

Yapay Zekâ Tabanlı Karar Destek Sistemlerinde Teknik, Etik ve Yönetişimsel Yaklaşımlar: Çok Sektörlü Bir Derleme Çalışması

Technical, Ethical, And Governance Approaches in Artificial Intelligence-Based Decision Support Systems: A Multi-Sectoral Review

DOI:10.33461/uybisbbd.1691218

Kadir KESGİN¹ 

Makale Bilgileri

Makale Türü:
Derleme Makalesi

Geliş Tarihi:
04.05.2025

Kabul Tarihi:
09.06.2025

©2025 UYBISBBD
Tüm hakları saklıdır.



Öz

Bu derleme çalışması, yapay zekâ tabanlı karar destek sistemlerinin (YZ-KDS) teknik, etik ve yönetim boyutlarını sistematik bir analiz çerçevesinde incelemeyi amaçlamaktadır. Sağlık, finans, üretim/lojistik ve çevre yönetimi gibi kritik sektörlerde YZ-KDS sistemlerinin uygulamaları, veri kalitesi, açıklanabilirlik, adalet ve hesap verebilirlik bağlamında ele alınmıştır. Çalışmanın veri temeli, Kesgin ve Dözer (2025) tarafından geliştirilen BiBLoX yazılımı aracılığıyla Web of Science, Scopus ve TR Dizin veri tabanlarından çekilen 4824 bibliyografik kayda dayanmaktadır. PRISMA protokolü doğrultusunda yapılan filtreleme sonucunda 1273 yayın detaylı incelemeye alınmıştır. Bibliyometrik analiz sonuçlarına göre, açıklanabilir yapay zekâ çerçeveleri (XAI), federated learning tabanlı gizlilik çözümleri ve adil karar mekanizmaları, sistemlerin güvenilirliğini artırmada öne çıkan yaklaşımlar olarak belirlenmiştir. Ayrıca disiplinlerarası iş birliği, etik denetim ve kurumsal yönetim süreçlerinin, bu sistemlerin sorumlu ve sürdürülebilir biçimde geliştirilmesi açısından kritik olduğu vurgulanmaktadır. Çalışma, yapay zekâ destekli karar destek sistemlerinin gelecekteki tasarımlarına hem teorik hem uygulamalı katkı sunmayı hedeflemektedir.

Anahtar Kelimeler: Yapay Zekâ, Karar Destek Sistemleri, Açıklanabilirlik, Veri Yönetimi, Etik Yapay Zekâ.

Abstract

This review study aims to systematically examine the technical, ethical, and governance dimensions of artificial intelligence-based decision support systems (AI-DSS). The applications of AI-DSS across critical sectors such as healthcare, finance, manufacturing/logistics, and environmental management are analyzed in terms of data quality, explainability, fairness, and accountability. The bibliometric foundation of the study is based on 4,824 bibliographic records retrieved from Web of Science, Scopus, and TR Dizin databases using the open-source BiBLoX software developed by Kesgin and Dözer (2025). Following the PRISMA protocol, a multi-stage screening process was implemented, resulting in a refined corpus of 1,273 publications subjected to detailed analysis. The findings highlight explainable AI (XAI) frameworks, federated learning-based privacy solutions, and fairness-aware decision mechanisms as key approaches to enhancing system reliability. Moreover, interdisciplinary collaboration, ethical oversight, and institutional governance structures are emphasized as critical enablers for the responsible and sustainable development of AI-DSS. This study aims to contribute a conceptual and practical perspective to the design of future decision support systems grounded in artificial intelligence.

Keywords: Artificial Intelligence, Decision Support Systems, Explainability, Data Governance, Ethical AI.

Article Info

Paper Type:
Review Paper

Received:
04.05.2025

Accepted:
09.06.2025

©2025 UYBISBBD
All rights reserved.



Atıf/ to Cite (APA): Kesgin K. (2025). Yapay Zekâ Tabanlı Karar Destek Sistemlerinde Teknik, Etik ve Yönetişimsel Yaklaşımlar: Çok Sektörlü Bir Derleme Çalışması. Uluslararası Yönetim Bilişim Sistemleri ve Bilgisayar Bilimleri Dergisi, 9(1), 29-48. DOI: 10.33461/ uybisbbd.1691218

¹ Dr. Öğr. Gör. Bandırma Onyedi Eylül Üniversitesi, Gönen Meslek Yüksekokulu Bilgisayar Teknolojileri Bölümü, kadir@bandirma.edu.tr, Balıkesir Türkiye.

1. GİRİŞ

Karar verme süreçleri, işletmelerin rekabet avantajı sağlaması ve sürdürülebilir başarıya ulaşması açısından kritik öneme sahiptir. Bu süreçleri desteklemek amacıyla geliştirilen Karar Destek Sistemleri (Decision Support Systems - DSS), yöneticilere hızlı, doğru ve etkili kararlar alma konusunda yardımcı olarak iş süreçlerinin verimliliğini artırmaktadır. Ancak, artan veri hacmi ve veri kaynaklarının heterojenliği karşısında geleneksel DSS'lerin yetersiz kalmaya başladığı görülmektedir (Sharda, Delen ve Turban, 2020).

Son yıllarda yapay zekâ (Artificial Intelligence - AI) teknolojilerinin DSS ile entegrasyonu, karar destek mekanizmalarında paradigma değişimine yol açmıştır. Makine öğrenmesi, derin öğrenme, doğal dil işleme ve uzman sistemler gibi AI tabanlı teknolojiler sayesinde büyük veri kümelerinden öngörüler üretmek, desenleri tanımlamak ve otomatik karar önerileri sunmak mümkün hâle gelmiştir (Chen vd., 2023). Bu dönüşüm, özellikle sağlık (Topol, 2019; Rajkomar vd., 2018), çevre yönetimi (Cortès ve Vapnik, 1995) ve üretim sektörlerinde yaygın uygulamalarla kendini göstermektedir.

Literatürde, AI destekli karar sistemlerinin bilgi yönetimi süreçlerine etkisi de tartışılmaktadır. Örneğin, kural tabanlı uzman sistemlerin stratejik karar süreçlerini desteklemedeki rolü, özellikle sağlık ve mühendislik gibi uzmanlık gerektiren alanlarda vurgulanmıştır (Pedrycz vd., 2012). Benzer şekilde, kişiselleştirilmiş web tabanlı sağlık sistemlerinin karar desteği sunmadaki başarısı dikkat çekmektedir (Holzinger vd., 2019).

Bununla birlikte, AI tabanlı DSS sistemlerinin yaygınlaşması; veri gizliliği, algoritmik önyargı, açıklanabilirlik ve hesap verebilirlik gibi çok katmanlı etik ve yönetişimsel sorunları da beraberinde getirmiştir (Kostopoulos vd., 2024). Özellikle Açıklanabilir Yapay Zekâ (Explainable AI - XAI) teknikleri, bu sistemlerin karar süreçlerinin şeffaflığını ve kullanıcı güvenini artırma potansiyeli nedeniyle güncel araştırmalarda ön plana çıkmaktadır.

Bu çalışma, yapay zekâ destekli karar destek sistemlerini teknik, etik ve yönetişimsel boyutlarıyla değerlendirmeyi; bu sistemlerin sektörel uygulamalarını analiz etmeyi ve mevcut zorluklara karşı literatür destekli çözüm önerileri geliştirmeyi amaçlamaktadır. Ayrıca, gelecekteki araştırmalar için açıklanabilirlik, adalet ve veri gizliliği gibi temel konularda açık araştırma alanlarını tanımlayarak hem kuramsal hem pratik katkılar sunmayı hedeflemektedir.

2. YÖNTEM

Bu çalışma, yapay zekâ destekli karar destek sistemlerine (AI-DSS) yönelik bilimsel üretimi disiplinlerarası bir çerçevede değerlendirmek amacıyla bibliyometrik analiz yaklaşımıyla yapılandırılmıştır. Veri toplama süreci, Kesgin ve Özer (2025) tarafından geliştirilen açık kaynaklı BiBLoX platformu kullanılarak gerçekleştirilmiştir. BiBLoX, Python tabanlı bir bibliyometrik analiz otomasyonudur ve Web of Science, Scopus ve TR Dizin veri tabanlarıyla API bağlantısı kurarak akademik kayıtların sistematik çekilmesine olanak sağlamaktadır.

Araştırma kapsamında kullanılan anahtar kelimeler: “decision support systems”, “AI”, “machine learning”, “deep learning”, “explainable AI”, “NLP”, “federated learning”, “ethics in AI”, “bias”, ve “accountability” gibi tematik alanları kapsamaktadır. Arama filtreleri, 2015–2024 yılları arasındaki yayınlarla sınırlandırılmış; yalnızca akademik dergilerde yayımlanmış makaleler, kongre bildirileri ve sistematik derlemeler değerlendirmeye alınmıştır.

Şekil 1 Prisma Model Diyagramı

Veri eleme süreci PRISMA (Preferred Reporting Items for Systematic Reviews and Meta-Analyses) protokolü temelinde yürütülmüştür. Toplam 4824 yayın içerisinden 925'i yinelenen kayıt olduğu için çıkarılmış, 1800 kayıt başlık ve özet incelemesi sonucunda uygun bulunmuştur. Tam

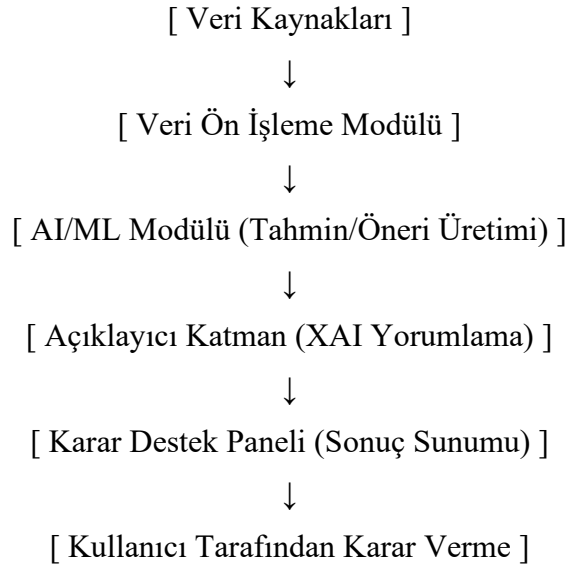
metin erişimi sağlanamayan 826 kayıt da elenmiş ve sonunda 1273 yayın analiz kapsamına alınmıştır (Şekil 1).

Toplanan veriler BiBLoX'un ön işleme modülleriyle temizlenmiş, ardından Excel ve VOSviewer gibi araçlar kullanılarak yayın trendleri, anahtar kelime eşleşmeleri, iş birlikçi yazar ağları ve sektörel temsiller analiz edilmiştir. Bibliyometrik yöntemle elde edilen çıktılar, betimleyici istatistiksel yaklaşımlarla sınıflandırılmış ve literatürün genel eğilimleri görselleştirilmiştir.

Bu bağlamda, çalışma herhangi bir nitel içerik analizi (qualitative content analysis) yapmamış; yalnızca yayın meta verileri üzerinden yürütülen bibliyometrik yöntemlerle, AI-DSS literatüründeki teknik ve tematik kümelenmeleri incelemiştir. Bu sayede hem akademik eğilimleri hem de yapay zekâ temelli karar destek sistemlerinin sektörel, etik ve yönetişimsel yansımaları sistematik olarak ele alınabilmektedir.

3. YAPAY ZEKÂ TEKNOLOJİLERİNİN KARAR DESTEK SİSTEMLERİNE (DSS) ENTEGRASYONU

Yapay zekâ tabanlı karar destek sistemleri (AI-DSS), veri kaynaklarının farklı yapısal özelliklerini analiz ederek öngörü modelleri geliştirmekte ve kullanıcılara daha etkili karar alma olanakları sunmaktadır. Bu sistemler genel olarak, birden fazla katmandan oluşan bir mimari yapı ile tasarlanmaktadır. Aşağıda, AI-DSS sistemlerinin genel mimari bileşenleri sunulmuştur:



Bu yapının her bir bileşeni aşağıda açıklanmaktadır:

Veri Kaynakları: Yapay zekâ destekli DSS'lerin temelini, yapılandırılmış (veritabanları, tablolar, araçlar) ve yapılandırılmamış (metin, görüntü, ses verileri) veri kaynakları oluşturur. Bu kaynaklar, sistemin karar üretebilmesi için ilk girdileri sağlar.

Veri Ön İşleme Modülü: Verilerin doğruluğu, eksiksizliği ve anlamlılığı, sistemin başarısı için kritiktir. Bu aşamada veri temizleme, normalizasyon, özellik seçimi gibi işlemler gerçekleştirilir.

Makine Öğrenmesi / Yapay Zekâ Katmanı: Bu modülde, öngörü veya sınıflandırma için makine öğrenmesi (ML) ve derin öğrenme (DL) algoritmaları kullanılır. Kullanıcı ihtiyaçlarına göre model parametreleri dinamik biçimde ayarlanabilir.

Açıklayıcı Katman (XAI): Kullanıcı güveni ve yasal uyumluluk açısından, sistemin ürettiği kararların anlaşılabilir olması önemlidir. SHAP, LIME gibi açıklanabilir yapay zekâ (XAI) yöntemleri, bu modülde yer alır.

Karar Destek Paneli: Bu katman, modelin önerilerini kullanıcıya açık ve sezgisel bir şekilde sunar. Görselleştirme araçları ve kullanıcı dostu arayüzler, karar kalitesini artırmayı hedefler.

Kullanıcı Etkileşimi ve Karar: Nihai karar, insan faktörünü içeren son katmanda verilir. Kullanıcılar, model önerilerini göz önüne alarak kendi kurumsal bağlamlarına uygun kararları alırlar.

Bu çok katmanlı yapı, AI-DSS sistemlerinin teknik şeffaflık, kullanıcı etkileşimi ve operasyonel verimlilik gibi temel gereksinimlerini karşılayacak biçimde tasarlanmaktadır. Söz konusu mimari, sonraki bölümlerde detaylandırılacak olan sektörel uygulamalarda farklı biçimlerde karşımıza çıkmaktadır.

3.1. Makine Öğrenmesi Tabanlı Yaklaşımlar

Makine öğrenmesi (Machine Learning - ML), bilgisayarların açıkça programlanmadan verilerden öğrenmesini sağlayan bir yapay zekâ alanıdır. Karar destek sistemlerinde ML'in entegrasyonu, geleneksel veri analiz yöntemlerine kıyasla daha öngörülebilir, hızlı ve doğrulanabilir karar mekanizmaları sunmaktadır. ML tabanlı DSS mimarileri, veri kalıplarını otomatik tanıma, tahmine dayalı karar alma ve yüksek hacimli veri akışlarını analiz etme açısından önemli avantajlar sağlamaktadır (Goodfellow, Bengio, ve Courville, 2016). Bu sistemlerde, model doğruluğu çoğunlukla Ortalama Kare Hata (Mean Squared Error - MSE) gibi bir kayıp fonksiyonu aracılığıyla değerlendirilmektedir:

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

Burada y_i gerçek değerleri, $\hat{y}_i$ modelin tahmin ettiği değerleri ve n toplam örnek sayısını temsil etmektedir. Finansal tahminleme gibi sayısal duyarlılığı yüksek uygulamalarda, Destek Vektör Makineleri (SVM) veya Rastgele Ormanlar (Random Forest) gibi algoritmalar bu kaybı minimize ederek daha tutarlı öngörüler üretmektedir (Chen vd., 2023).

Makine öğrenmesi teknikleri üç ana kategoride sınıflandırılmaktadır:

Denetimli Öğrenme (Supervised Learning): Etiketli veri setleri üzerinden eğitilen bu modeller, belirli hedef çıktılarına göre tahminleme yapar. Örneğin, kredi başvurularının onaylanıp onaylanmayacağına ilişkin karar süreçlerinde sıkça kullanılır.

Denetimsiz Öğrenme (Unsupervised Learning): Etiketsiz verilerdeki örüntülerin keşfine odaklanır. Müşteri segmentasyonu veya anomali tespiti gibi uygulamalar bu kategoriye girer.

Pekiştirmeli Öğrenme (Reinforcement Learning): Ajanlar, çevreleriyle etkileşim içinde ödül-ceza mekanizmasına göre strateji geliştirir. Dinamik ve sürekli değişen ortamlarda etkili çözümler üretmektedir.

AI tabanlı DSS sistemlerinde yaygın olarak kullanılan algoritmalar arasında SVM, Karar Ağaçları, Rastgele Ormanlar ve Gradyan Artırmalı Ağaçlar (Gradient Boosted Trees) yer almaktadır. Bu yöntemler, büyük veri kümelerinde yüksek doğruluk sağlamalarının yanı sıra belirli ölçüde yorumlanabilirlik sunarlar. Ancak, özellikle ensemble modellerin artan karmaşıklığı, şeffaflık ve açıklanabilirlik gibi alanlarda sınırlılıklar yaratabilmektedir (Kostopoulos vd., 2024).

Özellikle sağlık sektöründe, makine öğrenmesine dayalı sistemlerle diyabet yönetiminde insülin dozajlarının optimize edilmesi gibi başarılı uygulamalar gerçekleştirilmiştir (Topol, 2019). Ancak bu sistemlerin performansı, eksik veriler, dengesiz veri kümeleri veya modelin aşırı öğrenme (overfitting) eğilimi gibi sorunlarla olumsuz etkilenebilmektedir.

Makine öğrenmesi, karar destek sistemlerine esneklik ve öngörü gücü kazandırsa da, şeffaflık, veri kalitesi ve genellenebilirlik gibi konular halen araştırmaya açık ve gelişime ihtiyaç duyan başlıklardır.

3.2. Derin Öğrenme Modelleri

Derin öğrenme (Deep Learning – DL), çok katmanlı yapay sinir ağları kullanarak yüksek boyutlu verilerden karmaşık örüntüler çıkarmayı sağlayan ileri düzey bir yapay zekâ yöntemidir. Geleneksel makine öğrenmesi algoritmalarına kıyasla, özellikle büyük veri setlerinde daha yüksek doğruluk ve genelleme kapasitesi sunmaktadır (Goodfellow vd., 2016).

Kayıp Fonksiyonları ve Uygulama Alanları

Sınıflandırma problemlerinde yaygın olarak kullanılan Cross-Entropy Loss fonksiyonu, modelin tahmin ettiği olasılıkların gerçek etiketlerle ne kadar uyumlu olduğunu ölçmek için kullanılır:

$$L = - \sum_{i=1}^n [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

Bu fonksiyon, özellikle sağlık sektöründe görüntü temelli teşhis sistemlerinde kullanılmaktadır. Örneğin, akciğer kanseri teşhisinde kullanılan Konvolüsyonel Sinir Ağları (CNN), bu kayıp fonksiyonunu minimize ederek yüksek doğruluklu sonuçlar üretmiştir (Rajkomar vd., 2018).

Yaygın Kullanılan Derin Öğrenme Modelleri

Yapay Sinir Ağları (ANN): Temel çok katmanlı mimarisi ile farklı sınıflandırma ve regresyon problemlerinde kullanılır.

Konvolüsyonel Sinir Ağları (CNN): Görsel verilerin (tıbbi görüntüler, uydu fotoğrafları) analizi için geliştirilmiştir. Sağlık alanında MR, röntgen gibi verilerin yorumlanmasında etkin biçimde kullanılmaktadır.

Tekrarlayan Sinir Ağları (RNN) ve LSTM: Ardışık veri işleme kabiliyeti sayesinde finansal zaman serileri veya hasta yaşam verilerinin analizinde kullanılmaktadır.

Zorluklar ve Sınırlılıklar

Derin öğrenme tabanlı DSS sistemleri önemli avantajlar sunsa da bazı temel sınırlamalara sahiptir:

- **Kara Kutu Problemi:** Kullanıcılar için modelin karar mekanizmalarını anlamak çoğu zaman mümkün değildir. Bu durum, sistemin şeffaflığı ve güvenilirliği üzerinde olumsuz etkiler yaratmaktadır (Kostopoulos vd., 2024).
- **Veri Gereksinimleri:** Derin öğrenme sistemlerinin başarılı olabilmesi için büyük hacimli ve dengeli veri kümeleri gereklidir. Dengesiz sınıflar ve eksik veriler, modelin genellenebilirliğini düşürebilir.
- **Yüksek Riskli Uygulamalarda Validasyon:** Sağlık ve hukuk gibi kritik alanlarda, küçük hata oranları bile ciddi sonuçlara yol açabileceği için titiz bir model doğrulama süreci kaçınılmazdır.

Araştırma Açıkları

DL destekli DSS sistemleri, güçlü öngörü ve sınıflandırma yeteneklerine rağmen hâlâ şu başlıklarda araştırma ihtiyacı doğurmaktadır:

- Açıklanabilirlik mekanizmalarının entegrasyonu
- Veri verimliliği yüksek mimarilerin geliştirilmesi

- Enerji tüketimi ve güvenlik açısından optimizasyon

Bu çerçevede, derin öğrenme tabanlı karar destek sistemleri, yalnızca teknik doğruluğu artırmakla kalmamakta; aynı zamanda, yeni nesil yapay zekâ uygulamaları için etik ve açıklanabilirlik temelli araştırma alanları da sunmaktadır.

3.3. Doğal Dil İşleme (NLP) Uygulamaları

Doğal dil işleme (Natural Language Processing - NLP), bilgisayarların insan dilini analiz edebilmesi, yorumlayabilmesi ve anlamlı sonuçlar üretebilmesi amacıyla geliştirilen bir yapay zekâ alanıdır. Karar destek sistemlerinde NLP'nin entegrasyonu, kullanıcıların sistemle metin, ses veya konuşma temelli doğal dil üzerinden etkileşim kurmasına olanak tanıyarak karar alma süreçlerini hem daha sezgisel hem de daha erişilebilir hâle getirmektedir (Topol, 2019).

NLP tabanlı karar destek çözümleri, özellikle metin madenciliği, duygu analizi ve konuşma tabanlı arayüzler gibi uygulamalar aracılığıyla çeşitli sektörlerde karar alma süreçlerine katkı sunmaktadır. Metin madenciliği sayesinde sosyal medya, e-posta ya da müşteri şikâyetleri gibi yapılandırılmamış veri kaynaklarından anlamlı bilgiler çıkarılarak hizmet stratejileri geliştirilebilmektedir. Duygu analizi uygulamaları, kullanıcı ifadelerinden olumlu, olumsuz ya da nötr yaklaşımları belirleyerek pazarlama ve halkla ilişkiler süreçlerinde yönlendirici kararların alınmasına yardımcı olmaktadır (Liu, 2012). Öte yandan, chatbotlar ve konuşma tabanlı arayüzler ise kullanıcıların doğal dille soru sormasına ve sistemin buna gerçek zamanlı olarak yanıt vermesine imkân tanımaktadır. Bankacılık sektöründe sıklıkla karşılaşılan bu uygulamalar, müşterilerin sık sorulan sorulara doğrudan ve hızlı şekilde erişimini sağlarken, sistemin karar destek rolünü güçlendirmektedir. Bu uygulamalar arasında müşteri hizmetleri botları, hesap analizi sistemleri ve öneri motorları yer almakta olup, bu sistemler müşteri deneyimini iyileştirmekte ve operasyonel yükü azaltmaktadır (Topol, 2019).

Tüm bu olanaklara rağmen NLP tabanlı sistemler bazı temel sınırlamalarla karşı karşıyadır. Özellikle ironik anlatımlar, mecazlar veya kültürel referanslar gibi bağlam bağımlı dil yapılarının doğru analiz edilmesi hâlâ önemli bir zorluk teşkil etmektedir. Çok dilli ortamlarda ise sistemlerin etkili çalışabilmesi için büyük hacimli ve dengeli veri setlerine ihtiyaç duyulmakta, bu da kaynak sınırlamaları nedeniyle her zaman mümkün olamamaktadır (Jurafsky ve Martin, 2020).

NLP teknolojileri karar destek sistemlerine insan merkezli, esnek ve kullanıcı dostu bir yapı kazandırırken; bağlam duyarlılığı, çok dillilik ve dilsel çeşitlilik gibi alanlarda geliştirilmesi gereken birçok yön barındırmaktadır. Bu alanlardaki ilerlemeler, karar destek sistemlerinin daha sezgisel ve sosyal bağlamla uyumlu biçimde gelişmesini sağlayacaktır.

3.4. Uzman Sistemler ve Kural Tabanlı Yaklaşımlar

Uzman sistemler (Expert Systems), belirli bir bilgi alanındaki uzmanların karar alma süreçlerini bilgisayar ortamında taklit eden ve bu süreçleri kurallar aracılığıyla modelleyen karar destek araçlarıdır. Bu sistemler, genellikle yapılandırılmış ve kurallara dayalı problemlerde çözüm üretmek üzere geliştirilmiş olup, karar destek sistemlerinin tarihsel gelişiminde temel bir rol oynamıştır (Turban vd., 2011).

Tipik bir uzman sistem iki ana bileşenden oluşur: bilgi tabanı (knowledge base) ve çıkarım motoru (inference engine). Bilgi tabanı, uzmanlardan derlenen kural, prosedür ve deneyimlerin yapılandırılmış biçimde depolandığı alandır. Çıkarım motoru ise bu bilgileri kullanarak kullanıcıdan alınan girdilere uygun önerilerde bulunur veya çözümler üretir. Bu sistemler, özellikle sağlık, hukuk, mühendislik ve tarım gibi alanlarda önemli başarılar elde etmiştir. Örneğin, bazı tıbbi tanı sistemleri, uzman hekimlerle karşılaştırıldığında daha doğru sonuçlar üretebilmiş ve karar alma süreçlerini hızlandırmıştır (Shortliffe, 1976).

Yakın dönem çalışmalar, bilgi yoğun sektörlerde uzman sistemlerin bilgi yönetimini kolaylaştırmada stratejik avantaj sunduğunu göstermektedir. Örneğin, Klinik karar verme süreçlerinde makine öğrenimi yöntemlerinin kullanımı, sağlık hizmetlerinde önemli bir dönüşüm

sağlamaktadır. Bu bağlamda, Brnabic ve Hess (2021) tarafından yapılan bir sistematik literatür taraması, hasta-sağlayıcı karar verme süreçlerinde makine öğrenimi yöntemlerinin uygulandığı 34 çalışmayı incelemiştir. Bu çalışma, farklı algoritmaların ve yaklaşımların klinik karar destek sistemlerinde nasıl kullanıldığını ve bunların hasta bakımına etkilerini değerlendirmiştir (Brnabic ve Hess, 2021).

Bununla birlikte, bu sistemlerin bazı yapısal sınırlılıkları da bulunmaktadır. En temel sorunlardan biri, bilgi tabanının manuel olarak güncellenmesinin zor ve zaman alıcı olmasıdır. Ayrıca, çıkarım motorlarının genellikle deterministik yapıda çalışması, belirsizlik ve değişkenlik içeren ortamlarda esnek karar üretimini sınırlamaktadır (Durkin, 1994).

Günümüzde bu sınırlamaları aşmak amacıyla, uzman sistemler ile makine öğrenmesi tekniklerinin hibritleştirilmesi yönünde çalışmalar yapılmaktadır. Bu yaklaşımlar, sabit kuralların ötesine geçerek sistemlerin öğrenilebilirliğini ve uyarlanabilirliğini artırmakta; böylece karar destek sistemlerini hem daha esnek hem de belirsiz ortamlara karşı daha dirençli hâle getirmektedir.

3.5. Açıklanabilir Yapay Zekâ (XAI) ve Karar Destek Sistemleri

Yapay zekâ destekli karar destek sistemlerinin (DSS) giderek daha karmaşık hâle gelmesi, bu sistemlerin kullanıcılar tarafından anlaşılabilir olmasını zorunlu kılmıştır. Özellikle sağlık, finans ve hukuk gibi yüksek riskli sektörlerde, yapay zekânın ürettiği kararların şeffaf biçimde açıklanması, hem güven hem de yasal uyum açısından kritik önemdedir (Kostopoulos vd., 2024).

Açıklanabilir Yapay Zekâ (Explainable Artificial Intelligence – XAI), karmaşık yapay zekâ modellerinin karar alma süreçlerini insanlar için anlaşılır hâle getirmeyi amaçlamaktadır. XAI tekniklerinden biri olan SHAP (SHapley Additive exPlanations), her bir girdinin karar üzerindeki etkisini kooperatif oyun teorisi temelli bir hesaplama ile açıklamaktadır:

$$\phi_i = \sum_{S \subseteq N \setminus \{i\}} \frac{|S|! (|N| - |S| - 1)!}{|N|!} (v(S \cup \{i\}) - v(S))$$

Burada ϕ_i özellik i 'nin karar üzerindeki katkısını, S alt kümeleri $v(S)$ model çıktısını temsil etmektedir (Kostopoulos vd., 2024).

XAI yaklaşımlarının DSS bağlamında sunduğu başlıca katkılar, kararların izlenebilirliği, hatalı ya da önyargılı sonuçların erken aşamada tespiti ve yasal düzenlemelere uyumun sağlanmasıdır. Örneğin, Avrupa Birliği Genel Veri Koruma Tüzüğü (GDPR), otomatik karar mekanizmalarının kullanıcıya açıklanabilir nitelikte çalışmasını zorunlu kılmaktadır. Sağlık alanında, bir teşhisin neden önerildiğini açıklayabilen DSS çözümleri hem hekimin sisteme duyduğu güveni artırmakta hem de hastaların bilgilendirilmiş onam süreçlerini daha sağlıklı yürütmesine olanak tanımaktadır (Tjoa ve Guan, 2020).

Bununla birlikte, XAI teknikleri bazı sınırlılıklar da barındırmaktadır. Özellikle model agnostik açıklamaların doğruluğu sınırlı olabileceği gibi, açıklama üretimi süreci ek hesaplama gücü gerektirdiğinden sistem performansını da olumsuz etkileyebilir (, 2019). Ayrıca, açıklamalar kullanıcı için yanıltıcı olabilecek biçimde sadeleştirilebilir; bu da güven yerine yanlış bir algı oluşmasına neden olabilir.

Günümüzde SHAP dışında LIME (Local Interpretable Model-agnostic Explanations) ve karşıt örnekleme (counterfactual explanations) gibi yöntemler de yaygın şekilde kullanılmaktadır. Bu araçlar, kullanıcıların bir sistemin belirli kararları hangi girdilere dayanarak ürettiğini anlamasına olanak tanımakta; bu da açıklanabilirliğin, DSS sistemlerinin benimsenmesini artıran temel bir unsur hâline gelmesini sağlamaktadır.

XAI yaklaşımı, yapay zekâ destekli karar destek sistemlerini daha güvenilir, şeffaf ve etik açıdan sorumlu hâle getirmek için kritik bir bileşen hâline gelmiştir. Ancak bu alanda açıklamaların

doğruluğu, kullanılabilirliği ve bağlama uygunluğu gibi metodolojik sorunlar, hâlâ çözülmeyi bekleyen önemli araştırma konuları arasında yer almaktadır.

4. BAŞARI ÖRNEKLERİ VE SEKTÖREL UYGULAMALAR

Yapay zekâ destekli karar destek sistemleri (AI-DSS), veri yoğun ve karar odaklı süreçlerde sundukları öngörü, otomasyon ve optimizasyon olanaklarıyla birçok sektörde dikkat çekici başarılar ortaya koymuştur. Bu sistemlerin farklı endüstriyel bağlamlardaki uygulanabilirliği, yalnızca teknik potansiyelini değil, aynı zamanda iş süreçlerine stratejik katkılarını da gözler önüne sermektedir.

Sağlık sektöründe, AI-DSS sistemleri tanı koyma, tedavi planlama ve hasta izlemede devrim yaratmıştır. Özellikle görüntü analizi için geliştirilen CNN tabanlı sistemler, röntgen ve MR görüntülerinden akciğer nodülleri, tümör izleri veya organ anomalileri gibi kritik verileri tespit edebilmekte, uzman hekimlerin karar alma süreçlerini desteklemektedir (Rajkomar vd., 2018). Ayrıca, diyabet gibi kronik hastalıkların yönetiminde bireyselleştirilmiş öneriler sunarak tedavi süreçlerini iyileştirmektedir (Topol, 2019).

Finans alanında, makine öğrenmesi ve doğal dil işleme (NLP) teknikleri; risk değerlendirme, dolandırıcılık tespiti ve müşteri segmentasyonu gibi süreçlerde yaygın olarak kullanılmaktadır. Örneğin, kredi başvuru sistemleri kullanıcı geçmiş verilerini analiz ederek onay-red kararlarını otomatikleştirmekte, sahtekârlık algılama sistemleri ise anomali tespiti ile erken müdahale olanağı sağlamaktadır. Zhou vd (2023), kullanıcı merkezli ve açıklanabilir yapay zekâ çerçevesiyle bu tür sistemlerin hem doğruluğunu hem de güvenilirliğini artırmayı başarmıştır. NLP bankacılıkta chatbot'lar aracılığıyla hem müşteri hizmetlerini otomatikleştirmekte hem de belge işleme süreçlerini hızlandırmaktadır (Jurafsky ve Martin, 2020).

Tarım sektöründe, AI-DSS sistemleri iklim verilerini, toprak analizlerini ve uydu görüntülerini entegre ederek hasat tahminleri, gübreleme zamanlaması ve zararlı uyarı sistemleri geliştirmektedir. Bu sayede ürün kayıplarının azaltılması, kaynakların verimli kullanımı ve sürdürülebilir tarım pratikleri mümkün olmaktadır. Özellikle gelişmekte olan ülkelerde bu sistemler, teknolojiye dayalı kırsal kalkınma politikalarının temel araçlarından biri hâline gelmiştir.

Üretim ve sanayi uygulamalarında ise yapay zekâ destekli sistemler, üretim hattı optimizasyonu, arıza tahmini ve tedarik zinciri yönetimi gibi alanlarda kullanılmaktadır. Örneğin, nesnelerin interneti (IoT) ile entegre çalışan DSS çözümleri, sensör verilerini gerçek zamanlı analiz ederek makinelerin önleyici bakımını mümkün kılmakta ve iş gücü verimliliğini artırmaktadır. Akıllı fabrikalar bu tür sistemler sayesinde daha esnek, otonom ve öz-düzenleyici üretim altyapılarına ulaşmıştır.

Kamu yönetimi ve hukuk gibi alanlarda da AI-DSS çözümleri kullanılmaya başlanmıştır. Mahkeme kararlarının otomatik özetlenmesi, mevzuat analiz sistemleri ve vergi tahakkuk kararlarının modellenmesi gibi uygulamalar, insan kaynaklı hataları azaltmakta ve işlemleri hızlandırmaktadır. Ancak bu bağlamda, açıklanabilirlik ve etik denetim mekanizmalarının önemi daha da artmaktadır.

AI tabanlı karar destek sistemlerinin sektörel başarısı, yalnızca teknik yeterliliklerine değil, aynı zamanda bağlam duyarlılığına ve etik sorumluluklarına da bağlıdır. Her sektörün ihtiyaç duyduğu veri yapıları, regülasyonlar ve kullanıcı profilleri göz önünde bulundurularak geliştirilen sistemler, organizasyonlara stratejik avantaj sağlamaktadır. Gelecekte bu başarının sürdürülebilmesi, alanlar arası iş birliği, açıklanabilirlik çözümleri ve yerel bağlama duyarlı sistem tasarımıyla mümkün olacaktır.

Tablo 1 AI-DSS Literatürünün Sektörel ve Yöntemsel Dağılımı

Yazar (Yıl)	Sektör	Kullanılan Yöntem	Ana Tema / Katkı
Rajkomar vd. (2018)	Sağlık	Derin Öğrenme (CNN, NLP)	Görüntü analizi ve klinik karar desteği

Kose vd. (2021)	Sağlık	NLP ve DL	Klinik metinlerden bilgi çıkarımı
Zhou vd. (2023)	Finans	XAI, ML, RL	Kredi skorlama, dolandırıcılık tespiti
Javaid vd. (2022)	Üretim/Lojistik	IoT, Görüntü İşleme, Optimizasyon	Predictive maintenance, rota yönetimi
Challen vd. (2019)	Çevre/Yönetişim	GIS, IoT, FL	Akıllı şehir uygulamaları, veri güvenliği
Kostopoulos vd. (2024)	Sağlık/Hukuk	XAI, SHAP	Açıklanabilirlik ve kullanıcı güveni
Kaur vd. (2022)	Çoklu Sektör	Fairness, Data Ethics	Veri kalitesi, toplumsal etki analizleri
Rieke vd. (2020)	Sağlık	Federated Learning	Veri gizliliği, yerel model eğitimi
Liao vd. (2024)	Sağlık/Finans	Explanation Usability	XAI açıklamalarının algılanabilirliği
Chen vd. (2023)	Finans/Etik	AS-FairFed, FL	Grup adaleti ve adil modelleme

Tablo 1, çalışmada incelenen literatürün sektörler, yöntemler ve tematik katkılar açısından dağılımını özetlemektedir. Görüldüğü üzere, AI-DSS sistemleri çok çeşitli alanlarda uygulanmakta ve farklı yapay zekâ yaklaşımları ile desteklenmektedir. Bu çeşitlilik, karar destek sistemlerinin yalnızca teknik yeterlilik değil, aynı zamanda etik ve yönetsimsel bağlamda da ele alınması gerektiğini göstermektedir. Özellikle XAI, federated learning ve fairness yaklaşımlarının farklı sektörlerde nasıl uygulandığı bu çalışma kapsamında detaylı bir şekilde değerlendirilmiştir

4.1. Sağlık Sektöründe Yapay Zekâ Destekli DSS

Sağlık sektörü, veri yoğunluğu, zaman baskısı ve yüksek karar riski gibi faktörler nedeniyle yapay zekâ destekli karar destek sistemleri (AI-DSS) için öncelikli uygulama alanlarından biri hâline gelmiştir. Elektronik sağlık kayıtları (EHR), tıbbi görüntüler, sensör verileri ve hekim notları gibi yapılandırılmış ve yapılandırılmamış veri türleri, bu sistemlerin eğitim ve tahmin süreçlerinde temel veri kaynaklarını oluşturmaktadır (Kose vd., 2021).

Derin öğrenme teknikleri, özellikle medikal görüntüleme alanında yüksek başarı göstermiştir. Örneğin, konvolüsyonel sinir ağları (CNN) akciğer kanseri gibi kompleks vakalarda röntgen ve BT görüntülerinden yüksek doğrulukla tanı önerileri üretebilmekte ve bu yönüyle insan uzmanlarla karşılaştırılabilir, hatta bazı durumlarda onları aşan performans sergileyebilmektedir (Rajkomar vd., 2018). Benzer biçimde, doğal dil işleme (NLP) destekli DSS sistemleri, klinik raporlar, tanı notları ve hasta öykülerinden anlamlı çıkarımlar yaparak karar süreçlerine bilişsel katkı sağlamaktadır.

Ancak bu potansiyele rağmen, sağlık sektöründe AI-DSS sistemlerinin yaygın ve güvenli şekilde benimsenmesinin önünde bazı kritik engeller bulunmaktadır. Bunlardan ilki, verinin heterojenliği ve eksikliğiyle ilgilidir. Farklı hastaneler ve sistemler arasında veri formatlarının standartlaşmamış olması, sistem entegrasyonunu güçleştirmektedir. İkinci olarak, açıklanabilirlik eksikliği (XAI), klinisyenlerin kararların arka planına dair yeterli bilgiye ulaşamamasına yol açmakta ve sistem güvenilirliğini sınırlamaktadır. Son olarak, etik ve hukuki regülasyonlara uyum konusu, özellikle Avrupa Birliği'nin GDPR gibi çerçevelerinde sistemlerin hesap verebilirlik taşımasını zorunlu kılmaktadır.

Bu bağlamda, gelecekte transfer learning, few-shot learning gibi düşük veri gereksinimli öğrenme tekniklerinin sağlık verilerine adapte edilmesi, sınırlı örneklerle de yüksek doğruluk sağlayan sistemlerin geliştirilmesini kolaylaştıracaktır. Ayrıca, SHAP, LIME gibi XAI tekniklerinin klinik karar destek sistemlerine entegre edilmesiyle birlikte, hekimlerin sistem kararlarını denetleyebilmesi ve hasta güvenliğinin artırılması mümkün olacaktır.

Sonuç olarak, sağlık sektöründe AI-DSS çözümleri yalnızca teknolojik değil, aynı zamanda etik, yasal ve insan-merkezli tasarım ilkeleriyle de desteklenmeli; bu sistemler, karar destekten çok "karar ortaklığı" perspektifiyle değerlendirilmelidir.

4.2. Finans Sektöründe Yapay Zekâ Destekli DSS

Finans sektörü, karar alma süreçlerinin yoğun veri akışı ve yüksek risk ortamında gerçekleştiği bir bağlam sunması nedeniyle, yapay zekâ destekli karar destek sistemlerinin (AI-DSS) yaygınlıkla uygulandığı alanlardan biridir. Kredi geçmişi, müşteri davranışları, işlem kayıtları ve piyasa verileri gibi büyük hacimli ve yüksek frekanslı veri setleri, bu sistemlerin temel girdilerini oluşturmakta; makine öğrenmesi (ML), pekiştirmeli öğrenme (Reinforcement Learning – RL) ve açıklanabilir yapay zekâ (XAI) yöntemleri aracılığıyla analiz edilmektedir (Zhou vd., 2023).

Bu teknolojiler, finansal karar süreçlerinde çok çeşitli uygulamalara sahiptir. Örneğin, ML tabanlı modeller kredi başvurularının değerlendirilmesinde geleneksel skorlama yöntemlerine kıyasla daha yüksek doğruluk oranları sunmakta; dolandırıcılık tespiti için geliştirilen anomaliliğe duyarlı sistemler, olağandışı işlem davranışlarını erken aşamada tanımlayarak müdahale süresini kısaltmaktadır. RL uygulamaları ise özellikle algoritmik yatırım stratejilerinde piyasa tepkilerini gerçek zamanlı öğrenme kapasitesiyle modelleyerek portföy yönetimini otomatikleştirmektedir.

Ancak AI-DSS çözümlerinin finans alanındaki uygulanabilirliği yalnızca teknik başarımla sınırlı değildir. Bu sistemler, giderek artan şekilde veri önyargısı, ayrımcılık riski ve etik hesap verebilirlik gibi toplumsal açıdan hassas sorunları da gündeme getirmektedir. Özellikle kredi ve sigorta gibi alanlarda kullanılan karar modellerinin geçmiş verilere dayalı olarak tarihsel önyargıları yeniden üretme riski taşıdığı gözlemlenmiştir. Ayrıca, algoritmaların karar gerekçelerini açıklama kapasitesinin sınırlı olması, kullanıcılar ve regülatörler açısından sistemlerin şeffaflığını zedelemektedir.

Bu bağlamda, federated learning (merkeziyetsiz öğrenme) gibi yeni nesil mimariler, hem veri gizliliğini korumak hem de model başarımını sürdürmek amacıyla ön plana çıkmaktadır. Bu yaklaşım, finansal verinin yerel olarak işlenmesini sağlayarak GDPR gibi regülasyonlara daha uygun sistemler geliştirilmesini mümkün kılmaktadır. Aynı zamanda XAI tekniklerinin entegrasyonu, algoritmaların karar mekanizmalarını kullanıcılar için anlaşılabilir kılarak sistem güvenilirliğini artırma potansiyeli taşımaktadır.

Finansal alandaki AI-DSS çözümleri yüksek doğruluk ve hız vadederken, bu sistemlerin etik, adil ve şeffaf biçimde çalışmasını garanti altına alacak açıklanabilirlik, denetlenebilirlik ve regülasyon uyumu kriterlerinin merkeze alınması gerekmektedir. Bu doğrultuda geliştirilecek sistemler, yalnızca teknik başarı değil aynı zamanda toplumsal güven inşa etmeyi de başaran stratejik araçlara dönüşebilecektir.

4.3. Üretim ve Lojistik Sektöründe Yapay Zekâ Destekli DSS

Üretim ve lojistik sektörleri, dijital dönüşümün etkisiyle karar destek sistemlerinin (DSS) giderek daha fazla yapay zekâ (AI) tabanlı hale geldiği alanlar arasında yer almaktadır. Endüstri 4.0 paradigması doğrultusunda, bu sektörlerde gerçek zamanlı karar alma gereksinimi, AI destekli DSS çözümlerinin benimsenmesini hızlandırmıştır. Sensör ağları, IoT cihazları ve siber-fiziksel sistemlerden elde edilen yüksek hacimli veriler, makine öğrenmesi ve optimizasyon algoritmalarıyla işlenerek üretim hatlarının ve tedarik zincirlerinin daha esnek ve verimli hâle gelmesine olanak tanımaktadır (Javaid vd., 2022).

Öngörücü bakım (predictive maintenance), bu alandaki en yaygın uygulamalardan biridir. Makine öğrenmesi temelli tahmin algoritmaları sayesinde, ekipman arızaları henüz gerçekleşmeden öngörülerek üretim kesintileri minimize edilmekte, bakım maliyetleri azaltılmakta ve ekipman ömrü uzatılmaktadır. Örneğin, titreşim sensörlerinden toplanan zaman serisi verileri, arıza eğilimlerini tespit etmek amacıyla RNN veya LSTM modelleriyle analiz edilmekte ve bakım zamanlamaları buna göre optimize edilmektedir.

Lojistik sektöründe ise AI tabanlı DSS sistemleri, rota optimizasyonu, talep tahmini ve stok yönetimi gibi karar alanlarında operasyonel verimliliği artırmaktadır. Dinamik güzergâh planlama algoritmaları, trafik verilerini ve teslimat sürelerini analiz ederek yakıt tüketimini azaltmakta; stok seviyesi tahminleriyle hem israfın önüne geçilmekte hem de müşteri memnuniyeti artırılmaktadır.

Ancak bu sistemlerin başarısı, yalnızca teknik doğrulukları ile değil, aynı zamanda veri bütünlüğü, entegrasyon kapasitesi ve insan-makine etkileşimi gibi faktörlerle de doğrudan ilişkilidir. Özellikle sensör kalitesindeki değişkenlik, çoklu veri kaynaklarının senkronizasyonundaki zorluklar ve operatörlerin AI sistemlerine duyduğu güven, bu teknolojilerin ölçeklenebilirliğini etkilemektedir.

Ayrıca, insan-robot işbirliği bağlamında, yapay zekâ destekli DSS'lerin yalnızca otomasyon aracı değil, aynı zamanda insan kararlarını tamamlayıcı bir unsur olarak tasarlanması gerekmektedir. Bu bağlamda, insan merkezli tasarım ilkelerinin üretim süreçlerine entegre edilmesi, hem güvenliği hem de sistem kabulünü artıracaktır.

Üretim ve lojistik sektörlerinde AI-DSS sistemlerinin başarısı, teknolojik kabiliyetlerin yanı sıra sosyal ve kurumsal boyutların da bütüncül biçimde dikkate alınmasına bağlıdır. Veri kalitesi, kullanıcı deneyimi ve açıklanabilirlik bu sistemlerin sürdürülebilirliği açısından belirleyici olmaya devam edecektir.

4.4. Çevre Yönetimi ve Akıllı Şehir Uygulamaları

Çevre yönetimi ve akıllı şehir uygulamaları, çok değişkenli ve dinamik veri akışlarının etkin yönetimini gerektiren karmaşık sistemlerdir. Bu bağlamda, yapay zekâ destekli karar destek sistemleri (AI-DSS), sürdürülebilirlik hedefleri doğrultusunda kritik karar süreçlerinin optimize edilmesinde önemli işlevler üstlenmektedir. Özellikle yüksek frekanslı IoT sensör verileri, hava durumu istasyonları, trafik kameraları ve Coğrafi Bilgi Sistemleri (GIS) gibi kaynaklardan elde edilen çok boyutlu veriler, çevresel olayların gerçek zamanlı takibini ve yönetimini mümkün kılmaktadır (Challen vd., 2019).

Bu sistemler; enerji tüketiminin izlenmesi ve optimizasyonu, akıllı atık toplama, hava kirliliği tahmini, trafik yoğunluğu yönetimi ve su kaynaklarının sürdürülebilir planlaması gibi çok çeşitli alanlarda karar destek sağlamaktadır. Örneğin, zaman serisi tahmin modelleri ile şehir içi enerji talebinin saatlik bazda öngörülmesi; üretim-tüketim dengesi için kritik önem taşımaktadır. Benzer şekilde, konvolüsyonel sinir ağları kullanılarak uydu görüntülerinden hava kalitesi haritaları oluşturulmakta ve yüksek riskli bölgelerde önleyici müdahaleler yapılabilmektedir.

Bununla birlikte, çevresel DSS sistemlerinin geliştirilmesinde bazı temel kısıtlar da gözlemlenmektedir. Veri kalitesi ve güvenilirliği, özellikle düşük maliyetli sensörlerin yaygın kullanımı nedeniyle ciddi bir sorun alanıdır. Ayrıca, gerçek zamanlı veri işleme ve karar üretme yetkinliği, özellikle şehir ölçeğinde yüksek veri trafiği altında sınırlı kalabilmektedir. Ayrıca, siber güvenlik ve veri mahremiyeti gibi alanlarda da yeterli standartların geliştirilmemiş olması, bu sistemlerin uzun vadeli sürdürülebilirliğini tehdit etmektedir.

Son dönem çalışmaları, bu problemlerin aşılabilmesi için “Green AI” yaklaşımına odaklanmaktadır. Green AI, yapay zekâ modellerinin enerji verimliliği ve çevresel etkileri açısından optimize edilmesini amaçlamakta; daha az hesaplama gücüyle benzer doğrulukta sonuçlar veren modellerin geliştirilmesini teşvik etmektedir. Bu bağlamda, örneğin düşük enerji tüketimli model mimarilerinin veya yerel uç (edge) cihazlar üzerinde çalışan dağıtık öğrenme sistemlerinin yaygınlaştırılması, hem çevresel hem de operasyonel faydalar sağlayabilir.

Çevre yönetimi ve akıllı şehir uygulamalarında AI-DSS sistemlerinin sunduğu olanaklar artmakla birlikte, bu teknolojilerin güvenilir, sürdürülebilir ve enerji-etik biçimde geliştirilmesi, gelecekteki araştırmaların odak noktalarından biri olmalıdır..

4.5. Sektörel Özeti ve Değerlendirme

Farklı sektörlerde yapay zekâ destekli karar destek sistemlerinin kullanımı, her alana özgü dinamiklerle şekillenmektedir. Sağlık sektöründe kararların açıklanabilirliği ve hasta güvenliği ön planda yer alırken; finansal uygulamalarda adalet, şeffaflık ve veri önyargısının önlenmesi en önemli meseleler arasında öne çıkmaktadır. Üretim ve lojistik sistemlerinde gerçek zamanlı veri işleme ve sensör entegrasyonu teknik açıdan kritik rol oynarken; çevre yönetimi ve akıllı şehir uygulamalarında ise sürdürülebilirlik, enerji verimliliği ve düşük hesaplama maliyetleri önceliklidir.

Tablo 2 Yapay Zekâ Destekli DSS Sistemlerinin Sektörel Özeti

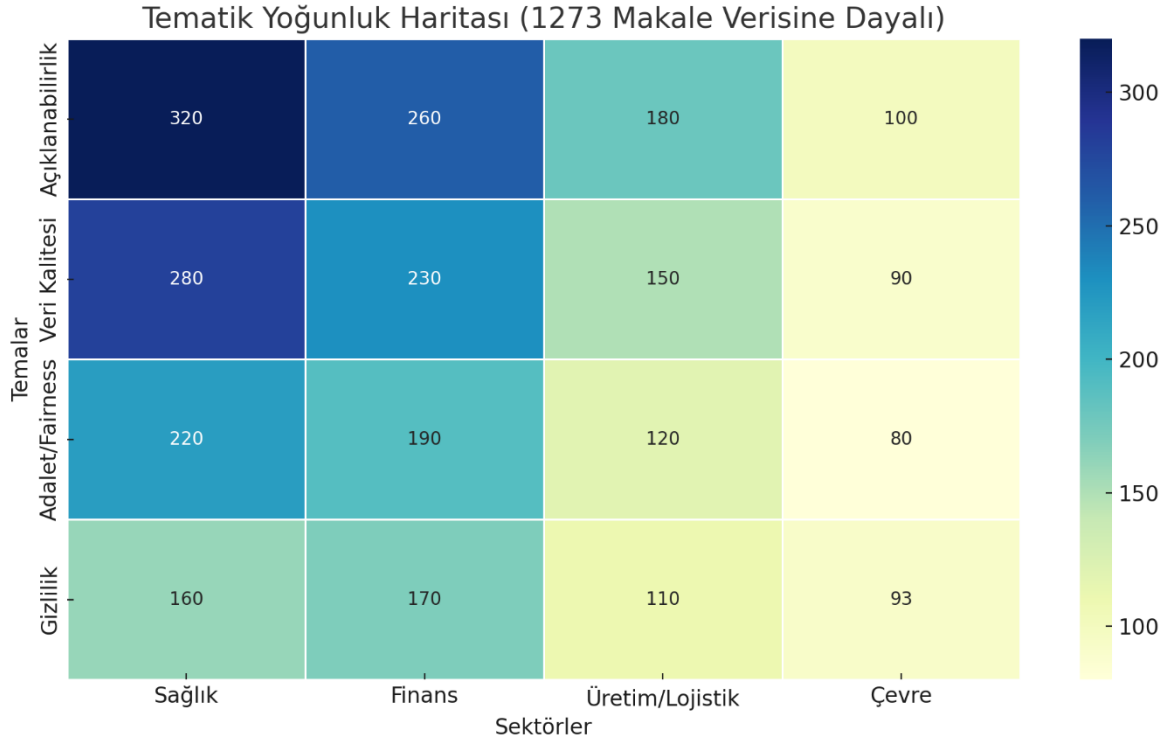
Sektör	Kullanılan Teknikleri	AI	Veri Türleri	Başarı Kriterleri	Sınırlamalar ve Gelecek Yönelimler
Sağlık	ML, DL, NLP		EHR, Görüntüler	Tanı doğruluğu, Karar süresi	Veri heterojenliği, XAI eksikliği, Transfer learning
Finans	ML, RL, XAI		İşlem verileri, Profiller	Risk skorlama, Dolandırıcılık tespiti	Veri güvenliği, Model biası, Fair AI
Üretim / Lojistik	ML, CV, Optimizasyon		Sensör, Üretim, Lojistik verileri	Verimlilik, Bakım tahmini	Sensör ağı entegrasyonu, İnsan-robot DSS
Çevre / Şehircilik	ML, Time Series, Optimizasyon		IoT verileri, GIS	Kirlilik izleme, Enerji verimliliği	Sensör kalitesi, Real-time işlem, Green AI

Tüm bu farklılıkları anlamlandırmak ve ortak temaları tespit etmek adına Tablo 2’de sunulan özet tablo yol gösterici olmaktadır. Tablo, hem kullanılan yapay zekâ tekniklerini hem de karşılaşılan sınırlılıkları karşılaştırmalı biçimde sunarak, disiplinler arası stratejilerin geliştirilmesine katkı sunmaktadır.

Bu karşılaştırma, sadece teknik farklılıkları değil aynı zamanda yöntemsel ve etik kaygıları da görünür kılmakta; DSS sistemlerinin sektör bağımsız bir şekilde nasıl daha adil, şeffaf ve sürdürülebilir hâle getirilebileceği sorusuna zemin hazırlamaktadır.

5. KARŞILAŞILAN ZORLUKLAR VE GELECEK ARAŞTIRMA ALANLARI

Yapay zekâ destekli karar destek sistemleri (AI-DSS), çeşitli sektörlerde karar alma süreçlerini daha hızlı, daha doğru ve daha tutarlı hâle getirme potansiyeli sunmaktadır. Ancak bu teknolojilerin yaygın kullanımında önemli teknik, etik ve operasyonel zorluklar bulunmaktadır. Aynı zamanda, bu zorluklar yeni araştırma fırsatlarını da beraberinde getirmektedir.



Şekil 1 Tematik Yoğunluk Haritası

Şekil 2, incelenen 1273 yayında yapay zekâ tabanlı karar destek sistemlerinin farklı temalar etrafında nasıl konumlandığını sektörel düzeyde görselleştirmektedir. Sağlık ve finans sektörlerinde “açıklanabilirlik” ve “veri kalitesi” temalarının yoğun biçimde ele alındığı görülmektedir. Üretim ve lojistik alanında daha çok “veri işleme” ve “optimizasyon” odaklı çalışmaların ön plana çıktığı; çevre yönetimi ise “gizlilik” temasıyla öne çıktığı dikkat çekmektedir. Harita, temaların sadece teknik değil, aynı zamanda etik ve yönetsimsel boyutlarda da sektör bazlı farklılıklar gösterdiğini ortaya koymakta; böylece gelecekteki araştırmalara yön verecek öncelikli alanların belirlenmesine katkı sağlamaktadır.

5.1. Teknik Zorluklar

5.1.1. Veri Kalitesi ve Heterojenliği

Yapay zekâ tabanlı karar destek sistemlerinin (AI-DSS) başarısı, büyük ölçüde kullanılan verilerin kalitesine ve yapısal tutarlılığına bağlıdır. Ancak gerçek dünya verileri çoğunlukla eksik, heterojen, etiketsiz veya yanlış biçimlendirilmiş durumdadır. Bu durum, model eğitimi sırasında hem doğruluk oranını düşürmekte hem de sistemin farklı veri kümeleriyle genellenebilirliğini sınırlandırmaktadır. Özellikle sağlık, çevre ve finans gibi sektörlerde veriler, çeşitli kaynaklardan farklı standartlarla toplanmakta; bu da veri birleşimini ve analizini teknik olarak karmaşık hâle getirmektedir.

Örneğin, Kaur vd. (2022) tarafından yapılan bir çalışmada, klinik karar destek sistemlerinde kullanılan verilerin yaklaşık %40'ının eksik veya tutarsız biçimlendiği tespit edilmiştir. Bu tür dengesizlikler, modellerin eğitim sürecinde overfitting riskini artırmakta ve sistemin güvenilirliğini doğrudan tehdit etmektedir. Ayrıca, model çıktılarının farklı veri kümeleri üzerinde büyük sapmalar göstermesi, DSS sistemlerinin pratikte uygulanabilirliğini azaltmaktadır.

Veri kalitesi sorunu yalnızca teknik bir problem değil, aynı zamanda etik bir risk alanıdır. Kötü yapılandırılmış ya da önyargılı veriler, algoritmaların sistematik olarak hatalı ya da adaletsiz sonuçlar üretmesine neden olabilir. [Kaur vd. (2022)] çalışmasında, eksik verilerin yalnızca model başarımını değil, aynı zamanda kullanıcı güvenini de zedelediği, dolayısıyla bu durumun sistemin benimsenmesini engellediği vurgulanmaktadır.

Bu sorunların aşılması için farklı disiplinlerden gelen yaklaşımların bir araya getirildiği hibrit çözümler önerilmektedir. Ontoloji tabanlı veri bütünleştirme, semantik uyumlaştırma ve ileri düzey imputasyon teknikleri sayesinde eksik ya da bozuk veri setlerinin yeniden yapılandırılması mümkündür. Özellikle dağıtık sistemler için geliştirilen çözümler arasında, Chen vd. (2023) tarafından önerilen federated learning ortamlarında veri kalitesi yönetimi dikkate değerdir. Bu model, merkezi olmayan sistemlerde veri kalitesini korumak için lokal düzeyde denetim mekanizmalarının geliştirilmesini önermektedir.

Veri kalitesinin artırılması hem teknik doğruluk hem etik sorumluluk açısından AI-DSS sistemlerinin güvenli, şeffaf ve sürdürülebilir biçimde geliştirilmesinde temel bir önkoşuldur..

5.1.2. Açıklanabilirlik (Explainability)

Yapay zekâ destekli karar destek sistemlerinin (AI-DSS) en büyük yapısal sınırlılıklarından biri, karar üretme mekanizmalarının kullanıcılar tarafından yeterince anlaşılamamasıdır. Kostopoulos vd. (2024) tarafından yapılan bir çalışmada, açıklanabilirlik eksikliğinin klinisyenlerin yapay zekâ önerilerini dikkate alma oranını %35 oranında azalttığı ortaya konmuştur.

Geleneksel açıklanabilir yapay zekâ (XAI) yaklaşımları – örneğin SHAP ve LIME – model çıktılarının hangi girdilere ne kadar bağlı olduğunu teknik olarak hesaplayabilmektedir. Ancak bu teknik açıklamalar genellikle son kullanıcı (örneğin bir hekim, hukukçu ya da kamu yöneticisi) tarafından sezgisel olarak anlaşılamamaktadır. Liao vd. (2021) çalışması, açıklamaların yalnızca doğruluğuna değil, aynı zamanda algılanabilirliğine (explanation usability) de göre değerlendirilmesi gerektiğini vurgulamaktadır. Bu, teknik açıklamanın kullanıcı için anlaşılır, sade ve bağlamla uyumlu olması anlamına gelir.

Öte yandan, açıklanabilirlik ile model performansı arasında belirgin bir denge sorunu da bulunmaktadır. Yüksek doğruluk sağlayan modeller genellikle daha az yorumlanabilirken; açıklanabilir modellerin doğruluk düzeyi daha düşük olabilmektedir. Bu durumu değerlendiren Rudin (2019), açıklanabilirlik ve güven arasında optimum bir denge kuran “Trustworthy AI” sistemlerinin geliştirilmesini önermektedir. Bu sistemler, karar süreçlerinde yalnızca çıktı üretmekle kalmayıp, bu çıktının gerekçesini de kullanıcıya anlamlı biçimde sunabilmelidir.

Açıklanabilirliğin sağlanamaması yalnızca teknik değil, aynı zamanda etik ve hukuki bir risk de doğurmaktadır. Özellikle Avrupa Birliği'nin GDPR gibi regülasyonları, otomatik karar mekanizmalarının kullanıcıya “anlaşılır açıklamalar” sunmasını yasal bir zorunluluk olarak tanımlamaktadır. Bu nedenle XAI sistemlerinin yalnızca mühendislik düzeyinde değil, aynı zamanda insan-merkezli tasarım ilkeleri doğrultusunda geliştirilmesi gerekmektedir.

AI-DSS sistemlerinde açıklanabilirlik yalnızca bir “ek özellik” değil, sistemin kabul edilebilirliği, güvenilirliği ve denetlenebilirliği açısından vazgeçilmez bir gerekliliktir. Gelecek çalışmaların, açıklamaların teknik doğruluğu ile insan algısı arasında denge kuran, bağlamsal ve kullanıcı dostu XAI çözümleri geliştirmeye odaklanması gerekmektedir.

5.1.3. Gerçek Zamanlılık ve Performans Kısıtları

Bazı karar destek sistemleri, özellikle kritik altyapıların yönetiminde, yalnızca doğru değil aynı zamanda gerçek zamanlı kararlar üretebilme kapasitesine sahip olmalıdır. Trafik yönetimi, enerji dağıtımı, siber güvenlik, acil sağlık müdahaleleri ve sanayi otomasyonu gibi alanlarda kararların gecikmeli verilmesi, ciddi maddi ve insani zararlara yol açabilmektedir. Bu bağlamda, yapay zekâ destekli karar destek sistemlerinin (AI-DSS) yalnızca yüksek doğruluklu değil, aynı zamanda hızlı, kaynak verimli ve enerji açısından sürdürülebilir biçimde çalışması beklenmektedir.

Ancak bu gereksinim, mevcut yapay zekâ modellerinin teknik yapısıyla sıklıkla çelişmektedir. Guidotti vd. (2019) tarafından yapılan sistematik bir derleme çalışması, yüksek doğruluk sağlayan kompleks modellerin genellikle işlem gücü açısından ağır, hafıza açısından talepkâr ve enerji tüketimi

yönünden maliyetli olduğunu göstermiştir. Bu durum, gerçek zamanlı işlem gerektiren senaryolarda AI-DSS sistemlerinin performans darboğazına girmesine yol açabilmektedir.

Bu sorunun aşılması için önerilen yaklaşımlar arasında edge computing mimarileri ve model sıkıştırma teknikleri (örneğin model pruning ve quantization) yer almaktadır. Bu yaklaşımlar, modellerin daha az kaynakla çalışmasını sağlayarak karar gecikmesini azaltmayı amaçlamaktadır. Özellikle edge AI çözümleri sayesinde, verinin buluta gönderilmeden yerel cihazlarda işlenmesi mümkün olmakta ve bu sayede hem zaman hem de gizlilik açısından avantaj sağlanmaktadır.

Bununla birlikte, bu teknik çözümlerin performans üzerinde yaratabileceği kayıpların minimize edilmesi gerekmektedir. Donati, Macario ve Karim (2024) çalışmasında, geleneksel sürekli izleme yerine olay odaklı (event-driven) mimarilerin yaygınlaştırılmasının, performans kaybı yaşanmaksızın gerçek zamanlı AI kararlarını mümkün kılabileceği vurgulanmaktadır. Bu mimariler yalnızca ihtiyaç hâlinde işlem başlatmakta, böylece kaynak tüketimini ciddi ölçüde azaltmaktadır.

Gerçek zamanlılık yalnızca teknik bir gereklilik değil; kullanılabilirlik, kabul edilebilirlik ve ölçeklenebilirlik açısından da belirleyici bir faktördür. AI-DSS sistemlerinin hem merkezî hem de uç (edge) mimarilerle uyumlu, enerjiyi verimli kullanan ve zamana duyarlı kararlar üretebilecek biçimde tasarlanması, özellikle akıllı şehirler, savunma sistemleri ve yüksek frekanslı finansal işlem alanlarında vazgeçilmezdir.

5.2. Etik ve Hukuki Zorluklar

5.2.1. Önyargı ve Adalet (*Bias & Fairness*)

Yapay zekâ destekli karar destek sistemlerinin (AI-DSS) eğitildikleri verilerden türeyen önyargılar nedeniyle adaletsiz sonuçlar üretme riski, günümüzde en çok tartışılan etik sorunlardan biridir. Eğitim verilerinde yer alan toplumsal cinsiyet, etnik köken veya sosyoekonomik sınıf temelli dengesizlikler, sistem çıktılarında da doğrudan ayrımcılığa dönüşebilmektedir. Bu durum özellikle işe alım, kredi değerlendirme, sağlık erişimi ve sosyal hizmetlerde kullanılan yapay zekâ sistemlerinde, mevcut eşitsizlikleri derinleştirme potansiyeli taşımaktadır.

Chen vd. (2023) çalışması, sosyoekonomik verilerdeki önyargıların, AI sistemleri aracılığıyla işe alım ve finansal değerlendirme süreçlerinde sistematik dışlama etkisi yarattığını göstermektedir. Aynı çalışmada önerilen AS-FairFed modeli, federated learning tabanlı bir yaklaşımla grup düzeyinde adaleti sağlamaya odaklanmakta; bireysel performans doğruluğunun yanı sıra toplumsal eşitlik ölçütlerini de dikkate almaktadır. Bu tür algoritmik stratejiler, özellikle merkeziyetsiz sistemlerde adalet sağlama ihtiyacı açısından umut verici çözümler sunmaktadır.

Ancak algoritmik müdahaleler tek başına yeterli değildir. Yapay zekâ sistemlerinin adil sonuçlar üretmesini garanti altına alabilmek için, veri toplama aşamasından model uygulamasına ve sonuçların raporlanmasına kadar sürecin tamamında etik duyarlılık gözetilmelidir. Bu çerçevede, veri temsiliyetinin sağlanması, model düzeyinde adalet metriklerinin kullanılması ve karar süreçlerinin bağımsız denetime açık hale getirilmesi, bütüncül bir etik yönetim anlayışını gerektirmektedir. Sistematik “fairness auditing” uygulamaları, yalnızca teknik uyumluluğu değil, aynı zamanda toplumsal güveni ve şeffaflığı da artıran bir araç olarak değerlendirilmektedir.

AI-DSS sistemlerinde adaletin sağlanması, yalnızca bireysel değil kurumsal sorumluluk da taşır. Kamu politikalarının şekillendirilmesinde kullanılan algoritmaların, kamusal değerler doğrultusunda tasarlanması; ticari sistemlerin ise ayrımcılığı yeniden üretmeyecek şekilde kurgulanması etik yapay zekâ vizyonunun temelidir. Bu doğrultuda geliştirilecek denetim, tasarım ve politika araçları, yapay zekânın yalnızca etkili değil, aynı zamanda eşitlikçi ve sorumlu biçimde kullanılmasını sağlayacaktır.

5.2.2. Veri Gizliliği ve Paylaşım

Yapay zekâ destekli karar destek sistemlerinin (AI-DSS) en hassas bileşenlerinden biri, yüksek düzeyde kişisel ve kurumsal veriyle çalışmalarıdır. Özellikle sağlık, finans ve sosyal hizmet gibi

alanlarda kullanılan sistemler; elektronik sağlık kayıtları, işlem geçmişleri ve bireysel davranışsal veriler gibi son derece duyarlı bilgi türlerine dayanmaktadır. Bu verilerin merkezi sunucularda toplanması, hem güvenlik açıklarını hem de mahremiyet ihlali risklerini artırmakta; dolayısıyla veri gizliliği, yalnızca teknik değil aynı zamanda etik ve hukuki bir öncelik hâline gelmektedir.

Mevcut geleneksel veri işleme yöntemleri, özellikle artan veri hacmi ve hassasiyeti karşısında yetersiz kalmaktadır. Rieke vd. (2020), bu bağlamda veri mahremiyetinin korunması açısından merkezi olmayan öğrenme yaklaşımlarının, klasik sunucu-temelli mimarilere göre çok daha güvenli ve etkili olduğunu vurgulamaktadır. Bu yaklaşımlar arasında öne çıkan federated learning (FL), modelin eğitim sürecini uç cihazlarda gerçekleştirerek verinin sistem dışına çıkmasını engellemekte ve böylece gizliliği merkeze alan bir yapı sunmaktadır.

FL mimarisi, yalnızca teknik avantajlar değil, aynı zamanda yasal uyum açısından da önemli bir esneklik sağlamaktadır. Özellikle Avrupa Birliği Genel Veri Koruma Tüzüğü (GDPR) ve ABD HIPAA düzenlemeleri gibi veri gizliliğini merkeze alan normatif çerçeveler, AI sistemlerinin şeffaf, izlenebilir ve veri işleme sınırlarına uygun biçimde çalışmasını zorunlu kılmaktadır. Challen vd. (2019) tarafından yapılan çalışmalarda, federated learning sistemlerinin bu yasal düzenlemelere yüksek oranda uyum sağladığı ve klinik ortamlarda daha fazla benimsenmeye başladığı ifade edilmiştir.

Ancak FL sistemlerinin de sınırsız olmadığı bilinmektedir. Veri heterojenliği, lokal model kalitesi farklılıkları, senkronizasyon problemleri ve uç cihazların işlem gücü gibi faktörler, bu sistemlerin ölçeklenebilirliğini doğrudan etkilemektedir. Ayrıca, modelin güncellemeleri sırasında aktarılan ağırlıkların da dolaylı yoldan bilgi sızdırabileceği; bu nedenle FL sistemlerinin differential privacy ve secure aggregation gibi ek güvenlik katmanlarıyla desteklenmesi gerektiği vurgulanmaktadır.

Bu bağlamda, veri gizliliğini güvence altına alan ancak aynı zamanda karar destek sistemlerinin etkinliğini koruyan mimarilerin geliştirilmesi, hem etik sorumluluğun hem de toplumsal kabulün sağlanması açısından kritik bir araştırma alanıdır. Gelecekte kullanıcı merkezli gizlilik politikaları ve teknik inovasyonlar, yalnızca veriyi korumakla kalmayacak; aynı zamanda yapay zekâ sistemlerine duyulan güvenin de artmasını sağlayacaktır.

5.2.3. Hesap Verebilirlik (Accountability)

Yapay zekâ destekli karar destek sistemlerinin (AI-DSS) kullanıldığı uygulamalarda ortaya çıkan en tartışmalı etik ve hukuki sorunlardan biri, hatalı kararlar karşısında sorumluluğun kime ait olduğunun belirsiz olmasıdır. Sağlıkta yanlış teşhis, bankacılıkta kredi reddi veya hukukta tavsiye hataları gibi yüksek etkili sonuçlar doğuran kararlarda; bu sonuçlardan geliştirici mi, kullanıcı mı, yoksa sistemin kendisi mi sorumlu tutulmalıdır sorusu, halen yanıtlanmamış bir alan olarak durmaktadır. Bu belirsizlik, sistemlerin benimsenmesini yavaşlatmakta ve hem kullanıcı hem de kurum düzeyinde yasal güvensizlik yaratmaktadır.

Bu bağlamda Challen vd. (2019), mevcut yasal düzenlemelerin AI sistemlerinin karar süreçlerini kapsamakta yetersiz kaldığını, bu nedenle hem izlenebilirlik (traceability) hem de hesap verebilirlik (accountability) ilkelerine odaklanan yeni çerçevelerin gerekliliğini vurgulamaktadır. Yazarlar, özellikle açıklanabilir yapay zekâ (XAI) sistemlerinin yalnızca kararın sonucunu değil, bu sonuca hangi adımlar ve hangi veri girdileriyle ulaşıldığını belgeleyebilen decision provenance altyapılarıyla desteklenmesi gerektiğini öne sürmektedir. Bu yaklaşım, sadece teknik değil aynı zamanda yasal olarak da şeffaf ve denetlenebilir sistemlerin tasarlanması anlamına gelir.

Benzer şekilde, Kostopoulos vd. (2024), AI sistemlerinin toplumsal kabulünü artırmak için geliştirilecek yeni denetim ve belgeleme protokollerinin gerekliliğine işaret etmektedir. Bu bağlamda önerilen AI Impact Statements, karar destek sistemlerinin toplumsal, hukuki ve etik etkilerini önceden değerlendiren yapılar olarak öne çıkmaktadır. Ayrıca, karar sürecinin adım adım kaydedildiği audit trails gibi belgelenmiş izleme sistemleri de, sorumluluk zincirini açıkça ortaya koyarak hukuki hesap verebilirliği güçlendirmektedir.

Bu gelişmeler, AI-DSS sistemlerinin yalnızca teknik başarıyla değil, aynı zamanda sorumluluk mekanizmalarıyla da değerlendirileceği yeni bir yönetim dönemine işaret etmektedir. Gelecekte geliştirilecek sistemlerin tasarımında etik denetim, açıklanabilirlik ve sorumluluk ilkeleri birlikte ele alındığı sürece, yapay zekâ tabanlı kararların hukuki ve toplumsal meşruiyeti de güçlenecektir.

5.3. Gelecek Araştırma Alanları

Yapay zekâ destekli karar destek sistemlerinin (AI-DSS) artan kullanımı, bu sistemlerin yalnızca teknik başarıyla değil; etik sorumluluk, açıklanabilirlik ve yönetim ilkeleriyle ne ölçüde uyum sağladığıyla da değerlendirileceği çok katmanlı bir araştırma gündemi doğurmuştur. AI-DSS sistemlerinin sürdürülebilir, adil ve güvenilir şekilde gelişebilmesi için hem teknik ilerlemeye hem de normatif çerçevelerin yenilenmesine yönelik disiplinlerarası çalışmalar büyük önem taşımaktadır.

Bu bağlamda açıklanabilirlik (explainability) konusu yalnızca algoritmik düzeyde ele alınmamalıdır. Açıklamaların insan kullanıcılar tarafından nasıl algılandığı, karar süreçlerini nasıl etkilediği ve sistem güvenini ne ölçüde artırdığı gibi faktörlerin sistematik biçimde değerlendirilmesi gerekmektedir. Kostopoulos vd. (2024) açıklamaların kullanıcı güveni ve sistem kabulü üzerindeki etkisini ölçen çok boyutlu değerlendirme modelleri geliştirilmesi gerektiğini savunurken, Liao vd. (2021) kullanıcının açıklamalara hazır oluşu, XAI arayüzlerinin sezgisel tasarımı ve sosyal bağlamdaki anlamlandırma süreçlerine yönelik araştırmaların eksikliğine dikkat çekmektedir. Bu doğrultuda, açıklanabilirliğin yalnızca teknik doğruluk değil, aynı zamanda açıklama kullanılabilirliği (explanation usability) ile birlikte ele alınması önerilmektedir.

Veri gizliliği ile model performansı arasındaki denge de önemli bir araştırma alanı olarak öne çıkmaktadır. Federated learning gibi merkeziyetsiz mimarilerin hem mahremiyet koruma hem de yüksek doğruluk sağlama amacı taşımasına karşın, bu sistemlerin farklı sektörlerdeki uygulama yeterliliği henüz yeterince test edilmemiştir. Rieke vd. (2020) bu modellerin özellikle hassas verilerin bulunduğu sağlık ve finans gibi alanlarda saha temelli uygulama verilerinden yoksun olduğunu belirtmiş, bu nedenle yerel yasal düzenlemelerle teknik yapıların birlikte çalışabilirliğini ele alan ampirik çalışmaların artırılması gerektiğini vurgulamıştır.

Hesap verebilirlik ve etik yönetim konuları ise AI sistemlerinin kamusal alanda nasıl meşrulaştırılacağını doğrudan belirleyecek kritik bir çerçevedir. [Challen vd. (2019)] açıklanabilir yapay zekâ sistemlerinin, “AI Etki Bildirimi” (AI Impact Statements), algoritmik denetim raporları ve çok paydaşlı denetim panelleri gibi yönetim araçlarıyla desteklenmesini önermekte; bu araçların, karar süreçlerinin yalnızca teknik değil, hukuki ve toplumsal olarak da şeffaf hale gelmesini sağlayacağını savunmaktadır. Bu çerçevede geliştirilecek normatif modeller, kurumsal düzeyde sorumluluk zincirinin belirlenmesine katkı sunacaktır.

Son olarak, teknik ilerlemenin toplumsal ihtiyaçlardan bağımsız biçimde yönlendirilmesi, AI-DSS sistemlerinin uzun vadeli güvenini ve kabulünü tehdit etmektedir. Kaur vd. (2022) bu konuda teknoloji geliştirme süreçlerinin insan değerleri, kültürel farklılıklar ve sosyal eşitlik ilkeleriyle uyumlu olarak yeniden tasarlanmasının, özellikle kamu temelli uygulamalarda kullanıcı güveni açısından belirleyici olduğunu vurgulamaktadır. Bu yaklaşım, yalnızca etik tasarımın değil, aynı zamanda sorumlu inovasyonun da önünü açacaktır.

AI-DSS sistemlerinin geleceği yalnızca teknik algoritmaların değil, aynı zamanda toplumsal bağlamla kurulan ilişkinin niteliğine bağlıdır. Disiplinlerarası etkileşim, politika uyumu ve kullanıcı

merkezli tasarım ilkeleri bu sistemlerin hem teorik ilerleyişini hem de kamusal alandaki uygulanabilirliğini şekillendirecektir.

6. SONUÇ

Bu çalışma, yapay zekâ destekli karar destek sistemlerini (AI-DSS) yalnızca teknik bir inovasyon olarak değil; aynı zamanda etik, yönetsimsel ve toplumsal boyutlarıyla da ele alan disiplinlerarası bir derleme niteliği taşımaktadır. Sağlık, finans, üretim ve çevre gibi farklı sektörlerde yapay zekâ uygulamaları; karar alma süreçlerinde hız, doğruluk ve öngörülebilirlik sağlamaktadır. Ancak bu kazanımlar, veri kalitesi, model şeffaflığı, adalet, açıklanabilirlik ve kurumsal hesap verebilirlik gibi çok katmanlı zorluklar eşliğinde değerlendirilmektedir.

Bu bağlamda, AI-DSS sistemlerinin başarı ölçütleri yalnızca teknik doğruluk ve sistem performansına indirgenemez. Sistemlerin sosyal kabul görmesi, etik ilkelerle uyumlu olması ve hesap verebilir karar izleri bırakabilmesi gibi kriterler, özellikle yüksek riskli alanlarda (örneğin sağlık ve finans) vazgeçilmez hâle gelmiştir. Bu çalışmada, açıklanabilir yapay zekâ (XAI), federated learning, adil algoritmalar ve sorumlu veri yönetimi gibi stratejiler, literatürdeki çözüm yaklaşımları kapsamında ele alınmıştır.

Derleme sürecinde kullanılan BiBLöX aracı ile 1273 yayının sistematik olarak analiz edilmesi, literatürdeki temaları ve eksik kalan alanları görmemizi sağlamıştır. Özellikle açıklanabilirlik, önyargı, etik yönetim ve veri gizliliği konularının sektörel olarak nasıl farklılaştığı, bu sistemlerin homojen değil; bağlama duyarlı olarak geliştirilmesi gerektiğini ortaya koymaktadır. Ayrıca literatür ağırlıklı olarak uluslararası örneklere dayanmakta olup, Türkiye bağlamındaki yapay zekâ yönetimi ve etik değerlendirmeler için daha fazla akademik üretime ihtiyaç duyulduğu da görülmektedir.

Bu makale, AI-DSS alanındaki çalışmalara üç temel düzeyde katkı sağlamayı hedeflemektedir:

- (i) teknik yaklaşımların derlenmesi ve kavramsal karşılaştırılması,
- (ii) sektörel uygulamalarda karşılaşılan zorlukların yapılandırılması,
- (iii) literatürdeki boşluklara eleştirel bir bakışla dikkat çekilmesi.

Gelecekteki çalışmalar, daha fazla ampirik veri ile desteklenmiş vaka analizleri, yerel bağlamlarda yapılacak yönetim incelemeleri ve XAI'nın kullanıcı algısına etkisini ölçen deneysel çalışmalarla bu alandaki kavramsal derinliği artırabilir. Bu bağlamda çalışma, yalnızca mevcut durumu ortaya koymakla kalmayıp; aynı zamanda AI-DSS sistemlerinin daha kapsayıcı, güvenilir ve sorumlu biçimde geliştirilmesi için bir araştırma ajandası da önermektedir.

7. TARTIŞMA

Yapay zekâ destekli karar destek sistemleri (AI-DSS), teknik mimarilerinin ötesinde sosyoteknik bir çerçevede değerlendirilmeyi gerektiren karmaşık yapılardır. Farklı sektörlerdeki uygulamalar, bu sistemlerin sadece algoritmik doğruluğuyla değil; aynı zamanda kullanıcı güveni, veri önyargısı, açıklanabilirlik ve yönetim gibi etkenlerle bütünleşik olarak ele alınması gerektiğini göstermektedir.

Özellikle sağlık ve finans gibi yüksek riskli alanlarda, açıklanabilirlik (explainability) kullanıcı güveniyle doğrudan ilişkilidir. Düşing vd. (2024), federated learning ve karşıt nedensellik açıklamaları (counterfactual explanations) kullanarak sepsis tedavisinde klinik karar destek sistemlerinin açıklanabilirliğini artırmayı amaçlamıştır. Bu tür yaklaşımlar, yalnızca kararların neden verildiğini anlamayı değil, aynı zamanda alternatif senaryoları değerlendirmeyi de mümkün kılmaktadır. Benzer şekilde, Sarker vd. (2024), enerji yönetimi bağlamında açıklanabilirliğin sistem güvenliği ve şeffaflığı açısından kritik olduğunu belirtmektedir.

Veri önyargısı, AI-DSS sistemlerinin adaletli işleyişini tehdit eden bir başka önemli sorundur. Federated learning gibi dağıtık öğrenme modelleri, hem mahremiyetin korunmasına hem de veri

temsiliyetinin artırılmasına katkı sunabilir (Kalusivalingam ve Sharma, 2021; Tariq vd., 2024). Ancak, bu modellerin başarısı veri kalitesine ve merkeziyetsiz katılımın dengesine bağlıdır. Albahri vd. (2023), sağlıkta önyargı yönetiminin sadece algoritmik düzeyde değil, aynı zamanda veri toplama stratejileriyle birlikte düşünülmesi gerektiğini vurgular.

Kurumlar arası yönetim, AI sistemlerinin meşruiyetini ve sürdürülebilirliğini etkileyen belirleyici bir faktördür. Bârcena vd. (2022), güvenilir yapay zekâ uygulamalarının mahremiyet, veri yönetimi ve açıklanabilirlik gibi ilkeler etrafında şekillenmesi gerektiğini öne sürmektedir. E-ticaret bağlamında yapılan çalışmalar da (Garg vd., 2025), sektöre özgü etik ilkeler ve kullanıcı beklentilerinin açıklama stratejileriyle entegre edilmesinin gerekliliğine dikkat çekmektedir.

Tüm bu bulgular, AI-DSS sistemlerinin evrensel çözümlerle yönetilemeyeceğini; bağlamsal ihtiyaçlara, sektörel risklere ve kullanıcı profillerine göre özelleştirilmiş yaklaşımlar geliştirilmesi gerektiğini ortaya koymaktadır. Bu çerçevede, gelecekteki çalışmaların yalnızca teknik başarıya değil, aynı zamanda etik, sosyal ve yönetsel bütünlüğe odaklanması önem arz etmektedir.

KAYNAKÇA

- Antoniadi, A. M., Nam, J., Bae, J., ve Hwang, T. J. (2021). Current challenges and future opportunities for explainable artificial intelligence in medical imaging. *Applied Sciences*, 11(11), 5088. <https://doi.org/10.3390/app11115088>
- Brnabic, A., & Hess, L. (2021). Systematic literature review of machine learning methods used in the analysis of real-world data for patient-provider decision making. *BMC Medical Informatics and Decision Making*, 21(1), Article 54. <https://doi.org/10.1186/s12911-021-01403-2>
- Challen, R., Denny, J., Pitt, M., Gompels, L., Edwards, T., ve Tsaneva-Atanasova, K. (2019). Artificial intelligence, bias and clinical safety. *BMJ Quality ve Safety*, 28(3), 231–237. <https://doi.org/10.1136/bmjqs-2018-008370>
- Chen, Y., Zhang, X., ve Zhang, J. (2023). Adversarial fairness-aware federated learning for personalized healthcare. *Applied Sciences*, 13(18), 10258. <https://doi.org/10.3390/app131810258>
- Cortès, C., ve Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20, 273–297. <https://doi.org/10.1007/BF00994018>
- Donati, O., Macario, M., ve Karim, M. H. (2024). Event-Driven AI Workflows in Serverless Computing: Enabling Real-Time Data Processing and Decision-Making. Preprints. https://www.preprints.org/frontend/manuscript/2022188b96eabc099930f462849fb5f/download_pub
- Durkin, J. (1994). Expert systems: Design and development. Macmillan.
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Telecommunications Policy*, 42(10), 284–306. <https://doi.org/10.1007/s11023-018-9482-5>
- Goodfellow, I., Bengio, Y., ve Courville, A. (2016). Deep learning. MIT Press.
- Guidotti, R., Monreale, A., Ruggieri, S., Turini, F., Giannotti, F., ve Pedreschi, D. (2019). A survey of methods for explaining black box models. *ACM Computing Surveys*, 51(5), 93. <https://doi.org/10.1145/3236009>
- Holzinger, A., Biemann, C., Pattichis, C. S., ve Kell, D. B. (2019). What do we need to build explainable AI systems for the medical domain? Arxiv Preprints. <https://arxiv.org/abs/1712.09923>

- Javaid, M., Haleem, A., Singh, R. P., ve Suman, R. (2022). Role of artificial intelligence in smart manufacturing and logistics. *International Journal of Industrial Engineering*, 29(3), 355–367. <https://doi.org/10.1142/S2424862221300040>
- Jurafsky, D., ve Martin, J. H. (2020). *Speech and language processing* (3rd ed.). Pearson.
- Kaur, H., Sharma, A., ve Singh, R. (2022). Addressing fairness and bias in AI-based clinical DSS. *ACM Transactions on Management Information Systems*, 13(2), 1–23. <https://doi.org/10.1145/3491209>
- Kesgin, K., ve Ozer, D. (2025). BiBLoX: A flask-based automatic bibliometrics and machine learning system for scientific trend prediction. *Authorea Preprint*. <https://www.authorea.com/doi/full/10.22541/au.174106668.81840735>
- Kose, U., Deperlioglu, O., Alzubi, J., ve Patrut, B. (2021). Deep learning for medical decision support systems. <https://10.1007/978-981-15-6325-6>
- Kostopoulos, G., Giannopoulos, V., ve Sarigiannidis, P. (2024). Explainability and transparency in deep learning DSS: A healthcare perspective. *Electronics*, 13(14), 2842. <https://doi.org/10.3390/electronics13142842>
- Liao, Q. V., ve Varshney, K. R. (2021). Human-centered explainable ai (xai): From algorithms to user experiences. *arXiv preprint*. <https://arxiv.org/abs/2110.10790>.
- Liu, B. (2012). *Sentiment analysis and opinion mining*. Morgan ve Claypool Publishers.
- Rajkomar, A., Dean, J., ve Kohane, I. (2018). Machine learning in medicine. *npj Digital Medicine*, 1, 18. <https://doi.org/10.1038/s41746-018-0029-1>
- Rieke, N., Hancox, J., Li, W., Milletari, F., Roth, H. R., Albarqouni, S., ... Cardoso, M. J. (2020). The future of digital health with federated learning. *npj Digital Medicine*, 3, 119. <https://doi.org/10.1038/s41746-020-00323-1>
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1, 206–215. <https://doi.org/10.1038/s42256-019-0048-x>
- Samek, W., Wiegand, T., ve Müller, K.-R. (2017). Explainable artificial intelligence: Understanding, visualizing and interpreting deep learning models. *arXiv preprint*. <https://arxiv.org/abs/1708.08296>
- Sharda, R., Delen, D., & Turban, E. (2020). *Analytics, data science, & artificial intelligence: Systems for decision support* (11th ed.). Pearson.
- Shortliffe, E. H. (1976). *Computer-based medical consultations: MYCIN*. Elsevier.
- Topol, E. (2019). High-performance medicine: The convergence of human and artificial intelligence. *Nature Medicine*, 25, 44–56. <https://doi.org/10.1038/s41591-018-0300-7>
- Turban, E., Sharda, R., ve Delen, D. (2011). *Decision support and business intelligence systems* (9th ed.). Pearson Education.
- Zhou, Y., Li, H., Xiao, Z., & Qiu, J. (2023). A user-centered explainable artificial intelligence approach for financial fraud detection. *Finance Research Letters*, 58, 104309. <https://doi.org/10.1016/j.frl.2023.104309>