplandırılmıştır. Tematik kodlama, belirli temalar etrafında analiz yapmayı içermektedir ve bu süreç, içeriklerin niteliksel bir şekilde sınıflandırılmasını sağlamaktadır.

Tematik analizde, Amazon'un algoritmasının cinsiyet temelli ayrımcılık oluşturması, veri seti önyargısı, algoritmanın şeffaflık eksiklikleri ve hesap verebilirlik gibi ana temalar belirlenmiştir.

Tablo 3.1. Tematik Kodlama ve Bulgular

Tema	Bulgular	Yorum	Veri Kaynağı
Cinsiyet Temelli Ayrımcılık	Kadınlarla ilişkilendirilen kelimeler negatif puanlandı.	Algoritma, cinsiyet önyargılarını pekiştirerek kadınlar hakkında olumsuz çağrışımlara neden olmuştur.	Reuters, 2018
Veri Seti Önyargısı	10 yıllık erkek ağırlıklı verilerle eğitildi.	Cinsiyet temsili dengesiz olan veri seti, kadınları yeterince yansıtmadığı için adil olmayan sonuçlara yol açmıştır.	Reuters 2018

Şeffaflık Eksikliği	Algoritmanın nasıl çalıştığı tam açıklanmadı.	Şeffaf olmayan sistemler, dış denetimi zorlaştırmakta ve kamu güvenini zedelemektedir.	Reuters 2018
Hesap Verebilirlik	Proje sonlandırıldı, ancak sorumluluk açıkça ifade edilmedi.	Sorumluluğun üstlenilmemesi, benzer sorunların tekrar yaşanma riskini artırmakta ve kurumsal etik açısından olumsuz bir örnek oluşturmaktadır.	Reuters 2018

Bu tablo, Amazon'un işe alım algoritmasının içeriğinde yer alan etik ihlalleri ve bu ihlallerin nasıl tespit edildiğini özetlemektedir.

Etik İlkelere Göre Değerlendirme

Amazon'un algoritması etik açıdan birkaç önemli ihlale yol açmıştır. Bu ihlaller, adalet, şeffaflık ve hesap verebilirlik gibi etik ilkelerle doğrudan ilişkilidir.

Tablo 3.2. Etik İlkelere Göre Değerlendirme

Etik İlke	İhlal Durumu	Açıklama
Adalet	Evet	Kadın adaylar sistematik olarak daha düşük puan aldı.
Şeffaflık	Evet	Algoritmanın çalışma mantığı kamuya açık değildi.
Hesap Verebilirlik	Kısmen	Amazon sistemi sonlandırdı, ancak etik sorumluluklar açıkça ifade edilmedi.

Bu tablo, algoritmanın ne şekilde etik ihlallere yol açtığını ve hangi etik ilkelere aykırı olduğunu sunmaktadır.

3.7.2. Vaka II: Irk Temelli Ayrımcılık: COMPAS Vakası

Northpointe, Inc. tarafından geliştirilen COMPAS algoritması, suçlu adaylarının gelecekteki suç işleme olasılıklarını tahmin etmek için kullanılmaktadır. Florida'nın Broward County bölgesindeki 7.000'den fazla suçlu adayının verileri üzerinden yapılan bir inceleme, algoritmanın suçluların yeniden suç işleme risklerini doğru tahmin etme oranının %61 olduğunu, ancak şiddetli suçların yeniden işlenmesi konusunda ise yalnızca %20 oranında doğru tahmin yaptığını göstermiştir. Siyah sanıkların, beyaz sanıklara göre daha yüksek riskli oldukları öngörülmüş, ancak aslında bu tahminler sıklıkla yanlış çıkmıştır. Beyaz sanıklar ise, suç işledikleri halde daha düşük riskli olarak değerlendirilmiştir (ProPublica, 2016)

Northpointe, Inc.'in geliştirdiği COMPAS (Correctional Offender Management Profiling for Alternative Sanctions) adlı algoritma, suçlu adaylarının tekrar suç işleyip işlemeyeceğini tahmin etme konusunda belirli gruplara karşı önyargılı olabilmektedir.

Bu analizde, COMPAS vakasından elde edilen veriler, ProPublica tarafından yayımlanan veriler, çeşitli medya haber siteleri ve akademik literatür gibi kaynaklardan derlenmiş ve tematik kodlama yoluyla sınıflandırılmıştır.

Tematik analizde, COMPAS algoritmasının ırk temelli ayrımcılık oluşturmaları, veri seti önyargısı, algoritmanın şeffaflık eksiklikleri ve hesap verebilirlik gibi ana temalar belirlenmiştir.

Tablo 3.3. COMPAS Algoritmasının Yanılma Oranları (ProPublica, 2016)

Sanık Tipi	Yüksek Risk Olarak Yanlış Değerlendirilenler (%)	Düşük Risk Olarak Yanlış Değerlendirilenler (%)
Siyah Sanıklar	45	28
Beyaz Sanıklar	23	48

Tablodan, siyah sanıkların, beyaz sanıklara göre yüksek risk olarak yanlış değerlendirilme oranının daha fazla olduğu açıkça görülebilmektedir.

Tablo 3.4. Tematik Kodlama ve Bulgular

Tema	Bulgular	Yorum	Veri Kaynağı
İrk Temelli Ayrımcılık	Siyah sanıklar %45 oranında suç işlememelerine rağmen yüksek riskli olarak değerlendirildi.	Algoritma, ırksal önyargıları yeniden üreterek adalet sisteminde yapay taraflılığa neden olmuş ve siyah bireyler aleyhine sonuçlar doğurmuştur.	ProPublica 2016
Veri Seti Önyargısı	Geçmişe dayalı, ırksal ayrımcılık içeren ceza adalet verileri ile eğitildi.	Tarihsel olarak önyargılı veriler, algoritmanın bu önyargıları yeniden üretmesine sebep olmuş ve sonuçlarda sistematik ırkçılık etkisi doğmuştur.	ProPublica 2016
Şeffaflık Eksikliği	COMPAS algoritmasının nasıl çalıştığı kamuya açıklanmadı.	Değerlendirilmesi gereken algoritmanın karar mekanizmaları belirsizdir, bu da şeffaflığı ve itiraz hakkını ortadan kaldırmaktadır.	ProPublica 2016
Hesap Verebilirlik	Hatalı tahminlere rağmen, algoritmanın kararlarından doğan sorumluluk açık biçimde üstlenilmedi.	Algoritmanın hatalı sonuçları için sorumluluk kabul edilmemesi, hesap verebilirlik ilkesinin ihlali ve kurumsal etik sorunu yaratmaktadır.	ProPublica 2016

Bu tabloda ProPublica (2016) verilerine dayanan bulgular, algoritmanın siyah sanıkları sistematik olarak daha yüksek riskli değerlendirdiğini, bu durumun tarihsel olarak önyargılı verilerden kaynaklandığını ve sistemin karar süreçlerinin hem şeffaf hem de hesap verebilir olmadığını göstermektedir. Bu durum, adalet, eşitlik ve etik sorumluluk ilkelerinin ihlal edildiğine işaret etmektedir.

Etik İlkeler Göre Değerlendirme

Tablo 3.5. Etik İlkeler Göre Değerlendirme

Etik İlke	İhlal Durumu	Açıklama
Adalet	Evet	Siyah sanıklar, suç işlememelerine rağmen daha yüksek riskli olarak değerlendirilmiştir; bu durum sistematik ayrımcılığa işaret etmektedir.
Şeffaflık	Evet	COMPAS algoritmasının karar mekanizması kamuya açık değildir; formüller ve değerlendirme kriterleri açıklanmamaktadır.
Hesap Verebilirlik	Kısmen	Hatalı kararların sorumluluğu şirket ya da yargı organlarınca açık biçimde üstlenilmemiştir; bu da etik hesap verebilirlik açısından sorun teşkil etmektedir.

Bu tablo, algoritmanın ne şekilde etik ihlallere yol açtığını ve hangi etik ilkelere aykırı olduğunu sunmaktadır.

3.7.3. Vaka III: MIT ve Stanford Araştırması : "Yüz Tanıma Sistemlerinde Cinsiyet ve Irk Temelli Algoritmik Önyargı"

MIT Media Lab'den Joy Buolamwini ve Stanford Üniversitesi'nden Timnit Gebru'nun yaptığı araştırma, ticari yüz tanıma yazılımlarındaki algoritmik önyargıları incelemiştir. Araştırma, üç ticari yüz tanıma programının cilt tonu ve cinsiyetle ilgili önyargılarını ortaya koymuştur. Çalışmada, bu sistemlerin açıkça cilt tonu ve cinsiyetle ilgili önemli hatalar yaptığı gözlemlenmiştir. Araştırmacılar, üç programın açık tenli erkekleri tanımlamada hata oranlarının %0,8'i geçmediğini, ancak özellikle koyu tenli kadınlarda hata oranının %20-34 arasında değişirken, veri setinde bulunan en koyu tenli kadınlarda %46'ya kadar çıktığını bulmuşlardır. Araştırmaya göre, bu tür yüz tanıma

sistemlerinin doğruluk oranları, özellikle sistemlerin eğitim verilerinin demografik çeşitliliğiyle doğrudan ilişkilidir. Çalışmaya katılan araştırmacılar, ABD'nin önde gelen teknoloji firmalarından birinin geliştirdiği yüz tanıma sisteminin %97 doğruluk oranına sahip olduğunu iddia ettiğini belirtmişlerdir. Ancak bu verilerde, %77'si erkek ve %83'ü beyaz olan bir örneklem kullanılmıştır. Çalışma, yüz tanıma algoritmalarının eğitim veri setlerinin çoğunlukla beyaz ve erkek bireyler içerdiğini ve bu dengesizliklerin ırksal ve cinsiyet temelli ayrımcılığa yol açtığını ortaya koymuştur (MIT, 2018).

Bu analizde, yüz tanıma sistemlerinde cinsiyet ve ırk temelli algoritmik önyargı konusunda ele alınan vakada elde edilen veriler, MIT Haber Ofisi tarafından yayımlanan veriler, çeşitli medya haber siteleri ve akademik literatür gibi kaynaklardan derlenmiş ve tematik kodlama yoluyla gruplandırılmıştır.

Tematik analizde, üç ticari yüz tanıma programı algoritmasının, cinsiyet ve ırk temelli ayrımcılık oluşturmaları, veri seti önyargısı, algoritmanın şeffaflık eksiklikleri ve hesap verebilirlik gibi ana temalar belirlenmiştir.

Tablo 3.6. Tematik Kodlama ve Bulgular

Tema	Bulgular	Yorum	Veri Kaynağı
Cinsiyet Temelli Ayrımcılık	Koyu tenli kadınlar için cinsiyet tanıma doğruluk oranı %20-34 arasında değişiyor.	Yüz tanıma algoritmaları, koyu tenli kadınlar üzerinde daha yüksek hata oranları göstermekte, bu da cinsiyet temelli ayrımcılığa işaret etmektedir.	MIT, 2018
İrk Temelli Ayrımcılık	Koyu tenli bireylerde hata oranı %46'ya kadar çıkıyor, özellikle siyah kadınlar üzerinde hatalar daha belirgin.	Algoritmalarda ırksal önyargılar bulunmakta; koyu tenli bireyler, özellikle siyah kadınlar, daha sık hatalı tanımlanmaktadır.	MIT, 2018

Veri Seti Önyargısı	Eğitim veri setlerinde %77 erkek ve %83 beyaz bireyler vardı.	Veri setinin cinsiyet ve ırk açısından dengesiz olması, algoritmaların bu grupları doğru tanımamasına yol açmaktadır.	MIT, 2018
Şeffaflık Eksikliği	Algoritmaların nasıl eğitildiği ve test verilerinin nasıl seçildiği hakkında yeterli bilgi verilmedi.	Şeffaflık eksiklikleri, algoritmaların adil olup olmadığı konusunda kamu güvenini zedelemekte ve dış denetim süreçlerini zorlaştırmaktadır.	MIT, 2018
Hesap Verebilirlik	Hatalı sonuçların sorumluluğu net bir şekilde belirtilmedi.	Sorumluluğun üstlenilmemesi, önyargıların devam etme riskini artırmakta ve kurumların etik sorumluluklarını yerine getirmemelerine yol açmaktadır.	MIT, 2018

Bu tablo, yüz tanıma sistemlerinde cinsiyet ve ırk temelli önyargılarla birlikte veri seti önyargısı, şeffaflık ve hesap verebilirlik eksikliklerini ortaya koymaktadır. Özellikle koyu tenli kadınlarda yüksek hata oranları dikkat çekmektedir.

Etik İlkelere Göre Değerlendirme

Tablo 3.7. Etik İlkelere Göre Değerlendirme

Etik İlke	İhlal Durumu	Açıklama
Adalet	Evet	Koyu tenli kadınlarda yüksek hata oranları, algoritmanın belirli demografik gruplara karşı sistematik ayrımcılık içerdiğini göstermektedir.

Şeffaflık	Evet	Algoritmaların nasıl eğitildiği, hangi verilerin kullanıldığı ve performans kriterleri hakkında kamuya açık bilgi verilmemektedir.
Hesap Verebilirlik	Kısmen	Hatalı sınıflandırmalara dair açık bir sorumluluk üstlenilmemiş, kurumlar etik sonuçlar karşısında yeterince hesap vermemiştir.

Bu tablo, yüz tanıma teknolojilerinde kullanılan algoritmaların uygulanma süreçlerinde temel etik ilkelerin nasıl ihlal edildiğini özetlemektedir.

3.7.4. Vaka IV: Google Fotoğraflarda "Goril" Etiketleme Skandalı

2015 yılında, yazılım geliştiricisi Jacky Alcine, Google Fotoğraflar uygulamasına yüklediği bir fotoğrafın otomatik etiketleme sistemi tarafından "goril" olarak sınıflandırıldığını fark etti. Bu durum, sosyal medyada büyük yankı uyandırdı ve Google, kullanıcıların siyahi bireylerini "goril" olarak etiketleyen algoritmasını hızla devre dışı bıraktı. Google, bu hatayı "tamamen kabul edilemez" olarak nitelendirerek özür diledi ve "goril", "şempanze" ve "maymun" gibi terimleri etiketleme sisteminden kalıcı olarak kaldırdı. Google temsilcisi, yüz etiketleme teknolojisinin hala yeteri kadar gelişmediğini söyledi (The Verge, 2018).

Bu analizde, Google Fotoğraflar uygulamasının otomatik etiketleme algoritmasının önyargılı olduğunu kanıtlayan bu vakada elde edilen veriler, temel olarak The Verge tarafından yayımlanan veriler ve çeşitli medya haber siteleri gibi kaynaklardan derlenmiş ve tematik kodlama yoluyla sınıflandırılmıştır.

Tematik analizde, Google Fotoğraflar uygulamasının ırk temelli ayrımcılık oluşturma, veri seti önyargısı, algoritmanın şeffaflık eksiklikleri ve hesap verebilirlik gibi ana temalar belirlenmiştir.

Tablo 3.8. Tematik Kodlama ve Bulgular

Tema	Bulgular	Yorum	Veri Kaynağı
İrk Temelli Ayrımcılık	Siyah bir çiftin fotoğrafı, "goriller" olarak etiketlendi.	Google Fotoğraflar algoritması, ırk temelli ayrımcılık yaparak siyahi bireyleri yanlış bir şekilde "goriller" olarak sınıflandırmıştır. Bu, algoritmalarındaki ırkçı önyargıyı ortaya koymaktadır.	The Verge 2018
Veri Seti Önyargısı	Fotoğraf etiketleme algoritması, çeşitliliği eksik bir veri seti ile eğitildi.	Algoritmanın hatalı sonuçları, kullanılan veri setindeki çeşitlilik eksikliklerinden kaynaklanmış ve bu, algoritmanın önyargılı sonuçlar üretmesine yol açmıştır.	The Verge 2018
Şeffaflık Eksikliği	Algoritmanın nasıl çalıştığı ve nasıl kararlar verdiği konusunda yeterli açıklama yapılmadı.	Google, algoritmasının çalışma şekli ve veri kullanımı konusunda şeffaf olamayarak, kullanıcıların algoritmaya olan güvenini zedelemiştir.	The Verge 2018
Hesap Verebilirlik	Google, hatayı fark ettikten sonra özür diledi ve algoritmayı düzeltme kararı aldı, ancak algoritma hala insan denetimi gerektiriyor.	Google, algoritmanın hatalı çalıştığını kabul etmesine rağmen, teknolojinin geliştirilmesi ve gelecekteki hataların önlenmesi konusunda daha fazla hesap verebilirlik gerekmektedir.	The Verge 2018

Tablo, Google Fotoğraflar algoritmasındaki ırkçı önyargıyı ve etik sorunları vurgulamaktadır. Özellikle, siyah bireylerin "goriller" olarak etiketlenmesi, veri seti önyargısı, şeffaflık eksikliği ve hesap verebilirlik yetersizlikleri gibi sorunlar algoritmanın güvenilirliğini ve adaletini sorgulamaktadır.

Etik İlkelere Göre Değerlendirme

Tablo 3.9. Etik İlkelere Göre Değerlendirme

Etik İlke	İhlal Durumu	Açıklama
Adalet	Evet	Siyahi bireylerin yanlış etiketlenmesi, algoritmanın adil olmadığını göstermektedir.
Şeffaflık	Evet	Algoritmanın işleyişi ve veri seti hakkında bilgi eksikliği, şeffaflık ilkesinin ihlalidir.
Hesap Verebilirlik	Kısmen	Hatalı etiketlemenin sorumluluğu açıkça belirtilmemiştir; ancak Google, hatayı kabul ederek düzeltme yapmıştır.

Bu tablo, Google Fotoğraflar uygulamasında kullanılan otomatik etiketleme algoritmasının ne şekilde etik ihlallere yol açtığını ve hangi etik ilkelere aykırı olduğunu göstermektedir.

3.7.5. Vaka V: Sağlık Kararlarını Yönlendiren Algoritmada Irksal Önyargı

ABD'de 70 milyondan fazla insanı etkileyen bir algoritma, hastaların gelecekteki karmaşık sağlık ihtiyaçlarını tahmin ederek onları özel bir "bakım yönetimi" programına dahil edip etmeme konusunda karar vermiştir. Bu program, hastalara ekstra sağlık kaynakları sağlayarak daha iyi izlenmelerini ve tedavilerinin daha etkili bir şekilde koordine edilmesini amaçlamaktadır. Bu kararlar, büyük ölçüde ticarileştirilmiş ve yaygın olarak kullanılan bir algoritma tarafından oluşturulan risk skorlarına dayanmaktadır. Ancak yapılan bir çalışmada, bu algoritmanın ciddi düzeyde irksal önyargı içerdiği tespit edilmiştir. Siyah ve beyaz hastalar aynı risk skoruna sahip

olduklarında, gerçek sağlık durumları önemli ölçüde farklıdır. Özellikle otomatik olarak programa dahil edilen siyah hastalar, aynı risk skoruna sahip beyaz hastalara göre daha fazla kronik hastalığa sahiptir. Kan şekeri (HbA1c) düzeylerinden tansiyon kontrolüne kadar birçok biyobelirteçte bu fark gözlemlenmiştir. Bu önyargının temel nedeni, algoritmanın sağlık ihtiyacını sağlık harcamalarıyla ölçmesidir. Ancak, sağlık harcamaları ile gerçek sağlık durumu arasında doğrudan bir ilişki bulunmamaktadır. Siyah bireyler, aynı hastalık düzeyine sahip olsalar bile daha az harcama yapılmasıyla sonuçlanan yapısal nedenlerden ötürü algoritma tarafından daha az "riskli" olarak değerlendirilmişlerdir. Bu da onları gerekli bakım programından mahrum bırakmaktadır (Obermeyer, Powers, Vogeli, Mullainathan, 2019).

Bu vaka, algoritmanın hedef fonksiyonunun sağlık değil maliyet olması nedeniyle sistematik bir ırksal önyargı ürettiğini ortaya koymaktadır. Dolayısıyla, bu önyargı sadece teknik bir mesele değil, aynı zamanda etik, toplumsal ve yapısal bir sorundur.

Bu analizde, Obermeyer ve arkadaşlarının (2019) sağlık algoritmasının ırk temelli önyargı ürettiğini kanıtlayan bu vakada elde edilen veriler, temel olarak akademik makaleler, çeşitli medya haber siteleri gibi kaynaklardan derlenmiştir.

Tematik analizde, algoritmanın ırk temelli ayrımcılık oluşturması, veri seti önyargısı, algoritmanın şeffaflık eksiklikleri ve hesap verebilirlik gibi ana temalar belirlenmiştir. Bu temalar, algoritmanın siyah hastaları beyazlarla aynı sağlık koşullarına sahip olsalar bile daha düşük riskli görmesi ve sağlık hizmetlerine eşit erişim sağlayamamaları gibi etik sorunları ortaya koymaktadır.

Tablo 3.10. Tematik Kodlama ve Bulgular

Tema	Bulgular	Yorum	Veri Kaynağı
İrk Temelli Ayrımcılık	Siyah ve beyaz hastalar aynı risk skoruna sahip olduklarında sağlık durumları önemli ölçüde farklılık gösteriyor.	Algoritma, eşit skorlar verdiği halde siyah hastalar daha ağır sağlık sorunları yaşamakta, bu da sistematik ırksal ayrımcılığa işaret etmektedir.	Obermeyer, Powers, Vogeli, Mullainathan 2019

Veri Seti Önyargısı	Siyah bireyler, aynı sağlık durumuna sahip olsalar bile sistem tarafından daha düşük maliyetli görülüyor ve daha az bakım alıyor.	Veri seti, siyah bireylerin sağlık sisteminden daha az hizmet aldığı gerçeğini yansıtıyor; bu durum algoritmada yapısal önyargıya neden oluyor.	Obermeyer, Powers, Vogeli, Mullainathan 2019
Şeffaflık Eksikliği	Algoritma, “maliyet” tahminine dayalı olarak tasarlanmış, ancak bu durum karar alıcılar tarafından yeterince anlaşılmamış olabilir.	Şeffaf olmayan hedef fonksiyon, algoritmanın neden bu şekilde davrandığını belirsiz hale getirmekte, bu da denetim ve etik değerlendirmeyi zorlaştırmaktadır.	Obermeyer, Powers, Vogeli, Mullainathan 2019
Hesap Verebilirlik	Sorunun verilerde değil, algoritmanın optimizasyon amacıyla (maliyet) olduğu tespit edildi.	Sorumluluğun belirsizliği, algoritma üreticilerinin ve uygulayıcılarının etik sonuçlardan kaçınmasına yol açabilmektedir.	Obermeyer, Powers, Vogeli, Mullainathan 2019

Bu tablo, sağlık alanında kullanılan bir algoritmanın yapısal ve sistematik ırk temelli ayrımcılık ürettiğini ortaya koymaktadır. Siyah bireyler, aynı sağlık durumuna sahip olsalar bile algoritma tarafından daha düşük riskli görülerek gerekli bakımdan mahrum bırakılmaktadır. Bu durum, veri seti önyargısı, şeffaflık eksikliği ve hesap verebilirliğin olmayışıyla birleşerek ciddi etik sorunlara yol açmaktadır.

Etik İlkeler Göre Değerlendirme

Tablo 3.11. Etik İlkeler Göre Değerlendirme

Etik İlke	İhlal Durumu	Açıklama
Adalet	Evet	Siyah hastalar, aynı sağlık durumuna rağmen beyazlara göre daha az sağlık hizmeti almakta; bu durum adalet ve eşitlik ilkelerinin açıkça ihlal edildiği anlamına gelmektedir.
Şeffaflık	Evet	Algoritmanın sadece maliyet tahminine göre çalıştığı net olarak belirtilmemiştir. Bu durum, karar vericilerin algoritmayı yanlış yorumlamasına yol açmaktadır.
Hesap Verebilirlik	Evet	Algoritmanın sağlık yerine maliyet odaklı optimize edilmesinin sonuçları hakkında açık bir sorumluluk alınmamıştır. Bu da hesap verebilirliğin eksik olduğunu göstermektedir.

Bu tablo, algoritmanın etik ilkeler açısından ciddi sorunlar barındırdığını vurgulamaktadır. Özellikle adalet ve şeffaflık ilkeleri açıkça ihlal edilmekte, hesap verebilirlik ise sağlanamamaktadır. Kısacası, algoritmanın hem adil çalışmadığı hem de bu adaletsizliğin sorumluluğunun üstlenilmediği gösterilmektedir.

3.7.6. Vaka VI: Sahne Tanıma Algoritmalarında Sosyoekonomik Önyargı

ABD’de çeşitli şehirlerden toplanan yaklaşık bir milyon görsel üzerinden yapılan analizlerde, aynı yapısal özelliklere sahip evlerin yalnızca bulundukları mahallelerin sosyoekonomik statülerine göre farklı etiketlendiği tespit edilmiştir. Örneğin, düşük gelirli bölgelerde yer alan bir ev, derin öğrenme algoritması tarafından "harabe", "gecekondu" veya "kullanım dışı bina" gibi olumsuz çağrışımlarla sınıflandırılırken; benzer yapısal özelliklere sahip bir ev yüksek gelirli bir bölgede yer aldığında "lüks iç mekân", "modern ev" ya da "estetik yapı" gibi olumlu etiketlerle tanımlanmıştır (Greene, Josyula, Si, Hart, 2024).

Bu örüntü, görüntülerin sınıflandırılmasında kullanılan yapay zekâ algoritmalarının yalnızca görsel içerik değil, aynı zamanda bu görsellerin geldiği bağlamsal sosyoekonomik verileri dolaylı olarak işlediğini ve bu nedenle ayrımcılığı pekiştiren sonuçlar doğurabileceğini göstermektedir. Bu önyargılı sınıflandırmalar, şehir planlaması, dijital haritalama, gayrimenkul değerlendirme ve otomatik reklam yerleştirme

gibi pek çok alanda ciddi toplumsal ve ekonomik eşitsizliklere neden olabilecek riskler taşımaktadır.

Bu analizde, Greene ve arkadaşlarının (2024) çalışmasına konu olan sahne tanıma algoritmasının sosyoekonomik önyargılar taşıdığını kanıtlayan bu vakada elde edilen bulgular, temel olarak akademik makaleler ve çeşitli medya haber kaynaklarından derlenmiştir.

Tematik analizde, algoritmanın sosyoekonomik temelli ayrımcılık üretmesi, veri seti önyargısı, algoritmanın şeffaflık eksiklikleri ve hesap verebilirlik yetersizliği gibi ana temalar belirlenmiştir. Bu temalar, düşük gelirli bölgelerdeki yapıları sistematik olarak olumsuz biçimde etiketleyen algoritmaların, sosyal eşitsizlikleri yeniden üretme riskini taşıdığını ve etik açıdan ciddi sorunlar doğurduğunu ortaya koymaktadır.

Tablo 3.12. Tematik Kodlama ve Bulgular

Tema	Bulgular	Yorum	Veri Kaynağı
Sosyoekonomik Temelli Ayrımcılık	Düşük gelirli bölgelerden gelen görseller sıklıkla "harabe", "gecekondu" gibi negatif etiketlerle eşleşti.	Derin öğrenme modelleri, sosyoekonomik statüye göre görüntüleri farklı şekilde sınıflandırarak ayrımcılığı pekiştirmektedir.	Greene, Josyula, Si, Hart 2024
Veri Seti Önyargısı	Eğitim verilerinin büyük çoğunluğu yüksek gelirli bölgelerden geldiği için düşük gelirli bölgeler az temsil edildi.	Veri dengesizliği, düşük gelirli bölgeleri temsil eden görsellerin sistem tarafından doğru etiketlenmesini engellemiştir.	Greene, Josyula, Si, Hart 2024

Şeffaflık Eksikliği	Algoritmanın nasıl etiketleme yaptığı açıklanmamıştır. Bu nedenle sınıflandırma kriterleri dış gözlemciler tarafından anlaşılamamaktadır.	Şeffaflık eksikliği algoritma davranışlarının değerlendirilmesini ve hataların düzeltilmesini zorlaştırmaktadır.	Greene, Josyula, Si, Hart 2024
Hesap Verebilirlik	Algoritmik etiketlerin yanlış sonuçlar doğurmasına rağmen sistem geliştiricileri sorumluluk üstlenmemiştir.	Geliştiricilerin veya kullanıcı kurumların hesap verme mekanizması olmaması, etik sorunların devamını teşvik etmektedir.	Greene, Josyula, Si, Hart 2024

Bu tablo, sahne tanıma algoritmalarının düşük gelirli bölgeleri olumsuz etiketleyerek sosyoekonomik önyargıyı pekiştirdiğini vurgulamaktadır. Veri seti önyargısı, şeffaflık eksikliği ve hesap verebilirlik sorunu, etik riskleri artırmaktadır.

Etik İlkelere Göre Değerlendirme

Tablo 3.13. Etik İlkelere Göre Değerlendirme

Etik İlke	İhlal Durumu	Açıklama
Adalet	Evet	Algoritma, düşük gelirli bölgeleri sistematik olarak olumsuz tanımlayarak sosyal adaleti ihlal etmektedir.
Şeffaflık	Evet	Algoritmanın sınıflandırma süreci açık değildir; karar mekanizması dış paydaşlara kapalıdır.
Hesap Verebilirlik	Evet	Yanıltıcı etiketlerin doğurduğu sonuçlara ilişkin hiçbir taraf sorumluluk üstlenmemiştir.

Bu tablo, sahne tanıma algoritmasının etik ilkelere uyumsuzluğunu vurgulamaktadır. Özellikle adaletin ihlali, şeffaflık eksikliği ve hesap verebilirlik sorunu algoritmanın sosyal eşitsizlikleri artırma potansiyeline sahip olduğunu göstermektedir.

3.8. BULGULARIN DEĞERLENDİRİLMESİ

Bu çalışmada elde edilen bulgular, yapay zekâ sistemlerinin insan karar süreçlerinin yerine geçebilecek kadar yaygınlaştığı günümüzde, bu sistemlerin yalnızca teknik değil, aynı zamanda toplumsal ve etik yapılarla da derin bağlara sahip olduğunu açık biçimde ortaya koymuştur. Algoritmik karar mekanizmaları, görünüşte tarafsız veri işleme yöntemleriyle çalışsa da, arka planda toplumsal önyargıların dijital biçimde yeniden üretildiği mekanizmalar hâline gelmiştir. Elde edilen her bulgu, bu teknolojilerin salt işlevsel araçlar değil, aynı zamanda ideolojik ve kültürel kodları yansıtan sistemler olduğunu göstermektedir.

Yapay zekâ algoritmalarındaki önyargı sorunu, farklı uygulama alanlarında kendini gösteren ve derinlemesine incelenmesi gereken kritik bir olgudur. Sunulan vaka analizleri bu önyargıların sistemik, yaygın ve çeşitli toplumsal kesimleri olumsuz etkileyebilecek nitelikte olduğunu açıkça ortaya koymaktadır. İşe alım süreçlerinden adalet sistemine, görüntü tanımadan sağlık hizmetlerine ve sosyoekonomik sınıflandırmalara kadar geniş bir yelpazede gözlemlenen bu durum, yapay zekâ teknolojilerinin tarafsız ve objektif olduğu yönündeki yaygın kanıyı sorgulamaktadır.

Amazon'un işe alım algoritmasının kadın adaylara sistematik olarak dezavantaj sağlaması, yalnızca teknolojik bir eksiklik değil; toplumsal cinsiyet eşitsizliğinin dijital kodlara sinmiş bir yansımasıdır. Bu bulgu, iş yaşamında süregelen ayrımcılık biçimlerinin, veri temelli sistemler aracılığıyla daha az görünür fakat daha kalıcı biçimde yeniden üretilebileceğini ortaya koymuştur. Diğer yandan, COMPAS algoritmasının Afro-Amerikan bireylere yönelik önyargılı risk skorları oluşturması, yapay zekânın "sayısal adalet" sağlama iddiasının ciddi biçimde sorgulanmasına neden olmuştur. Buradaki değerlendirme, veri temelli karar süreçlerinin, mevcut yargı sisteminin eşitsizliklerini meşrulaştırabileceğini ve bireysel özgürlükler üzerinde ciddi hak ihlallerine neden olabileceğini göstermiştir.

Yüz tanıma sistemlerine ilişkin MIT ve Stanford çalışması ise, yapay zekâ modellerinin eğitim verileri yeterince kapsayıcı olmadığında nasıl hatalı çıktılar ürettiğini gözler önüne sermektedir. Teknolojinin gelişkinliği, önyargıların ortadan kalkacağı anlamına gelmemektedir. Aksine, sistemin öğrenme süreci, kendisine sunulan örneklerle sınırlı olduğu sürece, çeşitliliğe ve eşitliğe dair sorunlar da bu modellerin içine gömülü

kalmaktadır. Google’ın siyah bireyleri “goril” olarak etiketlemesi, yalnızca teknolojik başarısızlık değil, aynı zamanda sistematik gözetim uygulamalarında marjinal grupların nasıl kolayca kurbanlaştırılabileceğine dair ciddi bir uyarıdır.

Sağlık sistemlerinde kullanılan algoritmaların siyahi hastaların tedavi önceliklerini düşürmesi, yapay zekânın yaşam hakkı gibi temel etik alanlarda dahi ne denli riskli sonuçlar doğurabileceğini göstermiştir. Bu, yalnızca bir hata meselesi değil; daha derin bir adalet problemidir. Sağlık hizmetlerine erişim gibi en temel insan haklarının dahi dijital önyargılarla şekillenmesi, teknolojinin etik denetimden bağımsız geliştirilemeyeceğinin somut göstergesidir. Sosyoekonomik düzeyi düşük bölgelerin yapay zekâ sistemleri tarafından “harabe” ya da “gecekondü” olarak sınıflandırılması ise, gelir eşitsizliklerinin dijital sistemler aracılığıyla nasıl yeniden üretildiğini gösteren çarpıcı örneklerden biridir.

Bu önyargıların temelinde, algoritmaların eğitildiği veri setlerindeki eksiklikler, dengesizlikler ve tarihsel önyargılar merkezi bir rol oynamaktadır. Yetersiz temsil, yanlış etiketleme veya geçmişteki ayrımcı uygulamaları yansıtan veriler, algoritmaların bu önyargıları öğrenmesine ve yeniden üretmesine neden olmaktadır. Birçok vakada, algoritmaların karar alma süreçlerinin şeffaflık ve anlaşılabilirlik düzeyinin düşük olduğu görülmektedir. Bu durum, önyargıların tespit edilmesini, anlaşılmasını ve giderilmesini zorlaştırmaktadır. Aynı zamanda, hatalı veya ayrımcı sonuçlar ortaya çıktığında sorumluluğun kimde olduğu ve nasıl hesap sorulacağı gibi kritik soruların cevapsız kalmasına yol açmaktadır.

Analizler, yapay zekâ algoritmalarının toplumsal boşlukta işlemeyip, mevcut eşitsizlikleri veri aracılığıyla özümlediğini ve karar alma süreçlerine yansıttığını göstermektedir. Dahası, önyargılı algoritmaların uygulanması, bu eşitsizlikleri daha da derinleştirme ve yeni ayrımcılık biçimleri yaratma potansiyeline sahiptir. Vakalar, yapay zekâ teknolojilerinin geliştirilmesi ve uygulanmasında etik ilkelerin ve bağımsız denetim mekanizmalarının ne kadar kritik olduğunu vurgulamaktadır. Teknik doğruluk tek başına yeterli olmayıp, algoritmaların adalet, eşitlik, şeffaflık ve hesap verebilirlik gibi etik değerlere uygun olarak tasarlanması ve işletilmesi gerekmektedir. Ayrıca, algoritmaların, geliştirildikleri ve uygulandıkları sosyokültürel ve ekonomik bağlamı dikkate alması

gerekmektedir. Evrensel çözümler üretme iddiasındaki algoritmaların, yerel dinamikleri ve potansiyel eşitsizlikleri göz ardı etmesi, istenmeyen sonuçlara yol açabilmektedir.

Sonuç olarak, yapay zekâ algoritmalarındaki önyargı sorunu çok boyutlu ve karmaşık bir yapıya sahiptir. Bu sorun, yalnızca teknik bir iyileştirme meselesi olmanın ötesine geçerek, derin etik, sosyal ve hukuki yansımaları bulunmaktadır. Algoritmaların yaygınlaştığı günümüz dünyasında, bu önyargıların potansiyel zararlı etkilerini en aza indirmek ve daha adil, eşitlikçi ve güvenilir yapay zekâ sistemleri inşa etmek için kapsamlı ve disiplinlerarası bir yaklaşım benimsemek zorunludur.



SONUÇ

Yapay zekâ, büyük bir teknolojik değişimi ve dönüşümü beraberinde getirmiştir. Bu durumun sağladığı avantajlarla birlikte, bir dizi önemli etik sorun ve toplumsal etki de ortaya çıkmıştır.

Bu çalışmada, yapay zekânın toplumsal dönüşüm sürecinde ortaya çıkan etik ihlallerden algoritmik önyargının temel bir sorun olduğu, seçilmiş örnek olaylar bağlamında derinlemesine incelenmiştir. Gerçek dünyadan derlenen vakalar, algoritmik önyargının salt teknik bir hata değil, köklü bir etik meseleyi temsil ettiğini ortaya koymuştur.

İncelemeler, yapay zekâ sistemlerinin karar mekanizmalarında toplumsal cinsiyet, ırk ve sosyoekonomik konum gibi etkenlere dayalı önyargıları yeniden ürettiğini göstermiştir. Bu durum yalnızca bireylerin yaşamını etkilemekle kalmayıp, toplumun genel adalet anlayışını da zedelemektedir. Kısacası, algoritmalar geçmişteki eşitsizlikleri ortadan kaldırmak bir yana, mevcut toplumsal yapıyı daha da güçlendirme eğilimi göstermektedir. Örnek olay analizleri, algoritmik önyargının farklı uygulama alanlarında benzer biçimde tezahür ettiğini ortaya koymuştur.

Cinsiyet temelli ayrımcılık Amazon'un işe alım algoritmasında; ırk temelli ayrımcılık ise COMPAS gibi adalet odaklı uygulamalarda gözlemlenmiştir. Görüntü işleme teknolojileri de benzer sorunlar barındırmaktadır: MIT ve Stanford araştırması, yüz tanıma sistemlerinin kadınlar ve farklı etnik köken grupları üzerinde hatalı performans gösterdiğini; Google Fotoğraflar uygulamasının ise siyah bireyleri "goril" olarak etiketlediğini rapor etmiştir. Örneğin sağlık sektöründe kullanılan bir algoritmanın, ırksal özelliklere dayalı modeli nedeniyle bazı hasta gruplarını ötelediği bildirilmiştir; benzer biçimde, yakın zamanda gerçekleştirilen bir çalışmada sahne tanıma yazılımlarının düşük gelirli bölgelere ait nesneleri tanımada başarısız olduğu ve bu sistemlerin sosyoekonomik eşitsizlikleri yeniden üretebileceği tespit edilmiştir. Bu bulgular, yapay zekâ uygulamalarının tasarımı ve kullanımı sırasında farklı toplumsal kimlikler üzerinde ayrımcı sonuçlar oluşturabileceğini işaret etmektedir.

Araştırma ayrıca, yapay zekânın bu etik ihlalleri üretmesinde insan kaynaklı faktörlerin belirleyici rol oynadığını göstermiştir. Yapay zekâ modelleri çoğunlukla geçmişteki toplumsal önyargıları taşıyan büyük veri kümeleriyle beslenmektedir. Bu verilerdeki cinsiyetçi, ırkçı ve sosyoekonomik temelli ayrımcılık eğilimleri, algoritmalar tarafından tekrarlanmakta ve sonuçlara yansımaktadır. Başka bir deyişle, algoritmalar teknik olarak tarafsız görünse de eğitim veri kümesindeki önyargıları sonuçlarına yansıtarak mevcut eşitsizlikleri yeniden üretebilmektedir. Ayrıca derin öğrenme tabanlı sistemlerdeki “kara kutu” problemi, algoritmaların çalışma mantığını şeffaf olmayan bir hale getirerek alınan sonuçların meşruiyetini sorgulamayı güçleştirmektedir. Tüm bu dinamikler, algoritmik karar verme süreçlerinde insan faktörünün yapısal önyargılarla iç içe geçtiğini ve bu sürecin etik ihlallerin kaynağı olabileceğini ortaya koymaktadır. Bu bağlamda elde edilen bulgular, yapay zekânın toplumsal adalet ve eşitlik ilkelerini ihlal edebileceğini açıkça göstermektedir.

Tarafsızlık iddiasıyla geliştirilen teknolojiler, içerdiği önyargılı çıktılar nedeniyle kurumlara ve kamuoyuna duyulan güveni zedeleyebilir. Örneğin, adalet sisteminde kullanılan öngörücü algoritmaların özellikle dezavantajlı grupları hedef alması, haksız cezalara yol açmakla kalmayıp toplumsal güven duygusunu da sarsmaktadır. Benzer şekilde, işe alım süreçlerinde cinsiyet temelli tercihler iş dünyasında fırsat eşitliğini zayıflatmaktadır. Kısacası, yapay zekâ uygulamalarındaki bu etik ihlaller, toplumsal kutuplaşmayı ve güç dengesini yeniden üretecek şekilde kültürel ve sosyo-ekonomik adaletsizliklere zemin hazırlayabilir. Çalışmanın sınırlılıkları göz önüne alındığında, elde edilen sonuçların belirli bir bağlamla sınırlı olduğu görülmektedir.

Araştırma, öncelikli olarak akademik ve medya kaynaklarında rapor edilmiş birkaç somut vaka üzerinden kurgulanmış nitel bir analizdir; bu nedenle bulguların bu örneklerle sınırlı kalmış olması muhtemeldir. Ele alınan vakaların çoğu belli coğrafi ve kültürel bağlamlara dayanmaktadır; bu bakımdan sonuçları farklı bölgelerdeki yapay zekâ uygulamalarına genellemek güçtür. Zaman sınırlılığı da bu çalışmanın bir diğer kısıtlayıcı unsurudur; 2014–2024 arası vakalar incelendiği için hızla değişen teknoloji sahnesinde yeni etik sorunlar ortaya çıkabilir. Ek olarak, nitel içerik analizinin doğası gereği elde edilen bulgular yoruma açık niteliktedir; daha geniş örneklem ve nicel yöntemlerle gerçekleştirilecek ileriki çalışmaların mevcut sonuçları güçlendirici nitelikte olması

beklenir. Tüm bu noktalar dikkate alındığında, mevcut bulgular temkinli değerlendirilmelidir, ancak elde edilen veriler yapay zekânın etik riskleri konusunda önemli bir farkındalık sunmaktadır.

Uzun vadede, etik ilkelerin somut uygulamalara nasıl yansıtılabileceği ve yapay zekâ teknolojisinin sürdürülebilir şekilde toplum yararına hizmet etmesi için hangi politika ve düzenlemelerin uygulanması gerektiği önemli araştırma konuları olmaya devam edecektir. Bütün bu bulgular ışığında, yapay zekâ ve toplumsal dönüşüm ilişkisinde etik ilkelerin öncelikli rolü açıkça görülmektedir.

Teknolojinin sunduğu imkânlar ancak etik kuralların rehberliğinde kullanıldığında dengeli ve sürdürülebilir toplumsal faydalar yaratabilir. Bu çalışmada ulaşılan sonuçlar, yapay zekâ sistemlerinin geliştirilmesi ve uygulanması süreçlerinin etik ve toplumsal değerlerden ayrılamayacağını göstermektedir. Gelecekte, teknolojik yeniliklerin toplumsal ihtiyaçlarla uyumlu biçimde yönlendirilmesi ve yapay zekânın adaletli kullanımı için disiplinlerarası işbirlikleri ile etik temelli politikalara ihtiyaç duyulacaktır. Bu yaklaşımlar benimsendiğinde, yapay zekânın potansiyel toplumsal yararı maksimize edilecek ve dijital dönüşüm süreci insan odaklı bir anlayışla ilerleyecektir.

Yapay zekâ algoritmalarındaki önyargı sorununu gidermek ve daha adil sistemler inşa etmek için temel öneriler şunlardır:

1- Veri kalitesini ve temsiliyetini artırmak amacıyla algoritmaların eğitildiği veri setlerinin çeşitliliği sağlanmalı, yanlılıklar giderilmeli ve veri toplama süreçlerinde etik ilkeler titizlikle gözetilmelidir.

2- Algoritmik şeffaflığı ve anlaşılabilirliği temin etmek üzere yapay zekâ algoritmalarının karar alma süreçleri şeffaflaştırılmalı, açıklanabilir modeller geliştirilmeli ve özellikle hassas kararlarda gerekçeler, ilgili paydaşlarla paylaşılmalıdır.

3- Etik denetim ve hesap verebilirliği güçlendirmek için yapay zekâ sistemlerinin geliştirme ve uygulama süreçlerine bağımsız etik kurullar dâhil edilmeli, potansiyel önyargılar değerlendirilmeli ve hatalı sonuçlarda sorumluluk mekanizmaları tesis edilmelidir.

4- Teknik çözümler üretmek adına veri setlerindeki ve modeldeki önyargıları otomatik olarak tespit edip giderebilen araçlar geliştirilmeli ve algoritmalara entegre edilmelidir.

5- İnsan kaynağını geliştirmek için yapay zekâ geliştiricileri ve ilgili personel, algoritmik önyargı, etik ilkeler ve sorumlu geliştirme konusunda düzenli ve kapsamlı eğitimler almalıdır. Ayrıca disiplinlerarası işbirliğini teşvik etmek amacıyla yapay zekâ geliştiricileri, sosyal bilimciler, hukukçular ve etik uzmanları arasında sürekli ve etkin bir işbirliği ortamı yaratılmalıdır.

6- Devlet politikaları çerçevesinde, yapay zekâ sistemlerinin geliştirilmesi ve uygulanmasına yönelik etik ilkeleri ve standartları belirleyen yasal ve düzenleyici çerçeveler oluşturulmalı, şeffaflık, hesap verebilirlik ve ayrımcılık yasağı gibi temel prensipler güvence altına alınmalıdır. Ayrıca, yapay zekâ araştırmalarına ve geliştirme süreçlerine etik ilkelerin entegrasyonunu teşvik eden teşvik mekanizmaları ve fonlama öncelikleri belirlenmelidir.

7- Algoritmaların sürekli iyileştirilmesi için performansları ve önyargıları düzenli olarak izlenmeli, değerlendirilmeli ve değişen toplumsal normlara uyum sağlayacak şekilde güncellenmelidir.

8- Toplumsal farkındalığı artırmak için yapay zekâ algoritmalarındaki önyargı sorununun etkileri konusunda kamuoyu bilinçlendirilmeli ve eğitim programları desteklenmelidir.

Bu bütüncül yaklaşım, yapay zekânın faydalarını maksimize ederken olası zararlarını en aza indirerek daha adil ve eşitlikçi bir geleceğe katkıda bulunacaktır.

KAYNAKÇA

- Ağgöl, B. (2021). *Derin Öğrenme Kullanılarak Sahte Plakalı Araç Tespit Sistemi Geliştirilmesi*.
- Alanka, D. (2024). Nitel bir araştırma yöntemi olarak içerik analizi: Teorik bir çerçeve. *Kronotop İletişim Dergisi*, 1(1), 64.
- Atalay, M., & Çelik, E. (2017). Büyük Veri Analizinde Yapay Zekâ Ve Makine Öğrenmesi Uygulamaları. *Mehmet Akif Ersoy Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, 9(22), 158-162.
- Barocas, S., & Selbst, A. D. (2016). Big data's disparate impact. *California Law Review*, 104(3), 671–732. <https://doi.org/10.2139/ssrn.2477899>
- Bayır, F. (2006). Yapay Sinir Ağları ve Tahmin Modellemesi Üzerine Bir Uygulama (s.8). [Yüksek Lisans Tezi, İstanbul Üniversitesi, Sosyal Bilimler Enstitüsü, İşletme Anabilim Dalı, Sayısal Yöntemler Bilim Dalı].
- Bozüyük, T., Yağcı, C., Gökçe, İ., & Akar, G. (2005). Yapay Zekâ Teknolojilerinin Endüstrideki Uygulamaları (s.21). *Marmara Üniversitesi, Teknik Bilimler Meslek Yüksekokulu, Elektrik Programı*.
- Cengiz, Y. E. (2023). Yapay Zekâ Teknikleri Kullanarak Mekansal Analizler [Yüksek Lisans Tezi, Harran Üniversitesi, Fen Bilimleri Enstitüsü].
- Coşkun, F., & Gülleroğlu, H. D. (2021). Yapay Zekânın Tarih İçindeki Gelişimi Ve Eğitimde Kullanılması. *Ankara Üniversitesi Eğitim Bilimleri Fakültesi Dergisi*, 54(3), 947–966. <https://doi.org/10.30964/auebfd.916220>
- Çağırkan, B. (2024). Yapay Zeka Ve Dijital Teknoloji Kullanımının Yalnızlık Üzerine Etkisi. *Yapay Zeka, Toplum Ve Kültür* (Ss. 233–234). Ütopya Yayınevi.
- Çalışkan, E., & Acar, H. H. (2006). Yapay Zeka Tekniklerinin Odun Hammaddesi Üretiminde Kullanılması Üzerine Bir Değerlendirme. *Kafkas Üniversitesi Artvin Orman Fakültesi Dergisi*, 7(1), 51–59.

- Darı, A. B., & Koçyiğit, A. (2024). Yapay Zekâ ve Etik: Yeni Medyanın Dönüşümünde Sorumluluk ve Sınırlar. *The Journal Of Communication And Social Studies*, 4(2), 246–261. <https://doi.org/10.59534/jcss.1492948>
- Demirci, V. G. (2022). Dijital Reklamcılıkta Makine Öğrenmesi ve Veri Gizliliği. *Kent Akademisi Dergisi*, 15(3), 1456–1460.
- Erhandı, B. (2020). Derin Öğrenme İle Metin Özetleme.
- Geeksforgeeks. (2021, Temmuz 7). Veri Madenciliğinde YSA'nın Avantajları ve Dezavantajları. <https://www.geeksforgeeks.org/advantages-and-disadvantages-of-ann-in-data-mining/>
- Greene, M. R., Josyula, M., Si, W., & Hart, J. A. (2024). Digital Divides İn Scene Recognition: Uncovering Socioeconomic Biases İn Deep Learning Systems. arXiv. <https://arxiv.org/abs/2401.13097>
- Gülpınar Demirci, V. (2022). Dijital reklamcılıkta makine öğrenmesi ve veri gizliliği. *Kent Akademisi*, 15(3), 1455–1474. <https://doi.org/10.35674/kent.1145325>
- Güvercin, C.H. (2020) “Tıpta Yapay Zekâ ve Etik”. (Ekmekçi PE, editör). *Yapay Zekâ ve Tıp Etiği*. Türkiye Klinikleri.
- Karaca, A. (2021). Derin Öğrenme ile Türkçe Haber Metinlerine Başlık Üretme.
- Koçak, H. (2024). Yapay Zekâda Etik: Riskler, İhlaller ve Güncel Tartışma Konuları. *Sosyal Bilimlerde Güncel Araştırma ve İncelemeler VII* (s. 69).
- Köçeri, K. (2023). Yapay Zekânın Siyasi, Etik ve Toplumsal Açından Dezenformasyon Tehdidi. *İletişim ve Diplomasi*, (11), 251.
- Maltarollo, V. G., Honório, K. M., & Ferreira Da Silva, A. B. (2013). Applications of Artificial Neural Networks in Chemical Problems. In K. Suzuki (Ed.), *Artificial Neural Networks: Methods and Applications* (s. 44). InTechOpen. <https://doi.org/10.5772/51275>
- Mccarthy, J. (2007). “From Here To Human-Level Aı”. *Artificial Intelligence*, 171(18), 1174-1182. <https://doi.org/10.1016/J.Artint.2007.10.09>.

- Obermeyer, Z., Powers, B., Vogeli, C., & Mullainathan, S. (2019). Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations. *Science*, 366(6464), 447–453. <https://doi.org/10.1126/science.aax2342>
- Öztemel, E. (2020). Yapay Zekâ ve İnsanlığın Geleceği. *Bilişim Teknolojileri ve İletişim: Birey ve Toplum Güvenliği* (ss. 100–105). <https://doi.org/10.53478/tuba.2020.011>
- Öztürk, M. (2024). Yenilikten Tartışmaya: Yapay Zekâ ve Deepfake Çalışmalarının Web of Science Üzerinden Bibliyometrik Analizi. *Akdeniz İletişim*, (46 - Yapay Zekâ ve İletişim Özel Sayısı), 78.
- Pirim, H. (2006). Yapay Zeka. *Journal of Yaşar University*, 1(1), 83–87.
- Rouhiainen, L. (2020). Yapay Zeka Geleceğimizle İlgili Bugün Bilmemiz Gereken 101 Şey (ss. 2, 277, 279). Pegasus Yayınları.
- Sargın, A., & Göçen, A. (2020). Çocuklar İçin Yapay Zeka (s. 12).
- Serhatlıoğlu, S., Ozan, A. T., & Hardalaç, F. (2005). Yapay Zeka Teknikleri ve Radyolojiye Uygulanması. *Pivolka Dergisi*, 4(18), 4.
- Sheikhi, M. (2022). Yapay Zekâ Kullanımının İş Piyasasına Etkisi. *Journal of Economics and Political Sciences*. Medipol Üniversitesi, Sosyal Bilimler Meslek Yüksekokulu, Dış Ticaret Programı. 2(1), 102–111. ORCID: 0000-0003-3493-8307.
- Taşçı, G., & Çelebi, M. (2020). Eğitimde Yeni Bir Paradigma: “Yükseköğretimde Yapay Zekâ”. *Uluslararası Toplum Araştırmaları Dergisi*, 16(29), 2351–2353. <https://doi.org/10.26466/opus.747634>
- Tosunoğlu, E., Yılmaz, R., Özeren, E., & Sağlam, Z. (2021). Eğitimde Makine Öğrenmesi: Araştırmalardaki Güncel Eğilimler Üzerine İnceleme. *Ahmet Keleşoğlu Eğitim Fakültesi Dergisi*, 3(2), 180–183.
- Tuğrul, B., & Ahmed, A. S. A. (2022). Makine Öğrenme Yöntemleri ile Ağ Trafik Analizi. *Ankara Üniversitesi, Adli Bilişim Bölümü, Niğde Ömer Halisdemir Üniversitesi Mühendislik Bilimleri Dergisi (NÖHÜ J. Eng. Sci.)*, 11(4), 862–870.

Turan, T. (2024). Dijital Vicdan: Yapay Zekâ Çağında Etik (s. 29). YAZ Yayınları. ISBN: 978-625-6171-12-1.

Yavuz, U. (1995). Uzman Sistemler ve Parametrik Hipotez Testleri Üzerine Bir Uygulama. (s.7) [Doktora Tezi, Atatürk Üniversitesi, Sosyal Bilimler Enstitüsü, İşletme Anabilim Dalı].

İNTERNET KAYNAKLARI

<https://Fykmobile.Com/Articles/Brain-Tranning.Html>

<https://Www.Matematiksel.Org/Turing-Testi-İnsani-Robottan-Nasil-Ayirir/>

<https://Www.Turhost.Com/Blog/Makine-Ogrenmesi-Machine-Learning-Nedir/>

<https://Www.Turhost.Com/Blog/Makine-Ogrenmesi-Machine-Learning-Nedir/>

<https://www.yenicaggazetesi.com.tr/yapay-zeka-destekli-robotlar-guvenlik-tehdidi-olusturuyor-866735h.htm>

Angwin, J., Larson, J., Mattu, S., & Kirchner, L. (2016, May 23). Machine Bias: There's Software Used Across the Country to Predict Future Criminals. And It's Biased Against Blacks. ProPublica. <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing> e.t: 13.04.2025

Dastin, J. (2018, October 11). Insight - Amazon Scraps Secret AI Recruiting Tool that Showed Bias Against Women. Reuters. <https://www.reuters.com/article/us-amazon-com-jobs-automation-insight/amazon-scraps-secret-ai-recruiting-tool-that-showed-bias-against-women-idUSKCN1MK08G>

Hardesty, L. (2018, February 11). Study Finds Gender, Skin-Type Bias in Some Artificial-İntelligence Systems. MIT News. <https://news.mit.edu/2018/study-finds-gender-skin-type-bias-artificial-intelligence-systems-0212>

İş Gücünün Geleceği İçin Yapay Zekâya Erişim Fırsatları Çoğalmalı. (2024). <https://hrdergi.com/is-gucunun-gelecegi-icin-yapay-zekaya-erisim-firsatlari-cogalmali#:~:text=Yapay%20zeka%2C%20engelli%20bireyler%20i%C3%A7in>

[,%39'da%20kald%C4%B1%C4%9F%C4%B1n%C4%B1%20g%C3%B6steriyor](#), (E.T: 21.12.2024)

Kabalcı, E. (2015). Yapay Sinir Ağları- (Artificial Neural Networks) (s. 17-19).

<https://ekblc.files.wordpress.com/2013/09/ysa.pdf>

Kayıhan, B. (2014). Yapay Zekâ ve Etik Sorunlar: İnsanlık, Teknolojinin Kontrolünü Kaybediyor Mu?. <https://bilimgenc.tubitak.gov.tr/makale/yapay-zeka-ve-etik-sorunlar-insanlik-teknolojinin-kontrolunu-kaybediyor-mu?> (E.T:15.12.2024).

Larson, J., Mattu, S., Kirchner, L., & Angwin, J. (2016, May 23). How We Analyzed the COMPAS Recidivism Algorithm. ProPublica.

<https://www.propublica.org/article/how-we-analyzed-the-compas-recidivism-algorithm>

Savunma Sanayii Yapay Zekâ Platformu. (2022, May 30) Pekiştirmeli Öğrenme ile Sürü Stratejisi Geliştirme Yarışması Başlıyor. (E.T: 20.10.2024).

<https://ssyz.org.tr/>

Türe, S., & Topuz, S. (2019). Yapay Zekâ ve Askeri Uygulamalar. MilSOFT Yazılım Teknolojileri AŞ. https://www.milsoft.com.tr/wp-content/uploads/2020/08/Yapay-Zeka-ve-Askeri-Uygulamalar_v2.2.pdf

Van Rijmenam, M. (2023, October 20). Privacy in the Age of AI: Risks, Challenges and Solutions. The Digital Speaker. <https://www.thedigitalspeaker.com/privacy-age-ai-risks-challenges-solutions/>

Vincent, J. (2018, January 12). Google 'Fixed' its Racist Algorithm by Removing Gorillas from its Image-Labeling Tech. The Verge.

<https://www.theverge.com/2018/1/12/16882408/google-racist-gorillas-photo-recognition-algorithm-ai>

Vizyoner Genç, (2019, February 26). Gelecekte Çok Duyacağız: Yapay Sinir Ağları ve Çalışma Mantıkları. <https://vizyonergenc.com/icerik/gelecekte-cok-duyacagiz-yapay-sinir-aglari-ve-calisma-mantiklari> (E.T: 15.10.2024)

Wigmore, I. (2013, 30 Temmuz). What is Natural Language Generation (NLG)?

TechTarget. <https://www.techtarget.com/searchenterpriseai/definition/Natural-Language-Generation-NLG>

Yapay Zeka (AI) Teknolojileri ve Etik Sorunları, (2024, May 6), Bilim & Teknik.

<https://www.hayatayonelik.com/yapay-zeka-teknolojileri-ve-etik-sorunlari/>

(E.T: 12.04.2025)

Yapay Zeka Ömer Civalek'le Söyleşi. (2003). Türkiye Mühendislik Haberleri, Sayı

423. <http://web.bilecik.edu.tr/murat-ozalp/files/2016/12/%C3%96mer-Civalekle-S%C3%B6yle%C5%9Fi.pdf> (E.T: 18.10.2024).

Yapay Zekâ, Etik ve Sosyal Sorumluluk. (2023, January 24).

<https://pcb.com/article/artificial-intelligence-ethics-and-social-responsibility>.

(E.T: 10.10.2024)

