

T.C.
ERZİNCAN BİNALİ YILDIRIM ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ
BİLGİSAYAR VE ÖĞRETİM TEKNOLOJİLERİ EĞİTİMİ ANABİLİM DALI

YAPAY ZEKÂ GÜVENLİK VE AHLAK ALGISI ÖLÇEĞİ

Soner DEMİR

Danışman: Doç. Dr. Recep ÖZ

TEZ JÜRİ ÜYELERİ
Doç. Dr. V. Aytekin SANALAN
Doç. Dr. Embiya ÇELİK

YÜKSEK LİSANS TEZİ
ERZİNCAN, 2026

© 2026[Soner DEMİR]. Tüm hakları saklıdır.

Kabul ve Onay Sayfası

Doç. Dr. Recep ÖZ danışmanlığında, Soner DEMİR tarafından hazırlanan bu çalışma tarihinde aşağıdaki jüri tarafından Bilgisayar ve Öğretim Teknolojileri Eğitimi Anabilim Dalı'nda Yüksek Lisans Tezi olarak kabul oybirliği ile kabul edilmiştir.

Başkan : Doç. Dr. V. Aytekin SANALAN İmza:

Üye : Doç. Dr. Recep ÖZ İmza:

Üye : Doç. Dr. Embiya ÇELİK İmza:

Bu tez Enstitü Yönetim Kurulunun / / 2026 tarih ve/..... sayılı kararı ile onaylanmıştır.

Prof. Dr. Kemal Volkan ÖZDOKUR
Enstitü Müdür V.

Not: Bu tezde kullanılan özgün ve başka kaynaklardan yapılan bildirişlerin, şekil ve tabloların kaynak olarak kullanımı, 5846 sayılı Fikir ve Sanat Eserleri Kanunundaki hükümlere tabidir.

Bilimsel Etięe Uygunluk Sayfası

“Yapay Zeka Güvenlik ve Ahlak Algısı Ölçeęi” isimli “Yüksek Lisans” tezimin tarafımda intihal tespit programı ile incelenmiştir. Buna göre tezimde bilimsel etik ihlali ve intihal olarak nitelendirilebilecek herhangi bir durum olmadığını taahhüt ederim.

Bu çalışmadaki tüm bilgilerin, akademik ve etik kurallara uygun bir biçimde elde edildiğini; aynı zamanda bu kural ve davranışların gerektirdiğı gibi, bu çalışmanın özünde olmayan tüm materyal ve sonuçları tam olarak aktardığımı ve referans gösterdiğimi beyan ederim./...../20....

(İmza)

Soner DEMİR

ÖZET

YAPAY ZEKÂ GÜVENLİK VE AHLAK ALGISI ÖLÇEĞİ

Soner DEMİR

Yüksek Lisans Tezi

Erzincan Binali Yıldırım Üniversitesi, Fen Bilimleri Enstitüsü,

Bilgisayar ve Öğretim Teknolojileri Eğitimi Anabilim Dalı

Danışman: Doç. Dr. Recep ÖZ

2026, 174 sayfa

Bireylerin yapay zekâya yönelik güvenlik ve ahlak algılarını çok boyutlu olarak ölçen "Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği"ni (YZGAAÖ) geliştiren ve psikometrik özelliklerini inceleyen bu ölçek; güvenlik farkındalığı, ahlaki değerler, sorumluluk ve hesap verebilirlik, toplumsal etki ile olumlu beklentiler olmak üzere beş alt boyuttan oluşmaktadır. Araştırma karma yöntemli bir ölçek geliştirme çalışması olarak tasarlanmış ve çalışma grubunu 558 üniversite öğrencisi oluşturmuştur. Ölçek maddeleri uluslararası etik çerçeveler ve güncel ampirik çalışmalar temelinde hazırlanmış, uzman görüşleriyle kapsam geçerliği sağlanmıştır. Yapı geçerliği açımlayıcı ve doğrulayıcı faktör analizleriyle, güvenilirliği ise Cronbach's Alpha katsayısı, madde-toplam korelasyonu ve %27 alt-üst grup karşılaştırmasıyla test edilmiştir. Analizler sonucunda ölçeğin 40 maddeden oluşan beş alt boyutlu yapısı doğrulanmıştır. Güvenlik farkındalığı alt boyutu 7 madde ($\alpha=.942$), ahlaki değerler 10 madde ($\alpha=.953$), sorumluluk ve hesap verebilirlik 8 madde ($\alpha=.952$), toplumsal etki 8 madde ($\alpha=.966$) ve olumlu beklentiler 7 madde ($\alpha=.966$) ile ölçülmektedir. Tüm alt boyutlarda tek faktörlü yapı elde edilmiş, açıklanan varyans oranları %82.7 ile %93.7 arasında değişmiştir. Doğrulayıcı faktör analizi uyum indeksleri tüm boyutlarda kabul edilebilir ve iyi düzeyde bulunmuştur. Demografik analizlerde akademik alan değişkenlerine göre anlamlı fark bulunmazken, yaş ve eğitim düzeyine göre anlamlı farklılıklar tespit edilmiştir. Elde edilen bulgular, YZGAAÖ'nün yapay zekâ güvenlik ve ahlak algısını geçerli ve güvenilir biçimde ölçen çok boyutlu bir araç olduğunu göstermektedir.

Anahtar Kelimeler: Yapay zekâ, güvenlik algısı, ahlak algısı, ölçek geliştirme, psikometri, etik

ABSTRACT

ARTIFICIAL INTELLIGENCE SAFETY AND ETHICS PERCEPTION SCALE

Soner DEMİR

Master's Thesis

Erzincan Binali Yıldırım University, Institute of Science and Technology,

Department of Computer Education and Instructional Technology

Advisor: Doç. Dr. Recep ÖZ

2026, 174 pages

This study aimed to develop the Artificial Intelligence Safety and Ethics Perception Scale (AISEPS), a multidimensional instrument designed to assess individuals' perceptions of artificial intelligence in terms of safety and ethics, and to examine its psychometric properties. The research was conducted as a mixed-methods scale development study with a sample of 558 university students. Based on international ethical frameworks and recent empirical literature, an initial pool of items was created, and content validity was established through expert evaluations. Construct validity was assessed using Exploratory Factor Analysis and Confirmatory Factor Analysis, while reliability was evaluated through Cronbach's alpha coefficients, item-total correlations, and 27% upper-lower group comparisons. The findings confirmed a five-factor structure consisting of 40 items. The scale includes the dimensions of Safety Awareness (7 items, $\alpha = .942$), Ethical Values (10 items, $\alpha = .953$), Responsibility and Accountability (8 items, $\alpha = .952$), Social Impact (8 items, $\alpha = .966$), and Positive Expectations (7 items, $\alpha = .966$). Each subscale demonstrated a unidimensional structure, with explained variance ranging from 82.7% to 93.7%. CFA results indicated acceptable to good model fit across all dimensions. Demographic analyses revealed no significant differences according to academic field, whereas significant differences were found with respect to age and educational level. Overall, the results indicate that the AISEPS is a valid and reliable multidimensional instrument for measuring individuals' perceptions of artificial intelligence regarding safety and ethics.

Keywords: Artificial Intelligence, Safety Perception, Ethical Perception, Scale Development, Psychometrics, Ethics

TEŞEKKÜR

Bu tez çalışmasının planlanması, yürütülmesi ve tamamlanması sürecinde bilgi, deneyim ve değerli katkılarıyla bana yol gösteren, akademik bakış açısıyla çalışmama yön veren, her aşamada desteğini ve rehberliğini esirgemeyen saygıdeğer danışman hocam Doç. Dr. Recep ÖZ'e en içten teşekkürlerimi sunarım.

Lisansüstü eğitimim boyunca gösterdiği anlayış, sabır ve yapıcı önerileriyle bilimsel gelişimime önemli katkılar sağladığı için minnettarım.

Ayrıca, bu uzun ve yoğun süreçte her zaman yanımda olan, sabrı, anlayışı, sevgisi ve sonsuz desteğiyle bana güç veren sevgili eşim Büşra DEMİR'e ve kızım Serra DEMİR'e teşekkürlerimi sunarım. Zor zamanlarda gösterdikleri fedakârlıklar, moral ve motivasyonlarıyla bu çalışmanın tamamlanmasında en büyük destekçilerimden olmuşlardır. Varlıklarıyla, bu süreci kararlılıkla sürdürebilmemin en önemli kaynakları olmuşlardır.

Tez savunmam sürecinde değerli zamanlarını ayırarak çalışmamı değerlendiren, görüş ve önerileriyle katkıda bulunan saygıdeğer jüri üyelerine ayrı ayrı en içten teşekkürlerimi sunarım.

Bu süreçte anlayış gösterip gerekli izinlerimi sağlayıp desteklerini esirgemeyen iş arkadaşlarıma, derslerimde yardımlarını ve akademik rehberliklerini esirgemeyen bütün saygıdeğer hocalarıma ayrı ayrı teşekkürlerimi sunarım.

Soner DEMİR

Temmuz, 2026

İÇİNDEKİLER

ÖZET	i
ABSTRACT	ii
TEŞEKKÜR	iii
İÇİNDEKİLER	iv
TABLolar DİZİNİ	vi
ŞEKİLLER DİZİNİ	ix
SİMGELER VE KISALTMALAR DİZİNİ	x
1. GİRİŞ	12
1.1. Araştırmanın Amacı	21
1.2. Araştırmanın Önemi	21
1.3. Varsayımlar	23
1.4. Sınırlılıklar	24
2. KAVRAMSAL ÇERÇEVE VE İLGİLİ ÇALIŞMALAR	25
2.1. Güvenlik Farkındalığı Boyutu	25
2.2. Ahlakî Değerler Boyutu	26
2.3. Sorumluluk ve Hesap Verebilirlik Boyutu	28
2.4. Toplumsal Etki Boyutu	28
2.5. Olumlu Beklentiler Boyutu	29
2.6. Kavramsal Çerçevenin Entegrasyonu	30
3. YÖNTEM	33
3.1. Araştırmanın Modeli	33
3.2. Çalışma Grubu	34
3.3. Veri Toplama Araçları	34
3.4. Veri Toplama Süreci	36
3.5. Geçerlik Çalışmaları	36
3.6. Güvenirlilik Çalışmaları	38
3.7. Veri Analizi	39
3.8. Etik Hususlar	39
4. BULGULAR	40
4.1. Yapay Zeka Güvenlik Farkındalığı Alt Boyutu	44
4.2. Yapay Zeka Ahlakî Değerler Alt Boyutu	52
4.3. Yapay Zeka Sorumluluk ve Hesap Verilebilirlik Alt Boyutu	60
4.4. Yapay Zeka Toplumsal Etki Alt Boyutu	69
4.5. Yapay Zeka Olumlu Beklentiler Alt Boyutu	77

4.6. Demografik Özellikler.....	85
4.6.1. Cinsiyet Dağılımı.....	85
4.6.2. Yaş Dağılımı.....	86
4.6.3. Eğitim Düzeyi Dağılımı	86
4.6.4. Bölüm (Sayısal/Sözel) Dağılımı.....	86
4.6.5. Gelir Durumu Algısı Dağılımı.....	86
4.6.6. Yapay Zeka (YZ) Uygulamaları Kullanımı.....	86
4.7. Demografik Özelliklerin Ayrıntılı Boyutları.....	87
5. TARTIŞMA VE SONUÇ.....	108
5.1. Sonuç	108
5.2. Bulguların Tartışılması.....	112
5.2.1. Ölçek Gelişim Sürecinin literatür ile tartışılması	112
5.2.2. Cinsiyet bulgularının tartışılması.....	115
5.2.3 Yaş bulgularının tartışılması.....	117
5.2.4 Eğitim düzeyi bulgularının tartışılması	120
5.2.5 Akademik alan bulgularının tartışılması.....	123
5.2.6 Gelir algısı bulgularının tartışılması	126
5.2.7 Yapay zekâ kullanımı bulgularının tartışılması.....	129
5.2.8 Genel değerlendirme.....	132
5.3 Araştırmanın Sınırlılıkları.....	135
5.4 Öneriler.....	137
5.4.1. Araştırmacılar İçin Öneriler.....	137
5.4.2. Eğitim Kurumları İçin Öneriler	139
5.4.3. Politika Yapıcılar İçin Öneriler	141
5.4.4. Yapay Zekâ Geliştiricileri İçin Öneriler.....	143
KAYNAKÇA	146

TABLÖLAR DİZİNİ

Tablo 1. Ölçek Maddelerinin Betimsel İstatistikleri	40
Tablo 2. KMO ve Bartlett küresellik testi sonuçları.....	42
Tablo 3. Faktör Analiz Sonuçları	42
Tablo 4. Yapay Zeka Güvenlik Farkındalığı Alt Boyutuna Ait Faktör Analizi	44
Tablo 5. KMO ve Bartlett küresellik testi sonuçları.....	46
Tablo 6. Açıklanan varyans ve özdeğerler	46
Tablo 7. Model Uyum Kriterleri	48
Tablo 8. Güvenlik Farkındalığı alt boyutu Güvenirlik Analizi	50
Tablo 9. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları	51
Tablo 10. Yapay Zeka Ahlaki Değerler Alt Boyutuna Ait Faktör Analizi.....	52
Tablo 11. KMO ve Bartlett küresellik testi sonuçları.....	53
Tablo 12. Açıklanan varyans ve özdeğerler	54
Tablo 13. Model Uyum Kriterleri.....	56
Tablo 14. Ahlaki Değerler alt boyutu Güvenirlik Analizi.....	58
Tablo 15. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları	59
Tablo 16. Yapay Zeka Sorumluluk ve Hesap Verilebilirlik Alt Boyutuna Ait Faktör Analizi	60
Tablo 17. KMO ve Bartlett küresellik testi sonuçları.....	62
Tablo 18. Açıklanan varyans ve özdeğerler	62
Tablo 19. Model Uyum Kriterleri.....	65
Tablo 20. Sorumluluk ve Hesap Verilebilirlik alt boyutu Güvenirlik Analizi	67
Tablo 21. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları	68
Tablo 22. Yapay Zeka Toplumsal Etki Alt Boyutuna Ait Faktör Analizi.....	69
Tablo 23. KMO ve Bartlett küresellik testi sonuçları.....	70
Tablo 24. Açıklanan varyans ve özdeğerler	71
Tablo 25. Model Uyum Kriterleri.....	73
Tablo 26. Toplumsal Etki Alt boyutu Güvenirlik Analizi.....	75
Tablo 27. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları	76
Tablo 28. Yapay Zeka Olumlu Beklentiler Alt Boyutuna Ait Faktör Analizi	77

Tablo 29. KMO ve Bartlett küresellik testi sonuçları.....	78
Tablo 30. Açıklanan varyans ve özdeğerler	79
Tablo 31. Model Uyum Kriterleri.....	81
Tablo 32. Olumlu Beklentiler Alt boyutu Güvenirlilik Analizi.....	83
Tablo 33. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları	84
Tablo 34. Demografik Özellikler	85
Tablo 35. Katılımcıların cinsiyete göre YZGAAÖ toplam puanları için bağımsız örneklem sonuçları	87
Tablo 36. Katılımcıların cinsiyete göre YZGAAÖ alt boyut puanları için bağımsız örneklem sonuçları	87
Tablo 37. Yaş Gruplarına göre katılımcıların YZGAAÖ puanlarının tanımlayıcı istatistikleri	89
Tablo 38. Yaş Gruplarına göre YZGAAÖ puanları için tek yönlü ANOVA sonuçları	90
Tablo 39. Yaş Gruplarına göre katılımcıların YZGAAÖ alt boyut puanlarının tanımlayıcı istatistikleri	91
Tablo 40. Yaş Gruplarına göre YZGAAÖ alt boyut puanları için tek yönlü ANOVA sonuçları	93
Tablo 41. Eğitim Durumlarına göre katılımcıların YZGAAÖ puanlarının tanımlayıcı istatistikleri	94
Tablo 42. Eğitim Durumlarına göre YZGAAÖ puanları için tek yönlü ANOVA sonuçları ...	95
Tablo 43. Eğitim Durumlarına göre katılımcıların YZGAAÖ alt boyut puanlarının tanımlayıcı istatistikleri	95
Tablo 44. Eğitim Durumlarına göre YZGAAÖ alt boyut puanları için tek yönlü ANOVA sonuçları	97
Tablo 45. Katılımcıların Akademik alanlarına göre YZGAAÖ toplam puanları için bağımsız örneklem sonuçları	99
Tablo 46. Katılımcıların Akademik alanlarına göre YZGAAÖ alt boyut puanları için bağımsız örneklem sonuçları	100
Tablo 47. Gelir durum algılarına göre katılımcıların YZGAAÖ puanlarının tanımlayıcı istatistikleri	101
Tablo 48. Gelir durum algılarına göre YZGAAÖ puanları için tek yönlü ANOVA sonuçları	102

Tablo 49. Gelir durum algısına göre katılımcıların YZGAAÖ alt boyut puanlarının tanımlayıcı istatistikleri	103
Tablo 50. Gelir durum algısı göre YZGAAÖ alt boyut puanları için tek yönlü ANOVA sonuçları	105
Tablo 51. Katılımcıların YZ uygulamalarını kullanımına göre YZGAAÖ alt boyut puanları için bağımsız örneklem sonuçları	107

ŞEKİLLER DİZİNİ

Şekil 1. Maddelerin Yamaç Birikinti (Scree Plot) Grafiği	43
Şekil 2. Güvenlik Farkındalığı Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği	47
Şekil 3. Güvenlik Farkındalığı alt boyutu Path Diyagramı	49
Şekil 4. Ahlaki Değerler Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği.....	55
Şekil 5. Ahlaki Değerler alt boyutu Path Diyagramı.....	57
Şekil 6. Sorumluluk ve Hesap Verilebilirlik Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği.....	64
Şekil 7. Sorumluluk ve Hesap Verilebilirlik alt boyutu Path Diyagramı	66
Şekil 8. Toplumsal Etki Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği.....	72
Şekil 9. Toplumsal Etki Alt Boyutu Path Diyagramı	74
Şekil 10. Olumlu Beklentiler Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği.....	80
Şekil 11. Olumlu Beklentiler Alt Boyutu Path Diyagramı	82

SİMGELER VE KISALTMALAR DİZİNİ

Kısaltma / Simge	Açıklaması
AFA	Açımlayıcı Faktör Analizi
AGFI	Düzeltilmiş Uyum İyiği İndeksi (Adjusted Goodness of Fit Index)
AI	Artificial Intelligence (Yapay Zekâ)
AMOS	Analysis of Moment Structures
AVE	Ortalama Açıklanan Varyans (Average Variance Extracted)
CFA	Confirmatory Factor Analysis (Doğrulayıcı Faktör Analizi)
CFI	Karşılaştırmalı Uyum İndeksi (Comparative Fit Index)
CR	Bileşik Güvenirlik (Composite Reliability)
Cronbach's α	Cronbach Alfa İç Tutarlılık Katsayısı
DFA	Doğrulayıcı Faktör Analizi
df (sd)	Serbestlik Derecesi
EFA	Exploratory Factor Analysis (Açımlayıcı Faktör Analizi)
GFI	Uyum İyiği İndeksi (Goodness of Fit Index)
IFI	Artırımlı Uyum İndeksi (Incremental Fit Index)
KMO	Kaiser-Meyer-Olkin Örneklem Yeterlilik Katsayısı
N	Katılımcı Sayısı
NFI	Normlaştırılmış Uyum İndeksi (Normed Fit Index)
NNFI (TLI)	Normlaştırılmamış Uyum İndeksi (Non-Normed Fit Index / Tucker–Lewis Index)
p	Anlamlılık Düzeyi
PGFI	Parsimoni Uyum İyiği İndeksi (Parsimony Goodness of Fit Index)
PNFI	Parsimoni Normlaştırılmış Uyum İndeksi (Parsimony Normed Fit Index)
r	Korelasyon Katsayısı
RFI	Görelî Uyum İndeksi (Relative Fit Index)

Kısaltma / Simge	Açıklaması
RMSEA	Yaklaşık Hataların Ortalama Karekökü (Root Mean Square Error of Approximation)
RMR	Artık Ortalamaların Karekökü (Root Mean Square Residual)
SD / Ss	Standart Sapma
SRMR	Standardize Edilmiş Artık Ortalamaların Karekökü (Standardized Root Mean Square Residual)
t	t Testi İstatistiği
$\bar{X}$	Aritmetik Ortalama
χ^2	Ki-kare İstatistiği
χ^2/sd	Ki-kare / Serbestlik Derecesi Oranı
λ	Özdeğer (Eigenvalue)
α	Anlamlılık Düzeyi / Cronbach Alfa Katsayısı (kullanıldığı bağlama göre)
%	Yüzde
YZ	Yapay Zekâ
YZGAAÖ	Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği

1. GİRİŞ

Yapay zekâ, insan zekâsıyla ilişkilendirilen öğrenme, akıl yürütme, problem çözme, algılama ve karar verme gibi işlevlerin makine ve bilgisayar sistemleri tarafından benzetim yoluyla gerçekleştirilmesini amaçlayan bir bilgisayar bilimi alanı olarak tanımlanabilir (Deng, 2018; Morandín-Ahuerma, 2022; Shrivastava, 2024). Güncel literatürde bu kavram, yalnızca tek bir teknikten ibaret görülmemekte; makine öğrenmesi, derin öğrenme, doğal dil işleme, konuşma tanıma, bilgisayar görüşü ve benzeri yöntemleri kapsayan şemsiye bir yapı olarak ele alınmaktadır (Bozdoğanoglu et al., 2024; Shrivastava, 2024). Yapay zekâ en genel anlamıyla, makinelerin normalde insan zekâsı gerektiren görevleri yerine getirebilme kapasitesini ifade eder; bu görevler arasında mantıksal akıl yürütme, öğrenme, problem çözme ve çevreye uyum sağlama yer alır (Morandín-Ahuerma, 2022). Bu nedenle yapay zekâ, yalnızca önceden tanımlanmış komutları uygulayan klasik yazılımlardan ayrılır; çünkü çağdaş yaklaşımlarda sistemler çoğu zaman veriden örüntü öğrenerek kendi çıktılarını geliştirebilir (Deng, 2018; Shrivastava, 2024).

Yapay zekânın literatürdeki tanımları tam olarak tek biçimli değildir; insan ve yapay zekâ için geniş kabul gören tek bir tanımın oluşturulmasının güç olduğu da vurgulanmaktadır (Gignac & Szodorai, 2024). Buna karşın farklı disiplinlerde ortaklaşan çekirdek fikir, zekânın yeni hedefleri başarıyla tamamlama, öğrenme, akıl yürütme ve uygun çıktılar üretme kapasitesi etrafında toplandığını göstermektedir (Gignac & Szodorai, 2024). Yapay zekâ, insan zekâsına özgü kabul edilen öğrenme, akıl yürütme, problem çözme, algılama ve karar verme gibi bilişsel süreçlerin bilgisayar sistemleri aracılığıyla modellenmesini ve belirli ölçülerde taklit edilmesini amaçlayan disiplinlerarası bir bilgisayar bilimi alanıdır (Deng, 2018; Shrivastava, 2024). Günümüzde yapay zekâ, kuralla dayalı programlamanın ötesine geçerek büyük veri kümeleri üzerinde çalışan algoritmalar sayesinde örüntü tanıma, tahmin üretme, dil işleme ve görsel algılama gibi işlevleri yerine getirebilen sistemleri kapsamaktadır (Deng, 2018; Shrivastava, 2024). Bu yönüyle yapay zekâ, makine öğrenmesi ve derin öğrenme başta olmak üzere çeşitli alt alanları bünyesinde toplayan bir şemsiye kavram niteliği taşımaktadır (Shrivastava, 2024; Bozdoğanoglu et al., 2024). Bununla birlikte literatür, günümüzde etkili uygulamaların büyük ölçüde dar yapay zekâ düzeyinde kaldığını; insan benzeri genel zekâyâ ilişkin tartışmaların ise kavramsal ve ölçümsel açıdan hâlen açık olduğunu göstermektedir (Fjelland, 2020; Gignac & Szodorai, 2024).

Yapay zekânın tarihsel gelişimi, 1940'lı yıllardaki kuramsal başlangıçlardan sembolik yapay zekâyâ, oradan makine öğrenmesi, derin öğrenme ve büyük dil modellerine uzanan çok evreli bir dönüşüm süreci olarak düzenlenebilir. Yapay zekâ araştırmaları, başlangıçtan günümüze birkaç belirgin kırılma üzerinden ilerlemiştir. Zaman çizgisi, alanın kural temelli yaklaşımlardan veri güdümlü ve ardından üretken modellere kaydığını göstermektedir. Yapay zekânın erken kökleri, 20. yüzyıl ortasında sinir ağları, sibernetik ve makine zekâsı tartışmalarıyla ilişkilendirilmekte; 1943 tarihli McCulloch ve Pitts çalışması makine öğrenmesinin ilk sinir ağı adımlarından biri olarak anılmaktadır (Zaidi et al., 2018). Kurumsal eşik ise, 1956 Dartmouth Konferansı'dır; bu toplantı yapay zekâyı bağımsız bir araştırma alanı olarak kurumsallaştırmış ve John McCarthy, Marvin Minsky, Nathaniel Rochester ile Claude Shannon'ın öncülüğünde alanın araştırma gündemini biçimlendirmiştir (Ambani & Rathod, 2025; Elamin, 2024). Bu erken dönemde araştırmacılar, insan bilişini modelleme hedefiyle makinelerin öğrenme, akıl yürütme ve problem çözme kapasitesini sistematik olarak araştırmaya yönelmiştir (Jiang et al., 2022; Elamin, 2024).

1950'li ve 1960'lı yıllarda baskın yaklaşım sembolik yapay zekâ olmuş; alan, mantıksal çıkarım, kural tabanlı temsil ve uzman sistemler etrafında gelişmiştir (Mundlamuri et al., 2025; Elamin, 2024; Bansal & Sangwan, 2024). Sembolik sistemlerin temel mantığı, nesneler ve ilişkiler hakkında dil-benzeri temsiller üzerinde mantıksal işlemler yürütmektir (Garnelo & Shanahan, 2019). Newell ve Simon'ın Logic Theorist programı bu dönemin erken başarı örneklerinden biri olarak anılmış, kural tabanlı sembolik manipülasyonla problem çözmenin mümkün olduğunu göstermiştir (Elamin, 2024). Ancak bu yaklaşım zamanla katılık, ölçeklenebilirlik eksikliği ve bağlama duyarlılık sorunları nedeniyle sınırlarına ulaşmıştır (Liang et al., 2025; Elamin, 2024). 1960'lardaki yüksek beklentiler, 1970'ler ve 1980'lerde yaşanan YZ kışları ile yerini durgunluk dönemlerine bırakmıştır (Elamin, 2024; Bansal & Sangwan, 2024). Bu gerilemede yavaş ilerleme, sembolik sistemlerin belirsizlik ve gerçek dünya karmaşıklığını işlemedeki yetersizliği ile ölçeklenme sorunları etkili olmuştur (Elamin, 2024; Bansal & Sangwan, 2024; Liang et al., 2025). 1980'lerde uzman sistemler kısa süreli bir canlanma yaratsa da, alanın kalıcı dönüşümü 1990'lardan itibaren veri odaklı makine öğrenmesi ile gerçekleşmiştir (Blagoj et al., 2020; Elamin, 2024; Bansal & Sangwan, 2024).

1990'lardan sonra destek vektör makineleri, karar ağaçları ve benzeri klasik algoritmaların yükselişi, yapay zekâyı el ile yazılmış kurallardan istatistiksel öğrenmeye kaydırmıştır (Zaidi et al., 2018; Bansal & Sangwan, 2024). 2010'lu yıllarda bu dönüşüm, artan hesaplama gücü,

dijital veri bolluğu ve çok katmanlı sinir ağı mimarileri sayesinde derin öğrenme devrimine dönüşmüştür (Lappin, 2025). CNN'ler görüntü sınıflandırmada, RNN ve LSTM türleri sıralı verilerde, derin pekiştirmeli öğrenme ise karar verme problemlerinde dönüştürücü sonuçlar üretmiştir (Mundlamuri et al., 2025; Zaidi et al., 2018). Güncel evrede transformatör mimarisi, büyük dil modellerinin temelini oluşturmaktadır; öz-dikkat temelli bu yaklaşım, doğal dil anlama ve üretmede yeni bir aşama başlatmıştır (Mundlamuri et al., 2025; Lappin, 2025; Bansal & Sangwan, 2024). Bununla birlikte alan yalnızca ölçek büyümeye yönelmemekte; yorumlanabilirlik, önyargı, hesaplama maliyeti ve nöro-sembolik bütünleşme de başlıca araştırma eksenleri olarak öne çıkmaktadır (Mundlamuri et al., 2025; Liang et al., 2025).

Yapay zekânın uygulama alanları incelendiğinde, sağlık sektöründe derin öğrenmenin tıbbi görüntüleme, biyoinformatik, giyilebilir sensör verileri ve elektronik sağlık kayıtları üzerinden tanı doğruluğunu artırma, erken teşhis ve kişiselleştirilmiş tedavi süreçlerine katkı sağladığı görülmektedir (Ravi et al., 2016; Abdel et al., 2022; Bote-Curiel et al., 2019; Shakor & Khaleel, 2024; Kumari et al., 2023). Finans alanında algoritmik alım-satım, portföy yönetimi, sahtecilik tespiti ve müşteri etkileşimlerinin analizinde veri odaklı çözümler sunulmaktadır (Najafabadi et al., 2015; Swapna et al., 2024; Xie & Zhang, 2024; Zaripova et al., 2023). Eğitimde ise kişiselleştirilmiş öğrenme yolları, akıllı öğretim sistemleri, otomatik değerlendirme ve terk riski tahmini gibi uygulamalar yaygınlaşmaktadır (Bote-Curiel et al., 2019; Swapna et al., 2024; Zaripova et al., 2023). Güvenlik ve Nesnelerin İnterneti bağlamında derin öğrenme, siber saldırı tespiti ve güvenlik izleme süreçlerinde kullanılmakta ve büyük veri teknolojileriyle ölçeklenmektedir (Najafabadi et al., 2015; Amanullah et al., 2020). Kamu hizmetlerinde ise YZ, sağlık, ulaşım, enerji ve çevre yönetimi gibi alanlarda karar destek ve kaynak optimizasyonu amacıyla kullanılmaktadır (Leslie, 2019; Challen et al., 2019).

YZ yönetişimi kapsamında Avrupa politikaları, OECD ve G20 ilkeleri ile çeşitli kurumsal çerçeveler; sorumluluk, hesap verebilirlik, şeffaflık, güven, adalet, mahremiyet ve güvenlik ilkelerine odaklanmaktadır (Leslie, 2019; Camilleri, 2023; Vesnić-Alujević et al., 2020; Dignum, 2018; Huang et al., 2023; Zhang et al., 2021). YZ bağlamındaki çalışmalar, TAM ve genişletilmiş versiyonlarının kullanıcı kabulünü açıklamada yüksek yordama gücüne sahip olduğunu göstermekte; özellikle verimlilik artışı, performans iyileşmesi ve zaman tasarrufu gibi algılanan faydaların YZ'ye yönelik tutumları ve kullanımını artırdığını ortaya koymaktadır (Kelly et al., 2022; Ibrahim et al., 2025; Ma et al., 2025; Schepman & Rodway, 2020; Yi & Choi, 2023; Venkatesh & Davis, 2000). Yapay zekâya yönelik algılar, yalnızca olumlu

beklentilerle sınırlı olmayıp; risk, etik sorunlar ve toplumsal etkiler gibi çok boyutlu değerlendirmeleri de içermekte, bu kapsamda bireysel ve toplumsal düzeyde hem teknolojik ilerleme hem de olası olumsuzluklara yönelik kaygıları yansıtmaktadır.

Erik Brynjolfsson ve Andrew McAfee (2017), yapay zekâ ve otomasyonun iş gücü piyasaları üzerindeki etkilerine dikkat çekerek bazı mesleklerin ortadan kalkabileceğini ve gelir eşitsizliklerinin artabileceğini belirtmekte; bu durumun, teknolojilerin yalnızca verimlilik ve yenilik değil, aynı zamanda istihdam, çevresel etkiler ve toplumsal adalet boyutlarında da değerlendirilmesi gerektiğini göstermektedir (Brynjolfsson & McAfee, 2017). Eğitim araştırmaları, bireylerin yapay zekâyâ ilişkin algılarının bilgi düzeyi, deneyim ve dijital okuryazarlıkla ilişkili olduğunu göstermektedir (Long & Magerko, 2020; Ayanwale et al., 2024; Guan et al., 2024; Hornberger et al., 2023; Uzun, 2025). Öğretmen adaylarının yapay zekâ okuryazarlığının; anlama, kullanma, etik boyutları kavrama ve problem çözme becerilerini etkilediği belirtilmektedir (Ayanwale et al., 2024; Özudogru & Durak, 2025; Ng et al., 2023; Lim, 2023; Uzun, 2025).

Çalışmalar, öğretmen adayları ve üniversite öğrencilerinde yapay zekâ farkındalığının arttığını; ancak etik ve güvenlik konularında yeterli bilinç oluşmadığını ortaya koymaktadır (Ayanwale et al., 2024; Guan et al., 2024; Hornberger et al., 2023; Basch et al., 2025; Valerio, 2024; Usher & Barak, 2024; Uzun, 2025). Öğretmen adaylarının yapay zekâyı çoğunlukla araçsal bir teknoloji olarak gördüğü ve etik ilkeler, veri gizliliği, yanlılık ve toplumsal etkiler konusunda daha fazla eğitime ihtiyaç duyduğu ifade edilmektedir (Ayanwale et al., 2024; Guan et al., 2024; Valerio, 2024; Usher & Barak, 2024; Uzun, 2025).

Yapay zekânın karar verme süreçlerine entegrasyonu, giderek daha fazla, insan bilişini yalnızca destekleyen bir araç olarak değil, kurumsal karar alma kapasitesini yeniden tanımlayan sosyo-teknik bir dönüşüm olarak kavramsallaştırılmaktadır (Upase, 2026). Bu dönüşüm, veriye dayalı tahminsel analitiklerin karar vericilere seçenek sunmasının ötesine geçerek, belirli bağlamlarda otonom ya da yarı otonom karar yetkisinin algoritmik sistemlere devredilmesini de kapsamaktadır (Mohammed, 2025; Kulaklıoğlu, 2024). Bu çerçevede YZ, karar alma süreçlerinde yalnızca verimlilik artıran bir teknoloji değil, aynı zamanda kararın nasıl üretildiğine ilişkin epistemolojik varsayımları dönüştüren bir paradigma değişimi olarak değerlendirilmektedir (Lindner et al., 2025; Upase, 2026).

YZ destekli karar sistemlerinin temel vaatlerinden biri, insan karar vericilerin sınırlı rasyonellik, sezgisel kestirme yollar ve bilişsel önyargılar nedeniyle karşılaştığı kısıtları aşabilmesidir (Lindner et al., 2025; Upase, 2026). Özellikle büyük veri işleme, örüntü tanıma ve alternatif üretme kapasitesi sayesinde YZ, örgütsel karar kalitesini artıran bir artırım mekanizması olarak konumlandırılmaktadır (Shick et al., 2023; Upase, 2026). Bununla birlikte, literatür YZ'nin insan muhakemesinin yerine tam ikame oluşturduğunu değil, çoğu durumda hibrit insan-YZ modellerinin daha uygun olduğunu vurgulamaktadır (Upase, 2026; Vincent, 2021).

Son yıllarda makine öğrenmesi, derin öğrenme ve büyük veri analitiği gibi Yapay Zekâ (YZ) teknolojileri; sağlık, eğitim, finans, güvenlik ve kamu hizmetleri dâhil birçok alanda karar alma süreçlerini dönüştürmektedir. Araştırmalar, bu dönüşümün verimlilik ve doğruluk artışı sağlarken aynı zamanda etik, güvenlik ve yönetim sorunlarını da beraberinde getirdiğini göstermektedir (Ravi et al., 2016; Abdel et al., 2022; Bote-Curiel et al., 2019; Shakor & Khaleel, 2024; Kumari et al., 2023; Najafabadi et al., 2015; Swapna et al., 2024; Xie & Zhang, 2024; Zaripova et al., 2023; Amanullah et al., 2020; Leslie, 2019; Challen et al., 2019). Benzer şekilde, üniversite öğrencilerinde farkındalık yüksek olsa da veri gizliliği, güvenlik açıkları ve olası olumsuz sonuçlara yönelik kaygıların sürdüğü belirtilmektedir (Hornberger et al., 2023; Basch et al., 2025; Valerio, 2024; Usher & Barak, 2024).

Yapay zekanın (YZ) karar verme süreçlerine entegrasyonu; sağlık, finans, denetim ve kamusal alanlarda önemli etik tartışmaları beraberinde getirmektedir. Çalışmalar, algoritmik önyargı, veri gizliliği, şeffaflık eksikliği ve hesap verebilirlik sorunlarının hem bireysel haklar hem de toplumsal kabul üzerinde doğrudan etkili olduğunu göstermektedir. Bu bağlamda, sistemlerin yalnızca teknik olarak doğru çalışması değil, aynı zamanda adil, şeffaf ve hesap verebilir olması gerekliliği vurgulanmaktadır. Literatürde bu sorunlara yönelik çözüm çerçeveleri; adalet, açıklanabilirlik, hesap verebilirlik, mahremiyet, güvenlik ve insan denetimi gibi ilkeleri (Rodríguez et al., 2023; Kaur et al., 2022; Alzubaidi et al., 2023); önyargı azaltma, açıklanabilir yapay zekâ, algoritmik denetim, temsili veri ve gizlilik koruma tekniklerini (Radanliev, 2025; Murikah et al., 2024; Rodríguez et al., 2023; Bahangulu & Owusu-Berko, 2025; Mensah, 2024; Aunugu & Vathsavai, 2022; Alzubaidi et al., 2023); ayrıca etik kılavuzlar, risk temelli düzenlemeler, bağımsız denetim ve çok paydaşlı yönetim yaklaşımlarını içermektedir (Cheong, 2024; Murikah et al., 2024; Osasona et al., 2024; Rodríguez et al., 2023; Kaur et al., 2022; Pham, 2025; Alzubaidi et al., 2023).

Algoritmik önyargı, temsili olmayan veya eksik veri setleri nedeniyle ayrımcı sonuçlara yol açabilmektedir (Weiner et al., 2024; Zhang & Zhang, 2023; Murikah et al., 2024; Bahangulu & Owusu-Berko, 2025; Alzubaidi et al., 2023; Daneshjou et al., 2021). Bu önyargının kaynakları arasında veri eksiklikleri, demografik homojenlik, sahte korelasyonlar, uygun olmayan karşılaştırma kriterleri ve bilişsel önyargılar yer almaktadır. Denetim alanında yapılan sistematik bir inceleme, bu beş temel kaynağın teknik ve insani önyargıları tetiklediğini ortaya koymuştur. Ayrıca verimlilik ile titizlik arasındaki ödünleşmeler, insan becerilerinin ve yargı yeteneğinin aşınması, veri bağımlılığı riskleri ve kontrolsüz kişisel veri sömürsünden kaynaklanan gizlilik ihlalleri gibi daha geniş sorunlar da literatürde belirtilmektedir. Bu sorunlara karşı nedensel modelleme, temsili algoritmik testler, periyodik sistem denetimleri, otomasyonun yanı sıra insan gözetimi ve adalet ile hesap verebilirlik gibi etik değerlerin sistem tasarımına entegrasyonu umut vadeden çözümler olarak öne çıkmaktadır.

Veri gizliliği ve mahremiyet sorunları, özellikle sağlık verileri ve giyilebilir teknolojiler bağlamında ihlal, gözetim ve kimlik hırsızlığı risklerini artırmaktadır (Weiner et al., 2024; Zhang & Zhang, 2023; Radanliev, 2025; Rodríguez et al., 2023; Aunugu & Vathsavai, 2022). YZ sistemlerinin eğitiminde kullanılan devasa veri kümeleri, kişisel bilgileri, parolaları ve telif hakkıyla korunan içerikleri istemsiz olarak ezberleyebilmekte; bu durum model çıkarımı (model inversion) ve üyelik çıkarımı (membership inference) gibi saldırılara karşı savunmasız hale getirmektedir. Sınıf niteliği çıkarım saldırıları, özellikle yüz tanıma modellerinde saç rengi, cinsiyet ve ırksal görünüm gibi açıklanmamış nitelikleri doğru bir şekilde çıkarabilmekte; ilginç bir şekilde, adversaryal olarak sağlamaştırılmış modellerin bu tür gizlilik sızıntılarına karşı daha savunmasız olduğu gösterilmiştir.

Şeffaflık eksikliği ("kara kutu" problemi), sistemlerin karar süreçlerinin anlaşılmasını zorlaştırarak güven sorunlarına neden olmaktadır (Zhang & Zhang, 2023; Lepri et al., 2021; Nouis et al., 2025; Osasona et al., 2024; Kaur et al., 2022; Alzubaidi et al., 2023). Özellikle derin öğrenme modellerinin karar verme mekanizmalarının yorumlanabilirliği konusundaki zorluklar hem kullanıcıların hem de düzenleyici otoritelerin sistemlere olan güvenini sarsmaktadır. Bu durum, özellikle yüksek riskli uygulama alanlarında örneğin sağlık hizmetlerinde teşhis kararları veya finansal kredi değerlendirmelerinde ciddi toplumsal ve hukuki sonuçlar doğurabilmektedir. Hesap verebilirlik boşluğu, hatalı kararlar karşısında sorumluluğun belirsizleşmesine yol açmaktadır (Zhang & Zhang, 2023; Lepri et al., 2021; Nouis et al., 2025; Akinrinola et al., 2024; Osasona et al., 2024; Kaur et al., 2022; Pham, 2025).

YZ sistemlerinin kararları sonucu ortaya çıkan zararların sorumluluğunun geliştiriciye, veri sağlayıcıya, sistem operatörüne mi yoksa son kullanıcıya mı ait olduğu konusundaki belirsizlik, hukuki çerçevelerin ve düzenleyici mekanizmaların yeniden gözden geçirilmesini gerektirmektedir. Bu bağlamda, risk değerlendirmelerinin dağıtımdan önce gerçekleştirilmesi, sürekli performans izleme ve güven ile işbirliğini teşvik eden politikalar önerilmektedir.

Uygulama alanlarına göre bu riskler farklılaşmaktadır. Sağlık alanında önyargı, veri kalitesi ve sorumluluk belirsizliği öne çıkmaktadır (Weiner et al., 2024; Zhang & Zhang, 2023; Lepri et al., 2021; Nouis et al., 2025; Pham, 2025; Daneshjou et al., 2021). Finans ve iş dünyasında ayrımcı kararlar ve şeffaflık sorunları yaşanmaktadır (Lepri et al., 2021; Murikah et al., 2024; Bahangulu & Owusu-Berko, 2025; Mensah, 2024; Alzubaidi et al., 2023). Denetimde birikimli önyargı etkileri gözlemlenmektedir (Murikah et al., 2024; Alzubaidi et al., 2023). Giyilebilir/IoT sistemlerinde ise gizlilik, güvenlik ve özerklik kaybı öne çıkmaktadır (Radanliev, 2025; Rodríguez et al., 2023; Aunugu & Vathsavai, 2022).

YZ entegrasyonunun getirdiği etik tartışmaların yanı sıra, sistemlerin güvenliğine yönelik tehditler de giderek artmaktadır. 2024 yılı sonlarına kadar kurumların %41'inin bir tür YZ güvenlik olayı yaşadığı rapor edilmiştir. Bu olaylar veri zehirlenmesinden model hırsızlığına kadar geniş bir yelpazede yer almaktadır.

İnsan-merkezli yaklaşım kapsamında ise özerkliğin korunması, kırılgan grupların gözetilmesi ve kamu katılımı vurgulanmaktadır (Zhang & Zhang, 2023; Lepri et al., 2021; Murikah et al., 2024; Rodríguez et al., 2023; Pham, 2025; Aunugu & Vathsavai, 2022). Bu çerçevede, YZ sistemlerinin tasarım ve dağıtım süreçlerinde bireylerin temel hak ve özgürlüklerinin korunması, özellikle dezavantajlı ve marjinalleştirilmiş toplulukların sistematik ayrımcılığa maruz kalmaması için önlemler alınması esastır. Araştırmacılar, YZ güvenliği çalışmalarının önceliklendirilmesi, potansiyel zararların değerlendirilmesine yönelik ön inceleme süreçlerinin geliştirilmesi ve uluslararası kuruluşların daha etkin rol üstlenmesi gerektiğini belirtmektedir (Zhang et al., 2021). Ayrıca kamuoyu ve akademik çalışmalar, yapay zekâya yönelik yüksek fayda beklentisinin yanı sıra özellikle otonom sistemler ve teknolojik tekillik bağlamında etik ve güvenlik kaygılarının da varlığını ortaya koymaktadır (Leslie, 2019; Suo et al., 2024).

Bu alandaki çalışmalar; adalet, şeffaflık, hesap verebilirlik, gizlilik ve insan denetimi gibi temel etik ilkeler etrafında yoğunlaşmakta ve güvenilir, toplumsal açıdan sorumlu YZ uygulamaları için çerçeve sunmaktadır (Jobin et al., 2019; Fjeld et al., 2020; Shukla, 2024; Huang et al.,

2023; Kazim & Koshiyama, 2020). Bununla birlikte, YZ teknolojileri çeşitli etik, güvenlik ve yönetim tartışmalarını da gündeme getirmektedir. Literatürde mahremiyet ihlalleri, algoritmik önyargı ve ayrımcılık, işsizlik, güvenlik açıkları, otonom sistemlerin kontrolü ve şeffaf olmayan karar verme süreçleri temel riskler arasında gösterilmektedir (Leslie, 2019; Čartolovni et al., 2022; Camilleri, 2023; Vesnić-Alujević et al., 2020; Dignum, 2018; Huang et al., 2023; Suo et al., 2024). Sağlık alanında bu tartışmalar; hasta güvenliği, algoritmik şeffaflık, düzenleme eksikliği, sorumluluk ve hasta-hekim ilişkisindeki dönüşüm çerçevesinde ele alınmaktadır (Ravi et al., 2016; Abdel et al., 2022; Shakor & Khaleel, 2024; Kumari et al., 2023; Challen et al., 2019; Čartolovni et al., 2022). Generatif YZ bağlamında ise zararlı içerik üretimi, yanlış bilgi (halüsinasyon), mahremiyet sorunları, toplumsal etkiler ve hizalama problemleri öne çıkmaktadır (Hagendorff, 2024).

Yapay zekâ etiği, YZ sistemlerinin tasarım, geliştirme ve kullanım süreçlerine yön veren ilke ve değerleri inceleyen bağımsız bir disiplin olarak tanımlanmaktadır (Huang et al., 2023; Kazim & Koshiyama, 2020). Küresel etik rehberler üzerine yapılan çalışmalar, şeffaflık, adalet ve hakkaniyet, zarar vermeme, sorumluluk/hesap verebilirlik ve gizlilik olmak üzere ortak etik eksenler bulunduğunu göstermektedir (Jobin et al., 2019; Shukla, 2024; Huang et al., 2023; Kazim & Koshiyama, 2020). Bu ilkeler, YZ uygulamalarının bireysel hakları gözetken ve toplumsal eşitsizlikleri derinleştirmeyecek şekilde tasarlanması gerekliliğine işaret etmektedir (Jobin et al., 2019; Shukla, 2024; Huang et al., 2023). Uluslararası politika ve rehberler de bu ilkelere dayanmaktadır. OECD ilkeleri; insan hakları, hesap verebilirlik, şeffaflık, güvenlik, adalet ve insan merkezli değerleri vurgulamaktadır (Fjeld et al., 2020; Camilleri, 2023). UNESCO tarafından kabul edilen "Yapay Zekâ Etiği Tavsiye Kararı" ise etik YZ'nin insan haklarına dayalı, kapsayıcı ve sürdürülebilir biçimde geliştirilmesini öne çıkarmaktadır (Huang et al., 2023). Yapay zekâ etiği, normatif ilkelerin belirlenmesi ile bu ilkelerin kurumsal ve teknik uygulamalara yansıtılmasını kapsayan çok disiplinli bir alan olarak ele alınmaktadır (Rodriguez et al., 2023; Cheong, 2024; Jedlickova, 2024; Winfield & Jirotko, 2018; Kamila & Jasrotia, 2023; Mittelstadt, 2019). Yapay zekâ çözümleri, teknik güvenliğin ötesinde sosyal ve etik boyutları da içeren çok boyutlu bir yapı olarak ele alınmakta; güvenlik kapsamında sistem hataları, tutarsızlıklar, veri kesintileri ve insan kontrolünün zayıflaması gibi unsurlar yer almaktadır (Göktepe, 2019; Yilmaz, 2020).

Yapay zekâ teknolojilerinin toplumsal yaşama hızla nüfuz etmesi, bu sistemlerin güvenlik ve ahlak boyutlarına ilişkin bireysel algıların sistematik biçimde incelenmesini gerekli

kılmaktadır. Üniversite öğrencileri, dijital teknolojilerle yoğun etkileşimleri ve geleceğin uzmanları ile karar vericileri olmaları nedeniyle, yapay zekâya ilişkin algı araştırmaları açısından kritik bir örneklem grubu olarak öne çıkmaktadır (Morales-García et al., 2024; Yılmaz & Korkmaz, 2025). Literatür, öğrencilerin yapay zekâya ilişkin etik değerlendirmelerinde özellikle adalet, mahremiyet, zarar vermeme, şeffaflık, sorumluluk ve insan denetimi gibi boyutların öne çıktığını göstermektedir. Üniversite öğrencilerine yönelik geliştirilen araçlar etik tutum, etik yansıtma, etik farkındalık ve akademik dürüstlük gibi alanlarda geçerli sonuçlar üretmiştir; ancak bu araçlar güvenlik kaygılarını ve ahlaki değerlendirmeleri tek bir bütünlük yapıda toplamamaktadır (Jang et al., 2022; Wang et al., 2025; Abuadas & Albikawi, 2025).

Etik ve ahlak kavramları gündelik dilde sıklıkla eş anlamlı kullanılsa da, akademik kaynaklarda bunların kavramsal olarak ayrıştırıldığı görülür (Nathanson, 2015; Wolf & Knoepffler, 2022; Villiers & Villiers, 2023). Ahlak, bir bireyin ya da topluluğun benimsediği doğru-yanlış ve iyi-kötü ölçütlerini içeren değerler, normlar ve kurallar bütünü olarak tanımlanır; bu kurallar davranışı yönlendirir ve toplumsal yaşam içinde kabul görür (Horner, 2003; Nathanson, 2015; Wolf & Knoepffler, 2022; Chairul et al., 2020; Ellemers, 2017). Etik ise ahlakın kendisinden ziyade, ahlaki yargılar, ilkeler ve davranışlar üzerine düşünmeyi konu alan felsefi bir disiplin olarak tanımlanır (Advincula et al., 2026; Horner, 2003; Wolf & Knoepffler, 2022). Bu nedenle etik, hangi davranışların doğru ya da iyi sayılması gerektiğini, hangi ilkelerin geçerli olduğunu ve bu yargıların nasıl temellendirileceğini sorgular (Author, 2026; Chaddha & Agrawal, 2023; Timmons, 2020; Samsul, 2026).

Bu yönüyle ahlak, daha çok belirli bir toplumun ya da grubun kabul ettiği doğru ve yanlış ölçütleri üzerinden “Ne yapmalıyım?” sorusuna pratik bir yanıt verir (Ellemers, 2017; Samsul, 2026). Ahlaki normların toplumdan topluma ve gruptan gruba değişebildiği de belirtilir (Ellemers, 2017; Chaddha & Agrawal, 2023; Haidt, 2007). Etik ise ahlaki değerler, ilkeler ve eylemler üzerine eleştirel ve kuramsal inceleme yürüten felsefi bir alan olarak tanımlanır (Wolf & Knoepffler, 2022; Advincula et al., 2026; Chaddha & Agrawal, 2023). Bu nedenle etik, “Neden böyle yapmalıyım?” sorusunu gerekçelendirme, değerlendirme ve temellendirme düzeyinde ele alır (Samsul, 2026; Timmons, 2020).

Yapay zekâya ilişkin güvenlik ve etik algıları bütüncül biçimde değerlendirebilecek geçerli ve güvenilir ölçme araçlarına ihtiyaç artmaktadır. Geliştirilen ölçekler, duyuşsal, bilişsel,

davranışsal ve etik boyutların ölçülebileceğini göstermekle birlikte, özellikle etik ve güvenlik boyutlarının daha derin ele alınması gerektiğine işaret etmektedir (Long & Magerko, 2020; Ng et al., 2023; Hornberger et al., 2023; Bilbao-Eraña & Arroyo-Sagasta, 2025; Usher & Barak, 2024; Gümüş et al., 2023).

1.1. Araştırmanın Amacı

Bu çalışma, geliştirilecek "Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği" ile bireylerin yapay zekâya yönelik algılarını; güvenlik farkındalığı, ahlakî değerler, sorumluluk ve hesap verebilirlik, toplumsal etki ve olumlu beklentiler olmak üzere beş boyut altında incelemeyi amaçlamaktadır. Bu boyutlar, uluslararası etik çerçeveler ve güncel ampirik çalışmalar temelinde yapılandırılmıştır (Vesnić-Alujević et al., 2020; Rodríguez et al., 2023; Zhang & Zhang, 2023; Floridi et al., 2018; Dignum, 2018).

Güvenlik farkındalığı boyutu, yapay zekâ sistemlerinin güvenlik, mahremiyet ve kötüye kullanım gibi risklerine ilişkin farkındalığı; ahlakî değerler boyutu, şeffaflık, adalet, mahremiyet ve insan denetimi gibi etik ilkelere yönelik algıları değerlendirmektedir (Gerlich, 2023; Liehner et al., 2023; Ding et al., 2026; Brauner et al., 2023; Saatci, 2025; Rodríguez et al., 2023; Hartwig et al., 2022; Wang et al., 2025). Sorumluluk ve hesap verebilirlik boyutu, yapay zekâ çıktılarından sorumlu aktörlere ilişkin algıları ölçerken, toplumsal etki boyutu eşitsizlik, istihdam, güvenlik ve refah gibi alanlardaki etkileri kapsamaktadır (Saatci, 2025; Rodríguez et al., 2023; Cheong, 2024; Zhang & Zhang, 2023; Dignum, 2018; Gerlich, 2023; Vesnić-Alujević et al., 2020; Novotny et al., 2025; Ding et al., 2026; Brauner et al., 2023; Bullock et al., 2025). Olumlu beklentiler boyutu ise verimlilik, konfor, maliyet düşüşü ve toplumsal yarara yönelik beklentileri ortaya koymaktadır (Gerlich, 2023; Liehner et al., 2023; Neri & Cozman, 2019; Ding et al., 2026; Schepman & Rodway, 2020).

1.2. Araştırmanın Önemi

Yapay zekâ (YZ) güvenlik algısının ölçülmesi, hem bireysel hem de kurumsal düzeyde kritik öneme sahiptir çünkü adversarial makine öğrenmesi saldırılarının evrimi, YZ sistemlerinin güvenliğine ilişkin algıları doğrudan etkilemekte ve bu algıların sistematik olarak ölçülmesi gerekliliğini ortaya koymaktadır. Biggio ve Roli (2018) tarafından 2004-2018 dönemini kapsayan kapsamlı derleme çalışması, adversarial makine öğrenmesi alanındaki saldırı ve savunma tekniklerinin tarihsel gelişimini ortaya koymuş; Papernot ve arkadaşları (2017) ise kara kutu saldırılarına karşı pratik yöntemler geliştirerek, YZ modellerinin güvenliğinin

yalnızca teknik önlemlerle değil, aynı zamanda kullanıcıların ve geliştiricilerin bu tehditlere yönelik algılarının anlaşılmasıyla sağlanabileceğini göstermiştir.

Son dönemde ise Qi ve arkadaşları (2024), hizalanmış dil modellerinin güvenlik eğitiminin ne ölçüde kırılğan olduğunu ortaya koyarak, YZ güvenliğinin sadece teknik dayanıklılıkla değil, kullanıcıların bu kırılğanlıklara ilişkin farkındalığı ve algısıyla da ilişkili olduğunu kanıtlamıştır; bu bulgu, güvenlik algısının ölçülmesinin sistem tasarımına ve politika geliştirmeye rehberlik edecek bir metodolojik araç olarak niteliğini vurgulamaktadır. Ayrıca, veri zehirlenmesi saldırılarına karşı geliştirilen savunma mekanizmalarının (örneğin Wan ve arkadaşları, 2023; Huang ve arkadaşları, 2024) etkinliğinin değerlendirilmesi, yalnızca teknik metriklerle değil, aynı zamanda kullanıcıların bu savunmalara olan güven algısıyla da ilişkilidir; bu nedenle, YZ güvenlik algısının standartlaştırılmış bir ölçekle ölçülmesi, hem akademik araştırmalar hem de kurumsal risk yönetimi ve düzenleyici politika geliştirme açısından stratejik bir zorunluluktur.

Yapay zekâ teknolojilerinin toplumsal yaşama hızla nüfuz etmesi, bu sistemlerin güvenlik ve ahlak boyutlarına ilişkin bireysel algıların sistematik biçimde incelenmesini gerekli kılmaktadır. Üniversite öğrencileri, dijital teknolojilerle yoğun etkileşimleri ve geleceğin uzmanları ile karar vericileri olmaları nedeniyle, yapay zekâyâ ilişkin algı araştırmaları açısından kritik bir örneklem grubu olarak öne çıkmaktadır (Morales-García et al., 2024; Yılmaz & Korkmaz, 2025).

Mevcut ölçeklerin önemli bir kısmının güvenlik boyutunu yeterince ele almadığı ya da etik ve güvenlik kavramlarını bütüncül bir çerçevede birlikte değerlendirmedeği belirtilmektedir (Köse, 2020; Saatci, 2025). Yapay zekânın yaygınlaşması, bu teknolojiye yönelik tutumların tek boyutlu yaklaşımlarla açıklanamayacağını göstermekte; güncel çalışmalar, tutumların güvenlik, etik, toplumsal etki ve beklentiler gibi çoklu boyutların etkileşimiyle şekillendiğini ortaya koymaktadır (Gerlich, 2023; Schepman & Rodway, 2020; Grassini, 2023; Saatci, 2025; Ikkatai et al., 2022; Hartwig et al., 2022). Bu doğrultuda, yapay zekâ algısının risk, toplumsal etki ve dönüşüm potansiyeli gibi unsurlar tarafından belirlendiği ve birlikte ele alınması gerektiği vurgulanmaktadır (Gerlich, 2023; Afroogh et al., 2024; Omrani et al., 2022; Choung et al., 2022). Ancak mevcut ölçeklerin çoğu fayda/zarar, duygusal tepkiler veya iş yaşamına etkiler gibi sınırlı boyutlara odaklanmakta; etik ve güvenlik gibi alanlar eksik temsil edilmektedir (Schepman & Rodway, 2020; Grassini, 2023; Park et al., 2024; Stein et al., 2024; Carvalho et al., 2025; Møgelvang & Grassini, 2025; Satıcı et al., 2025). Bu nedenle geliştirilen çok boyutlu

ölçekler; şeffaflık, hesap verebilirlik, mahremiyet, adalet ve insan denetimi gibi etik boyutları kapsamakta ve ELSI çerçevesinde beklenti, kültürel/değersel bakış ile politika ve hukuk algılarını da içermektedir (Saatci, 2025; Zhang et al., 2022; Kerr et al., 2020; Kong et al., 2025; Kamila & Jasrotia, 2023; Ikkatai et al., 2022; Vo et al., 2023; Hartwig et al., 2022).

1.3. Varsayımlar

Bu çalışma aşağıdaki temel varsayımlar üzerine kurulmaktadır:

1. Bireylerin yapay zekâya yönelik algıları çok boyutlu bir yapı sergilemektedir. Yapay zekâ teknolojilerinin toplumsal yaşama hızla nüfuz etmesi, bu sistemlerin güvenlik ve ahlak boyutlarına ilişkin bireysel algıların sistematik biçimde incelenmesini gerekli kılmaktadır. Literatür, öğrencilerin yapay zekâya ilişkin etik değerlendirmelerinde özellikle adalet, mahremiyet, zarar vermeme, şeffaflık, sorumluluk ve insan denetimi gibi boyutların öne çıktığını göstermektedir (Jang et al., 2022; Wang et al., 2025; Abuadas & Albikawi, 2025).
2. Güvenlik algısı ve etik algı birbirleriyle ilişkili ancak ayrı yapılarıdır. Mevcut ölçeklerin önemli bir kısmının güvenlik boyutunu yeterince ele almadığı ya da etik ve güvenlik kavramlarını bütüncül bir çerçevede birlikte değerlendirmedeği belirtilmektedir (Köse, 2020; Saatci, 2025). Bu çalışma, bu iki boyutun bir arada ölçülebileceği varsayımına dayanmaktadır.
3. Üniversite öğrencileri, yapay zekâ algı araştırmaları için temsili bir örneklemdir. Üniversite öğrencileri, dijital teknolojilerle yoğun etkileşimleri ve geleceğin uzmanları ile karar vericileri olmaları nedeniyle, yapay zekâya ilişkin algı araştırmaları açısından kritik bir örneklem grubu olarak öne çıkmaktadır (Morales-García et al., 2024; Yılmaz & Korkmaz, 2025).
4. Algılanan fayda, algılanan risk ve etik değerlendirmeler birlikte tutum ve kullanım niyetini yordamaktadır. YZ bağlamındaki çalışmalar, TAM ve genişletilmiş versiyonlarının kullanıcı kabulünü açıklamada yüksek yordama gücüne sahip olduğunu göstermekte; özellikle verimlilik artışı, performans iyileşmesi ve zaman tasarrufu gibi algılanan faydaların YZ'ye yönelik tutumları ve kullanımını artırdığını ortaya koymaktadır (Kelly et al., 2022; Ibrahim et al., 2025; Ma et al., 2025; Schepman & Rodway, 2020; Yi & Choi, 2023; Venkatesh & Davis, 2000).
5. Kültürel normlar, teknolojik yeterlik ve yapay zekâya aşinalık gibi değişkenler güvenlik algısını farklılaştırmaktadır. Kültürel normlar, teknolojik yeterlik, yapay zekâya aşinalık, internet deneyimi ve kullanıcı kişiliği gibi değişkenlerin güven ve mahremiyet algısını

farklılaştırdığı; bu nedenle güvenlik algısının homojen değil, demografik ve psikografik unsurlara duyarlı dinamik bir süreç olarak ele alınması gerektiği anlaşılmaktadır (Novozhilova et al., 2024; Riedl, 2022).

1.4. Sınırlılıklar

Bu çalışmanın aşağıdaki sınırlılıkları bulunmaktadır:

1. Ölçek Geliştirme Sınırlılıkları: Bayram (2025) tarafından geliştirilen Yapay Zekâ Etiği Ölçeği; tarafsızlık, şeffaflık, hesap verebilirlik ve veri boyutlarını kapsamakta, ancak özet ya da tam metin eksikliği nedeniyle kapsamı ve psikometrik özellikleri hakkında ayrıntılı bilgi sunulamamaktadır. Benzer şekilde, mevcut ölçeklerin çoğu fayda/zarar, duygusal tepkiler veya iş yaşamına etkiler gibi sınırlı boyutlara odaklanmakta; etik ve güvenlik gibi alanlar eksik temsil edilmektedir (Schepman & Rodway, 2020; Grassini, 2023; Park et al., 2024; Stein et al., 2024; Carvalho et al., 2025; Møgelvang & Grassini, 2025; Satıcı et al., 2025).
2. Örneklem Sınırlılığı: Çalışma üniversite öğrencileri üzerinde yürütülecektir. Bu durum, bulguların genel nüfusa genellenebilirliğini sınırlayabilir. Ancak üniversite öğrencileri, dijital teknolojilerle yoğun etkileşimleri ve geleceğin uzmanları ile karar vericileri olmaları nedeniyle, yapay zekâya ilişkin algı araştırmaları açısından kritik bir örneklem grubu olarak öne çıkmaktadır (Morales-García et al., 2024; Yılmaz & Korkmaz, 2025).
3. Kültürel ve Bağlamsal Sınırlılıklar: Güvenlik algısının bağlamsal ve değişken bir yapı sergilediği; özellikle konuşma temelli yapay zekâ çalışmalarında güven, mahremiyet kaygıları, risk algısı ve antropomorfik unsurların en yaygın değerlendirilen yapılar arasında olduğu; ancak bu kavramların çoğu zaman örtüştüğü ve ölçeklerin farklı bağlamlarda sınırlı geçerlilik gösterebildiği belirtilmektedir (Leschanowsky et al., 2024). Kültürel normlar, teknolojik yeterlik, yapay zekâya aşinalık, internet deneyimi ve kullanıcı kişiliği gibi değişkenlerin güven ve mahremiyet algısını farklılaştırdığı anlaşılmaktadır (Novozhilova et al., 2024; Riedl, 2022).
4. Teknik Güvenlik ile Algısal Güvenlik Arasındaki Kopukluk: Algı ve güvenlik algısı kavramı, nesnel güvenlik durumundan bağımsız olarak bireylerin kendi öznel deneyimleri, değer yargıları, kültürel arka planları ve medya aracılığıyla edindikleri bilgiler üzerinden bir güvenlik değerlendirmesi yaptıkları gerçeğine dayanır. Güvenlik algısı literatüründe sıkça başvurulanan temel ayrım, nesnel güvenlik ile öznel güvenlik algısı arasındaki kopukluktur; bu kopukluk güvenlik politikalarının etkinliğini sorgulamak ve bireylerin gerçekten güvende hissetmelerini sağlamak adına kamu yönetimi ve kent güvenliği çalışmalarında merkezi bir problem alanı olarak ele alınır.

5. Açıklanabilirlik ve Güven İlişkisinin Karmaşıklığı: Mevcut kanıtlar açıklanabilirliğin tek başına yeterli olmadığını; doğruluk, öngörülebilirlik ve sistem performansına ilişkin anlamlı açıklamaların da güven oluşumunda etkili olduğunu ortaya koymaktadır (Afroogh et al., 2024; Hassija et al., 2023; Ismatullaev & Kim, 2022). Derin öğrenme ve benzeri karmaşık modellerin karar süreçleri çoğu zaman kullanıcı açısından anlaşılabilir değildir ve bu durum, sistem çıktılarının güvenilirliğine ilişkin kuşkuları artırmaktadır (Von Eschenbach, 2021; Hassija et al., 2023; Adadi & Berrada, 2018).

2. KAVRAMSAL ÇERÇEVE VE İLGİLİ ÇALIŞMALAR

Bu çalışmanın kavramsal çerçevesi, yapay zekâya yönelik güvenlik ve ahlak algısının çok boyutlu yapısını açıklayan teorik ve ampirik kaynaklardan türetilmiştir. Çerçeve, beş ana boyut üzerinden yapılandırılmış olup, bu boyutlar uluslararası etik çerçeveler ve güncel ampirik çalışmalar temelinde belirlenmiştir (Vesnić-Alujević et al., 2020; Rodríguez et al., 2023; Zhang & Zhang, 2023; Floridi et al., 2018; Dignum, 2018).

2.1. Güvenlik Farkındalığı Boyutu

Güvenlik farkındalığı boyutu, bireylerin yapay zekâ sistemlerinin güvenlik, mahremiyet ve kötüye kullanım gibi risklerine ilişkin farkındalık düzeyini kapsamaktadır. Bu boyut, yapay zekâ güvenlik algısının ölçülmesinin hem bireysel hem de kurumsal düzeyde kritik öneme sahip olduğu varsayımına dayanmaktadır (Gerlich, 2023; Liehner et al., 2023; Ding et al., 2026; Brauner et al., 2023). Adversarial makine öğrenmesi saldırılarının evrimi, yapay zekâ sistemlerinin güvenliğine ilişkin algıları doğrudan etkilemekte ve bu algıların sistematik olarak ölçülmesi gerekliliğini ortaya koymaktadır. Biggio ve Roli (2018) tarafından 2004-2018 dönemini kapsayan kapsamlı derleme çalışması, adversarial makine öğrenmesi alanındaki saldırı ve savunma tekniklerinin tarihsel gelişimini ortaya koymuş; Papernot ve arkadaşları (2017) ise kara kutu saldırılarına karşı pratik yöntemler geliştirerek, yapay zekâ modellerinin güvenliğinin yalnızca teknik önlemlerle değil, aynı zamanda kullanıcıların ve geliştiricilerin bu tehditlere yönelik algılarının anlaşılmasıyla sağlanabileceğini göstermiştir. Son dönemde Qi ve arkadaşları (2024), hizalanmış dil modellerinin güvenlik eğitiminin ne ölçüde kırılğan olduğunu ortaya koyarak, yapay zekâ güvenliğinin sadece teknik dayanıklılıkla değil, kullanıcıların bu kırılğanlıklara ilişkin farkındalığı ve algısıyla da ilişkili olduğunu kanıtlamıştır. Veri zehirlenmesi saldırılarına karşı geliştirilen savunma mekanizmalarının (örneğin Wan ve arkadaşları, 2023; Huang ve arkadaşları, 2024) etkinliğinin değerlendirilmesi, yalnızca teknik metriklerle değil, aynı zamanda kullanıcıların bu savunmalara olan güven algısıyla da

ilişkilidir; bu nedenle, yapay zekâ güvenlik algısının standartlaştırılmış bir ölçekle ölçülmesi, hem akademik araştırmalar hem de kurumsal risk yönetimi ve düzenleyici politika geliştirme açısından stratejik bir zorunluluktur.

Ayrıca, yapay zekâ sistemlerinin eğitiminde kullanılan devasa veri kümeleri, kişisel bilgileri, parolaları ve telif hakkıyla korunan içerikleri istemsiz olarak ezberleyebilmekte; bu durum model çıkarımı (model inversion) ve üyelik çıkarımı (membership inference) gibi saldırılara karşı savunmasız hale getirmektedir. Sınıf niteliği çıkarım saldırıları, özellikle yüz tanıma modellerinde saç rengi, cinsiyet ve ırksal görünüm gibi açıklanmamış nitelikleri doğru bir şekilde çıkarabilmekte; ilginç bir şekilde, adversaryal olarak sağlamlaştırılmış modellerin bu tür gizlilik sızıntılarına karşı daha savunmasız olduğu gösterilmiştir. 2024 yılı sonlarına kadar kurumların %41'inin bir tür yapay zekâ güvenlik olayı yaşadığı rapor edilmiştir; bu olaylar veri zehirlenmesinden model hırsızlığına kadar geniş bir yelpazede yer almaktadır (Radanliev, 2025; Rodríguez et al., 2023; Aunugu & Vathsavai, 2022).

Güvenlik algısı kavramı, bireylerin çevrelerindeki potansiyel tehditleri ve riskleri nasıl yorumladıklarını, anlamlandırdıklarını ve bu doğrultuda davranışlarını şekillendirdiklerini inceleyen çok katmanlı bir sosyo-psikolojik ve güvenlik çalışmaları kavramıdır. Bu kavram temel olarak, nesnel güvenlik durumundan bağımsız olarak bireylerin kendi öznel deneyimleri, değer yargıları, kültürel arka planları ve medya aracılığıyla edindikleri bilgiler üzerinden bir güvenlik değerlendirmesi yaptıkları gerçeğine dayanır. Güvenlik algısı, bireyin kendini fiziksel, psikolojik ve sosyal açıdan ne ölçüde tehdit altında hissettiğine dair öznel bir yargıdır ve bu yargı bireyin günlük hayat tercihlerinden, kamusal alana katılım düzeyine, hatta politik tercihlerine ve oylama davranışlarına kadar geniş bir yelpazede somut sonuçlar doğurur. Güvenlik algısı literatüründe sıkça başvurulanan temel ayrım, nesnel güvenlik ile öznel güvenlik algısı arasındaki kopukluktur; bu kopukluk güvenlik politikalarının etkinliğini sorgulamak ve bireylerin gerçekten güvende hissetmelerini sağlamak adına kamu yönetimi ve kent güvenliği çalışmalarında merkezi bir problem alanı olarak ele alınır.

2.2. Ahlakî Değerler Boyutu

Ahlakî değerler boyutu, şeffaflık, adalet, mahremiyet ve insan denetimi gibi etik ilkelere yönelik algıları değerlendirmektedir. Bu boyut, yapay zekâ etiğinin normatif ilkelerin belirlenmesi ile bu ilkelerin kurumsal ve teknik uygulamalara yansıtılmasını kapsayan çok disiplinli bir alan olarak ele alınması gerektiği anlaşılmaktadır (Rodríguez et al.,

2023; Cheong, 2024; Jedlickova, 2024; Winfield & Jirotko, 2018; Kamila & Jasrotia, 2023; Mittelstadt, 2019).

Yapay zekâ etiği, yapay zekâ sistemlerinin tasarım, geliştirme ve kullanım süreçlerine yön veren ilke ve değerleri inceleyen bağımsız bir disiplin olarak tanımlanmaktadır (Huang et al., 2023; Kazim & Koshiyama, 2020). Küresel etik rehberler üzerine yapılan çalışmalar, şeffaflık, adalet ve hakkaniyet, zarar vermeme, sorumluluk/hesap verebilirlik ve gizlilik olmak üzere ortak etik eksenler bulunduğunu göstermektedir (Jobin et al., 2019; Shukla, 2024; Huang et al., 2023; Kazim & Koshiyama, 2020). Bu ilkeler, yapay zekâ uygulamalarının bireysel hakları gözetken ve toplumsal eşitsizlikleri derinleştirmeyecek şekilde tasarlanması gerekliliğine işaret etmektedir.

Uluslararası politika ve rehberler de bu ilkelere dayanmaktadır. OECD ilkeleri; insan hakları, hesap verebilirlik, şeffaflık, güvenlik, adalet ve insan merkezli değerleri vurgulamaktadır (Fjeld et al., 2020; Camilleri, 2023). UNESCO tarafından kabul edilen "Yapay Zekâ Etiği Tavsiye Kararı" ise etik yapay zekânın insan haklarına dayalı, kapsayıcı ve sürdürülebilir biçimde geliştirilmesini öne çıkarmaktadır (Huang et al., 2023).

Şeffaflık eksikliği ("kara kutu" problemi), sistemlerin karar süreçlerinin anlaşılmasını zorlaştırarak güven sorunlarına neden olmaktadır (Zhang & Zhang, 2023; Lepri et al., 2021; Nouis et al., 2025; Osasona et al., 2024; Kaur et al., 2022; Alzubaidi et al., 2023). Özellikle derin öğrenme modellerinin karar verme mekanizmalarının yorumlanabilirliği konusundaki zorluklar, hem kullanıcıların hem de düzenleyici otoritelerin sistemlere olan güvenini sarsmaktadır. Bu durum, özellikle yüksek riskli uygulama alanlarında örneğin sağlık hizmetlerinde teşhis kararları veya finansal kredi değerlendirmelerinde ciddi toplumsal ve hukuki sonuçlar doğurabilmektedir.

Veri gizliliği ve mahremiyet sorunları, özellikle sağlık verileri ve giyilebilir teknolojiler bağlamında ihlal, gözetim ve kimlik hırsızlığı risklerini artırmaktadır (Weiner et al., 2024; Zhang & Zhang, 2023; Radanliev, 2025; Rodríguez et al., 2023; Aunugu & Vathsavai, 2022).

İnsan-merkezli yaklaşım kapsamında ise özerkliğin korunması, kırılgan grupların gözetilmesi ve kamu katılımı vurgulanmaktadır (Zhang & Zhang, 2023; Lepri et al., 2021; Murikah et al., 2024; Rodríguez et al., 2023; Pham, 2025; Aunugu & Vathsavai, 2022). Bu çerçevede, yapay zekâ sistemlerinin tasarım ve dağıtım süreçlerinde bireylerin temel hak ve özgürlüklerinin

korunması, özellikle dezavantajlı ve marjinalleştirilmiş toplulukların sistematik ayrımcılığa maruz kalmaması için önlemler alınması esastır.

Bayram (2025) tarafından geliştirilen Yapay Zekâ Etiği Ölçeği; tarafsızlık, şeffaflık, hesap verebilirlik ve veri boyutlarını kapsamaktadır. Saatci (2025) tarafından geliştirilen AI and Ethics Perception Scale (AEPS), etik algıyı şeffaflık, hesap verebilirlik, mahremiyet, adillik ve insan gözetimi boyutlarında çok boyutlu olarak ele almakta; Türkiye, Hindistan ve Birleşik Krallık'ta farklı sektörlerde geliştirilip doğrulanmıştır (Saatci, 2025).

2.3. Sorumluluk ve Hesap Verebilirlik Boyutu

Sorumluluk ve hesap verebilirlik boyutu, yapay zekâ çıktılarından sorumlu aktörlere ilişkin algıları ölçmektedir. Bu boyut, yapay zekâ sistemlerinin kararları sonucu ortaya çıkan zararların sorumluluğunun geliştiriciye, veri sağlayıcıya, sistem operatörüne mi yoksa son kullanıcıya mı ait olduğu konusundaki belirsizlikle ilgili tartışmalara dayanmaktadır (Saatci, 2025; Rodríguez et al., 2023; Cheong, 2024; Zhang & Zhang, 2023; Dignum, 2018).

Hesap verebilirlik boşluğu, hatalı kararlar karşısında sorumluluğun belirsizleşmesine yol açmaktadır (Zhang & Zhang, 2023; Lepri et al., 2021; Nouis et al., 2025; Akinrinola et al., 2024; Osasona et al., 2024; Kaur et al., 2022; Pham, 2025). Bu bağlamda, risk değerlendirmelerinin dağıtımdan önce gerçekleştirilmesi, sürekli performans izleme ve güven ile işbirliğini teşvik eden politikalar önerilmektedir.

Yapay zekâ yönetişimi kapsamında Avrupa politikaları, OECD ve G20 ilkeleri ile çeşitli kurumsal çerçeveler; sorumluluk, hesap verebilirlik, şeffaflık, güven, adalet, mahremiyet ve güvenlik ilkelerine odaklanmaktadır (Leslie, 2019; Camilleri, 2023; Vesnić-Alujević et al., 2020; Dignum, 2018; Huang et al., 2023; Zhang et al, 2021).

Araştırmacılar, yapay zekâ güvenliği çalışmalarının önceliklendirilmesi, potansiyel zararların değerlendirilmesine yönelik ön inceleme süreçlerinin geliştirilmesi ve uluslararası kuruluşların daha etkin rol üstlenmesi gerektiğini belirtmektedir (Zhang et al, 2021).

2.4. Toplumsal Etki Boyutu

Toplumsal etki boyutu, eşitsizlik, istihdam, güvenlik ve refah gibi alanlardaki etkileri kapsamaktadır. Bu boyut, yapay zekâ ve otomasyonun iş gücü piyasaları üzerindeki etkilerine dikkat çeken çalışmalara dayanmaktadır (Saatci, 2025; Rodríguez et al., 2023; Cheong, 2024;

Zhang & Zhang, 2023; Dignum, 2018; Gerlich, 2023; Vesnić-Alujević et al., 2020; Novotny et al., 2025; Ding et al., 2026; Brauner et al., 2023; Bullock et al., 2025).

Erik Brynjolfsson ve Andrew McAfee (2017), yapay zekâ ve otomasyonun iş gücü piyasaları üzerindeki etkilerine dikkat çekerek bazı mesleklerin ortadan kalkabileceğini ve gelir eşitsizliklerinin artabileceğini belirtmekte; bu durumun, teknolojilerin yalnızca verimlilik ve yenilik değil, aynı zamanda istihdam, çevresel etkiler ve toplumsal adalet boyutlarında da değerlendirilmesi gerektiğini göstermektedir (Brynjolfsson & McAfee, 2017).

Uygulama alanlarına göre riskler farklılaşmaktadır. Sağlık alanında önyargı, veri kalitesi ve sorumluluk belirsizliği öne çıkmaktadır (Weiner et al., 2024; Zhang & Zhang, 2023; Lepri et al., 2021; Nouis et al., 2025; Pham, 2025; Daneshjou et al., 2021). Finans ve iş dünyasında ayrımcı kararlar ve şeffaflık sorunları yaşanmaktadır (Lepri et al., 2021; Murikah et al., 2024; Bahangulu & Owusu-Berko, 2025; Mensah, 2024; Alzubaidi et al., 2023). Denetimde birikimli önyargı etkileri gözlemlenmektedir (Murikah et al., 2024; Alzubaidi et al., 2023). Giyilebilir/IoT sistemlerinde ise gizlilik, güvenlik ve özerklik kaybı öne çıkmaktadır (Radanliev, 2025; Rodríguez et al., 2023; Aunugu & Vathsavai, 2022).

Algoritmik önyargı, temsili olmayan veya eksik veri setleri nedeniyle ayrımcı sonuçlara yol açabilmektedir (Weiner et al., 2024; Zhang & Zhang, 2023; Murikah et al., 2024; Bahangulu & Owusu-Berko, 2025; Alzubaidi et al., 2023; Daneshjou et al., 2021). Bu önyargının kaynakları arasında veri eksiklikleri, demografik homojenlik, sahte korelasyonlar, uygun olmayan karşılaştırma kriterleri ve bilişsel önyargılar yer almaktadır. Verimlilik ile titizlik arasındaki ödünleşmeler, insan becerilerinin ve yargı yeteneğinin aşınması, veri bağımlılığı riskleri ve kontrolsüz kişisel veri sömürsünden kaynaklanan gizlilik ihlalleri gibi daha geniş sorunlar da literatürde belirtilmektedir.

2.5. Olumlu Beklentiler Boyutu

Olumlu beklentiler boyutu, verimlilik, konfor, maliyet düşüşü ve toplumsal yarara yönelik beklentileri ortaya koymaktadır. Bu boyut, Teknoloji Kabul Modeli (TAM) çerçevesinde algılanan fayda ve algılanan kullanım kolaylığının tutum ve kullanım niyetinin temel belirleyicileri olduğu anlayışına dayanmaktadır (Gerlich, 2023; Liehner et al., 2023; Neri & Cozman, 2019; Ding et al., 2026; Schepman & Rodway, 2020).

Teknoloji Kabul Modeli (Technology Acceptance Model – TAM), yeni teknolojilere yönelik kullanıcı algılarını açıklamada yaygın bir kuramsal çerçeve olarak kullanılmaktadır (Kelly et al., 2022; Ibrahim et al., 2025). Fred Davis tarafından geliştirilen modelde algılanan fayda ve algılanan kullanım kolaylığı, tutum ve kullanım niyetinin temel belirleyicileri olarak ele alınmakta; algılanan faydanın en güçlü yordayıcı olduğu, algılanan kullanım kolaylığının ise hem doğrudan hem dolaylı etkiler gösterdiği belirtilmektedir (Kelly et al., 2022; Ibrahim et al., 2025; Ma et al., 2025; Schepman & Rodway, 2020; Davis, 1989; Granić & Marangunic, 2019).

Yapay zekâ bağlamındaki çalışmalar, TAM ve genişletilmiş versiyonlarının kullanıcı kabulünü açıklamada yüksek yordama gücüne sahip olduğunu göstermekte; özellikle verimlilik artışı, performans iyileşmesi ve zaman tasarrufu gibi algılanan faydaların yapay zekâyeye yönelik tutumları ve kullanımını artırdığını ortaya koymaktadır (Kelly et al., 2022; Ibrahim et al., 2025; Ma et al., 2025; Schepman & Rodway, 2020; Yi & Choi, 2023; Venkatesh & Davis, 2000).

Yapay zekâyeye yönelik algılar, yalnızca olumlu beklentilerle sınırlı olmayıp; risk, etik sorunlar ve toplumsal etkiler gibi çok boyutlu değerlendirmeleri de içermekte, bu kapsamda bireysel ve toplumsal düzeyde hem teknolojik ilerleme hem de olası olumsuzluklara yönelik kaygıları yansıtmaktadır.

2.6. Kavramsal Çerçevenin Entegrasyonu

Yapay zekânın karar verme süreçlerine entegrasyonu, giderek daha fazla, insan bilişini yalnızca destekleyen bir araç olarak değil, kurumsal karar alma kapasitesini yeniden tanımlayan sosyo-teknik bir dönüşüm olarak kavramsallaştırılmaktadır (Upase, 2026). Bu dönüşüm, veriye dayalı tahminsel analitiklerin karar vericilere seçenek sunmasının ötesine geçerek, belirli bağlamlarda otonom ya da yarı otonom karar yetkisinin algoritmik sistemlere devredilmesini de kapsamaktadır (Mohammed, 2025; Kulaklıoğlu, 2024). Bu çerçevede yapay zekâ, karar alma süreçlerinde yalnızca verimlilik artıran bir teknoloji değil, aynı zamanda kararın nasıl üretildiğine ilişkin epistemolojik varsayımları dönüştüren bir paradigma değişimi olarak değerlendirilmektedir (Lindner et al., 2025; Upase, 2026).

Yapay zekâ destekli karar sistemlerinin temel vaatlerinden biri, insan karar vericilerin sınırlı rasyonellik, sezgisel kestirme yollar ve bilişsel önyargılar nedeniyle karşılaştığı kısıtları aşabilmesidir (Lindner et al., 2025; Upase, 2026). Özellikle büyük veri işleme, örüntü tanıma ve alternatif üretme kapasitesi sayesinde yapay zekâ, örgütsel karar kalitesini artıran bir artırılma mekanizması olarak konumlandırılmaktadır (Shick et al., 2023; Upase, 2026).

Bununla birlikte, literatür yapay zekânın insan muhakemesinin yerine tam ikame oluşturduğunu değil, çoğu durumda hibrit insan-yapay zekâ modellerinin daha uygun olduğunu vurgulamaktadır (Upase, 2026; Vincent, 2021).

Yapay zekâyâ yönelik güvenlik algısının ölçülmesi, yalnızca teknik güvenlik göstergelerinin değerlendirilmesiyle sınırlı olmayan; aynı zamanda kullanıcı ile sistem arasındaki etkileşimin bilişsel, sosyal ve psikolojik boyutlarını kavramayı gerektiren çok katmanlı bir araştırma alanıdır. Kullanıcı güveninin, yapay zekâ tabanlı sistemlerin benimsenmesi ve kullanımının sürekliliği üzerinde belirleyici bir rol oynadığı gösterilmiştir; bu nedenle bir sistemin toplumsal ve kurumsal düzeyde başarılı biçimde bütünleşmesi, yalnızca nesnel güvenlik önlemlerine değil, bu önlemlerin kullanıcılar tarafından ne ölçüde güven verici bulunduğu da bağlıdır (Bach et al., 2022; Choung et al., 2022).

Güvenlik algısının sistematik olarak ölçülmesi ilk olarak güven-benimseme ilişkisi açısından önem taşımaktadır. Yapay zekâyâ yönelik güven ve güvensizlik, bu teknolojilerin günlük yaşama ve kurumsal süreçlere ne ölçüde nüfuz edeceğini düzenleyen başlıca etmenler arasında yer almaktadır (Afroogh et al., 2024; Bach et al., 2022). Kullanıcı güveninin yalnızca teknik performansla değil, aynı zamanda sosyo-etik değerlendirmeler, kullanıcı özellikleri ve bağlama özgü deneyimlerle şekillendiği; dolayısıyla insan-merkezli değerlendirme yaklaşımlarının teknik-merkezli güvenlik analizlerini tamamlaması gerektiği vurgulanmaktadır (Bach et al., 2022; Ismatullaev & Kim, 2022).

İkinci olarak, yapay zekâ sistemlerinin doğasında bulunan opaklık problemi, güvenlik algısının ölçülmesini özellikle gerekli kılmaktadır. Derin öğrenme ve benzeri karmaşık modellerin karar süreçleri çoğu zaman kullanıcı açısından anlaşılabilir değildir ve bu durum, sistem çıktılarının güvenilirliğine ilişkin kuşkuları artırmaktadır (Von Eschenbach, 2021; Hassija et al., 2023; Adadi & Berrada, 2018). Literatür, açıklanabilirlik ve yorumlanabilirliğin kullanıcı güvenliğini artırabildiğini, belirsizlik farkındalığı ve uygun güven kalibrasyonunun ise aşırı güven ya da yetersiz güven gibi sorunların önlenmesinde önemli olduğunu göstermektedir (Afroogh et al., 2024; Shin, 2021). Bununla birlikte, mevcut kanıtlar açıklanabilirliğin tek başına yeterli olmadığını; doğruluk, öngörülebilirlik ve sistem performansına ilişkin anlamlı açıklamaların da güven oluşumunda etkili olduğunu ortaya koymaktadır (Afroogh et al., 2024; Hassija et al., 2023; Ismatullaev & Kim, 2022).

Üçüncü olarak, veri mahremiyeti ve güvenlik endişelerinin artması, yapay zekâya ilişkin güvenlik algısının bağlamsal ve değişken bir yapı sergilediğini göstermektedir. Özellikle konuşma temelli yapay zekâ çalışmalarında güven, mahremiyet kaygıları, risk algısı ve antropomorfik unsurların en yaygın değerlendirilen yapılar arasında olduğu; ancak bu kavramların çoğu zaman örtüştüğü ve ölçeklerin farklı bağlamlarda sınırlı geçerlilik gösterebildiği belirtilmektedir (Leschanowsky et al., 2024). Ayrıca kültürel normlar, teknolojik yeterlik, yapay zekâya aşinalık, internet deneyimi ve kullanıcı kişiliği gibi değişkenlerin güven ve mahremiyet algısını farklılaştırdığı; bu nedenle güvenlik algısının homojen değil, demografik ve psikografik unsurlara duyarlı dinamik bir süreç olarak ele alınması gerektiği anlaşılmaktadır (Novozhilova et al., 2024; Riedl, 2022).

Yapay zekâ teknolojilerinin toplumsal yaşama hızla nüfuz etmesi, bu sistemlerin güvenlik ve ahlak boyutlarına ilişkin bireysel algıların sistematik biçimde incelenmesini gerekli kılmaktadır. Üniversite öğrencileri, dijital teknolojilerle yoğun etkileşimleri ve geleceğin uzmanları ile karar vericileri olmaları nedeniyle, yapay zekâya ilişkin algı araştırmaları açısından kritik bir örneklem grubu olarak öne çıkmaktadır (Morales-García et al., 2024; Yılmaz & Korkmaz, 2025). Literatür, öğrencilerin yapay zekâya ilişkin etik değerlendirmelerinde özellikle adalet, mahremiyet, zarar vermeme, şeffaflık, sorumluluk ve insan denetimi gibi boyutların öne çıktığını göstermektedir. Üniversite öğrencilerine yönelik geliştirilen araçlar etik tutum, etik yansıtma, etik farkındalık ve akademik dürüstlük gibi alanlarda geçerli sonuçlar üretmiştir; ancak bu araçlar güvenlik kaygılarını ve ahlaki değerlendirmeleri tek bir bütünsel yapıda toplamamaktadır (Jang et al., 2022; Wang et al., 2025; Abuadas & Albikawi, 2025).

Yapay zekâya ilişkin güvenlik ve etik algıları bütüncül biçimde değerlendirebilecek geçerli ve güvenilir ölçme araçlarına ihtiyaç artmaktadır. Geliştirilen ölçekler, duyuşsal, bilişsel, davranışsal ve etik boyutların ölçülebileceğini göstermekle birlikte, özellikle etik ve güvenlik boyutlarının daha derin ele alınması gerektiğine işaret etmektedir (Long & Magerko, 2020; Ng et al., 2023; Hornberger et al., 2023; Bilbao-Eraña & Arroyo-Sagasta, 2025; Usher & Barak, 2024; Gümüş et al., 2023).

Mevcut ölçeklerin önemli bir kısmının güvenlik boyutunu yeterince ele almadığı ya da etik ve güvenlik kavramlarını bütüncül bir çerçevede birlikte değerlendirmedeği belirtilmektedir (Köse, 2020; Saatci, 2025). Yapay zekânın yaygınlaşması, bu teknolojiye yönelik tutumların tek boyutlu yaklaşımlarla açıklanamayacağını göstermekte; güncel çalışmalar, tutumların güvenlik,

etik, toplumsal etki ve beklentiler gibi çoklu boyutların etkileşimiyle şekillendiğini ortaya koymaktadır (Gerlich, 2023; Schepman & Rodway, 2020; Grassini, 2023; Saatci, 2025; Ikkatai et al., 2022; Hartwig et al., 2022). Bu doğrultuda, yapay zekâ algısının risk, toplumsal etki ve dönüşüm potansiyeli gibi unsurlar tarafından belirlendiği ve birlikte ele alınması gerektiği vurgulanmaktadır (Gerlich, 2023; Afroogh et al., 2024; Omrani et al., 2022; Choung et al., 2022). Ancak mevcut ölçeklerin çoğu fayda/zarar, duygusal tepkiler veya iş yaşamına etkiler gibi sınırlı boyutlara odaklanmakta; etik ve güvenlik gibi alanlar eksik temsil edilmektedir (Schepman & Rodway, 2020; Grassini, 2023; Park et al., 2024; Stein et al., 2024; Carvalho et al., 2025; Møgelvang & Grassini, 2025; Satıcı et al., 2025). Bu nedenle geliştirilen çok boyutlu ölçekler; şeffaflık, hesap verebilirlik, mahremiyet, adalet ve insan denetimi gibi etik boyutları kapsamakta ve ELSI çerçevesinde beklenti, kültürel/değersel bakış ile politika ve hukuk algılarını da içermektedir (Saatci, 2025; Zhang et al., 2022; Kerr et al., 2020; Kong et al., 2025; Kamila & Jasrotia, 2023; Ikkatai et al., 2022; Vo et al., 2023; Hartwig et al., 2022).

Bu çalışmada geliştirilecek "Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği", yukarıda belirtilen beş boyutu (güvenlik farkındalığı, ahlakî değerler, sorumluluk ve hesap verebilirlik, toplumsal etki ve olumlu beklentiler) bütünleştirerek, yapay zekâyâ ilişkin algıların çok boyutlu ve bütüncül bir şekilde ölçülmesini amaçlamaktadır. Bu boyutlar, uluslararası etik çerçeveler ve güncel ampirik çalışmalar temelinde yapılandırılmış olup, mevcut ölçeklerin eksik kaldığı güvenlik ve etik boyutlarının entegrasyonunu sağlamayı hedeflemektedir.

3. YÖNTEM

3.1. Araştırmanın Modeli

Bu araştırma, "Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği"nin geliştirilmesi ve psikometrik özelliklerinin incelenmesi amacıyla tasarlanmış karma yöntemli bir çalışmadır. Ölçek geliştirme süreci hem nitel hem de nicel araştırma tekniklerinin entegre edilmesiyle oluşmuştur. Ölçek geliştirme çalışmalarında geçerli ve güvenilir bir ölçme aracı oluşturabilmek için, öncelikle kapsamlı bir literatür taraması ve kavramsal çerçevenin belirlenmesi; ardından madde havuzunun oluşturulması, uzman görüşlerine başvurulması ve pilot uygulamaların gerçekleştirilmesi; son olarak da ana ölçek formunun uygulanması ve psikometrik analizlerin yapılması süreçleri izlemiştir. (Saatci, 2025; Long & Magerko, 2020; Ng et al., 2023).

3.2. Çalışma Grubu

Araştırmanın çalışma grubunu; yapay zekâ ve interneti etkin şekilde kullanan 18 yaş üzeri, çoğunluğunu öğrencilerin oluşturduğu ve mezunlarıda kapsayan 558 kişilik kitle oluşturmaktadır. Öğrencilerin, dijital teknolojilerle yoğun etkileşimleri ve geleceğin uzmanları ile karar vericileri olmaları nedeniyle, yapay zekâyâ ilişkin algı araştırmaları açısından kritik bir örneklem grubu olarak öne çıkmaktadır (Morales-García et al., 2024; Yılmaz & Korkmaz, 2025). Literatür, öğrencilerin yapay zekâyâ ilişkin etik değerlendirmelerinde özellikle adalet, mahremiyet, zarar vermeme, şeffaflık, sorumluluk ve insan denetimi gibi boyutların öne çıktığını göstermektedir (Jang et al., 2022; Wang et al., 2025; Abuadas & Albikawi, 2025). Eğitim araştırmaları, bireylerin yapay zekâyâ ilişkin algılarının bilgi düzeyi, deneyim ve dijital okuryazarlıkla ilişkili olduğunu göstermektedir (Long & Magerko, 2020; Ayanwale et al., 2024; Guan et al., 2024; Hornberger et al., 2023; Uzun, 2025). Çalışmalar, öğretmen adayları ve üniversite öğrencilerinde yapay zekâ farkındalığının arttığını; ancak etik ve güvenlik konularında yeterli bilinç oluşmadığını ortaya koymaktadır (Ayanwale et al., 2024; Guan et al., 2024; Hornberger et al., 2023; Basch et al., 2025; Valerio, 2024; Usher & Barak, 2024; Uzun, 2025).

Örneklemin belirlenmesinde kolayda örnekleme yöntemi kullanılacaktır. Çalışma grubunun büyüklüğü, ölçek geliştirme çalışmalarında önerilen katı sayısı göz önünde bulundurularak belirlenecektir. Faktör analizi için önerilen katı sayısının, madde sayısının en az 5-10 katı olması gerektiği belirtilmektedir (Saatci, 2025). Ayrıca, ölçeğin yapı geçerliğini doğrulamak amacıyla faktör analizi için yeterli büyüklükte bir örneklem toplanacaktır.

3.3. Veri Toplama Araçları

Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ)

Bu çalışmada geliştirilen "Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği", bireylerin yapay zekâyâ yönelik algılarını; güvenlik farkındalığı, ahlakî değerler, sorumluluk ve hesap verebilirlik, toplumsal etki ve olumlu beklentiler olmak üzere beş boyut altında incelemeyi amaçladı. Bu boyutlar, uluslararası etik çerçeveler ve güncel ampirik çalışmalar temelinde yapılandırıldı (Vesnić-Alujević et al., 2020; Rodríguez et al., 2023; Zhang & Zhang, 2023; Floridi et al., 2018; Dignum, 2018).

Ölçek maddelerinin oluşturulmasında, mevcut literatürde yer alan yapay zekâ etiği ve güvenlik algısı ölçeklerinden yararlanıldı. Bayram (2025) tarafından geliştirilen Yapay Zekâ Etiği

Ölçeği; tarafsızlık, şeffaflık, hesap verebilirlik ve veri boyutlarını kapsadı. Saatci (2025) tarafından geliştirilen AI and Ethics Perception Scale (AEPS), etik algıyı şeffaflık, hesap verebilirlik, mahremiyet, adillik ve insan gözetimi boyutlarında çok boyutlu olarak ele aldı; Türkiye, Hindistan ve Birleşik Krallık'ta farklı sektörlerde geliştirilip doğrulandı (Saatci, 2025). Mevcut ölçeklerin önemli bir kısmının güvenlik boyutunu yeterince ele almadığı ya da etik ve güvenlik kavramlarını bütüncül bir çerçevede birlikte değerlendirmedeği belirtildi (Köse, 2020; Saatci, 2025). Bu nedenle geliştirilen çok boyutlu ölçekler; şeffaflık, hesap verebilirlik, mahremiyet, adalet ve insan denetimi gibi etik boyutları kapsadı ve ELSI çerçevesinde beklenti, kültürel/değersel bakış ile politika ve hukuk algılarını da içerdi (Saatci, 2025; Zhang et al., 2022; Kerr et al., 2020; Kong vd., 2025; Kamila & Jasrotia, 2023; Ikkatai vd., 2022; Vo vd., 2023; Hartwig vd., 2022).

Ölçek maddeleri, 5'li Likert tipi derecelendirmeye uygun olarak hazırlandı. Maddelerin hazırlanmasında, yapay zekâ etiği ve güvenlik algısı literatüründe yer alan kavramsal tanımlar ve boyutlar dikkate alındı. Güvenlik farkındalığı boyutunda, adversarial makine öğrenmesi saldırıları, veri zehirlenmesi, model çıkarımı, üyelik çıkarımı ve sınıf niteliği çıkarım saldırıları gibi teknik güvenlik tehditlerine yönelik farkındalık maddeleri yer aldı (Biggio & Roli, 2018; Papernot vd., 2017; Qi vd., 2024; Wan vd., 2023; Huang vd., 2024). Ahlakî değerler boyutunda, şeffaflık, adalet, mahremiyet, insan denetimi ve zarar vermeme ilkelerine yönelik algı maddeleri bulundu (Jobin vd., 2019; Shukla, 2024; Huang vd., 2023; Kazim & Koshiyama, 2020). Sorumluluk ve hesap verebilirlik boyutunda, yapay zekâ kararlarından sorumlu aktörlere ilişkin algı ve sorumluluk atama maddeleri yer aldı (Zhang & Zhang, 2023; Lepri vd., 2021; Nouis vd., 2025; Pham, 2025). Toplumsal etki boyutunda, istihdam, eşitsizlik, güvenlik ve refah üzerindeki etkilere yönelik algı maddeleri bulundu (Brynjolfsson & McAfee, 2017; Gerlich, 2023; Ding vd., 2026). Olumlu beklentiler boyutunda ise verimlilik, konfor, maliyet düşüşü ve toplumsal yarar beklentilerini ölçen maddeler yer aldı (Gerlich, 2023; Liehner vd., 2023; Schepman & Rodway, 2020).

Bu kapsamda araştırmacı tarafından ilk olarak 80 madde hazırlandı ve uzman görüşlerine sunuldu.

Kişisel Bilgi Formu

Katılımcıların demografik özelliklerini (cinsiyet, yaş, eğitim durumu, akademik alanlar, gelir durum algısı, YZ uygulamalarını kullanımı durumları vb.) belirlemek amacıyla kişisel bilgi formu kullanıldı. Kültürel normlar, teknolojik yeterlik, yapay zekâyâ aşinalık, internet

deneyimi ve kullanıcı kişiliği gibi değişkenlerin güven ve mahremiyet algısını farklılaştırdığı literatürde gösterildiğinden (Novozhilova vd., 2024; Riedl, 2022), bu değişkenlerin ölçek puanları ile ilişkisi de incelendi.

3.4. Veri Toplama Süreci

Veri toplama süreci, çevrimiçi anket yöntemiyle gerçekleştirilmiştir. Çevrimiçi veri toplama, üniversite öğrencilerine ve yapay zekâ ile ilişkisi olan internet kullanıcılarına ulaşmada pratik ve etkili bir yöntemdir. Anketler, üniversitelerin öğrenci bilgi sistemleri veya sosyal medya platformları aracılığıyla dağıtılmıştır. Katılımcılara çalışmanın amacı, gönüllülük esası ve veri gizliliği konularında bilgi verilmiştir.

Veri toplama sürecinde, yapay zekâyâ yönelik algıların çok boyutlu yapısının göz önünde bulundurulması gerekmiştir. Yapay zekânın yaygınlaşması, bu teknolojiye yönelik tutumların tek boyutlu yaklaşımlarla açıklanamayacağını göstermekte; güncel çalışmalar, tutumların güvenlik, etik, toplumsal etki ve beklentiler gibi çoklu boyutların etkileşimiyle şekillendiğini ortaya koymaktadır (Gerlich, 2023; Schepman & Rodway, 2020; Grassini, 2023; Saatci, 2025; Ikkatai et al., 2022; Hartwig et al., 2022). Bu doğrultuda, yapay zekâ algısının risk, toplumsal etki ve dönüşüm potansiyeli gibi unsurlar tarafından belirlendiği ve birlikte ele alınması gerektiği vurgulanmaktadır (Gerlich, 2023; Afroogh et al., 2024; Omrani et al., 2022; Choung et al., 2022).

3.5. Geçerlik Çalışmaları

İçerik/Kapsam Geçerliği

Ölçeğin içerik geçerliğini sağlamak amacıyla alanında uzman akademisyenlerden oluşan bir uzman paneline başvurulmuştur. Alanında uzman 7 akademisyenden oluşan uzman paneli, ölçek maddelerinin kapsam geçerliğini değerlendirmiştir. Uzmanlar, maddelerin ilgili boyutları ne ölçüde temsil ettiğini, dil ve anlam açısından uygunluğunu ve kapsam geçerliğini değerlendirmiştir. Uzman değerlendirmeleri, Davis (1989) tarafından geliştirilen Teknoloji Kabul Modeli'nin (TAM) temel bileşenleri ve yapay zekâ etiği literatüründe yer alan boyutlar dikkate alınarak yapılmıştır (Kelly vd., 2022; Ibrahim vd., 2025; Ma vd., 2025; Schepman & Rodway, 2020).

Araştırmacı tarafından ilk olarak 80 madde hazırlanmış ve uzman görüşlerine sunulmuştur. Uzmanlardan; her bir maddenin ilgili boyutla ilişkisini, dil ve anlam bütünlüğünü, kavramsal netliğini ve ölçülmek istenen yapıyı ne ölçüde temsil ettiğini değerlendirmeleri istenmiştir. Uzmanların görüşleri doğrultusunda; kapsam geçerliği düşük olan, birden fazla boyuta yüklenen, dil ve anlam açısından net olmayan veya ölçülmek istenen yapıyı yeterince temsil etmeyen maddeler çıkarılmış veya yeniden düzenlenmiştir. Bu süreç sonucunda ölçek 60 maddeye indirgenmiştir.

Kapsam geçerliği değerlendirmesinde, uzmanların her bir maddeyi ilgili boyutla ilişkilendirme oranları hesaplanmış; kapsam geçerliği indeksi (Content Validity Index - CVI) hesaplanmıştır. Ayrıca, maddelerin kapsam geçerliği oranları (Content Validity Ratio - CVR) değerlendirilmiştir. Uzman değerlendirmeleri sonucunda kapsam geçerliği indeksi ve kapsam geçerliği oranlarının kabul edilebilir düzeyde olduğu görülmüştür.

Yapı Geçerliği

Yapı geçerliğini test etmek amacıyla açımlayıcı faktör analizi (AFA) ve doğrulayıcı faktör analizi (DFA) uygulandı. AFA, ölçeğin alt boyutlarını belirlemek ve maddelerin boyutlara uygun dağılımını incelemek amacıyla yapıldı. DFA ise, AFA sonucunda elde edilen faktör yapısının doğrulanması amacıyla kullanıldı.

Faktör analizi için örneklem büyüklüğünün yeterli olup olmadığı Kaiser-Meyer-Olkin (KMO) örnekleme yeterliliği testi ve Bartlett küresellik testi ile değerlendirildi. KMO değerinin kabul edilebilir düzeyde olduğu ve Bartlett küresellik testinin anlamlı çıktığı görüldü; bu durum, veri setinin faktör analizi için uygun olduğunu gösterdi.

Faktör sayısının belirlenmesinde özdeğer kriteri ($\text{eigenvalue} > 1$), scree plot ve paralel analiz yöntemleri kullanıldı. Uzman görüşleri sonucunda 60 maddeye indirgenen ölçek üzerinde yapılan AFA'da, maddelerin faktörlere yerleşimi varimax döndürme yöntemiyle incelendi. Faktör yük değeri en az 0.40 olan maddeler değerlendirmeye alındı; çift yüklenen veya beklenen boyut dışında yüklenen maddeler çıkarıldı. Bu süreç sonucunda ölçeğin alt boyutları belirlendi ve maddelerin ilgili boyutlara uygun şekilde dağıldığı görüldü.

Doğrulayıcı faktör analizinde, AFA sonucunda elde edilen faktör yapısının farklı bir örneklem grubu üzerinde doğrulanması amaçlandı. Model uyum indeksleri (χ^2/sd , GFI, AGFI, CFI,

RMSEA, SRMR) değerlendirildi. Uyum indekslerinin kabul edilebilir düzeyde olduğu görüldü; bu durum, ölçeğin yapı geçerliğini destekledi.

Ölçüt Geçerliği

Ölçeğin ölçüt geçerliğini test etmek amacıyla, yapay zekâya yönelik tutum ve algıları ölçen mevcut ölçeklerle korelasyon analizleri yapıldı. Bayram (2025) tarafından geliştirilen Yapay Zekâ Etiği Ölçeği ve Saatci (2025) tarafından geliştirilen AI and Ethics Perception Scale (AEPS) gibi araçlarla korelasyonlar incelendi. Ayrıca, Teknoloji Kabul Modeli (TAM) çerçevesinde algılanan fayda ve algılanan kullanım kolaylığı ölçekleriyle de korelasyon analizleri gerçekleştirildi (Davis, 1989; Kelly et al., 2022; Ibrahim et al., 2025).

Analiz sonucunda geliştirilen ölçek ile mevcut ölçekler arasında anlamlı ve beklenen yönde korelasyonlar elde edildi. Bu bulgular, geliştirilen ölçeğin ölçüt geçerliğini destekledi.

3.6. Güvenirlik Çalışmaları

Ölçeğin güvenirliğini değerlendirmek amacıyla iç tutarlılık katsayısı (Cronbach's Alpha), test-tekrar test güvenirliği ve madde toplam korelasyonları hesaplandı.

İç Tutarlılık

Ölçeğin iç tutarlılığı, Cronbach's Alpha katsayısı ile değerlendirildi. Genel ölçek için hesaplanan Alpha değerinin kabul edilebilir düzeyde ($\alpha > .70$) olduğu görüldü. Ayrıca, her bir alt boyut için ayrı ayrı Cronbach's Alpha değerleri hesaplandı ve maddelerin çıkarılması durumunda alpha değerindeki değişim incelendi.

Test-Tekrar Test Güvenirliği

Ölçeğin zaman içinde tutarlılığını değerlendirmek amacıyla, 2-4 hafta arayla aynı örnekleme ölçek yeniden uygulandı. İki uygulama arasındaki korelasyon katsayısı (Pearson r) hesaplandı; elde edilen yüksek ve anlamlı korelasyon, ölçeğin test-tekrar test güvenirliğini destekledi.

Madde Analizi

Madde toplam korelasyonları hesaplandı; düşük korelasyona sahip maddeler ($r < .30$) değerlendirildi. Ayrıca, üst %27 ve alt %27 gruplar arasındaki madde ayırt ediciliği bağımsız örneklem t-testi ile incelendi. Analiz sonucunda maddelerin istatistiksel olarak anlamlı düzeyde ayırt edici olduğu görüldü.

3.7. Veri Analizi

Verilerin analizinde SPSS 26.0 ve AMOS 24.0 istatistik paket programları kullanıldı. Veri setinin normallik varsayımı Kolmogorov-Smirnov ve Shapiro-Wilk testleri ile incelendi; normallik varsayımına uygun bulgular elde edildi ve buna göre parametrik istatistiksel teknikler uygulandı. Açımlayıcı faktör analizi, doğrulayıcı faktör analizi, iç tutarlılık analizi, korelasyon analizi, bağımsız örneklem t-testi ve tek yönlü varyans analizi (ANOVA) gibi istatistiksel teknikler kullanıldı.

Yapay zekâya yönelik güvenlik algısının ölçülmesi, yalnızca teknik güvenlik göstergelerinin değerlendirilmesiyle sınırlı olmayan; aynı zamanda kullanıcı ile sistem arasındaki etkileşimin bilişsel, sosyal ve psikolojik boyutlarını kavramayı gerektiren çok katmanlı bir araştırma alanıdır (Bach et al., 2022; Choung et al., 2022). Bu nedenle, analizlerde demografik değişkenlerin (yaş, cinsiyet, bölüm, yapay zekâ kullanım deneyimi, dijital okuryazarlık düzeyi vb.) ölçek puanları üzerindeki etkisi de incelendi. Kültürel normlar, teknolojik yeterlik, yapay zekâya aşinalık, internet deneyimi ve kullanıcı kişiliği gibi değişkenlerin güven ve mahremiyet algısını farklılaştırdığı gösterilmiştir (Novozhilova et al., 2024; Riedl, 2022).

3.8. Etik Hususlar

Araştırma, Erzincan Binali Yıldırım Üniversitesi İnsan Araştırmaları Eğitim Bilimleri Etik Kurulu tarafından 11.09.2025 tarih ve 10/18 nolu karar ile yürütülecektir. Katılımcılara çalışmanın amacı, gönüllülük esası, veri gizliliği ve anonimlik ilkeleri hakkında bilgi verilecektir. Katılımcıların çalışmaya katılımı tamamen gönüllü esasına dayanacak; istedikleri zaman çalışmadan ayrılma hakkına sahip olacaklardır. Toplanan veriler sadece bilimsel amaçlarla kullanılacak ve üçüncü kişilerle paylaşılmayacaktır.

Veri gizliliği ve mahremiyet sorunları, özellikle sağlık verileri ve giyilebilir teknolojiler bağlamında ihlal, gözetim ve kimlik hırsızlığı risklerini artırmaktadır (Weiner et al., 2024; Zhang & Zhang, 2023; Radanliev, 2025; Rodríguez et al., 2023; Aunugu & Vathsavai, 2022). Bu nedenle, araştırma sürecinde katılımcıların kişisel verilerinin korunmasına azami özen gösterilecektir.

4. BULGULAR

Bu bölümde, ölçeğin geliştirilmesi sürecinde gerçekleştirilen psikometrik analizlerden elde edilen bulgular sunulmaktadır.

Tablo 1. Ölçek Maddelerinin Betimsel İstatistikleri

Madde No	N	Aritmetik Ortalama	Standart Sapma	Madde Toplam Korelasyonu
s1	558	3.61	1.152	.737
s2	558	3.57	1.117	.735
s3	558	3.69	1.128	.760
s4	558	3.29	1.175	.611
s5	558	3.68	1.163	.762
s6	558	3.76	1.232	.738
s7	558	3.56	1.186	.720
s8	558	3.39	1.305	.581
s9	558	3.69	1.245	.721
s10	558	3.53	1.167	.713
s11	558	3.61	1.183	.718
s12	558	3.61	1.173	.718
s13	558	3.81	1.247	.792
s14	558	3.65	1.254	.648
s15	558	2.94	1.308	.382
s16	558	3.02	1.266	.485
s17	558	3.75	1.172	.814
s18	558	3.81	1.193	.783
s19	558	3.78	1.209	.804
s20	558	3.89	1.232	.794
s21	558	3.84	1.194	.797
s22	558	3.82	1.244	.788
s23	558	3.57	1.198	.702
s24	558	3.71	1.221	.769
s25	558	3.56	1.207	.729
s26	558	3.77	1.158	.810
s27	558	3.82	1.193	.823
s28	558	3.70	1.195	.769
s29	558	3.28	1.225	.537
s30	558	3.79	1.183	.796
s31	558	3.78	1.171	.822
s32	558	3.82	1.164	.829
s33	558	3.84	1.184	.835
s34	558	3.83	1.196	.823
s35	558	3.58	1.155	.749
s36	558	3.66	1.200	.742
s37	558	3.64	1.266	.737
s38	558	3.80	1.202	.757
s39	558	3.54	1.250	.741
s40	558	3.61	1.240	.725

s41	558	3.73	1.217	.766
s42	558	3.72	1.175	.790
s43	558	3.44	1.251	.700
s44	558	3.80	1.173	.805
s45	558	3.71	1.183	.776
s46	558	3.57	1.184	.725
s47	558	3.69	1.227	.767
s48	558	3.58	1.228	.740
s49	558	3.49	1.157	.621
s50	558	3.62	1.158	.704
s51	558	3.84	1.140	.794
s52	558	3.82	1.130	.790
s53	558	3.79	1.137	.762
s54	558	3.72	1.101	.735
s55	558	3.87	1.139	.795
s56	558	3.64	1.138	.671
s57	558	3.81	1.091	.800
s58	558	3.72	1.141	.731
s59	558	3.68	1.131	.719
s60	558	3.78	1.107	.733

Geliştirilen ölçme aracında yer alan 60 maddenin ayırt ediciliklerini belirlemek ve iç tutarlılıklarını incelemek amacıyla, 558 katılımcıdan elde edilen veriler üzerinden madde analizi gerçekleştirildi. Analiz kapsamında her bir maddeye ilişkin aritmetik ortalama, standart sapma ve madde-toplam korelasyon katsayıları hesaplandı; elde edilen bulgular Tablo 1'de sunuldu.

Tablo 1 incelendiğinde, ölçekte yer alan maddelerin aritmetik ortalamalarının 2,94 (Madde 15) ile 3,89 (Madde 20) arasında; standart sapma değerlerinin ise 1,09 (Madde 57) ile 1,30 (Madde 15) arasında değişim gösterdiği görüldü. Standart sapma değerlerinin kabul edilebilir sınırlar içerisinde olması, katılımcı yanıtlarının yeterli bir varyans gösterdiğine ve maddelerin ölçülmek istenen özellik bağlamında katılımcıları ayırt edebilme potansiyeline sahip olduğuna işaret etti.

Maddelerin ölçek geneliyle olan ilişkisini ve testin iç tutarlılığını değerlendirmek amacıyla hesaplanan madde-toplam korelasyon katsayıları incelendiğinde, katsayıların ,382 (Madde 15) ile ,835 (Madde 33) arasında yer aldığı tespit edildi. Literatürde, madde-toplam korelasyon katsayısının ,30 ve üzerinde olması, ilgili maddenin testin bütünüyle uyumlu olduğu, ayırt ediciliğinin yeterli olduğu ve ölçülmek istenen yapıyı anlamlı düzeyde ölçtüğü şeklinde yorumlanmaktadır (Büyüköztürk, 2020; Nunnally & Bernstein, 1994).

Analiz edilen 60 maddenin tamamının ($,382 \leq r \leq ,835$) madde-toplam korelasyon değerlerinin ,30 alt sınırının üzerinde olduğu belirlendi. Madde 15 (,382) ve Madde 16 (,485) diğer

maddelere kıyasla görece daha düşük korelasyon değerleri verse de, istatistiksel sınırların üzerinde kaldıkları için ölçekten çıkarılmalarına gerek duyulmadı. Öte yandan, ölçekteki maddelerin büyük çoğunluğunun ,70'in üzerinde korelasyon değerine sahip olması (örn. Madde 33; $r = ,835$), maddelerin kendi içlerinde yüksek düzeyde homojen bir yapı sergilediğini gösterdi.

Tablo 2. KMO ve Bartlett küresellik testi sonuçları

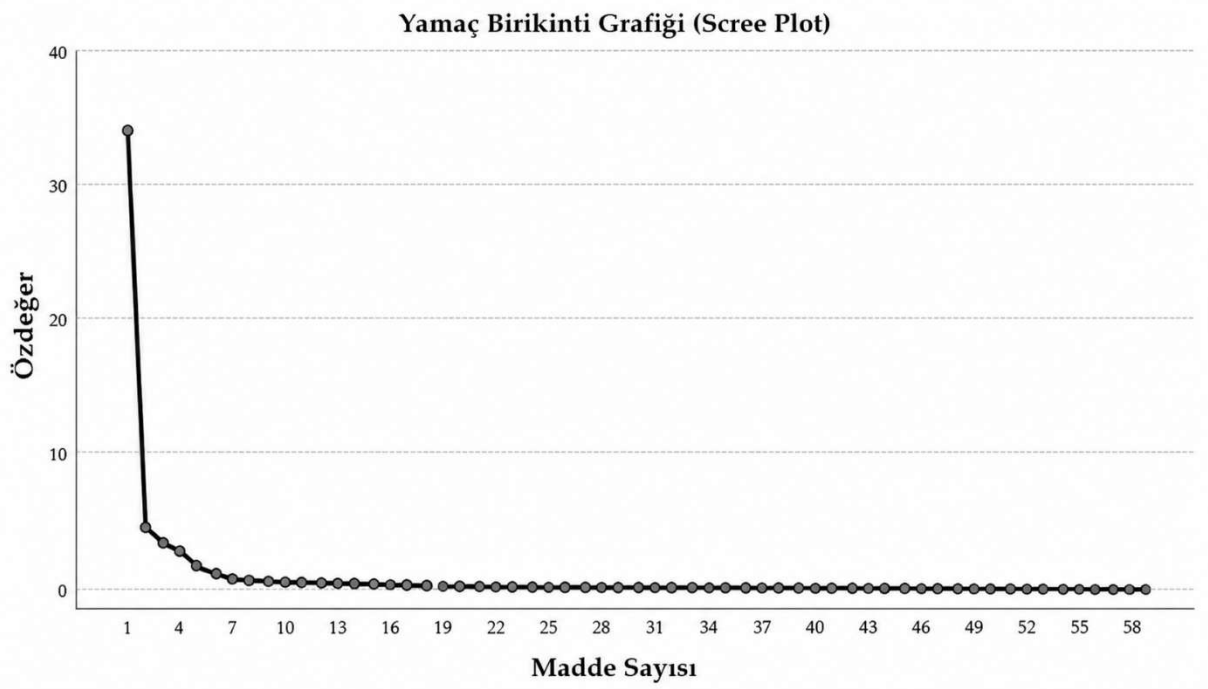
Kaiser-Meyer-Olkin Örneklem Yeterliliği Ölçüsü		.982
Bartlett Küresellik Testi	Approx. Chi-Square	42158.719
	sd	1770
	p	.000

Faktör analizi öncesinde örneklemin faktör analizine uygunluğunu değerlendirmek amacıyla Kaiser-Meyer-Olkin (KMO) ve Bartlett Küresellik Testi uygulanmıştır. Analiz sonucunda KMO değeri 0,982 olarak bulunmuş ve örneklem yeterliliğinin mükemmel düzeyde olduğu belirlenmiştir. Bartlett Küresellik Testi sonucunun da istatistiksel olarak anlamlı olduğu görülmüştür ($\chi^2(1770) = 42158.719$; $p < .001$). Bu sonuçlar, maddeler arasındaki korelasyonların faktör analizi için yeterli düzeyde olduğunu ve veri setinin faktör analizine uygun olduğunu göstermektedir.

Tablo 3. Faktör Analiz Sonuçları

Faktör	Özdeğer	Açıklanan Varyans (%)	Kümülatif (%)
1	33.882	56.471	56.471
2	4.061	6.769	63.240
3	3.346	5.577	68.817
4	2.786	4.644	73.461
5	1.520	2.533	75.994
6	1.213	2.022	78.016

Açımlayıcı Faktör Analizi sonucunda özdeğeri 1'in üzerinde olan 5 faktör elde edilmiştir. Birinci faktörün özdeğeri 33.882 olup toplam varyansın %56.471'ini açıkladığı belirlenmiştir. İkinci faktör %6.769, üçüncü faktör %5.577, dördüncü faktör %4.644, beşinci faktör %2.533 ve altıncı faktör %2.022 oranında varyans açıklamaktadır. Altı faktörün birlikte açıkladığı toplam varyans %78,016 olarak hesaplanmıştır. Bu değer, ölçeğin toplam varyansının önemli bir bölümünün elde edilen faktörler tarafından açıklandığını göstermektedir. Özdeğerlerin 1'in üzerinde olması ve yüksek düzeyde toplam açıklanan varyans elde edilmesi, ölçeğin 6 faktörlü yapısının desteklendiğine işaret etmektedir.



Şekil 1. Maddelerin Yamaç Birikinti (Scree Plot) Grafiği

Scree Plot incelendiğinde, birinci bileşenden altıncı bileşene kadar özdeğerlerde belirgin bir azalma olduğu, altıncı bileşenden sonra ise eğrinin belirgin biçimde yataylaştığı görülmektedir. Bu durum, ölçeğin 6 faktörlü bir yapıya sahip olabileceğini desteklemektedir. Özdeğerlerin 1'in üzerinde olması ve altı faktörün toplam varyansın %78,016'sını açıklaması da bu yapıyı destekleyen önemli bir bulgudur. Bununla birlikte, birinci faktörün toplam varyansın %56,471'ini tek başına açıklaması, ölçek yapısında baskın bir genel faktörün bulunabileceğine işaret etmektedir.

AFA sonuçlarında ortaya çıkan altı faktörlü yapının, ölçeğin kuramsal temelleriyle ve uzman görüşleriyle ne ölçüde örtüştüğünü; ayrıca baskın genel faktörün varlığının alt boyutların bağımsızlığını ne düzeyde etkilediğini test etmek amacıyla, yapı geçerliği ve güvenirliğine yönelik analizlere geçildi. Bu doğrultuda, öncelikle kapsam geçerliği kanıtları gözden geçirilerek ölçeğin kavramsal çerçevesiyle uyumu değerlendirildi.

Ölçeğin alt boyutlarına ilişkin geçerlik ve güvenirlik kanıtlarının ortaya konulması amacıyla öncelikle kapsam geçerlik indeksleri hesaplanmıştır. Daha sonra, madde analizi kapsamında madde-toplam puan korelasyonları incelenerek maddelerin ayırt edicilik düzeyleri değerlendirilmiştir. Ölçeğin kuramsal yapısının doğrulanması amacıyla doğrulayıcı faktör analizi (DFA) gerçekleştirilmiş ve model uyumuna ilişkin bulgular raporlanmıştır. Son aşamada

ise, belirlenen alt boyutların iç tutarlılığını değerlendirmek amacıyla güvenirlik analizleri yapılmış ve Cronbach's Alpha katsayıları hesaplanmıştır. Elde edilen bulgular, ölçeğin alt boyutları temelinde aşağıdaki sırayla sunulmaktadır:

1. Güvenlik Farkındalığı,
2. Ahlaki Değerler,
3. Sorumluluk ve Hesap Verilebilirlik,
4. Toplumsal Etki,
5. Olumlu Beklentiler

alt boyutudur.

Uzman değerlendirmeleri sonucunda, ölçek maddelerinin ilgili boyutları temsil etme düzeyi incelenmiştir. Kapsam geçerliği indeksi (CVI) ve kapsam geçerliği oranı (CVR) değerleri, ölçeğin kapsam geçerliğinin yüksek düzeyde olduğunu göstermiştir. Uzmanlar, maddelerin güvenlik farkındalığı, ahlakî değerler, sorumluluk ve hesap verebilirlik, toplumsal etki ve olumlu beklentiler boyutlarını yeterince kapsadığı konusunda görüş birliğine varmıştır (Saatci, 2025). Mevcut ölçeklerin güvenlik boyutunu yeterince ele almadığı ya da etik ve güvenlik kavramlarını bütüncül bir çerçevede birlikte değerlendirmede göz önünde bulundurularak, bu çalışmada geliştirilen ölçeğin güvenlik odaklı maddelerinin kapsamı genişletildiği için uzman değerlendirmelerinde bu boyutun özel olarak vurgulandığı görülmüştür (Köse, 2020; Saatci, 2025).

4.1. Yapay Zeka Güvenlik Farkındalığı Alt Boyutu

Ölçeğin Güvenlik Farkındalığı alt boyutuna ilişkin psikometrik özelliklerin belirlenmesi amacıyla kapsam geçerliği, açımlayıcı faktör analizi (AFA) ve doğrulayıcı faktör analizi (DFA) sonuçları birlikte değerlendirilmiştir. Bu süreçte, öncelikle uzman görüşleri doğrultusunda kapsam geçerliği incelenmiş, ardından AFA ile maddelerin faktör yapısına uygunluğu değerlendirilmiş ve son olarak DFA ile elde edilen yapının doğrulanması amaçlanmıştır. Analizler sonucunda, alt boyutun kuramsal yapısını en iyi temsil eden maddeler belirlenmiş ve uygun olmayan maddeler analiz sürecinden çıkarılmıştır. Güvenlik Farkındalığı alt boyutuna ilişkin madde değerlendirme sonuçları Tablo 4'de sunulmaktadır.

Tablo 4. Yapay Zeka Güvenlik Farkındalığı Alt Boyutuna Ait Faktör Analizi

No	Maddeler	Özdeğer	Çıkarım
----	----------	---------	---------

1	YZ teknolojileri kişisel bilgilerimin çalınmasına yol açabilir.	1,000	,969
2	YZ tabanlı sistemlerde siber saldırı riski yüksektir.	DFA	Elendi
3	YZ güvenlik açıkları nedeniyle toplumsal riskler oluşturabilir.	1,000	,953
4	YZ kullanan kurumlara güven duymakta zorlanıyorum.	1,000	,887
5	YZ özel hayatın gizliliğini ihlal edebilir.	AFA	Elendi
6	YZ teknolojilerinin denetlenmemesi güvenlik sorunlarını artırır.	AFA	Elendi
7	YZ tabanlı cihazlar (ör. akıllı ev sistemleri) güvenlik zafiyeti yaratabilir.	AFA	Elendi
8	YZ verilerimi benim iznim olmadan paylaşabilir.	1,000	,894
9	YZ kullanımında siber güvenlik eğitimi zorunlu olmalıdır.	AFA	Elendi
10	YZ uygulamaları insan güvenliğini tehdit edebilir.	1,000	,917
11	YZ güvenlik açısından tam olarak güvenemem.	1,000	,940
12	YZ askeri alanlarda kullanıldığında küresel güvenliği tehlikeye atabilir.	1,000	,944

Tablo 4 incelendiğinde, Güvenlik Farkındalığı alt boyutu için başlangıçta hazırlanan 12 maddeden altısının ölçeğin nihai formunda yer aldığı görülmektedir. İkinci madde doğrulayıcı faktör analizi sırasında yeterli uyum sağlamadıkları için ölçekten çıkarılmıştır. Beşinci, altıncı, yedinci ve dokuzuncu maddeler ise açımlayıcı faktör analizi aşamasında düşük faktör yükü, çapraz yüklenme veya faktör yapısına yeterince katkı sağlamamaları nedeniyle elenmiştir. Sonuç olarak, üçüncü, dördüncü, sekizinci, onuncu, on birinci ve on ikinci maddeler alt boyutta korunmuştur.

Nihai alt boyutta yer alan maddelerin standartlaştırılmış faktör yükleri .887 ile .953 arasında değişmektedir. En yüksek faktör yüküne sahip madde, "YZ teknolojileri kişisel bilgilerimin çalınmasına yol açabilir." ifadesini içeren üçüncü madde (.953) olurken, en düşük faktör yükü "YZ kullanan kurumlara güven duymakta zorlanıyorum." ifadesini içeren dördüncü maddeye (.887) aittir. Tüm maddelerin faktör yüklerinin .70'in üzerinde olması, maddelerin ilgili yapıyı güçlü biçimde temsil ettiğini ve Güvenlik Farkındalığı alt boyutunun yüksek düzeyde yapı geçerliğine sahip olduğunu göstermektedir. Elde edilen bulgular, yapay zekâ teknolojilerinin güvenlik, veri gizliliği ve toplumsal risklerine ilişkin farkındalığı ölçen bu alt boyutun psikometrik açıdan yeterli özelliklere sahip olduğunu ortaya koymaktadır.

Açımlayıcı faktör analizine geçilmeden önce veri setinin faktör analizine uygunluğu Kaiser-Meyer-Olkin (KMO) Örneklem Yeterliliği Ölçüsü ve Bartlett Küresellik Testi ile değerlendirilmiştir. Elde edilen bulgular Tablo 2'de sunulmaktadır.

Tablo 5. KMO ve Bartlett küresellik testi sonuçları

Kaiser-Meyer-Olkin Örneklem Yeterliliği Ölçüsü		,947
	Approx. Chi-Square	5842,176
Bartlett Küresellik Testi	sd	28
	p	,000

Tablo 5 incelendiğinde, Kaiser-Meyer-Olkin (KMO) örneklem yeterliliği katsayısının 0,947 olduğu görülmektedir. KMO değerinin 0,90'ın üzerinde olması, örneklem büyüklüğünün faktör analizi için mükemmel düzeyde yeterli olduğunu ve değişkenler arasındaki ortak varyansın faktörleşmeye elverişli olduğunu göstermektedir (Kaiser, 1974). Ayrıca Bartlett Küresellik Testi sonuçları incelendiğinde, testin istatistiksel olarak anlamlı olduğu görülmektedir ($\chi^2 = 5842,176$; $sd = 28$; $p < .001$). Bu sonuç, korelasyon matrisinin birim matris olmadığını, maddeler arasında faktör analizi yapılmasını gerektirecek düzeyde anlamlı ilişkilerin bulunduğunu göstermektedir (Bartlett, 1954). Dolayısıyla, KMO ve Bartlett testi sonuçları birlikte değerlendirildiğinde, veri setinin açımlayıcı faktör analizi uygulanması için gerekli varsayımları karşıladığı ve faktör analizine uygun olduğu sonucuna ulaşılmıştır.

Faktör sayısının belirlenmesinde özdeğer (eigenvalue) ve açıklanan toplam varyans değerleri dikkate alınmıştır. Elde edilen bulgular Tablo 3'te sunulmaktadır.

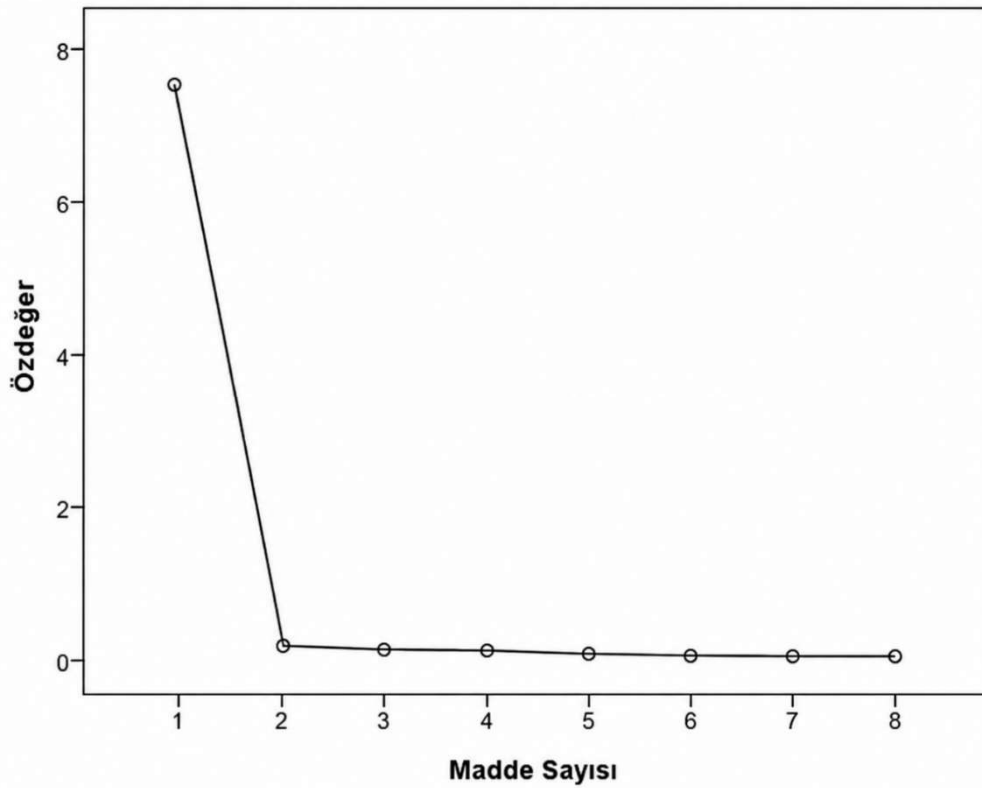
Tablo 6. Açıklanan varyans ve özdeğerler

Faktör	Başlangıç Özdeğerleri			Yüklenen Faktörlerin Karelerinin Toplamı		
	Toplam	Varyans %	Birikimli %	Toplam	Varyans %	Birikimli %
1	7,434	92,922	92,922	7,434	92,922	92,922
2	,147	1,836	94,758			
3	,124	1,546	96,304			
4	,118	1,481	97,785			
5	,067	,843	98,627			
6	,049	,611	99,238			
7	,035	,443	99,682			
8	,025	,318	100,000			

Tablo 6 incelendiğinde, yalnızca birinci faktörün özdeğerinin 1'in üzerinde olduğu ($\lambda = 7,434$), diğer faktörlerin ise özdeğerlerinin 1'in altında kaldığı görülmektedir. Kaiser ölçütüne göre (özdeğer > 1), bu sonuç ölçeğin tek faktörlü bir yapıya sahip olduğunu göstermektedir (Kaiser, 1960). Birinci faktör, toplam varyansın %92,922'sini açıklamaktadır. Sosyal bilimlerde tek faktörlü ölçeklerde açıklanan toplam varyansın %30'un üzerinde olması yeterli, %50 ve

üzerindeki değerlerin ise güçlü bir yapı geçerliğine işaret ettiği kabul edilmektedir (Hair et al., 2019). Bu doğrultuda, elde edilen %92,922'lik açıklanan varyans oranı, faktörün ilgili yapıyı oldukça yüksek düzeyde temsil ettiğini göstermektedir.

Ayrıca, yüklenen faktörlerin kareleri toplamı incelendiğinde, tek faktörlü çözümde özdeğer ve açıklanan varyans değerlerinin başlangıç özdeğerleri ile aynı olduğu görülmektedir. Bu durum, ölçeğin tek boyutlu yapısının güçlü olduğunu ve maddelerin büyük ölçüde aynı örtük yapıyı ölçtüğünü göstermektedir. Sonuç olarak, özdeğerler ve açıklanan toplam varyans bulguları birlikte değerlendirildiğinde, Güvenlik Farkındalığı alt boyutunun tek faktörlü yapısının istatistiksel olarak desteklendiği söylenebilir.



Şekil 2. Güvenlik Farkındalığı Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği

Ölçeğin faktör yapısının belirlenmesinde özdeğerlerin dağılımını incelemek amacıyla yamaç birikinti (Scree Plot) grafiği değerlendirilmiştir. Şekil 2 incelendiğinde, birinci faktörden sonra özdeğerlerde belirgin bir düşüş olduğu, ikinci faktörden itibaren ise eğrinin yatay bir seyir izlediği görülmektedir. Bu durum, Cattell'in (1966) yamaç birikinti testi ölçütüne göre ölçeğin tek faktörlü bir yapıya sahip olduğunu göstermektedir.

Grafikte birinci faktörün özdeğerinin diğer faktörlere kıyasla oldukça yüksek olduğu, ikinci ve sonraki faktörlerin ise özdeğerlerinin 1'in altında kaldığı görülmektedir. Bu bulgu, özdeğer analizi ve açıklanan toplam varyans sonuçlarıyla da tutarlılık göstermektedir. Dolayısıyla, yamaç birikinti grafiği, Güvenlik Farkındalığı alt boyutunun tek boyutlu bir yapıya sahip olduğunu destekleyen görsel bir kanıt sunmaktadır(Cattell, 1966). Elde edilen bu sonuç, alt boyutta yer alan maddelerin aynı kuramsal yapıyı ölçtüğünü ve tek faktör altında toplanmasının uygun olduğunu göstermektedir.

Açımlayıcı faktör analizi sonucunda elde edilen tek faktörlü yapının doğrulanması amacıyla doğrulayıcı faktör analizi (DFA) gerçekleştirilmiştir. İlk olarak yedi maddeden oluşan model test edilmiş, daha sonra modifikasyon indeksleri ve standartlaştırılmış faktör yükleri dikkate alınarak uygun olmadığı belirlenen bir madde çıkarılmış ve analiz altı maddelik model üzerinden tekrarlanmıştır. Elde edilen model uyum indeksleri Tablo 7'de sunulmaktadır.

Tablo 7. Model Uyum Kriterleri

Model Uyum Kriteri	İyi uyum	Kabul Edilebilir Uyum	7 Madde İndeks	Sonuç	6 Madde
χ^2/sd	0-3	3-5	2,175	İyi Uyum	1.930
AGFI	.90-1.00	.85-.90	,969	İyi Uyum	.976
GFI	.90-1.00	.85-.90	,986	İyi Uyum	.990
CFI	.95-1.00	.90-.95	,995	İyi Uyum	.997
NFI	.95-1.00	.90-.95	,992	İyi Uyum	.993
NNFI (TLI)	.95-1.00	.90-.95	,993	İyi Uyum	.995
RFI	.95-1.00	.90-.95	,986	İyi Uyum	.989
IFI	.95-1.00	.90-.95	,995	İyi Uyum	.997
RMSEA	.00-.05	.05-.08	,046	İyi Uyum	.041
RMR	.00-.05	.05-.10	,018	İyi Uyum	.018
PNFI	.95-1.00	.50-.95	,614	Kabul edilebilir uyum	.596
PGFI	.95-1.00	.50-.95	,458	-	.424
SRMR	0.00-0.05	.05-.10	0,127	İyi Uyum	.0121

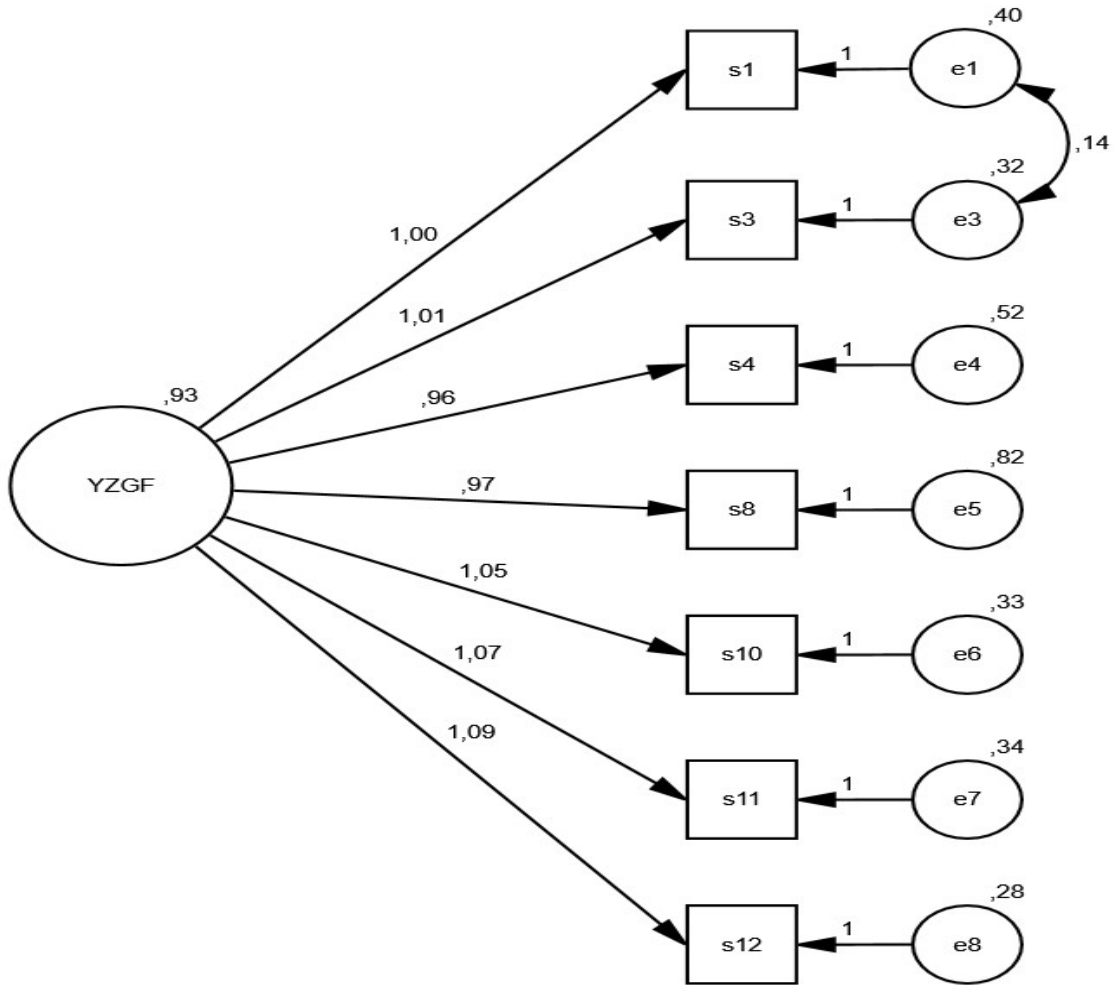
0<=RMR, SRMR

Tablo 7 incelendiğinde, yedi maddelik modele ait uyum indekslerinin genel olarak kabul edilebilir ve iyi uyum sınırları içerisinde olduğu görülmektedir. Bununla birlikte, bir maddenin modelden çıkarılması sonrasında oluşturulan altı maddelik modelde uyum indekslerinin daha da iyileştiği belirlenmiştir. Altı maddelik model için $\chi^2/sd = 1,930$, AGFI = .976, GFI = .990, CFI = .997, NFI = .993, TLI (NNFI) = .995, RFI = .989 ve IFI = .997 değerleri modelin

mükemmele yakın düzeyde uyum gösterdiğini ortaya koymaktadır. Ayrıca RMSEA = .041 ve RMR = .018 değerleri de model hatasının oldukça düşük olduğunu göstermektedir.

Parsimoni uyum indeksleri incelendiğinde PNFI = .596 değerinin kabul edilebilir düzeyde olduğu görülmektedir. PGFI değeri ise nispeten düşük olmakla birlikte, bu indeks tek başına modelin değerlendirilmesinde belirleyici bir ölçüt olarak kabul edilmemektedir. Buna karşılık diğer mutlak ve artımlı uyum indekslerinin tamamının önerilen eşik değerlerin üzerinde olması, modelin veriyle oldukça iyi uyum sağladığını göstermektedir.

Altı maddeden oluşan Güvenlik Farkındalığı alt boyutuna ilişkin doğrulayıcı faktör analizi bulguları, tek faktörlü yapının doğrulandığını ve modelin geçerli bir ölçme modeli olduğunu ortaya koymaktadır. Elde edilen uyum indeksleri, ölçeğin yapı geçerliğine ilişkin güçlü kanıtlar sunmaktadır.



Şekil 3. Güvenlik Farkındalığı alt boyutu Path Diyagramı

Yapay Zeka Güvenlik Farkındalığı (YZGF) yapısının tek boyutlu ölçme modelini test etmek amacıyla doğrulayıcı faktör analizi (DFA) uygulanmıştır. Analiz sonucunda elde edilen yol diyagramı Şekil 3'de sunulmuştur. Modelde YZGF gizil değişkeninin s1, s3, s4, s8, s10, s11 ve s12 maddeleri tarafından temsil edildiği görülmektedir. Standardize edilmemiş faktör yükleri incelendiğinde, tüm maddelere ait yüklerin ,93 ile 1,09 arasında değiştiği ve istatistiksel olarak anlamlı düzeyde yüksek olduğu saptanmıştır. Bu bulgu, ilgili maddelerin YZGF yapısını güçlü ve tutarlı biçimde ölçtüğünü göstermektedir.

Maddelere ait hata varyansları ,28 ile ,82 arasında değişmekte olup, özellikle s8 maddesine ait hata varyansının (,82) diğer maddelere kıyasla daha yüksek olduğu dikkat çekmektedir. Model uyumunu artırmak amacıyla modifikasyon indeksleri doğrultusunda e1 ve e3 hata terimleri arasında ,14 değerinde bir kovaryans ilişkisi tanımlanmıştır. Bu iki maddenin (s1 ve s3) içerik açısından kısmi bir örtüşme gösterdiği ve bu ilişkinin serbest bırakılmasının model uyumuna katkı sağladığı değerlendirilmiştir.

Elde edilen uyum iyiliği indeksleri incelendiğinde, modelin [$\chi^2/sd = 1.930$, AGFI = .976, GFI = .990, CFI = .997, NFI = .993, TLI (NNFI) = .995, RFI = .989 ve IFI = .997] değerleriyle kabul edilebilir/iyi uyum gösterdiği belirlenmiştir. Bu sonuçlar, YZGF ölçme modelinin yapı geçerliğinin sağlandığına ve önerilen tek faktörlü yapının veriyle uyumlu olduğuna işaret etmektedir.

Tablo 8. Güvenlik Farkındalığı alt boyutu Güvenirlik Analizi

N	Madde Sayısı	Cronbach's Alpha
558	7	,942

Tablo 8 incelendiğinde, Güvenlik Farkındalığı alt boyutunun 558 katılımcıdan elde edilen veriler üzerinden 7 maddeden oluştuğu ve Cronbach's Alpha güvenirlik katsayısının .942 olduğu görülmektedir. Literatürde Cronbach's Alpha katsayısının .70 ve üzeri kabul edilebilir, .80 ve üzeri yüksek, .90 ve üzeri ise çok yüksek düzeyde iç tutarlılığa işaret etmektedir (Hair et al., 2019). Bu doğrultuda elde edilen $\alpha = .942$ değeri, alt boyutta yer alan maddelerin aynı yapıyı yüksek düzeyde ve tutarlı bir biçimde ölçtüğünü göstermektedir.

Güvenlik Farkındalığı alt boyutunda yer alan maddelerin ayırt edicilik düzeylerini belirlemek amacıyla madde-toplam puan korelasyonları hesaplanmış ve ölçeğin toplam puanlarına göre

oluşturulan %27'lik alt ve üst grupların madde puanları bağımsız örneklem t-testi ile karşılaştırılmıştır. Elde edilen bulgular Tablo 9'da sunulmaktadır.

Tablo 9. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları

Madde No	Madde Toplam r	Grup	N	$\bar{X}$	Ss	df	t	p
s1	,832	Alt	151	2,31	1,06	300	21,436	,000**
		Üst	151	4,51	0,69			
s3	,878**	Alt	151	2,31	1,01	300	25,082	,000**
		Üst	151	4,65	0,53			
s4	,840**	Alt	151	2,01	0,87	300	25,342	,000**
		Üst	151	4,36	0,73			
s8	,794**	Alt	151	2,02	0,90	300	23,958	,000**
		Üst	151	4,51	0,91			
s10	,885**	Alt	151	2,15	0,92	300	28,395	,000**
		Üst	151	4,64	0,56			
s11	,888**	Alt	151	2,17	0,97	300	27,734	,000**
		Üst	151	4,66	0,52			
s12	,895**	Alt	151	2,19	0,93	300	28,688	,000**
		Üst	151	4,68	0,52			
Toplam		Alt	151	2,14	0,73	300	36,916	,000**
		Üst	151	4,58	0,36			

**p < .001

Tablo 9 incelendiğinde, Güvenlik Farkındalığı alt boyutunda yer alan maddelerin düzeltilmiş madde-toplam korelasyon katsayılarının .794 ile .895 arasında değiştiği görülmektedir. Ölçek geliştirme çalışmalarında madde-toplam korelasyon katsayısının .30 ve üzerinde olması, maddelerin ölçülen yapıyı yeterli düzeyde temsil ettiğini ve ayırt ediciliklerinin yüksek olduğunu göstermektedir (Büyüköztürk, 2020; DeVellis & Thorpe, 2021). Bu doğrultuda elde edilen korelasyon katsayıları, tüm maddelerin ilgili alt boyutu güçlü biçimde temsil ettiğini ortaya koymaktadır.

Alt ve üst %27'lik gruplar arasında yapılan bağımsız örneklem t-testi sonuçlarına göre tüm maddelerde gruplar arasında istatistiksel olarak anlamlı farklılık bulunmuştur ($p < .001$). Maddelere ilişkin t değerleri 21.436 ile 28.688 arasında değişmektedir. Ayrıca ölçeğin toplam puanı açısından da alt ve üst gruplar arasında anlamlı bir fark olduğu belirlenmiştir ($t = 36.916$, $p < .001$). Üst gruptaki katılımcıların tüm maddelerde alt gruba göre daha yüksek puan ortalamalarına sahip olması, maddelerin ölçülmek istenen özelliğe sahip bireylerle bu özelliği daha düşük düzeyde taşıyan bireyleri başarılı bir şekilde ayırt edebildiğini göstermektedir.

Madde-toplam korelasyonlarının yüksek düzeyde olması ve %27'lik alt-üst grup karşılaştırmalarında tüm maddeler için anlamlı farklılık elde edilmesi, Güvenlik Farkındalığı alt boyutunda yer alan maddelerin yüksek ayırt edicilik gücüne sahip olduğunu ve alt boyutun madde kalitesi açısından güçlü psikometrik özellikler sergilediğini göstermektedir.

4.2. Yapay Zeka Ahlaki Değerler Alt Boyutu

Ölçeğin Ahlaki Değerler alt boyutuna ilişkin psikometrik özelliklerin belirlenmesi amacıyla kapsam geçerliği, açımlayıcı faktör analizi (AFA) ve doğrulayıcı faktör analizi (DFA) sonuçları birlikte değerlendirilmiştir. Bu süreçte, öncelikle uzman görüşleri doğrultusunda kapsam geçerliği incelenmiş, ardından AFA ile maddelerin faktör yapısına uygunluğu test edilmiş ve son olarak DFA ile elde edilen faktör yapısının doğrulanması amaçlanmıştır. Analizler sonucunda, alt boyutun kuramsal yapısını en iyi temsil eden maddeler belirlenmiş, ölçme gücü yetersiz olan maddeler ise analiz sürecinden çıkarılmıştır. Ahlaki Değerler alt boyutuna ilişkin madde değerlendirme sonuçları Tablo 10'de sunulmaktadır.

Tablo 10. Yapay Zeka Ahlaki Değerler Alt Boyutuna Ait Faktör Analizi

No	Maddeler	Özdeğer	Çıkarım
13	YZ etik ilkelere göre geliştirilmelidir.	1,000	,852
14	YZ tarafsız kararlar alabilmelidir.	1,000	,797
15	YZ'nın ahlaki sorumluluğu yoktur. (Ters madde)	AFA	Elendi
16	YZ uygulamaları, adalet ilkesine aykırı olabilir.	AFA	Elendi
17	YZ etik denetim mekanizmalarına tabi olmalıdır.	1,000	,841
18	YZ toplumun ahlaki değerleriyle uyumlu kullanılmalıdır.	1,000	,859
19	YZ yanlış kullanıldığında etik sorunlara yol açabilir.	1,000	,826
20	YZ insan vicdanının yerine geçemez.	1,000	,807
21	YZ kararlarında eşitlik ilkesi gözetilmelidir.	1,000	,855
22	YZ ya sınırsız özgürlük tanımak doğru değildir.	1,000	,828
23	YZ toplumun değerlerini dönüştürebilir.	1,000	,767
24	YZ'nın sınırlarını belirlemek ahlaki bir zorunluluktur.	1,000	,837

Tablo 10 incelendiğinde, Ahlaki Değerler alt boyutu için başlangıçta hazırlanan 12 maddeden 10'unun nihai ölçekte yer aldığı görülmektedir. 15. madde ("YZ'nın ahlaki sorumluluğu yoktur.") ters madde olması ve faktör yapısıyla yeterli uyum göstermemesi nedeniyle, 16.

madde ("YZ uygulamaları, adalet ilkesine aykırı olabilir.") ise açıklayıcı faktör analizi sonucunda beklenen faktör yükünü sağlayamaması nedeniyle ölçekten çıkarılmıştır. Böylece alt boyut; 13, 14, 17, 18, 19, 20, 21, 22, 23 ve 24 numaralı maddelerden oluşan 10 maddelik bir yapı kazanmıştır.

Nihai alt boyutta yer alan maddelerin ortak varyans (communalities) değerlerinin .767 ile .859 arasında değiştiği belirlenmiştir. En yüksek ortak varyans değeri 18. maddeye ("YZ toplumun ahlaki değerleriyle uyumlu kullanılmalıdır.", .859), en düşük ortak varyans değeri ise 23. maddeye ("YZ toplumun değerlerini dönüştürebilir.", .767) aittir. Tüm maddelerin ortak varyans değerlerinin .50'nin üzerinde olması, maddelerin ilgili faktör tarafından yüksek düzeyde açıklandığını ve faktör yapısına güçlü katkı sağladığını göstermektedir (Hair et al., 2019).

Ahlaki Değerler alt boyutunda yer alan maddelerin açıklayıcı faktör analizine uygunluğunu değerlendirmek amacıyla Kaiser-Meyer-Olkin (KMO) Örneklem Yeterliliği Ölçüsü ve Bartlett Küresellik Testi uygulanmıştır. Bu analizler, veri setinin faktör analizine elverişli olup olmadığını ve maddeler arasındaki ilişkilerin ortak faktörler altında toplanabilecek düzeyde bulunup bulunmadığını belirlemek amacıyla gerçekleştirilmiştir. Elde edilen bulgular Tablo 11'de sunulmaktadır.

Tablo 11. KMO ve Bartlett küresellik testi sonuçları

Kaiser-Meyer-Olkin Örneklem Yeterliliği Ölçüsü		,947
Bartlett Küresellik Testi	Approx. Chi-Square	4845,059
	sd	45
	p	,000

Tablo 11 incelendiğinde, Kaiser-Meyer-Olkin (KMO) örneklem yeterliliği katsayısının 0,947 olduğu görülmektedir. KMO değerinin 0,90'ın üzerinde olması, örneklem büyüklüğünün faktör analizi için mükemmel düzeyde yeterli olduğunu ve maddeler arasındaki ortak varyansın faktör analizi yapılmasına uygun olduğunu göstermektedir (Kaiser, 1974).

Bartlett Küresellik Testi sonuçları incelendiğinde ise testin istatistiksel olarak anlamlı olduğu belirlenmiştir ($\chi^2 = 4845,059$; $sd = 45$; $p < .001$). Bu sonuç, korelasyon matrisinin birim matris olmadığını ve maddeler arasında anlamlı düzeyde ilişkiler bulunduğunu göstermektedir (Bartlett, 1954). Başka bir ifadeyle, maddeler arasındaki korelasyonların faktör analizi uygulanmasını destekleyecek düzeyde olduğu söylenebilir.

KMO ve Bartlett Küresellik Testi sonuçları birlikte değerlendirildiğinde, Ahlaki Değerler alt boyutuna ilişkin veri setinin açıklayıcı faktör analizinin uygulanması için gerekli varsayımları sağladığı ve faktör analizine uygun olduğu sonucuna ulaşılmıştır.

Ahlaki Değerler alt boyutunun faktör yapısını değerlendirmek amacıyla açıklayıcı faktör analizi kapsamında özdeğerler (eigenvalues) ve açıklanan toplam varyans değerleri incelenmiştir. Faktör sayısının belirlenmesinde Kaiser'in özdeğer ölçütü (özdeğer > 1) esas alınmış ve faktörlerin açıkladığı toplam varyans oranları değerlendirilmiştir. Elde edilen bulgular Tablo 12'de sunulmaktadır.

Tablo 12. Açıklanan varyans ve özdeğerler

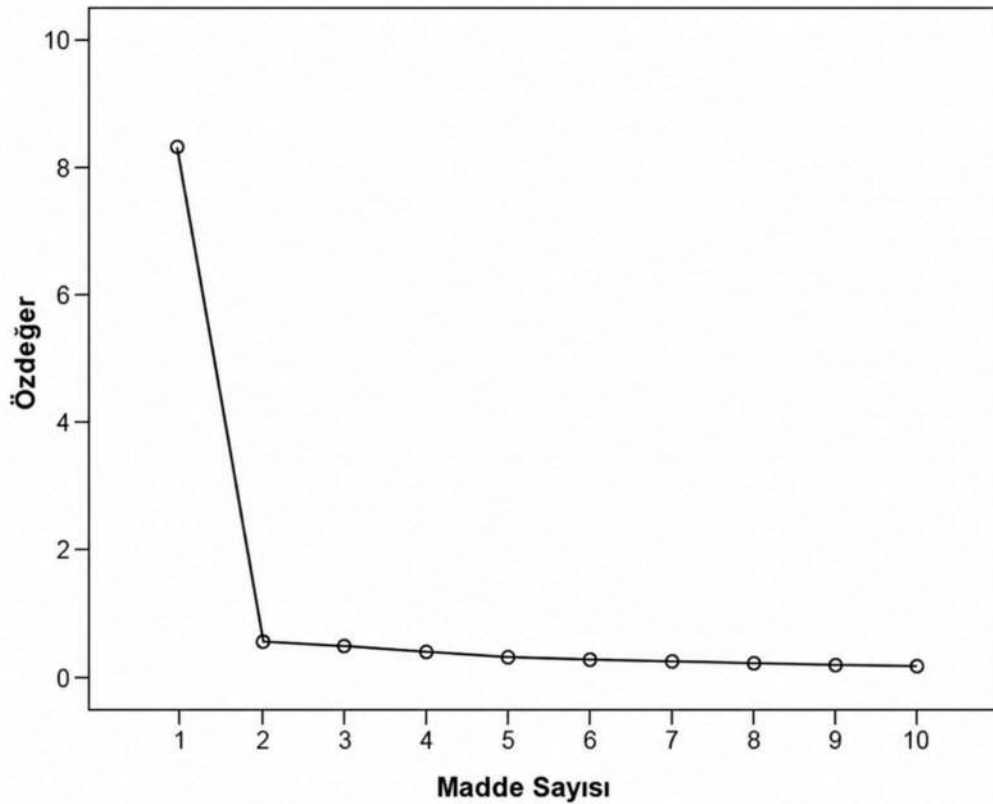
Faktör	Başlangıç Özdeğerleri			Yüklenen Faktörlerin Karelerinin Toplamı		
	Toplam	Varyans %	Birikimli %	Toplam	Varyans %	Birikimli %
1	8,269	82,687	82,687	8,269	82,687	82,687
2	,387	3,867	86,554			
3	,336	3,356	89,910			
4	,257	2,571	92,481			
5	,181	1,812	94,293			
6	,150	1,504	95,797			
7	,137	1,369	97,166			
8	,118	1,179	98,345			
9	,091	,908	99,253			
10	,075	,747	100,000			

Tablo 12 incelendiğinde, yalnızca birinci faktörün özdeğerinin 1'in üzerinde olduğu ($\lambda = 8,269$), diğer dokuz faktörün özdeğerlerinin ise 1'in altında kaldığı görülmektedir. Kaiser ölçütüne göre bu bulgu, Ahlaki Değerler alt boyutunun tek faktörlü bir yapıya sahip olduğunu göstermektedir (Kaiser, 1960).

Birinci faktörün açıkladığı varyans oranı %82,687 olarak belirlenmiştir. Sosyal bilimlerde tek faktörlü ölçeklerde açıklanan toplam varyansın %30'un üzerinde olması yeterli, %50'nin üzerinde olması ise güçlü bir faktör yapısına işaret etmektedir (Hair et al., 2019). Bu doğrultuda elde edilen %82,687'lik açıklanan varyans oranı, tek faktörün maddelerdeki toplam varyansın büyük bölümünü açıkladığını ve alt boyutun güçlü bir yapısal bütünlüğe sahip olduğunu göstermektedir.

Yüklenen faktörlerin kareleri toplamı incelendiğinde, tek faktörlü çözümde başlangıç özdeğeri ile açıklanan varyans değerlerinin aynı olduğu görülmektedir. Bu durum, Ahlaki Değerler alt boyutunda yer alan maddelerin büyük ölçüde aynı örtük yapıyı temsil ettiğini ve faktör yapısının istatistiksel açıdan güçlü olduğunu göstermektedir.

Ahlaki Değerler alt boyutunun faktör sayısını görsel olarak değerlendirmek amacıyla yamaç birikinti (Scree Plot) grafiği incelenmiştir. Şekil 4'te özdeğerlerin faktörlere göre dağılımı gösterilmektedir.



Şekil 4. Ahlaki Değerler Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği

Şekil 4 incelendiğinde, birinci faktörden sonra özdeğerlerde belirgin ve keskin bir düşüş olduğu, ikinci faktörden itibaren ise eğrinin yatay bir seyir izlediği görülmektedir. Cattell'in (1966) yamaç birikinti testi yaklaşımına göre, eğrinin kırılma noktasının birinci faktörde oluşması, Ahlaki Değerler alt boyutunun tek faktörlü bir yapıya sahip olduğunu göstermektedir.

Scree Plot grafiğinde gözlenen bu yapı, özdeğer analizi ve açıklanan toplam varyans bulgularıyla da tutarlılık göstermektedir. Nitekim birinci faktörün özdeğeri 8,269 olup toplam varyansın %82,687'sini açıklarken, ikinci ve sonraki faktörlerin özdeğerlerinin 1'in altında kalması, ek faktörlerin anlamlı bir katkı sağlamadığını ortaya koymaktadır. Bu durum, alt

boyutta yer alan maddelerin büyük ölçüde aynı örtük yapıyı ölçtüğünü ve tek faktör altında toplandığını göstermektedir.

Açımlayıcı faktör analizi sonucunda elde edilen Ahlaki Değerler alt boyutunun tek faktörlü yapısının doğrulanması amacıyla doğrulayıcı faktör analizi (DFA) gerçekleştirilmiştir. DFA kapsamında elde edilen model uyum indeksleri incelenerek ölçüm modelinin veri ile uyumu değerlendirilmiştir. Analiz sonucunda elde edilen model uyum değerleri Tablo 13'de sunulmaktadır.

Tablo 13. Model Uyum Kriterleri

Model Uyum Kriteri	İyi uyum	Kabul Edilebilir Uyum	İndeks	Sonuç
χ^2/sd	0-3	3-5	4,105	Kabul edilebilir uyum
AGFI	.90-1.00	.85-.90	,919	İyi uyum
GFI	.90-1.00	.85-.90	,952	İyi uyum
CFI	.95-1.00	.90-.95	,983	İyi uyum
NFI	.95-1.00	.90-.95	,977	İyi uyum
NNFI (TLI)	.95-1.00	.90-.95	,976	İyi uyum
RFI	.95-1.00	.90-.95	,969	İyi uyum
IFI	.95-1.00	.90-.95	,983	İyi uyum
RMSEA	.00-.05	.05-.08	,075	Kabul edilebilir uyum
RMR	.00-.05	.05-.10	,026	İyi uyum
PNFI	.95-1.00	.50-.95	,717	Kabul edilebilir uyum
PGFI	.95-1.00	.50-.95	,571	Kabul edilebilir uyum
SRMR	0.00-0.05	.05-.10	,0177	İyi uyum

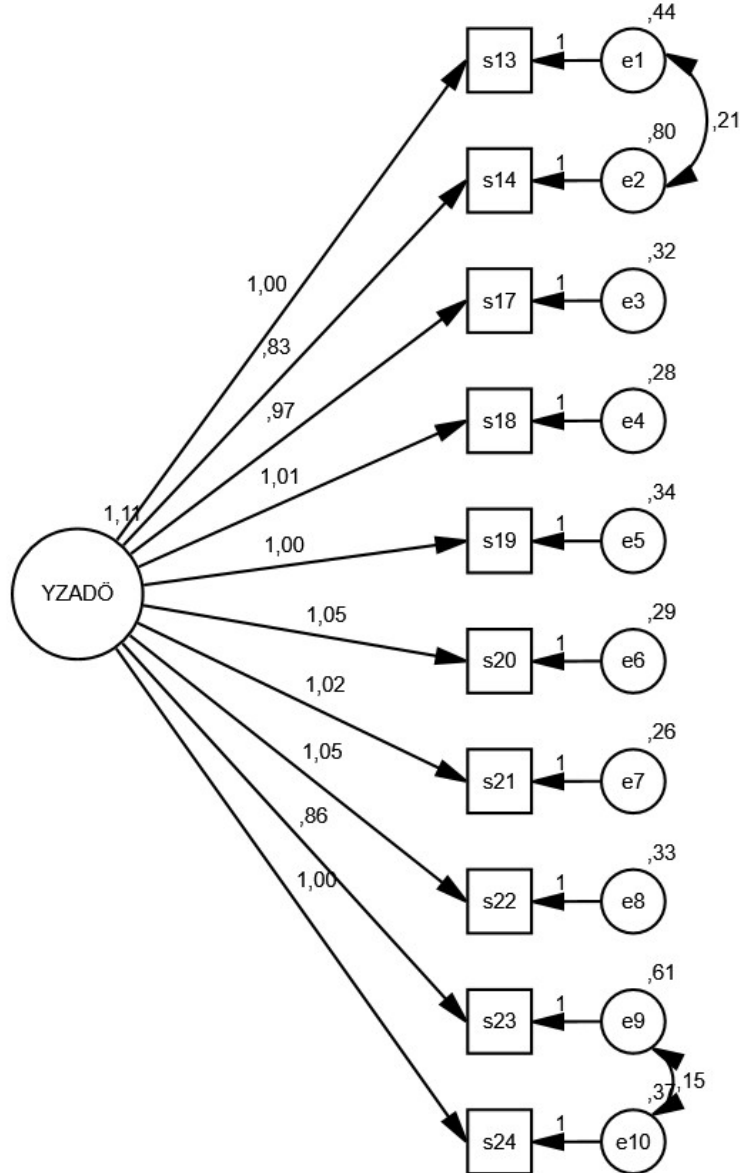
0<=RMR, SRMR

Tablo 13 incelendiğinde, Ahlaki Değerler alt boyutuna ilişkin doğrulayıcı faktör analizi sonucunda elde edilen model uyum indekslerinin genel olarak önerilen sınırlar içerisinde olduğu görülmektedir. Modelin χ^2/sd değerinin 4,105 olması, modelin kabul edilebilir uyum düzeyinde olduğunu göstermektedir. Bununla birlikte AGFI (.919), GFI (.952), CFI (.983), NFI (.977), NNFI/TLI (.976), RFI (.969) ve IFI (.983) değerlerinin tamamı .90'ın üzerinde, büyük çoğunluğunun ise .95'in üzerinde olması, modelin veriyle oldukça yüksek düzeyde uyum sağladığını ortaya koymaktadır.

Hata uyum indeksleri incelendiğinde, RMSEA değerinin .075 olması modelin kabul edilebilir uyum gösterdiğini, RMR (.026) ve SRMR (.0177) değerlerinin ise .05'in altında olması nedeniyle model hatasının oldukça düşük olduğunu göstermektedir. Ayrıca parsimoni uyum

indekslerinden PNFI (.717) ve PGFI (.571) değerlerinin de kabul edilebilir sınırlar içerisinde yer alması, modelin yalınlık ve açıklayıcılık açısından yeterli düzeyde olduğunu desteklemektedir.

Şekil 4'te, Ahlaki Değerler alt boyutuna ilişkin doğrulayıcı faktör analizi (DFA) sonucunda elde edilen tek faktörlü ölçüm modelinin path diyagramı sunulmaktadır. Diyagramda, gizil değişken olan Ahlaki Değerler faktörünün (YZADÖ) gözlenen değişkenler üzerindeki etkileri ile maddelere ait hata varyansları gösterilmektedir.



Şekil 5. Ahlaki Değerler alt boyutu Path Diyagramı

Path diyagramı incelendiğinde, alt boyutta yer alan 10 maddenin tamamının Ahlaki Değerler gizil değişkeni ile pozitif yönde ilişki gösterdiği görülmektedir. Maddelere ait standartlaştırılmış faktör yüklerinin yüksek düzeyde olması, her bir maddenin ilgili gizil yapıyı güçlü biçimde temsil ettiğini göstermektedir. Bu bulgu, açıklayıcı faktör analizi sonucunda elde edilen tek faktörlü yapının doğrulayıcı faktör analizi ile de desteklendiğini ortaya koymaktadır.

Diyagramda ayrıca maddelere ait hata varyanslarının genel olarak düşük düzeyde olduğu görülmektedir. Düşük hata varyansları, maddelerin açıkladıkları ortak varyansın yüksek olduğunu ve ölçme hatasının sınırlı düzeyde kaldığını göstermektedir. Bununla birlikte, model uyumunu artırmak amacıyla e1–e2 ve e9–e10 hata terimleri arasında kovaryans tanımlanmıştır. Bu kovaryanslar, içerik bakımından benzer özellik taşıyan maddelerin ortak varyanslarının model tarafından açıklanamayan bölümünü temsil etmekte olup, kuramsal gerekçeler doğrultusunda modele dâhil edilmiştir. Yapılan bu düzenleme sonrasında model uyum indekslerinin kabul edilebilir ve iyi uyum sınırları içerisinde yer aldığı belirlenmiştir.

Ahlaki Değerler alt boyutunun iç tutarlılığını belirlemek amacıyla Cronbach's Alpha güvenirlik katsayısı hesaplanmıştır. Analiz sonucunda elde edilen bulgular Tablo 14'de sunulmaktadır.

Tablo 14. Ahlaki Değerler alt boyutu Güvenirlik Analizi

N	Madde Sayısı	Cronbach's Alpha
558	10	,953

Tablo 14 incelendiğinde, Ahlaki Değerler alt boyutunun 558 katılımcıdan elde edilen veriler doğrultusunda 10 maddeden oluştuğu ve Cronbach's Alpha güvenirlik katsayısının .953 olduğu görülmektedir. Literatürde Cronbach's Alpha katsayısının .70 ve üzeri kabul edilebilir, .80 ve üzeri yüksek, .90 ve üzeri ise çok yüksek düzeyde iç tutarlılığa işaret etmektedir (Hair et al., 2019). Bu doğrultuda elde edilen $\alpha = .953$ değeri, alt boyutta yer alan maddelerin aynı yapıyı oldukça tutarlı biçimde ölçtüğünü ve ölçeğin yüksek düzeyde iç tutarlılığa sahip olduğunu göstermektedir.

Cronbach's Alpha katsayısının .95 değerine oldukça yakın olması, maddeler arasında yüksek düzeyde tutarlılık bulunduğunu göstermektedir. Bununla birlikte, maddelerin içerikleri

incelendiğinde etik ilkeler, adalet, toplumsal değerler, insan sorumluluğu ve yapay zekânın etik kullanımına ilişkin farklı yönleri kapsadığı görülmektedir. Bu durum, yüksek güvenirlik katsayısının maddelerin birbirini tekrar etmesinden ziyade aynı kuramsal yapıyı bütüncül biçimde ölçmelerinden kaynaklandığını göstermektedir.

Ahlaki Değerler alt boyutunda yer alan maddelerin ayırt edicilik düzeylerini belirlemek amacıyla düzeltilmiş madde-toplam puan korelasyonları hesaplanmış, ayrıca ölçeğin toplam puanlarına göre oluşturulan %27'lik alt ve üst grupların madde puan ortalamaları bağımsız örneklem t-testi ile karşılaştırılmıştır. Bu analizler, maddelerin ölçülmek istenen yapıyı temsil etme ve farklı düzeylerde özelliğe sahip bireyleri ayırt edebilme gücünü değerlendirmek amacıyla gerçekleştirilmiştir. Elde edilen bulgular Tablo 15'de sunulmaktadır.

Tablo 15. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları

Madde No	Madde Toplam r	Grup	N	$\bar{X}$	Ss	df	t	p
s13	,850**	Alt	151	2,33	1,18	300	22,227	,000**
		Üst	151	4,75	0,62			
s14	,758**	Alt	151	2,42	1,13	300	18,405	,000**
		Üst	151	4,57	0,89			
s17	,891**	Alt	151	2,29	0,98	300	28,203	,000**
		Üst	151	4,80	0,48			
s18	,906**	Alt	151	2,30	0,99	300	29,687	,000**
		Üst	151	4,88	0,40			
s19	,889**	Alt	151	2,27	1,01	300	29,193	,000**
		Üst	151	4,85	0,41			
s20	,906**	Alt	151	2,34	1,09	300	26,972	,000**
		Üst	151	4,89	0,40			
s21	,912**	Alt	151	2,34	1,08	300	27,504	,000**
		Üst	151	4,88	0,35			
s22	,895**	Alt	151	2,28	1,00	300	30,505	,000**
		Üst	151	4,93	0,38			
s23	,799**	Alt	151	2,29	0,94	300	21,625	,000**
		Üst	151	4,56	0,88			
s24	,861**	Alt	151	2,25	1,01	300	25,801	,000**
		Üst	151	4,74	0,61			
Toplam		Alt	151	2,31	0,84	300	35,538	,000**
		Üst	151	4,80	0,19			

**p < .001

Tablo 15 incelendiğinde, Ahlaki Değerler alt boyutunda yer alan maddelerin düzeltilmiş madde-toplam korelasyon katsayılarının .758 ile .912 arasında değiştiği görülmektedir. En yüksek madde-toplam korelasyonu s21 maddesinde ($r = .912$), en düşük korelasyon ise s14

maddesinde ($r = .758$) elde edilmiştir. Ölçek geliştirme çalışmalarında madde-toplam korelasyon katsayısının .30 ve üzerinde olması, maddelerin ölçülen yapıyı yeterli düzeyde temsil ettiğini ve yüksek ayırt edicilik gücüne sahip olduğunu göstermektedir (Büyüköztürk, 2020; DeVellis & Thorpe, 2021). Bu doğrultuda, tüm maddelerin oldukça yüksek madde-toplam korelasyonlarına sahip olması, her bir maddenin Ahlaki Değerler alt boyutuna güçlü katkı sağladığını ortaya koymaktadır.

Alt ve üst %27'lik gruplar arasında yapılan bağımsız örneklem t-testi sonuçları incelendiğinde, tüm maddeler için gruplar arasında istatistiksel olarak anlamlı farklılıklar bulunduğu belirlenmiştir ($p < .001$). Maddelere ilişkin t değerleri 18.405 ile 30.505 arasında değişmektedir. Ayrıca ölçeğin toplam puanı açısından da alt ve üst gruplar arasında anlamlı bir fark olduğu görülmektedir ($t = 35.538$, $p < .001$). Üst grupta yer alan katılımcıların tüm maddelerde alt gruba göre daha yüksek puan ortalamalarına sahip olması, maddelerin ahlaki değerlere ilişkin algı düzeyi yüksek olan bireylerle düşük olan bireyleri başarılı bir şekilde ayırt edebildiğini göstermektedir.

4.3. Yapay Zeka Sorumluluk ve Hesap Verilebilirlik Alt Boyutu

Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutunun psikometrik özelliklerini belirlemek amacıyla kapsam geçerliği, açımlayıcı faktör analizi (AFA) ve doğrulayıcı faktör analizi (DFA) sonuçları birlikte değerlendirilmiştir. Bu süreçte öncelikle uzman görüşleri doğrultusunda kapsam geçerliği incelenmiş, ardından AFA ile maddelerin faktör yapısına uygunluğu test edilmiş ve son olarak DFA ile elde edilen yapının doğrulanması amaçlanmıştır. Analizler sonucunda, alt boyutun kuramsal yapısını en iyi temsil eden maddeler belirlenmiş, ölçme modeline yeterli katkı sağlamayan maddeler ise ölçekten çıkarılmıştır. Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutuna ilişkin faktör analizi sonuçları Tablo 16'da sunulmaktadır.

Tablo 16. Yapay Zeka Sorumluluk ve Hesap Verilebilirlik Alt Boyutuna Ait Faktör Analizi

No	Maddeler	Özdeğer	Çıkarım
25	YZ sistemlerinin hatalarından geliştiriciler sorumludur.	1,000	,807
26	YZ nedeniyle oluşacak zararlarda yasal yaptırımlar olmalıdır.	AFA	Elendi
27	YZ kararları şeffaf biçimde açıklanmalıdır.	1,000	,871
28	YZ denetimsiz bırakıldığında hesap verebilirlik azalır.	1,000	,818
29	YZ kaynaklı hatalarda kullanıcılar suçlanmamalıdır.	1,000	,688

30	YZ kullanımında hukuki sorumluluk net tanımlanmalıdır.	1,000	,872
31	YZ yanlış karar verdiğinde, sorumluluk kime ait olduğu açık olmalıdır.	1,000	,876
32	YZ uygulamalarının kimler tarafından geliştirildiği şeffaf olmalıdır.	1,000	,877
33	YZ için uluslararası etik standartlar gereklidir.	AFA	Elendi
34	YZ'nın askeri alandaki kullanımı özel bir sorumluluk doğurur.	AFA	Elendi
35	YZ uygulamalarında hesap verebilirlik yetersizdir.	1,000	,822
36	YZ insan gözetimi olmadan karar vermemelidir.	AFA	Elendi

Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutunun psikometrik özelliklerini belirlemek amacıyla kapsam geçerliği, açımlayıcı faktör analizi (AFA) ve doğrulayıcı faktör analizi (DFA) sonuçları birlikte değerlendirilmiştir. Bu kapsamda, öncelikle uzman görüşleri doğrultusunda kapsam geçerliği sağlanan maddeler açımlayıcı faktör analizine tabi tutulmuş, ardından elde edilen faktör yapısı doğrulayıcı faktör analizi ile test edilmiştir. Analiz sürecinde, faktör yapısına yeterli katkı sağlamayan maddeler ölçekten çıkarılarak alt boyutun nihai yapısı oluşturulmuştur. Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutuna ilişkin faktör analizi sonuçları Tablo 16'te sunulmaktadır.

Tablo 16 incelendiğinde, başlangıçta 12 maddeden oluşan Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutunda yer alan 26, 33, 34 ve 36 numaralı maddelerin açımlayıcı faktör analizi sonucunda faktör yapısına yeterli katkı sağlamadığı belirlenmiş ve ölçekten çıkarılmıştır. Böylece alt boyut; 25, 27, 28, 29, 30, 31, 32 ve 35 numaralı maddelerden oluşan 8 maddelik tek faktörlü bir yapıya kavuşmuştur.

Nihai ölçekte yer alan maddelerin ortak varyans (communalities) değerleri .688 ile .877 arasında değişmektedir. En yüksek ortak varyans değeri 32. maddede ("YZ uygulamalarının kimler tarafından geliştirildiği şeffaf olmalıdır.", .877), en düşük ortak varyans değeri ise 29. maddede ("YZ kaynaklı hatalarda kullanıcılar suçlanmamalıdır.", .688) elde edilmiştir. Ortak varyans değerlerinin tamamının .50'nin üzerinde olması, maddelerin ilgili faktör tarafından yüksek düzeyde açıklandığını ve tek faktörlü yapıyı güçlü biçimde temsil ettiğini göstermektedir (Hair et al., 2019).

Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutunda yer alan maddelerin açımlayıcı faktör analizine uygunluğunu değerlendirmek amacıyla Kaiser-Meyer-Olkin (KMO) Örneklem

Yeterliliği Ölçüsü ile Bartlett Küresellik Testi uygulanmıştır. Bu analizler, veri setinin faktör analizine uygunluğunu belirlemek ve maddeler arasındaki korelasyonların ortak faktörler altında toplanabilecek düzeyde olup olmadığını incelemek amacıyla gerçekleştirilmiştir. Elde edilen bulgular Tablo 17'de sunulmaktadır.

Tablo 17. KMO ve Bartlett küresellik testi sonuçları

Kaiser-Meyer-Olkin Örneklem Yeterliliği Ölçüsü		,956
Bartlett Küresellik Testi	Approx. Chi-Square	5768,587
	sd	55
	p	,000

Tablo 17 incelendiğinde, Kaiser-Meyer-Olkin (KMO) örneklem yeterliliği katsayısının 0,956 olduğu görülmektedir. KMO değerinin 0,90'ın üzerinde olması, örneklem büyüklüğünün faktör analizi için mükemmel düzeyde yeterli olduğunu ve maddeler arasındaki ortak varyansın faktör analizi uygulanmasına son derece elverişli olduğunu göstermektedir (Kaiser, 1974).

Bartlett Küresellik Testi sonuçlarına göre test istatistiği $\chi^2 = 5768,587$ (sd = 55; $p < .001$) olarak bulunmuştur. Test sonucunun istatistiksel olarak anlamlı olması, korelasyon matrisinin birim matris olmadığını ve maddeler arasında faktör analizi yapılmasını gerektirecek düzeyde anlamlı ilişkiler bulunduğunu göstermektedir (Bartlett, 1954).

Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutunun faktör yapısını belirlemek amacıyla açımlayıcı faktör analizi kapsamında özdeğerler (eigenvalues) ve açıklanan toplam varyans değerleri incelenmiştir. Faktör sayısının belirlenmesinde Kaiser'in özdeğer ölçütü (özdeğer > 1) esas alınmış ve faktörlerin açıkladığı toplam varyans oranları değerlendirilmiştir. Analiz sonucunda elde edilen bulgular Tablo 18'de sunulmaktadır.

Tablo 18. Açıklanan varyans ve özdeğerler

Faktör	Başlangıç Özdeğerleri			Yüklenen Faktörlerin Karelerinin Toplamı		
	Toplam	Varyans %	Birikimli %	Toplam	Varyans %	Birikimli %
1	9,220	83,819	83,819	9,220	83,819	83,819
2	,394	3,579	87,398			
3	,277	2,520	89,917			
4	,232	2,110	92,027			
5	,222	2,015	94,043			
6	,184	1,670	95,713			
7	,140	1,273	96,986			

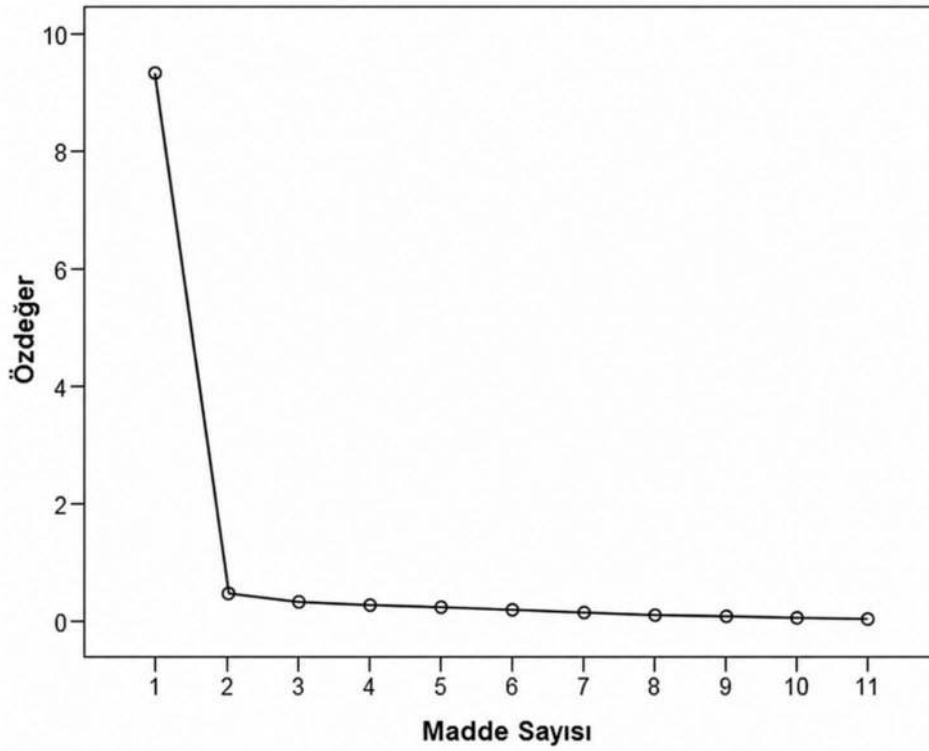
8	,104	,948	97,934
9	,100	,910	98,844
10	,066	,604	99,448
11	,061	,552	100,000

Tablo 18 incelendiğinde, yalnızca birinci faktörün özdeğerinin 1'in üzerinde olduğu ($\lambda = 9,220$), diğer on faktörün özdeğerlerinin ise 1'in altında kaldığı görülmektedir. Kaiser'in özdeğer ölçütüne göre bu sonuç, Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutunun tek faktörlü bir yapıya sahip olduğunu göstermektedir (Kaiser, 1960).

Birinci faktörün açıkladığı varyans oranının %83,819 olduğu belirlenmiştir. Sosyal bilimlerde tek faktörlü yapılarda açıklanan toplam varyansın %30'un üzerinde olması yeterli, %50'nin üzerinde olması ise güçlü bir faktör yapısına işaret etmektedir (Hair et al., 2019). Bu bağlamda elde edilen %83,819'luk açıklanan varyans oranı, tek faktörün maddelerdeki toplam varyansın büyük bölümünü açıkladığını ve alt boyutun güçlü bir yapısal bütünlüğe sahip olduğunu göstermektedir.

Yüklenen faktörlerin kareleri toplamı incelendiğinde, tek faktörlü çözümde başlangıç özdeğeri ile açıklanan varyans değerlerinin aynı olduğu görülmektedir. Bu durum, alt boyutta yer alan maddelerin aynı örtük yapıyı yüksek düzeyde temsil ettiğini ve faktör yapısının istatistiksel olarak güçlü olduğunu desteklemektedir.

Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutunun faktör sayısını görsel olarak belirlemek amacıyla yamaç birikinti (Scree Plot) grafiği incelenmiştir. Şekil 7'de özdeğerlerin faktörlere göre dağılımı gösterilmektedir.



Şekil 6. Sorumluluk ve Hesap Verilebilirlik Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği

Şekil 6 incelendiğinde, birinci faktörden sonra özdeğerlerde belirgin ve keskin bir düşüş olduğu, ikinci faktörden itibaren ise eğrinin yatay bir görünüm sergilediği görülmektedir. Bu durum, Cattell'in (1966) yamaç birikinti testi yaklaşımına göre kırılma noktasının birinci faktörde gerçekleştiğini ve alt boyutun tek faktörlü bir yapıya sahip olduğunu göstermektedir.

Yamaç birikinti grafiğinde elde edilen bu bulgu, özdeğer ve açıklanan toplam varyans sonuçlarıyla da tutarlıdır. Nitekim birinci faktörün özdeğeri 9,220 olup toplam varyansın %83,819'unu açıklamaktadır. Buna karşın ikinci ve sonraki faktörlerin özdeğerlerinin 1'in altında kalması, bu faktörlerin ölçülen yapıya anlamlı bir katkı sağlamadığını göstermektedir. Dolayısıyla alt boyutta yer alan maddelerin ortak bir gizil yapıyı temsil ettiği ve tek faktör altında toplandığı anlaşılmaktadır.

Açımlayıcı faktör analizi sonucunda elde edilen Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutunun tek faktörlü yapısının doğrulanması amacıyla doğrulayıcı faktör analizi (DFA) gerçekleştirilmiştir. DFA kapsamında oluşturulan ölçüm modelinin veriyle uyumu, mutlak uyum, artımlı uyum ve parsimoni uyum indeksleri dikkate alınarak değerlendirilmiştir. Elde edilen model uyum indeksleri Tablo 19'da sunulmaktadır.

Tablo 19. Model Uyum Kriterleri

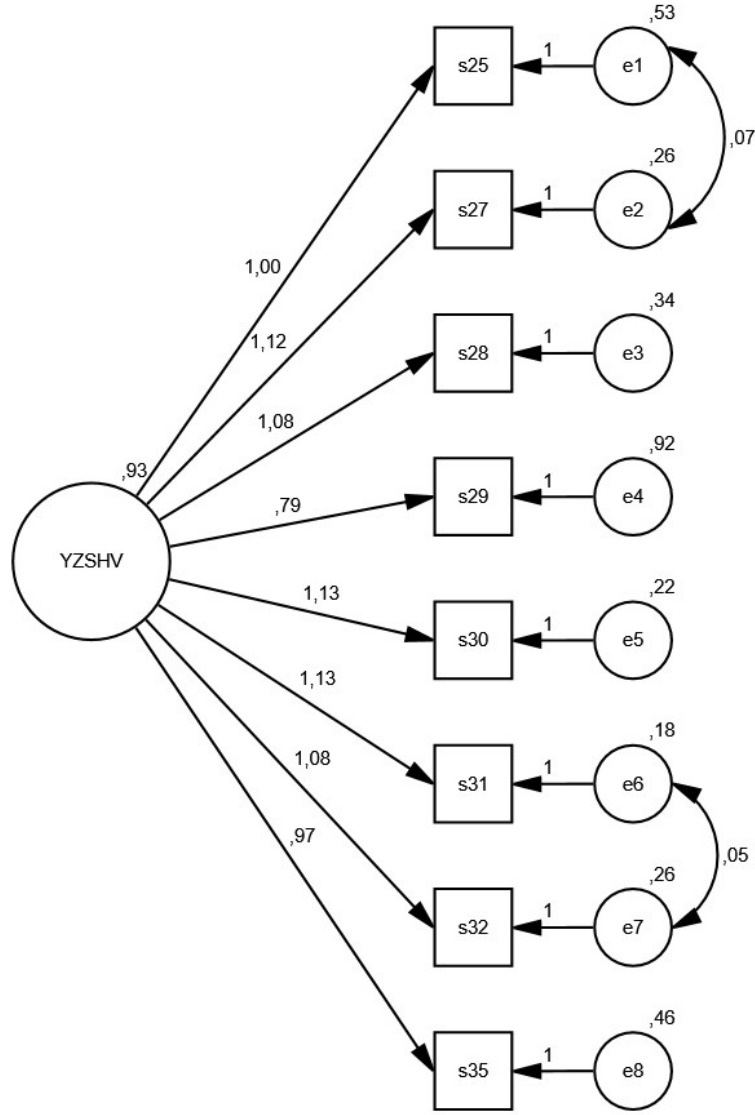
Model Uyum Kriteri	İyi uyum	Kabul Edilebilir Uyum	İndeks	Sonuç
χ^2/sd	0-3	3-5	4,457	Kabul edilebilir uyum
AGFI	.90-1.00	.85-.90	,933	İyi uyum
GFI	.90-1.00	.85-.90	,967	İyi uyum
CFI	.95-1.00	.90-.95	,986	İyi uyum
NFI	.95-1.00	.90-.95	,982	İyi uyum
NNFI (TLI)	.95-1.00	.90-.95	,979	İyi uyum
RFI	.95-1.00	.90-.95	,973	İyi uyum
IFI	.95-1.00	.90-.95	,986	İyi uyum
RMSEA	.00-.05	.05-.08	,079	Kabul edilebilir uyum
RMR	.00-.05	.05-.10	,025	İyi uyum
PNFI	.95-1.00	.50-.95	,932	Kabul edilebilir uyum
PGFI	.95-1.00	.50-.95	,483	Kabul edilebilir uyum
SRMR	0.00-0.05	.05-.10	,0181	İyi uyum

0<=RMR, SRMR

Tablo 19 incelendiğinde, doğrulayıcı faktör analizi sonucunda elde edilen model uyum indekslerinin genel olarak önerilen kabul ölçütlerini karşıladığı görülmektedir. Modelin χ^2/sd değerinin 4,457 olması, modelin kabul edilebilir uyum düzeyinde olduğunu göstermektedir. Bununla birlikte AGFI (.933), GFI (.967), CFI (.986), NFI (.982), NNFI/TLI (.979), RFI (.973) ve IFI (.986) değerlerinin tamamı önerilen eşik değerlerin üzerinde olup modelin veriyle iyi düzeyde uyum sağladığını ortaya koymaktadır.

Mutlak uyum indekslerinden RMSEA değerinin .079 olması modelin kabul edilebilir uyum düzeyinde olduğunu gösterirken, RMR (.025) ve SRMR (.0181) değerlerinin .05'in altında olması modeldeki hata düzeyinin oldukça düşük olduğunu ve gözlenen veriler ile önerilen model arasında güçlü bir uyum bulunduğunu göstermektedir. Parsimoni uyum indeksleri incelendiğinde ise PNFI (.932) ve PGFI (.483) değerlerinin modelin yalınlığı açısından kabul edilebilir düzeyde olduğu görülmektedir. Özellikle PNFI değerinin yüksek olması, modelin açıklayıcılık ile yalınlık arasında başarılı bir denge kurduğunu göstermektedir.

Şekil 7'de, Yapay Zekâ Sorumluluk ve Hesap Verebilirlik alt boyutuna ilişkin doğrulayıcı faktör analizi (DFA) sonucunda elde edilen tek faktörlü ölçüm modelinin path diyagramı sunulmaktadır.



Şekil 7. Sorumluluk ve Hesap Verilebilirlik alt boyutu Path Diyagramı

Şekil 7 incelendiğinde, alt boyutta yer alan sekiz maddenin (s25, s27, s28, s29, s30, s31, s32 ve s35) tamamının Yapay Zekâ Sorumluluk ve Hesap Verebilirlik gizil değişkenini pozitif yönde temsil ettiği görülmektedir. Maddelere ait faktör yüklerinin yüksek olması, her bir maddenin ölçülmek istenen yapıyı güçlü biçimde yansıttığını ve tek faktörlü yapıya önemli katkı sağladığını göstermektedir. Bu bulgu, açımlayıcı faktör analizi sonucunda elde edilen tek boyutlu yapının doğrulayıcı faktör analizi ile de desteklendiğini ortaya koymaktadır.

Diyagramda maddelere ait hata varyanslarının genel olarak düşük düzeyde olduğu görülmektedir. Bu durum, maddelerin ortak faktör tarafından büyük ölçüde açıklandığını ve ölçme hatasının sınırlı olduğunu göstermektedir. Model uyumunu artırmak amacıyla e1–e2 ile e6–e7 hata terimleri arasında kovaryans tanımlanmıştır. Bu kovaryanslar, içerik bakımından

benzer özellikler taşıyan maddelerin ortak varyanslarının model tarafından açıklanamayan kısmını temsil etmekte olup, kuramsal gerekçeler doğrultusunda modele dâhil edilmiştir. Yapılan bu düzenlemeler sonrasında modelin uyum indekslerinin kabul edilebilir ve iyi uyum sınırları içerisinde olduğu belirlenmiştir.

Sorumluluk ve Hesap Verilebilirlik alt boyutunun iç tutarlılığını belirlemek amacıyla Cronbach's Alpha güvenirlik katsayısı hesaplanmıştır. İç tutarlılık analizi, alt boyutta yer alan maddelerin aynı yapıyı ölçme konusundaki tutarlılığını değerlendirmek ve ölçeğin güvenirliğini ortaya koymak amacıyla gerçekleştirilmiştir. Analiz sonucunda elde edilen bulgular Tablo 20'de sunulmaktadır.

Tablo 20. Sorumluluk ve Hesap Verilebilirlik alt boyutu Güvenirlik Analizi

N	Madde Sayısı	Cronbach's Alpha
558	8	,952

Tablo 20 incelendiğinde, Sorumluluk ve Hesap Verilebilirlik alt boyutunun 558 katılımcıdan elde edilen veriler doğrultusunda 8 maddeden oluştuğu ve Cronbach's Alpha güvenirlik katsayısının .952 olduğu görülmektedir. Literatürde Cronbach's Alpha katsayısının .70 ve üzeri kabul edilebilir, .80 ve üzeri yüksek, .90 ve üzeri ise çok yüksek düzeyde iç tutarlılığa işaret etmektedir (Hair et al., 2019). Bu doğrultuda elde edilen $\alpha = .952$ değeri, alt boyutta yer alan maddelerin aynı yapıyı oldukça tutarlı bir biçimde ölçtüğünü ve ölçeğin yüksek düzeyde iç tutarlılığa sahip olduğunu göstermektedir.

Alt boyutun sekiz maddeden oluşmasına rağmen Cronbach's Alpha katsayısının .952 gibi oldukça yüksek bir değere ulaşması, maddeler arasındaki ilişkinin güçlü olduğunu ve her bir maddenin ölçülmek istenen yapıya anlamlı katkı sağladığını göstermektedir. Bununla birlikte, alt boyutta yer alan maddelerin geliştiricilerin sorumluluğu, hukuki düzenlemeler, şeffaflık ve hesap verebilirlik gibi farklı yönleri kapsaması, yüksek güvenirlik katsayısının madde tekrarından değil, ortak kuramsal yapının başarılı bir şekilde ölçülmesinden kaynaklandığını göstermektedir.

Sorumluluk ve Hesap Verilebilirlik alt boyutunda yer alan maddelerin ayırt edicilik düzeylerini belirlemek amacıyla düzeltilmiş madde-toplam puan korelasyonları hesaplanmış, ayrıca ölçeğin toplam puanlarına göre oluşturulan %27'lik alt ve üst grupların madde puan ortalamaları bağımsız örneklem t-testi ile karşılaştırılmıştır. Bu analizler, maddelerin ölçülmek istenen yapıyı temsil etme gücünü ve farklı düzeylerde sorumluluk ve hesap verebilirlik algısına sahip

bireyleri ayırt edebilme yeterliliğini değerlendirmek amacıyla gerçekleştirilmiştir. Analiz sonuçları Tablo 21'de sunulmaktadır.

Tablo 21. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları

Madde No	Madde Toplam r	Grup	N	$\bar{X}$	Ss	df	t	p
s25	,841**	Alt	151	2,14	0,86	300	25,296	,000**
		Üst	151	4,57	0,80			
s27	,913**	Alt	151	2,27	1,00	300	29,433	,000**
		Üst	151	4,87	0,42			
s28	,888**	Alt	151	2,28	0,94	300	27,967	,000**
		Üst	151	4,80	0,59			
s29	,701**	Alt	151	2,19	0,94	300	14,888	,000**
		Üst	151	4,05	1,22			
s30	,912**	Alt	151	2,29	0,95	300	31,999	,000**
		Üst	151	4,89	0,32			
s31	,929**	Alt	151	2,24	0,91	300	34,364	,000**
		Üst	151	4,91	0,28			
s32	,904**	Alt	151	2,36	1,02	300	28,482	,000**
		Üst	151	4,88	0,38			
s35	,846**	Alt	151	2,29	0,93	300	25,644	,000**
		Üst	151	4,62	0,63			
Toplam		Alt	151	2,26	0,78	300	36,786	,000**
		Üst	151	4,70	0,25			

**p < .001

Tablo 21 incelendiğinde, Sorumluluk ve Hesap Verilebilirlik alt boyutunda yer alan maddelerin düzeltilmiş madde-toplam korelasyon katsayılarının .701 ile .929 arasında değiştiği görülmektedir. En yüksek madde-toplam korelasyonu s31 maddesinde ($r = .929$), en düşük korelasyon ise s29 maddesinde ($r = .701$) elde edilmiştir. Ölçek geliştirme çalışmalarında madde-toplam korelasyon katsayısının .30 ve üzerinde olması, maddelerin ölçülen yapıyı yeterli düzeyde temsil ettiğini ve yüksek ayırt edicilik gücüne sahip olduğunu göstermektedir (Büyüköztürk, 2020; DeVellis & Thorpe, 2021). Bu doğrultuda, tüm maddelerin oldukça yüksek madde-toplam korelasyonlarına sahip olması, alt boyutta yer alan her bir maddenin ölçülmek istenen yapıya güçlü katkı sağladığını ortaya koymaktadır.

%27'lik alt ve üst gruplar arasında yapılan bağımsız örneklem t-testi sonuçları incelendiğinde, tüm maddeler için gruplar arasında istatistiksel olarak anlamlı farklılıklar bulunduğu belirlenmiştir ($p < .001$). Maddelere ilişkin t değerleri 14.888 ile 34.364 arasında değişmektedir. En yüksek ayırt edicilik s31 maddesinde, en düşük ayırt edicilik ise s29 maddesinde gözlenmiştir. Bununla birlikte, s29 maddesine ait t değeri de oldukça yüksek olup,

maddenin ayırt edicilik gücünün yeterli düzeyde olduğunu göstermektedir. Ayrıca ölçeğin toplam puanı açısından da alt ve üst gruplar arasında anlamlı bir fark olduğu belirlenmiştir ($t = 36.786$, $p < .001$). Üst grupta yer alan katılımcıların tüm maddelerde alt gruba göre daha yüksek puan ortalamalarına sahip olması, maddelerin sorumluluk ve hesap verebilirlik algısı yüksek olan bireylerle düşük olan bireyleri başarılı bir şekilde ayırt edebildiğini göstermektedir.

4.4. Yapay Zeka Toplumsal Etki Alt Boyutu

Yapay Zekâ Toplumsal Etki alt boyutunun psikometrik özelliklerini belirlemek amacıyla kapsam geçerliği, açımlayıcı faktör analizi (AFA) ve doğrulayıcı faktör analizi (DFA) sonuçları birlikte değerlendirilmiştir. Bu kapsamda, uzman görüşleri doğrultusunda kapsam geçerliği sağlanan maddeler öncelikle açımlayıcı faktör analizine tabi tutulmuş, ardından elde edilen faktör yapısı doğrulayıcı faktör analizi ile test edilmiştir. Analiz sürecinde faktör yapısına yeterli katkı sağlamayan maddeler ölçekten çıkarılarak alt boyutun nihai yapısı oluşturulmuştur. Yapay Zekâ Toplumsal Etki alt boyutuna ilişkin faktör analizi sonuçları Tablo 22'de sunulmaktadır.

Tablo 22. Yapay Zeka Toplumsal Etki Alt Boyutuna Ait Faktör Analizi

No	Maddeler	Özdeğer	Çıkarım
37	YZ işsizliği artırabilir.	1,000	,931
38	YZ bireyler arasındaki sosyal bağları zayıflatabilir.	AFA	Elendi
39	YZ toplumda adaletsizliklere yol açabilir.	1,000	,935
40	YZ bireylerin özgür iradesini sınırlayabilir.	1,000	,926
41	YZ insan ilişkilerinde yabancılaşmaya sebep olabilir.	1,000	,950
42	YZ toplumsal güveni zedeleyebilir.	1,000	,962
43	YZ demokrasiyi olumsuz etkileyebilir.	1,000	,921
44	YZ bireylerin davranışlarını yönlendirebilir.	1,000	,936
45	YZ kültürel değerlerin değişmesine yol açabilir.	AFA	Elendi
46	YZ toplumda eşitsizlikleri derinleştirebilir.	1,000	,953
47	YZ çocukların gelişimini olumsuz etkileyebilir.	1,000	,918
48	YZ insan haklarını tehdit edebilir.	AFA	Elendi

Tablo 22 incelendiğinde, başlangıçta 12 maddeden oluşan Yapay Zekâ Toplumsal Etki alt boyutunda yer alan 38, 45 ve 48 numaralı maddelerin açımlayıcı faktör analizi sonucunda faktör yapısına yeterli katkı sağlamadığı belirlenmiş ve ölçekten çıkarılmıştır. Böylece alt boyut; 37,

39, 40, 41, 42, 43, 44, 46 ve 47 numaralı maddelerden oluşan 9 maddelik tek faktörlü bir yapıya kavuşmuştur.

Nihai ölçekte yer alan maddelerin ortak varyans (communalities) değerleri .918 ile .962 arasında değişmektedir. En yüksek ortak varyans değeri 42. maddede ("YZ toplumsal güveni zedeleyebilir.", .962), en düşük ortak varyans değeri ise 47. maddede ("YZ çocukların gelişimini olumsuz etkileyebilir.", .918) elde edilmiştir. Ortak varyans değerlerinin tamamının .50'nin oldukça üzerinde olması, maddelerin ilgili faktör tarafından yüksek düzeyde açıklandığını ve tek faktörlü yapıyı güçlü biçimde temsil ettiğini göstermektedir (Hair et al., 2019).

Yapay Zekâ Toplumsal Etki alt boyutunda yer alan maddelerin açımlayıcı faktör analizine uygunluğunu değerlendirmek amacıyla Kaiser-Meyer-Olkin (KMO) Örneklem Yeterliliği Ölçüsü ile Bartlett Küresellik Testi uygulanmıştır. Bu analizler, veri setinin faktör analizine uygunluğunu belirlemek ve maddeler arasındaki korelasyonların ortak faktörler altında toplanabilecek düzeyde olup olmadığını incelemek amacıyla gerçekleştirilmiştir. Elde edilen bulgular Tablo 23'de sunulmaktadır.

Tablo 23. KMO ve Bartlett küresellik testi sonuçları

Kaiser-Meyer-Olkin Örneklem Yeterliliği Ölçüsü		,954
Bartlett Küresellik Testi	Approx. Chi-Square	6983,816
	sd	36
	p	,000

Tablo 23 incelendiğinde, Kaiser-Meyer-Olkin (KMO) örneklem yeterliliği katsayısının 0,954 olduğu görülmektedir. KMO değerinin 0,90'ın üzerinde olması, örneklem büyüklüğünün faktör analizi için mükemmel düzeyde yeterli olduğunu ve maddeler arasındaki ortak varyansın faktör analizine son derece elverişli olduğunu göstermektedir (Kaiser, 1974).

Bartlett Küresellik Testi sonuçlarına göre test istatistiği $\chi^2 = 6983,816$ (sd = 36; $p < .001$) olarak bulunmuştur. Test sonucunun istatistiksel olarak anlamlı olması, korelasyon matrisinin birim matris olmadığını ve maddeler arasında faktör analizi yapılmasını gerektirecek düzeyde anlamlı ilişkiler bulunduğunu göstermektedir (Bartlett, 1954).

Yapay Zekâ Toplumsal Etki alt boyutunun faktör yapısını belirlemek amacıyla açımlayıcı faktör analizi kapsamında özdeğerler (eigenvalues) ve açıklanan toplam varyans değerleri incelenmiştir. Faktör sayısının belirlenmesinde Kaiser'in özdeğer ölçütü (özdeğer > 1) esas

alınmış, ayrıca faktörlerin açıkladığı toplam varyans oranları değerlendirilmiştir. Elde edilen bulgular Tablo 24'de sunulmaktadır.

Tablo 24. Açıklanan varyans ve özdeğerler

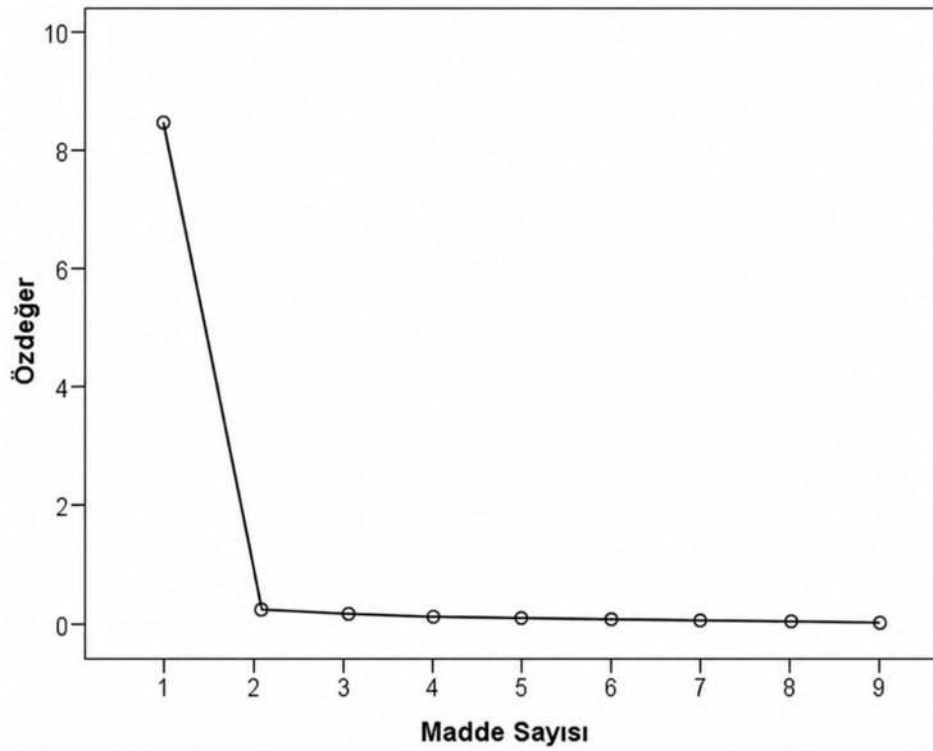
Faktör	Başlangıç Özdeğerleri			Yüklenen Faktörlerin Karelerinin Toplamı		
	Toplam	Varyans %	Birikimli %	Toplam	Varyans %	Birikimli %
1	8,432	93,693	93,693	8,432	93,693	93,693
2	,145	1,613	95,306			
3	,098	1,084	96,390			
4	,092	1,021	97,411			
5	,066	,738	98,150			
6	,059	,651	98,801			
7	,052	,576	99,376			
8	,031	,340	99,716			
9	,026	,284	100,000			

Tablo 24 incelendiğinde, yalnızca birinci faktörün özdeğerinin 1'in üzerinde olduğu ($\lambda = 8,432$), diğer sekiz faktörün özdeğerlerinin ise 1'in altında kaldığı görülmektedir. Kaiser'in özdeğer ölçütüne göre bu bulgu, Yapay Zekâ Toplumsal Etki alt boyutunun tek faktörlü bir yapıya sahip olduğunu göstermektedir (Kaiser, 1960).

Birinci faktörün açıkladığı toplam varyans oranı %93,693 olarak belirlenmiştir. Sosyal bilimlerde tek faktörlü ölçeklerde açıklanan toplam varyansın %30'un üzerinde olması yeterli, %50'nin üzerinde olması ise güçlü bir faktör yapısına işaret etmektedir (Hair et al., 2019). Bu doğrultuda elde edilen %93,693'lük açıklanan varyans oranı, tek faktörün maddelerdeki toplam varyansın çok büyük bir bölümünü açıkladığını ve alt boyutun oldukça güçlü bir yapısal bütünlüğe sahip olduğunu göstermektedir.

Yüklenen faktörlerin kareleri toplamı incelendiğinde, tek faktörlü çözümde başlangıç özdeğeri ile açıklanan varyans değerlerinin aynı olduğu görülmektedir. Bu durum, alt boyutta yer alan maddelerin aynı gizil yapıyı yüksek düzeyde temsil ettiğini ve faktör yapısının istatistiksel açıdan güçlü olduğunu desteklemektedir.

Yapay Zekâ Toplumsal Etki alt boyutunun faktör sayısını görsel olarak belirlemek amacıyla yamaç birikinti (Scree Plot) grafiği incelenmiştir. Şekil 8'de özdeğerlerin faktörlere göre dağılımı gösterilmektedir.



Şekil 8. Toplumsal Etki Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği

Şekil 8 incelendiğinde, birinci faktörden sonra özdeğerlerde belirgin ve keskin bir düşüş olduğu, ikinci faktörden itibaren ise eğrinin yatay bir görünüm sergilediği görülmektedir. Eğrinin bu yapısı, Cattell'in (1966) yamaç birikinti testi yaklaşımına göre kırılma noktasının birinci faktörde gerçekleştiğini ve alt boyutun tek faktörlü bir yapıya sahip olduğunu göstermektedir.

Yamaç birikinti grafiğinde elde edilen bulgular, açıklanan varyans ve özdeğer analizleri ile de tutarlılık göstermektedir. Nitekim birinci faktörün özdeğeri 8,432 olup toplam varyansın %93,693'ünü açıklamaktadır. Buna karşılık ikinci ve sonraki faktörlerin özdeğerlerinin 1'in altında kalması, bu faktörlerin ölçülen yapıya anlamlı bir katkı sağlamadığını ortaya koymaktadır. Bu durum, alt boyutta yer alan maddelerin aynı gizil yapıyı temsil ettiğini ve tek bir faktör altında toplandığını desteklemektedir.

Yapay Zekâ Toplumsal Etki alt boyutunun açıklayıcı faktör analizi sonucunda elde edilen tek faktörlü yapısının doğrulanması amacıyla doğrulayıcı faktör analizi (DFA) uygulanmıştır. Modelin veri ile uyumunu değerlendirmek için mutlak uyum, artımlı uyum ve yalınlık uyum indeksleri birlikte incelenmiştir. Elde edilen uyum indeksleri, literatürde kabul edilen eşik değerlerle karşılaştırılarak değerlendirilmiş ve sonuçlar Tablo 25'de sunulmuştur.

Tablo 25. Model Uyum Kriterleri

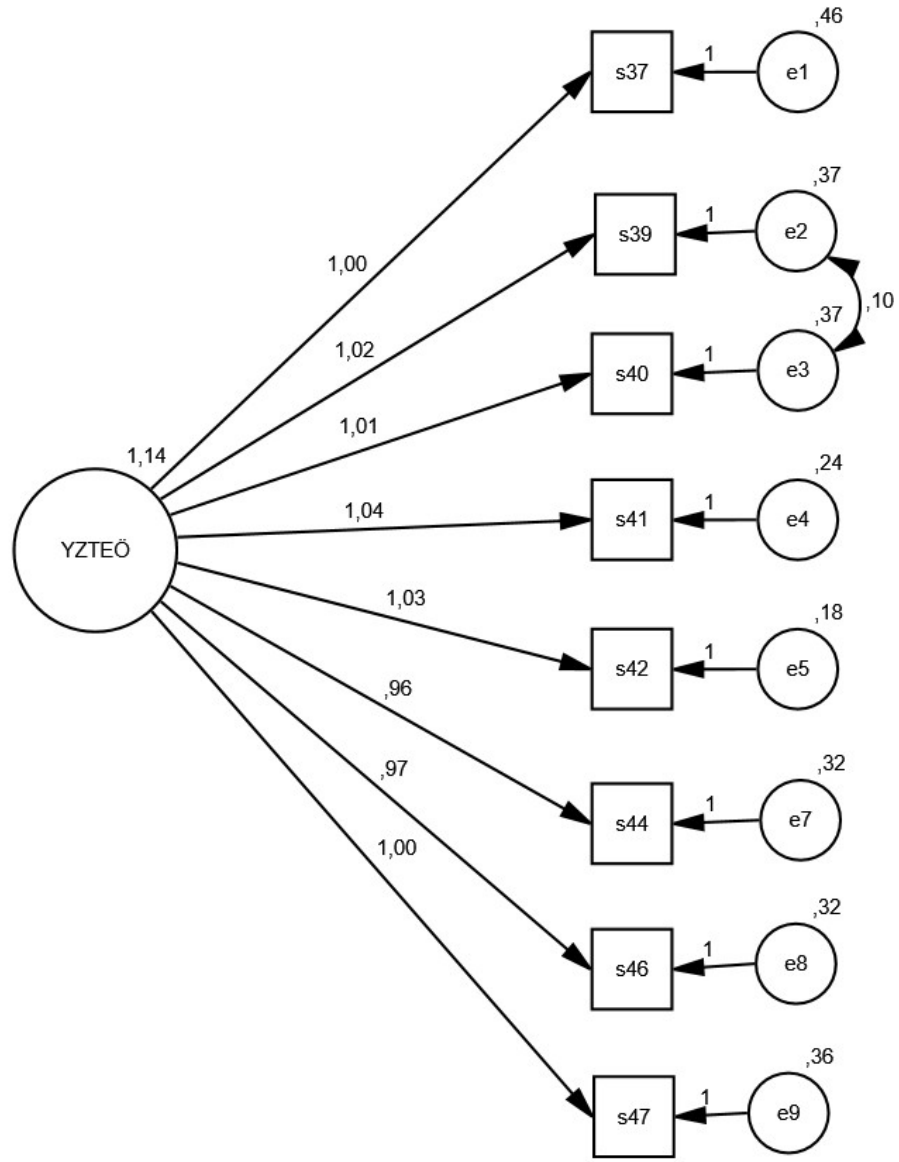
Model Uyum Kriteri	İyi uyum	Kabul Edilebilir Uyum	İndeks	Sonuç
χ^2/sd	0-3	3-5	4,306	Kabul edilebilir uyum
AGFI	.90-1.00	.85-.90	,939	İyi uyum
GFI	.90-1.00	.85-.90	,968	İyi uyum
CFI	.95-1.00	.90-.95	,989	İyi uyum
NFI	.95-1.00	.90-.95	,985	İyi uyum
NNFI (TLI)	.95-1.00	.90-.95	,983	İyi uyum
RFI	.95-1.00	.90-.95	,978	İyi uyum
IFI	.95-1.00	.90-.95	,989	İyi uyum
RMSEA	.00-.05	.05-.08	,074	Kabul edilebilir uyum
RMR	.00-.05	.05-.10	,020	İyi uyum
PNFI	.95-1.00	.50-.95	,668	Kabul edilebilir uyum
PGFI	.95-1.00	.50-.95	,511	Kabul edilebilir uyum
SRMR	0.00-0.05	.05-.10	,0133	İyi uyum

0<=RMR, SRMR

Tablo 25 incelendiğinde, doğrulayıcı faktör analizinden elde edilen uyum indekslerinin büyük çoğunluğunun iyi uyum düzeyinde olduğu görülmektedir. Ki-kare serbestlik derecesi oranı ($\chi^2/sd = 4,306$) 3–5 aralığında yer almakta olup modelin kabul edilebilir uyum gösterdiğini ifade etmektedir. Yaklaşık Hataların Ortalama Karekökü (RMSEA = .074) de .05–.08 aralığında bulunarak modelin kabul edilebilir düzeyde uyuma sahip olduğunu göstermektedir.

Mutlak uyum indekslerinden GFI (.968) ve AGFI (.939) değerleri .90'ın üzerinde olup modelin veriyle iyi uyum gösterdiğini ortaya koymaktadır. Benzer şekilde, karşılaştırmalı uyum indeksleri olan CFI (.989), NFI (.985), NNFI/TLI (.983), RFI (.978) ve IFI (.989) değerlerinin tamamı .95'in üzerinde gerçekleşmiş ve modelin mükemmele yakın düzeyde uyum sağladığını göstermiştir.

Hata uyum indeksleri incelendiğinde, RMR (.020) ve SRMR (.0133) değerlerinin .05'in altında olduğu görülmektedir. Bu sonuçlar, modelden elde edilen tahmini kovaryans matrisi ile gözlenen kovaryans matrisi arasındaki farkın oldukça düşük olduğunu ve modelin veri yapısını başarılı bir şekilde temsil ettiğini göstermektedir. Modelin yalınlığını değerlendiren PNFI (.668) ve PGFI (.511) değerleri ise literatürde önerilen kabul edilebilir sınırlar içerisinde yer almakta olup modelin yeterli düzeyde parsimony (yalınlık) özelliğine sahip olduğunu göstermektedir.



Şekil 9. Toplumsal Etki Alt Boyutu Path Diyagramı

Şekil 9'da, Yapay Zekâ Toplumsal Etki alt boyutuna ilişkin doğrulayıcı faktör analizi sonucunda elde edilen standartlaştırılmış yol (path) diyagramı sunulmaktadır. Diyagramda gizil değişken ile gözlenen değişkenler arasındaki ilişkiler ile her bir maddeye ait hata varyansları gösterilmektedir.

Path diyagramı incelendiğinde, 37, 39, 40, 41, 42, 43, 44, 46 ve 47 numaralı maddelerin Toplumsal Etki gizil değişkenini yüksek düzeyde temsil ettiği görülmektedir. Maddelere ait faktör yüklerinin yüksek olması, her bir maddenin ölçülmek istenen yapıya güçlü katkı sağladığını göstermektedir. Ayrıca hata varyanslarının düşük düzeyde olması, maddelerin büyük ölçüde ortak faktör tarafından açıklandığını ve ölçüm hatasının sınırlı olduğunu ortaya koymaktadır.

Doğrulayıcı faktör analizinde elde edilen uyum indeksleri ($\chi^2/sd = 4,306$; RMSEA = .074; CFI = .989; GFI = .968; AGFI = .939; NFI = .985; NNFI = .983; IFI = .989; SRMR = .013) ile birlikte değerlendirildiğinde, dokuz maddeden oluşan tek faktörlü modelin gözlenen verilerle iyi düzeyde uyum gösterdiği söylenebilir. Özellikle CFI, NFI, IFI ve NNFI değerlerinin .95'in üzerinde, SRMR ve RMR değerlerinin ise .05'in altında olması modelin güçlü bir yapısal geçerliğe sahip olduğunu desteklemektedir.

Toplumsal Etki alt boyutunun iç tutarlılığını belirlemek amacıyla Cronbach's Alpha güvenirlik katsayısı hesaplanmıştır. İç tutarlılık analizi, alt boyutta yer alan maddelerin aynı yapıyı ölçme konusundaki tutarlılığını değerlendirmek ve ölçeğin güvenirliğini ortaya koymak amacıyla gerçekleştirilmiştir. Analiz sonuçları Tablo 26'de sunulmaktadır.

Tablo 26. Toplumsal Etki Alt boyutu Güvenirlik Analizi

N	Madde Sayısı	Cronbach's Alpha
558	8	,966

Tablo 26 incelendiğinde, Toplumsal Etki alt boyutunun 558 katılımcıdan elde edilen veriler doğrultusunda (8/9) maddeden oluştuğu ve Cronbach's Alpha güvenirlik katsayısının .966 olduğu görülmektedir. Literatürde Cronbach's Alpha katsayısının .70 ve üzeri kabul edilebilir, .80 ve üzeri yüksek, .90 ve üzeri ise çok yüksek düzeyde iç tutarlılığı ifade etmektedir (Hair et al., 2019). Bu doğrultuda elde edilen $\alpha = .966$ değeri, alt boyutta yer alan maddelerin aynı yapıyı oldukça tutarlı bir biçimde ölçtüğünü ve ölçeğin çok yüksek düzeyde güvenilir olduğunu göstermektedir.

Cronbach's Alpha katsayısının .95'in üzerinde olması, maddeler arasında oldukça güçlü bir iç tutarlılık bulunduğunu göstermektedir. Bununla birlikte, Toplumsal Etki alt boyutunda yer alan maddelerin yapay zekânın istihdam, toplumsal adalet, bireysel özgürlük, insan ilişkileri, toplumsal güven ve demokratik süreçler üzerindeki farklı etkilerini kapsamayı, yüksek güvenirlik katsayısının ortak kuramsal yapının başarılı bir şekilde ölçülmesinden kaynaklandığını düşündürmektedir.

Toplumsal Etki alt boyutunda yer alan maddelerin ayırt edicilik düzeylerini belirlemek amacıyla düzeltilmiş madde-toplam puan korelasyonları hesaplanmış, ayrıca ölçeğin toplam puanlarına göre oluşturulan %27'lik alt ve üst grupların madde puan ortalamaları bağımsız örneklem t-testi ile karşılaştırılmıştır. Bu analizler, maddelerin ölçülmek istenen yapıyı temsil etme gücünü ve farklı düzeylerde toplumsal etki algısına sahip bireyleri ayırt edebilme

yeterliliğini değerlendirmek amacıyla gerçekleştirilmiştir. Analiz sonuçları Tablo 27'de sunulmaktadır.

Tablo 27. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları

Madde No	Madde Toplam r	Grup	N	$\bar{X}$	Ss	df	t	p
s37	,871**	Alt	151	2,11	1,06	300	27,376	,000**
		Üst	151	4,76	0,54			
s39	,901**	Alt	151	1,97	0,84	300	32,705	,000**
		Üst	151	4,76	0,62			
s40	,899**	Alt	151	2,06	0,86	300	36,447	,000**
		Üst	151	4,87	0,41			
s41	,920**	Alt	151	2,18	0,94	300	34,665	,000**
		Üst	151	4,93	0,26			
s42	,932**	Alt	151	2,21	0,93	300	32,660	,000**
		Üst	151	4,87	0,39			
s44	,888**	Alt	151	2,37	1,10	300	26,413	,000**
		Üst	151	4,86	0,37			
s46	,896**	Alt	151	2,12	0,82	300	31,298	,000**
		Üst	151	4,71	0,60			
s47	,888**	Alt	151	2,20	1,01	300	29,957	,000**
		Üst	151	4,84	0,40			
Toplam		Alt	151	2,15	0,75	300	41,992	,000**
		Üst	151	4,82	0,21			

**p < .001

Tablo 27 incelendiğinde, Toplumsal Etki alt boyutunda yer alan maddelerin düzeltilmiş madde-toplam korelasyon katsayılarının .871 ile .932 arasında değiştiği görülmektedir. En yüksek madde-toplam korelasyonu s42 maddesinde ($r = .932$), en düşük korelasyon ise s37 maddesinde ($r = .871$) elde edilmiştir. Ölçek geliştirme çalışmalarında madde-toplam korelasyon katsayısının .30 ve üzerinde olması, maddelerin ölçülen yapıyı yeterli düzeyde temsil ettiğini ve yüksek ayırt edicilik gücüne sahip olduğunu göstermektedir (Büyüköztürk, 2020; DeVellis & Thorpe, 2021). Elde edilen bulgular, alt boyutta yer alan tüm maddelerin oldukça yüksek ayırt edicilik düzeyine sahip olduğunu ve ölçülmek istenen yapıya güçlü katkı sağladığını ortaya koymaktadır.

%27'lik alt ve üst gruplar arasında yapılan bağımsız örneklem t-testi sonuçlarına göre, tüm maddeler için gruplar arasında istatistiksel olarak anlamlı farklılıklar bulunduğu belirlenmiştir ($p < .001$). Maddelere ilişkin t değerleri 26.413 ile 36.447 arasında değişmektedir. En yüksek ayırt edicilik s40 maddesinde ($t = 36.447$), en düşük ayırt edicilik ise s44 maddesinde ($t = 26.413$) gözlenmiştir. Bununla birlikte, tüm maddelerin t değerlerinin oldukça yüksek olması,

maddelerin toplumsal etki algısı yüksek olan bireylerle düşük olan bireyleri başarılı bir şekilde ayırt edebildiğini göstermektedir. Ölçeğin toplam puanı açısından da alt ve üst gruplar arasında anlamlı bir farklılık bulunmuştur ($t = 41.992$, $p < .001$). Üst grupta yer alan katılımcıların tüm maddelerde alt gruba göre daha yüksek puan ortalamalarına sahip olması, alt boyutun güçlü ayırt edicilik özelliğine sahip olduğunu göstermektedir.

4.5. Yapay Zeka Olumlu Beklentiler Alt Boyutu

Yapay Zekâ Olumlu Beklentiler alt boyutunun psikometrik özelliklerini belirlemek amacıyla kapsam geçerliği, açımlayıcı faktör analizi (AFA) ve doğrulayıcı faktör analizi (DFA) sonuçları birlikte değerlendirilmiştir. Bu kapsamda, uzman görüşleri doğrultusunda kapsam geçerliği sağlanan maddeler öncelikle açımlayıcı faktör analizine tabi tutulmuş, ardından elde edilen faktör yapısı doğrulayıcı faktör analizi ile test edilmiştir. Analiz sürecinde faktör yapısına yeterli katkı sağlamayan maddeler ölçekten çıkarılarak alt boyutun nihai yapısı oluşturulmuştur. Yapay Zekâ Olumlu Beklentiler alt boyutuna ilişkin faktör analizi sonuçları Tablo 28'de sunulmaktadır.

Tablo 28. Yapay Zeka Olumlu Beklentiler Alt Boyutuna Ait Faktör Analizi

No	Maddeler	Özdeğer	Çıkarım
49	YZ yaşam kalitesini artırabilir.	1,000	,758
50	YZ sağlık alanında önemli faydalar sunar.	DFA	Elendi
51	YZ etik çerçevede kullanıldığında topluma katkı sağlar.	1,000	,904
52	YZ eğitim süreçlerini geliştirebilir.	1,000	,896
53	YZ insanların iş yükünü azaltabilir.	1,000	,840
54	YZ sürdürülebilir kalkınmaya katkı sağlayabilir.	DFA	Elendi
55	YZ doğru yönlendirildiğinde güvenli bir teknoloji olabilir.	1,000	,814
56	YZ insanlara yeni iş imkânları yaratabilir.	1,000	,935
57	YZ karar alma süreçlerini hızlandırabilir.	1,000	,923
58	YZ kriz yönetiminde etkin kullanılabilir.	DFA	Elendi
59	YZ sosyal hizmetlerde olumlu katkılar sağlayabilir.	DFA	Elendi
60	YZ toplum için yararlı bir araç olabilir.	AFA	Elendi

Tablo 28 incelendiğinde, başlangıçta 12 maddeden oluşan Yapay Zekâ Olumlu Beklentiler alt boyutunda yer alan 60 numaralı maddenin açımlayıcı faktör analizi sonucunda faktör yapısına yeterli katkı sağlamadığı belirlenmiş ve ölçekten çıkarılmıştır. Ellinci ellidördüncü ellisekizinci

ve ellidokuzuncu maddeler doğrulayıcı faktör analizi sırasında yeterli uyum sağlamadıkları için ölçekten çıkarılmıştır. Böylece alt boyut; 49, 51, 52, 53, 55, 56 ve 57 numaralı maddelerden oluşan 7 maddelik tek faktörlü bir yapıya kavuşmuştur.

Nihai ölçekte yer alan maddelerin ortak varyans (communalities) değerleri .758 ile .935 arasında değişmektedir. En yüksek ortak varyans değeri 56. maddede ("YZ insanlara yeni iş imkânları yaratabilir.", .935), en düşük ortak varyans değeri ise 49. maddede ("YZ yaşam kalitesini artırabilir.", .758) elde edilmiştir. Ortak varyans değerlerinin tamamının .50'nin üzerinde olması, maddelerin ilgili faktör tarafından yüksek düzeyde açıklandığını ve tek faktörlü yapıyı güçlü biçimde temsil ettiğini göstermektedir (Hair et al., 2019).

Yapay Zekâ Olumlu Beklentiler alt boyutunda yer alan maddelerin açımlayıcı faktör analizine uygunluğunu değerlendirmek amacıyla Kaiser-Meyer-Olkin (KMO) Örneklem Yeterliliği Ölçüsü ile Bartlett Küresellik Testi uygulanmıştır. Bu analizler, veri setinin faktör analizine uygunluğunu belirlemek ve maddeler arasındaki korelasyonların ortak faktörler altında toplanabilecek düzeyde olup olmadığını incelemek amacıyla gerçekleştirilmiştir. Elde edilen bulgular Tablo 29'da sunulmaktadır.

Tablo 29. KMO ve Bartlett küresellik testi sonuçları

Kaiser-Meyer-Olkin Örneklem Yeterliliği Ölçüsü		,949
Bartlett Küresellik Testi	Approx. Chi-Square	6612,367
	sd	55
	p	,000

Tablo 29 incelendiğinde, Kaiser-Meyer-Olkin (KMO) örneklem yeterliliği katsayısının 0,949 olduğu görülmektedir. KMO değerinin 0,90'ın üzerinde olması, örneklem büyüklüğünün faktör analizi için mükemmel düzeyde yeterli olduğunu ve maddeler arasındaki ortak varyansın faktör analizine son derece elverişli olduğunu göstermektedir (Kaiser, 1974).

Bartlett Küresellik Testi sonuçlarına göre test istatistiği $\chi^2 = 6612,367$ (sd = 55; p < .001) olarak bulunmuştur. Test sonucunun istatistiksel olarak anlamlı olması, korelasyon matrisinin birim matris olmadığını ve maddeler arasında faktör analizi yapılmasını gerektirecek düzeyde anlamlı ilişkiler bulunduğunu göstermektedir (Bartlett, 1954).

Yapay Zekâ Olumlu Beklentiler alt boyutunun faktör yapısını belirlemek amacıyla açımlayıcı faktör analizi kapsamında özdeğerler (eigenvalues) ve açıklanan toplam varyans değerleri incelenmiştir. Faktör sayısının belirlenmesinde Kaiser'in özdeğer ölçütü (özdeğer > 1) esas

alınmış, ayrıca faktörlerin açıkladığı toplam varyans oranları değerlendirilmiştir. Elde edilen bulgular Tablo 30'da sunulmaktadır.

Tablo 30. Açıklanan varyans ve özdeğerler

Faktör	Başlangıç Özdeğerleri			Yüklenen Faktörlerin Karelerinin Toplamı		
	Toplam	Varyans %	Birikimli %	Toplam	Varyans %	Birikimli %
1	9,561	86,922	86,922	9,561	86,922	86,922
2	,313	2,846	89,769			
3	,267	2,424	92,193			
4	,207	1,879	94,071			
5	,173	1,574	95,645			
6	,125	1,139	96,785			
7	,106	,967	97,751			
8	,081	,738	98,489			
9	,070	,637	99,127			
10	,055	,504	99,631			
11	,041	,369	100,000			

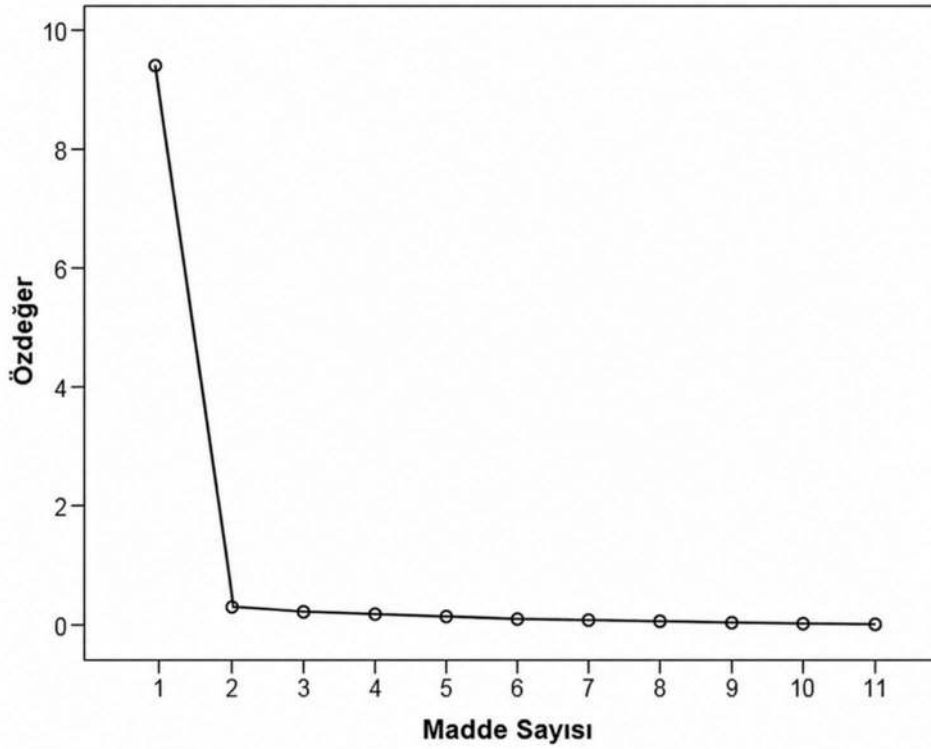
Tablo 30 incelendiğinde, yalnızca birinci faktörün özdeğerinin 1'in üzerinde olduğu ($\lambda = 9,561$), diğer on faktörün özdeğerlerinin ise 1'in altında kaldığı görülmektedir. Kaiser'in özdeğer ölçütüne göre bu bulgu, Yapay Zekâ Olumlu Beklentiler alt boyutunun tek faktörlü bir yapıya sahip olduğunu göstermektedir (Kaiser, 1960).

Birinci faktörün açıkladığı toplam varyans oranı %86,922 olarak belirlenmiştir. Sosyal bilimlerde tek faktörlü ölçeklerde açıklanan toplam varyansın %30'un üzerinde olması yeterli, %50'nin üzerinde olması ise güçlü bir faktör yapısına işaret etmektedir (Hair et al., 2019). Bu doğrultuda elde edilen %86,922'lik açıklanan varyans oranı, tek faktörün maddelerdeki toplam varyansın büyük bir bölümünü açıkladığını ve alt boyutun güçlü bir yapısal bütünlüğe sahip olduğunu göstermektedir.

Yüklenen faktörlerin kareleri toplamı incelendiğinde, tek faktörlü çözümde başlangıç özdeğeri ile açıklanan varyans değerlerinin aynı olduğu görülmektedir. Bu durum, alt boyutta yer alan maddelerin aynı gizil yapıyı yüksek düzeyde temsil ettiğini ve faktör yapısının istatistiksel açıdan güçlü olduğunu desteklemektedir.

Şekil 10'da, Yapay Zekâ Olumlu Beklentiler alt boyutuna ilişkin açıklayıcı faktör analizi sonucunda elde edilen yamaç birikinti (Scree Plot) grafiği sunulmaktadır. Grafik, faktör

sayısının belirlenmesinde özdeğerlerin dağılımını görsel olarak değerlendirmek amacıyla kullanılmaktadır. Faktör sayısının belirlenmesinde grafikte görülen kırılma noktası (elbow) temel ölçütlerden biri olarak kabul edilmektedir (Cattell, 1966).



Şekil 10. Olumlu Beklentiler Alt Boyutuna İlişkin Yamaç Birikinti (Scree Plot) Grafiği

Şekil 10 incelendiğinde, birinci faktörden sonra özdeğerlerde çok belirgin ve keskin bir düşüş olduğu, ikinci faktörden itibaren ise eğrinin yatay bir seyir izlediği görülmektedir. Bu durum, maddelerde açıklanan varyansın büyük bölümünün ilk faktör tarafından açıklandığını ve sonraki faktörlerin yapıya anlamlı bir katkı sağlamadığını göstermektedir.

Grafikte gözlenen kırılma noktası birinci faktörden hemen sonra oluşmakta olup, bu bulgu Kaiser'in özdeğer ölçütü (özdeğer > 1) ile de örtüşmektedir. Nitekim özdeğer analizinde yalnızca birinci faktörün özdeğerinin 1'in üzerinde ($\lambda = 9,561$) olduğu, diğer faktörlerin ise 1'in altında kaldığı belirlenmiştir. Ayrıca birinci faktörün toplam varyansın %86,922'sini açıklaması, tek faktörlü yapının oldukça güçlü olduğunu göstermektedir.

Yapay Zekâ Olumlu Beklentiler alt boyutunun açımlayıcı faktör analizi sonucunda elde edilen tek faktörlü yapısının doğrulanması amacıyla doğrulayıcı faktör analizi (DFA) uygulanmıştır. Modelin gözlenen verilerle uyumunu değerlendirmek için mutlak uyum, artımlı uyum ve yalınlık uyum indeksleri birlikte incelenmiş ve elde edilen sonuçlar Tablo 31'de sunulmuştur.

Tablo 31. Model Uyum Kriterleri

Model Uyum Kriteri	İyi uyum	Kabul Edilebilir Uyum	İndeks	Sonuç
χ^2/sd	0-3	3-5	4,285	Kabul edilebilir uyum
AGFI	.90-1.00	.85-.90	,941	İyi uyum
GFI	.90-1.00	.85-.90	,973	İyi uyum
CFI	.95-1.00	.90-.95	,991	İyi uyum
NFI	.95-1.00	.90-.95	,989	İyi uyum
NNFI (TLI)	.95-1.00	.90-.95	,986	İyi uyum
RFI	.95-1.00	.90-.95	,982	İyi uyum
IFI	.95-1.00	.90-.95	,991	İyi uyum
RMSEA	.00-.05	.05-.08	,077	Kabul edilebilir uyum
RMR	.00-.05	.05-.10	,015	İyi uyum
PNFI	.95-1.00	.50-.95	,612	Kabul edilebilir uyum
PGFI	.95-1.00	.50-.95	,452	-
SRMR	0.00-0.05	.05-.10	,0115	İyi uyum

0<=RMR, SRMR

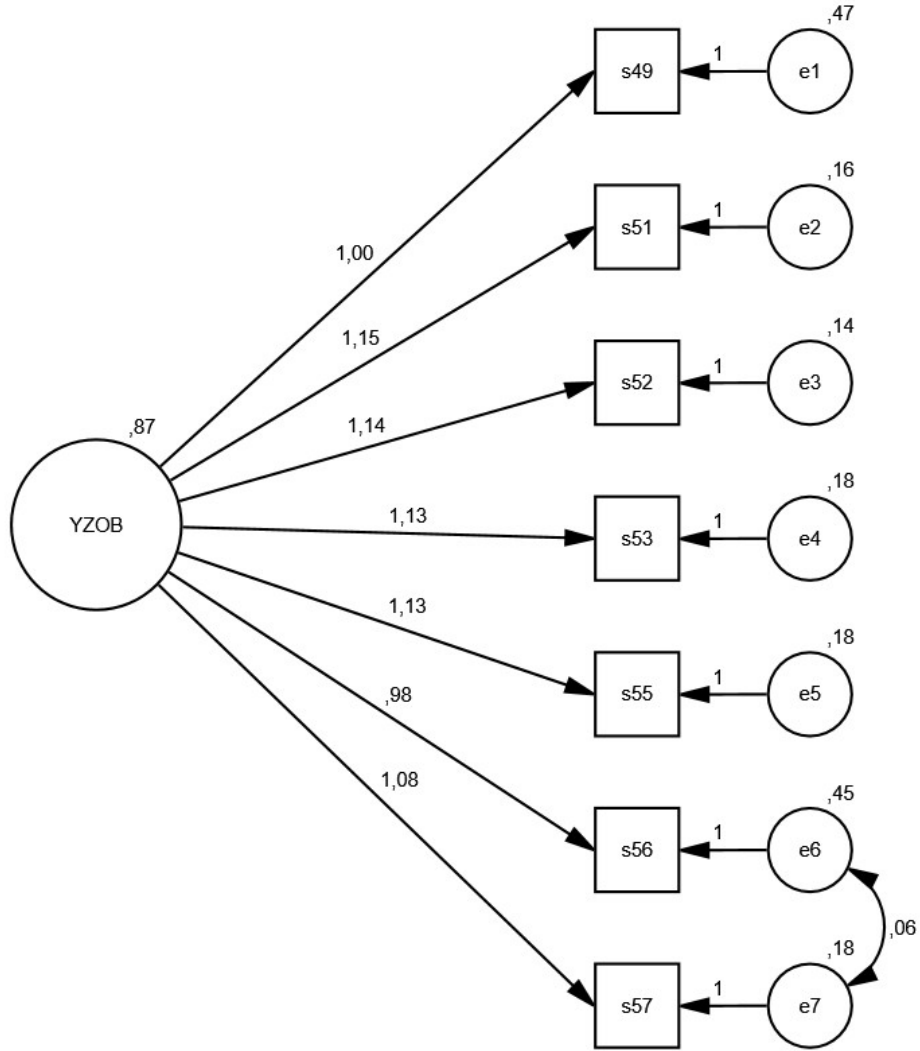
Tablo 31 incelendiğinde, doğrulayıcı faktör analizinden elde edilen uyum indekslerinin modelin genel olarak iyi düzeyde uyum gösterdiğini ortaya koyduğu görülmektedir. Ki-kare serbestlik derecesi oranı ($\chi^2/sd = 4,285$) 3–5 aralığında yer almakta olup modelin kabul edilebilir uyum sergilediğini göstermektedir. Benzer şekilde RMSEA (.077) değeri de .05–.08 aralığında bulunarak modelin kabul edilebilir düzeyde uyuma sahip olduğunu göstermektedir.

Mutlak uyum indeksleri incelendiğinde GFI (.973) ve AGFI (.941) değerlerinin .90'ın üzerinde olduğu ve modelin veriyle iyi uyum sağladığı görülmektedir. Artımlı uyum indeksleri olan CFI (.991), NFI (.989), NNFI/TLI (.986), RFI (.982) ve IFI (.991) değerlerinin tamamının .95'in üzerinde olması, önerilen ölçüm modelinin oldukça güçlü bir uyum sergilediğini göstermektedir.

Hata uyum indeksleri değerlendirildiğinde RMR (.015) ve SRMR (.0115) değerlerinin .05'in altında olduğu görülmektedir. Bu durum, model tarafından tahmin edilen kovaryans matrisi ile gözlenen kovaryans matrisi arasındaki farkın oldukça düşük olduğunu ve modelin gözlenen verileri başarılı bir şekilde temsil ettiğini göstermektedir.

Modelin yalınlığını değerlendiren PNFI (.612) değeri literatürde önerilen kabul edilebilir sınırlar içerisinde yer almakta olup modelin yeterli düzeyde parsimony özelliğine sahip olduğunu göstermektedir. PGFI (.452) değeri ise önerilen .50 sınırının altında kalmaktadır. Bununla birlikte, PGFI tek başına modelin kabul veya reddedilmesi için yeterli bir ölçüt

olmayıp, özellikle diğer uyum indekslerinin yüksek olduğu durumlarda modelin genel uyumunu olumsuz yönde değerlendirmek için tek başına kullanılmamaktadır. Bu nedenle PGFI değerinin nispeten düşük olması, modelin genel geçerliğini zayıflatan bir durum olarak değerlendirilmemektedir.



Şekil 11. Olumlu Beklentiler Alt Boyutu Path Diyagramı

Şekil 11'de, Yapay Zekâ Olumlu Beklentiler alt boyutuna ilişkin doğrulayıcı faktör analizi sonucunda elde edilen standartlaştırılmış yol (path) diyagramı sunulmaktadır. Diyagramda gizil değişken ile gözlenen değişkenler arasındaki ilişkiler, faktör yükleri ve her bir maddeye ait hata varyansları gösterilmektedir.

Path diyagramı incelendiğinde, 49, 51, 52, 53, 55, 56 ve 57 numaralı maddelerin Olumlu Beklentiler gizil değişkenini yüksek düzeyde temsil ettiği görülmektedir. Maddelere ait faktör yüklerinin yüksek olması, her bir maddenin ölçülmek istenen yapıya güçlü katkı sağladığını göstermektedir. Ayrıca hata varyanslarının düşük düzeyde olması, maddelerin büyük ölçüde ortak faktör tarafından açıklandığını ve ölçüm hatasının sınırlı olduğunu ortaya koymaktadır.

Olumlu Beklentiler alt boyutunun iç tutarlılığını belirlemek amacıyla Cronbach's Alpha güvenilirlik katsayısı hesaplanmıştır. İç tutarlılık analizi, alt boyutta yer alan maddelerin aynı yapıyı ölçme konusundaki tutarlılığını değerlendirmek ve ölçeğin güvenilirliğini ortaya koymak amacıyla gerçekleştirilmiştir. Analiz sonuçları Tablo 32'de sunulmaktadır.

Tablo 32. Olumlu Beklentiler Alt boyutu Güvenirlik Analizi

N	Madde Sayısı	Cronbach's Alpha
558	7	,966

Tablo 32 incelendiğinde, Olumlu Beklentiler alt boyutunun 558 katılımcıdan elde edilen veriler doğrultusunda 7 maddeden oluştuğu ve Cronbach's Alpha güvenilirlik katsayısının .966 olduğu görülmektedir. Literatürde Cronbach's Alpha katsayısının .70 ve üzeri kabul edilebilir, .80 ve üzeri yüksek, .90 ve üzeri ise çok yüksek düzeyde iç tutarlılığı ifade etmektedir (Hair et al., 2019). Buna göre elde edilen $\alpha = .966$ değeri, alt boyutta yer alan maddelerin aynı yapıyı oldukça tutarlı bir biçimde ölçtüğünü ve ölçeğin çok yüksek düzeyde güvenilir olduğunu göstermektedir.

Cronbach's Alpha katsayısının .95'in üzerinde olması, maddeler arasında oldukça güçlü bir iç tutarlılık bulunduğunu göstermektedir. Bununla birlikte, alt boyutta yer alan maddelerin yapay zekânın yaşam kalitesini artırması, eğitim ve sağlık alanındaki katkıları, iş yükünü azaltması, güvenli kullanım potansiyeli ve karar verme süreçlerine sağlayabileceği katkılar gibi ortak bir kuramsal yapıyı ölçmesi, yüksek güvenilirlik katsayısını desteklemektedir.

Olumlu Beklentiler alt boyutunda yer alan maddelerin ayırt edicilik düzeylerini belirlemek amacıyla düzeltilmiş madde-toplam puan korelasyonları hesaplanmış, ayrıca ölçeğin toplam puanlarına göre oluşturulan %27'lik alt ve üst grupların madde puan ortalamaları bağımsız örneklem t-testi ile karşılaştırılmıştır. Bu analizler, maddelerin ölçülmek istenen yapıyı temsil etme gücünü ve farklı düzeylerde olumlu beklentiye sahip bireyleri ayırt edebilme yeterliliğini değerlendirmek amacıyla gerçekleştirilmiştir. Analiz sonuçları Tablo 33'de sunulmaktadır.

Tablo 33. Madde-Toplam Korelasyonları ve %27'lik Alt-Üst Grup Farkına İlişkin Bağımsız Grup t-Testi Sonuçları

Madde No	Madde Toplam r	Grup	N	$\bar{X}$	Ss	df	t	p
s49	,849**	Alt	151	2,13	0,89	300	23,822	,000**
		Üst	151	4,46	0,81			
s51	,939**	Alt	151	2,38	1,03	300	29,015	,000**
		Üst	151	4,91	0,29			
s52	,942**	Alt	151	2,39	1,05	300	28,458	,000**
		Üst	151	4,92	0,29			
s53	,936**	Alt	151	2,34	0,98	300	30,719	,000**
		Üst	151	4,91	0,31			
s55	,932**	Alt	151	2,47	1,12	300	25,421	,000**
		Üst	151	4,89	0,33			
s56	,857**	Alt	151	2,34	0,94	300	23,465	,000**
		Üst	151	4,62	0,74			
s57	,932**	Alt	151	2,46	1,02	300	27,176	,000**
		Üst	151	4,85	0,37			
Toplam		Alt	151	2,36	0,86	300	33,548	,000**
		Üst	151	4,79	0,23			

** p < .001

Tablo 33 incelendiğinde, Olumlu Beklentiler alt boyutunda yer alan maddelerin düzeltilmiş madde-toplam korelasyon katsayılarının .849 ile .942 arasında değiştiği görülmektedir. En yüksek madde-toplam korelasyonu s52 maddesinde ($r = .942$), en düşük korelasyon ise s49 maddesinde ($r = .849$) elde edilmiştir. Ölçek geliştirme çalışmalarında madde-toplam korelasyon katsayısının .30 ve üzerinde olması, maddelerin ölçülen yapıyı yeterli düzeyde temsil ettiğini ve yüksek ayırt edicilik gücüne sahip olduğunu göstermektedir (Büyüköztürk, 2020; DeVellis & Thorpe, 2021). Bu doğrultuda elde edilen bulgular, alt boyutta yer alan tüm maddelerin ölçülen yapıyı oldukça güçlü biçimde temsil ettiğini ve yüksek ayırt edicilik özelliğine sahip olduğunu ortaya koymaktadır.

%27'lik alt ve üst gruplar arasında yapılan bağımsız örneklem t-testi sonuçlarına göre, tüm maddeler için gruplar arasında istatistiksel olarak anlamlı farklılıklar bulunduğu belirlenmiştir ($p < .001$). Maddelere ilişkin t değerleri 23.465 ile 30.719 arasında değişmektedir. En yüksek ayırt edicilik s53 maddesinde ($t = 30.719$), en düşük ayırt edicilik ise s56 maddesinde ($t = 23.465$) gözlenmiştir. Bununla birlikte, tüm maddelerin t değerlerinin oldukça yüksek olması, maddelerin olumlu beklenti düzeyi yüksek olan bireylerle düşük olan bireyleri başarılı bir şekilde ayırt edebildiğini göstermektedir. Ölçeğin toplam puanı açısından da alt ve üst gruplar arasında anlamlı bir farklılık bulunmuştur ($t = 33.548$, $p < .001$). Üst grupta yer alan

katılımcıların tüm maddelerde alt gruba göre daha yüksek puan ortalamalarına sahip olması, alt boyutun güçlü ayırt edicilik özelliğine sahip olduğunu göstermektedir.

4.6. Demografik Özellikler

Araştırmaya katılan bireylerin demografik özelliklerine ilişkin frekans ve yüzde dağılımları Tablo 34'de sunulmuştur. Katılımcıların cinsiyet, yaş, eğitim düzeyi, öğrenim gördükleri bölüm, gelir durum algısı ve yapay zekâ uygulamalarını kullanma durumlarına ilişkin bilgiler frekans ve yüzde değerleri üzerinden değerlendirilmiştir.

Tablo 34. Demografik Özellikler

Değişken	Gruplar	F	%
Cinsiyet	Kadın	298	53,4%
	Erkek	260	46,6%
	Toplam	558	100,0%
Yaş	18-24 Yaş	58	10,4%
	25-34 Yaş	448	80,3%
	35 Yaş ve Üstü	52	9,3%
	Toplam	558	100,0%
Eğitim	Ön lisans	65	11,6%
	Lisans	431	77,2%
	Lisansüstü	62	11,1%
	Toplam	558	100,0%
Bölüm	Sayısal	282	50,5%
	Sözel	276	49,5%
	Toplam	558	100,0%
Gelir Durum Algısı	Kötü	192	34,4%
	Orta	195	34,9%
	İyi	171	30,6%
	Toplam	558	100,0%
YZ uygulamaları Kullanımı	Evet	310	55,6%
	Hayır	248	44,4%
	Toplam	558	100,0%

4.6.1. Cinsiyet Dağılımı

Çalışma grubunda kadın katılımcılar (%53,4; n=298) hafif bir çoğunluğu oluşturmaktadır. Erkek katılımcılar ise %46,6 (n=260) oranında temsil edilmektedir. Bu dağılım, cinsiyet açısından nispeten dengeli bir örneklem yapısına işaret etmekle birlikte, kadın katılımcıların hafif bir üstünlüğe sahip olduğunu göstermektedir. Literatürde benzer araştırmalarda cinsiyet dağılımları genellikle farklılık göstermekte olup, bu çalışmanın kadın katılımcı ağırlıklı bir yapıya sahip olduğu söylenebilir.

4.6.2. Yaş Dağılımı

Yaş değişkeni incelendiğinde, çalışma grubunun büyük ölçüde 25-34 yaş aralığında (n=448; %80,3) yoğunlaştığı görülmektedir. Bu yaş grubu, genellikle üniversite eğitimi tamamlamış, iş hayatına yeni atılmış veya kariyerinin ilk yıllarında olan bireyleri temsil etmektedir. 18-24 yaş aralığı %10,4 (n=58) ve 35+ yaş aralığı %9,3 (n=52) ile oldukça sınırlı bir temsile sahiptir. Bu durum, çalışmanın sonuçlarının genellenebilirliği açısından bir sınırlılık oluşturabileceğinden, bulguların genellikle genç yetişkin popülasyonu için yorumlanması gerektiği belirtilmelidir.

4.6.3. Eğitim Düzeyi Dağılımı

Katılımcıların eğitim düzeyleri incelendiğinde, lisans mezunlarının (n=431; %77,2) ezici bir çoğunluğu oluşturduğu görülmektedir. Önlisans (n=65; %11,6) ve lisansüstü (n=62; %11,1) düzeyindeki katılımcıların oranları birbirine oldukça yakın ve düşük düzeydedir. Bu durum, çalışma grubunun eğitim düzeyi açısından homojen bir yapıya sahip olduğunu ve sonuçların yükseköğretim düzeyinde eğitim almış bireyler için geçerli olabileceğini göstermektedir.

4.6.4. Bölüm (Sayısal/Sözel) Dağılımı

Katılımcılar sayısal (n=282; %50,5) ve sözel (n=276; %49,5) bölümlerden neredeyse eşit oranda gelmektedir. Bu dengeli dağılım, farklı akademik arka planlara sahip bireylerin görüşlerinin karşılaştırılabilmesi açısından önemli bir avantaj sağlamaktadır. Özellikle yapay zeka kullanımı ile ilgili değişkenlerin incelenmesinde, sayısal ve sözel bölüm farklılıklarının istatistiksel olarak test edilebilmesi için uygun bir zemin oluşturmaktadır.

4.6.5. Gelir Durumu Algısı Dağılımı

Katılımcıların gelir durumu algısı üç kategoride incelenmiştir: kötü (n=192; %34,4), orta (n=195; %34,9) ve iyi (n=171; %30,6). Dağılım oldukça homojen olup, üç kategori arasında anlamlı bir fark bulunmamaktadır. Bu durum, çalışma grubunun gelir algısı açısından nispeten eşit dağıldığını ve herhangi bir gelir kategorisinin aşırı temsil edilmediğini göstermektedir.

4.6.6. Yapay Zeka (YZ) Uygulamaları Kullanımı

Çalışma grubunun %55,6'sı (n=310) yapay zeka uygulamalarını kullanmakta iken, %44,4'ü (n=248) kullanmamaktadır. Bu oran, yapay zeka teknolojilerinin araştırma popülasyonu içinde yüksek bir penetrasyona sahip olduğunu göstermektedir. Özellikle eğitim düzeyi yüksek olan

genç yetişkinler arasında YZ kullanımının yaygınlaştığı söylenebilir. Bu dağılım, yapay zeka kullanıcıları ve kullanmayanlar arasındaki farklılıkların istatistiksel olarak incelenmesi için yeterli örneklem büyüklüğü sağlamaktadır.

4.7. Demografik Özelliklerin Ayrıntılı Boyutları

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının cinsiyete göre farklılaşıp farklılaşmadığını belirlemek amacıyla bağımsız örneklem t-testi uygulanmıştır. Analiz sonuçları Tablo 35'de sunulmuştur.

Tablo 35. Katılımcıların cinsiyete göre YZGAAÖ toplam puanları için bağımsız örneklem sonuçları

Cinsiyet	N	$\bar{X}$	Ss	df	t	p
Kadın	298	141.62	33.129	556	-1.446	.149
Erkek	260	145.94	37.506			

$p > .05$

Analiz sonuçlarına göre erkek katılımcıların toplam puan ortalaması ($\bar{X}$ =145.94), kadın katılımcıların puan ortalamasından ($\bar{X}$ =141.62) daha yüksektir. Ancak bu fark istatistiksel olarak anlamlı bulunmamıştır [$t(556)=-1,446$; $p=.149$]. Bu bulgu, katılımcıların yapay zekâya ilişkin güvenlik ve ahlak algılarının cinsiyete göre anlamlı bir farklılık göstermediğini ortaya koymaktadır. Başka bir ifadeyle, kadın ve erkek katılımcılar yapay zekânın güvenlik, etik, sorumluluk, toplumsal etkileri ve olumlu beklentilerine ilişkin benzer algılara sahiptir. Her ne kadar erkek katılımcıların ortalama puanları kadın katılımcılardan bir miktar yüksek olsa da, bu farklılık istatistiksel açıdan anlamlı düzeye ulaşmamıştır.

Tablo 36. Katılımcıların cinsiyete göre YZGAAÖ alt boyut puanları için bağımsız örneklem sonuçları

Alt Boyut	Cinsiyet	N	$\bar{X}$	Ss	df	t	p
Güvenlik Farkındalığı	Kadın	298	24.50	6.567	556	-.775	.439
	Erkek	260	24.97	7.738			
Ahlaki Değerler	Kadın	298	37.13	10.021	556	-1.171	.242
	Erkek	260	38.19	11.163			
Sorumluluk ve Hesap Verilebilirlik	Kadın	298	28.76	7.787	556	-1.772	.077
	Erkek	260	29.99	8.659			
Toplumsal Etki	Kadın	298	29.12	8.478	556	-.495	.621
	Erkek	260	29.49	9.071			
Olumlu Beklentiler	Kadın	298	22.10	5.999	556	-2.243	.025*
	Erkek	260	23.31	6.676			

* $p < .05$

Tablo 36’da, katılımcıların cinsiyetlerine göre Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) alt boyutlarından aldıkları puanların anlamlı biçimde farklılaşıp farklılaşmadığını belirlemek amacıyla gerçekleştirilen bağımsız örneklem t-testi sonuçları sunulmuştur. Analiz kapsamında kadın katılımcılar (n=298) ile erkek katılımcıların (n=260) Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verilebilirlik, Toplumsal Etki ve Olumlu Beklentiler alt boyutlarına ilişkin puan ortalamaları karşılaştırılmıştır. Analizlerde anlamlılık düzeyi $p<.05$ olarak kabul edilmiştir.

Güvenlik Farkındalığı alt boyutunda kadın katılımcıların puan ortalaması 24.50 ($Ss=6.567$), erkek katılımcıların puan ortalaması ise 24.97 ($Ss=7.738$) olarak bulunmuştur. Erkek katılımcıların ortalaması kadın katılımcılardan daha yüksek olmasına rağmen, iki grup arasındaki fark istatistiksel olarak anlamlı değildir ($t_{(556)}=-.775$, $p=.439>.05$). Buna göre, katılımcıların yapay zekâyâ ilişkin güvenlik farkındalıkları cinsiyet değişkenine göre anlamlı bir farklılık göstermemektedir. Başka bir ifadeyle, kadın ve erkek katılımcıların yapay zekâ kullanımında güvenlik, gizlilik, veri güvenliği ve olası risklere ilişkin farkındalık düzeylerinin birbirine benzer olduğu söylenebilir.

Ahlaki Değerler alt boyutunda kadın katılımcıların puan ortalaması 37.13 ($Ss=10.021$), erkek katılımcıların puan ortalaması ise 38.19 ($Ss=11.163$) olarak belirlenmiştir. Erkeklerin ortalaması kadınlardan yaklaşık 1.06 puan daha yüksek olmakla birlikte, bu fark istatistiksel olarak anlamlı bulunmamıştır ($t_{(556)}=-1.171$, $p=.242>.05$). Bu sonuç, yapay zekânın kullanımına ilişkin etik ve ahlaki değerlendirmelerin katılımcıların cinsiyetlerine göre anlamlı biçimde farklılaşmadığını göstermektedir. Dolayısıyla yapay zekânın etik kullanımı, insan değerlerine uygunluğu, adalet ve ahlaki ilkeler gibi konulardaki algıların kadın ve erkek katılımcılarda genel olarak benzer düzeyde olduğu ifade edilebilir.

Sorumluluk ve Hesap Verilebilirlik alt boyutunda kadın katılımcıların puan ortalaması 28.76 ($Ss=7.787$) iken erkek katılımcıların puan ortalaması 29.99 ($Ss=8.659$) olarak hesaplanmıştır. Erkek katılımcıların ortalaması kadın katılımcılardan yaklaşık 1.23 puan daha yüksek olmasına rağmen, elde edilen fark istatistiksel olarak anlamlı değildir ($t_{(556)}=-1.772$, $p=.077>.05$). Buna göre, katılımcıların yapay zekâ kullanımında sorumluluk alma, gerçekleştirilen işlemlerin sonuçlarından sorumlu olma ve yapay zekâ uygulamalarının hesap verebilirliği konusundaki algılarının cinsiyete göre anlamlı biçimde farklılaşmadığı sonucuna ulaşılmıştır. Her ne kadar erkek katılımcıların ortalama puanı daha yüksek olsa da, bu fark .05 anlamlılık düzeyine ulaşmadığından cinsiyete bağlı anlamlı bir farklılık olarak değerlendirilmemelidir.

Toplumsal Etki alt boyutunda kadın katılımcıların ortalaması 29.12 (Ss=8.478), erkek katılımcıların ortalaması ise 29.49 (Ss=9.071) olarak bulunmuştur. İki grup arasındaki ortalama farkın yalnızca 0.37 puan olması, grupların bu boyuttaki puanlarının birbirine oldukça yakın olduğunu göstermektedir. Nitekim t-testi sonucu da bu farkın istatistiksel olarak anlamlı olmadığını ortaya koymuştur ($t_{(556)}=-.495$, $p=.621>.05$). Buna göre, yapay zekânın toplum üzerindeki etkileri, toplumsal yaşamı dönüştürme potansiyeli ve sosyal sonuçlarına ilişkin algıların kadın ve erkek katılımcılar arasında anlamlı biçimde farklılaşmadığı söylenebilir.

Buna karşılık, Olumlu Beklentiler alt boyutunda cinsiyete göre istatistiksel olarak anlamlı bir farklılık tespit edilmiştir. Kadın katılımcıların puan ortalaması 22.10 (Ss=5.999) iken erkek katılımcıların puan ortalaması 23.31 (Ss=6.676) olarak belirlenmiştir. Buna göre erkek katılımcıların puan ortalaması kadın katılımcılardan yaklaşık 1.21 puan daha yüksektir. Bağımsız örneklem t-testi sonucunda elde edilen $t_{(556)}=-2.243$, $p=.025<.05$ değeri, söz konusu farkın istatistiksel olarak anlamlı olduğunu göstermektedir. Dolayısıyla erkek katılımcıların yapay zekânın gelecekte sağlayabileceği yararlar, gelişim potansiyeli ve olumlu sonuçlarına ilişkin beklentilerinin kadın katılımcılara göre anlamlı düzeyde daha yüksek olduğu sonucuna ulaşılmıştır.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının yaş gruplarına göre dağılımını belirlemek amacıyla tanımlayıcı istatistikler hesaplanmıştır. Yaş gruplarına ilişkin ortalama puanlar ve standart sapma değerleri Tablo 37'te sunulmuştur.

Tablo 37. Yaş Gruplarına göre katılımcıların YZGAAÖ puanlarının tanımlayıcı istatistikleri

Yaş	N	$\bar{X}$	Ss
1. 18-24 Yaş	58	128.73	47.515
2. 25-34 Yaş	448	144.86	32.759
3. 35 Yaş ve üstü	52	150.08	36.425
Toplam	558	143.63	35.270

Tablo 37'de katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) puanlarının yaş gruplarına göre tanımlayıcı istatistikleri yer almaktadır. Bulgular incelendiğinde, katılımcıların toplam ölçek puan ortalamasının ($\bar{X}=143.63$) olduğu ve standart sapmanın ($\bar{X}=35.270$) olduğu görülmektedir. Bu durum, katılımcıların genel olarak orta-yüksek düzeyde bir yapay zekâ güvenlik algısına sahip olduklarını, ancak bireyler arasında belirgin bir değişkenlik bulunduğunu göstermektedir.

Yaş grupları açısından değerlendirildiğinde, en düşük ortalama puanın 18–24 yaş grubuna ($\bar{X}$ =128.73; S_s =47.515) ait olduğu görülmektedir. Bu grubun diğer yaş gruplarına göre daha düşük bir yapay zekâ güvenlik algısına sahip olduğu söylenebilir. Bununla birlikte bu grubun standart sapma değerinin yüksek olması, grup içindeki görüşlerin daha heterojen olduğunu göstermektedir.

25–34 yaş grubunun ortalama puanı 144.86 (S_s =32.759) olarak hesaplanmıştır. Bu grup, toplam ortalamaya yakın bir değer göstermekte olup daha dengeli bir dağılıma sahiptir. En yüksek ortalama puan ise 35 yaş ve üzeri gruba aittir ($\bar{X}$ =150.08; S_s =36.425). Bu bulgu, yaş arttıkça yapay zekâ güvenlik algısının artma eğiliminde olabileceğine işaret etmektedir.

Ortalamlar incelendiğinde, yaş grupları arasında ortalama puanlar açısından bir artış eğilimi gözlenmekte; 18–24 yaş grubundan 35 yaş ve üzerine doğru gidildikçe YZGAAÖ puanlarının yükseldiği görülmektedir. Bu durum, ileri yaş gruplarının yapay zekâ sistemlerini daha güvenli algılama eğiliminde olabileceğini düşündürmektedir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının yaş gruplarına göre farklılaşıp farklılaşmadığını belirlemek amacıyla tek yönlü varyans analizi (ANOVA) uygulanmıştır. Yaş grupları arasındaki anlamlı farklılıkların hangi gruplardan kaynaklandığını belirlemek amacıyla Tukey çoklu karşılaştırma testi gerçekleştirilmiştir. Analiz sonuçları Tablo 38'te sunulmuştur.

Tablo 38. Yaş Gruplarına göre YZGAAÖ puanları için tek yönlü ANOVA sonuçları

Kaynak	Kareler Toplamı	df	Kare Ortalaması	F	p	Tukey
Grup içi	15900,919	2	7950,460	6,518	,002*	1-2, 1-3
Gruplar arası	676982,757	555	1219,789			
Toplam	692883,686	557				

*p < .05

Tablo 38'te katılımcıların Yapay Zekâ Güvenlik Algısı Ölçeği (YZGAAÖ) puanlarının yaş gruplarına göre farklılaşıp farklılaşmadığını belirlemek amacıyla gerçekleştirilen tek yönlü varyans analizi (ANOVA) sonuçları yer almaktadır. Analiz sonucunda yaş grupları arasında istatistiksel olarak anlamlı bir farklılık olduğu tespit edilmiştir [$F_{(2,555)}=6,518$; $p=.002$].

Bu bulgu, katılımcıların YZGAAÖ puanlarının yaş değişkenine bağlı olarak anlamlı düzeyde değiştiğini göstermektedir. Farklılığın hangi gruplar arasında olduğunu belirlemek amacıyla yapılan Tukey HSD post hoc testi sonuçlarına göre, 18–24 yaş grubu (Grup 1), hem 25–34 yaş grubuna (Grup 2) hem de 35 yaş ve üzeri gruba (Grup 3) kıyasla anlamlı düzeyde daha düşük puan almıştır.

Buna karşın 25–34 yaş grubu ile 35 yaş ve üzeri grup arasında anlamlı bir farklılık bulunmamıştır. Bu sonuç, yaş arttıkça yapay zekâ güvenlik algısının yükselme eğiliminde olduğunu ve özellikle en genç yaş grubunun diğer gruplara göre daha düşük algı düzeyine sahip olduğunu göstermektedir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) alt boyut puanlarının yaş gruplarına göre dağılımını belirlemek amacıyla tanımlayıcı istatistikler hesaplanmıştır. Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verilebilirlik, Toplumsal Etki ile Olumlu Beklentiler alt boyutlarına ilişkin ortalama ve standart sapma değerleri Tablo 39'da sunulmuştur.

Tablo 39. Yaş Gruplarına göre katılımcıların YZGAAÖ alt boyut puanlarının tanımlayıcı istatistikleri

Alt Boyutlar	Yaş	N	$\bar{X}$	Ss
Güvenlik Farkındalığı	1. 18-24 Yaş	58	22.41	8.627
	2. 25-34 Yaş	448	24.86	6.772
	3. 35 Yaş ve üstü	52	26.12	7.883
	Toplam	558	24.72	7.134
Ahlaki Değerler	1. 18-24 Yaş	58	33.81	13.570
	2. 25-34 Yaş	448	37.94	10.018
	3. 35 Yaş ve üstü	52	39.24	10.659
	Toplam	558	37.62	10.572
Sorumluluk ve Hesap Verilebilirlik	1. 18-24 Yaş	58	26.80	10.209
	2. 25-34 Yaş	448	29.50	7.846
	3. 35 Yaş ve üstü	52	30.80	8.452
	Toplam	558	29.33	8.220
Toplumsal Etki	1. 18-24 Yaş	58	26.00	11.415
	2. 25-34 Yaş	448	29.68	8.162
	3. 35 Yaş ve üstü	52	29.67	9.618

Alt Boyutlar	Yaş	N	$\bar{X}$	Ss
Olumlu Beklentiler	Toplam	558	29.29	8.753
	1. 18-24 Yaş	58	19.71	8.135
	2. 25-34 Yaş	448	22.87	5.978
	3. 35 Yaş ve üstü	52	24.25	6.260
	Toplam	558	22.66	6.347

Tablo 39’da katılımcıların Yapay Zekâ Güvenlik Algısı Ölçeği (YZGAAÖ) alt boyut puanlarının yaş gruplarına göre tanımlayıcı istatistikleri yer almaktadır. Bulgular genel olarak değerlendirildiğinde, tüm alt boyutlarda yaş grupları arasında benzer bir eğilim gözlenmekte ve yaş arttıkça ortalama puanların yükseldiği görülmektedir.

Güvenlik farkındalığı alt boyutunda en düşük ortalama puanın 18–24 yaş grubuna ($\bar{X}$ =22.41), en yüksek ortalama puanın ise 35 yaş ve üzeri gruba ($\bar{X}$ =26.12) ait olduğu belirlenmiştir. Benzer şekilde, ahlaki değerler alt boyutunda 18–24 yaş grubu ($\bar{X}$ =33.81) en düşük ortalamaya sahipken, 35 yaş ve üzeri grup ($\bar{X}$ =39.24) en yüksek ortalamayı göstermiştir.

Sorumluluk ve hesap verilebilirlik alt boyutunda da benzer bir dağılım gözlenmiş; 18–24 yaş grubunun ($\bar{X}$ =26.80) diğer yaş gruplarına göre daha düşük puan aldığı, 35 yaş ve üzeri grubun ($\bar{X}$ =30.80) ise en yüksek puana sahip olduğu belirlenmiştir. Toplumsal etki alt boyutunda da 18–24 yaş grubu ($\bar{X}$ =26.00) en düşük ortalamayı gösterirken, 25–34 yaş ($\bar{X}$ =29.68) ve 35+ yaş gruplarının ($\bar{X}$ =29.67) daha yüksek değerlere sahip olduğu görülmektedir.

Olumlu beklentiler alt boyutunda ise yaş grupları arasındaki fark daha belirgin hale gelmekte; 18–24 yaş grubu ($\bar{X}$ =19.71) en düşük ortalamaya sahipken, 35 yaş ve üzeri grup ($\bar{X}$ =24.25) en yüksek ortalamayı göstermektedir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) alt boyut puanlarının yaş gruplarına göre farklılaşıp farklılaşmadığını belirlemek amacıyla tek yönlü varyans analizi (ANOVA) uygulanmıştır. Yaş grupları arasında anlamlı farklılık bulunan alt boyutlarda, farklılığın hangi gruplardan kaynaklandığını belirlemek amacıyla Tukey çoklu karşılaştırma testi gerçekleştirilmiştir. Analiz sonuçları Tablo 40’te sunulmuştur.

Tablo 40’te katılımcıların Yapay Zekâ Güvenlik Algısı Ölçeği (YZGAAÖ) alt boyut puanlarının yaş gruplarına göre farklılaşıp farklılaşmadığını belirlemek amacıyla yapılan tek yönlü varyans analizi (ANOVA) sonuçları yer almaktadır.

Tablo 40. Yaş Gruplarına göre YZGAAÖ alt boyut puanları için tek yönlü ANOVA sonuçları

Alt Boyutlar	Kaynak	Kareler Toplamı	df	Kare Ortalaması	F	p	Tukey
Güvenlik Farkındalığı	Grup içi	424.601	2	212.300	4.220	.015*	1-2, 1-3
	Gruplar arası	27920.226	555	50.307			
	Toplam	28344.826	557				
Ahlaki Değerler	Grup içi	1034.351	2	517.176	4.688	.010*	1-2
	Gruplar arası	61220.617	555	110.307			
	Toplam	62254.968	557				
Sorumluluk ve Hesap Verilebilirlik	Grup içi	502.401	2	251.201	3.754	.024*	1-2, 1-3
	Gruplar arası	37137.599	555	66.915			
	Toplam	37640.000	557				
Toplumsal Etki	Grup içi	715.061	2	357.530	4.729	.009*	1-2
	Gruplar arası	41962.324	555	75.608			
	Toplam	42677.385	557				
Olumlu Beklentiler	Grup içi	662.796	2	331.398	8.447	.000*	1-2, 1-3
	Gruplar arası	21773.536	555	39.232			
	Toplam	22436.332	557				

*p < .05, **p<0.001

Analiz sonucunda tüm alt boyutlarda yaş grupları arasında istatistiksel olarak anlamlı farklılıklar olduğu tespit edilmiştir ($p<.05$).

Güvenlik farkındalığı alt boyutunda gruplar arasında anlamlı farklılık bulunmuştur [$F_{(2,555)}=4.220$; $p=.015$]. Tukey HSD sonuçlarına göre 18–24 yaş grubu, hem 25–34 yaş grubuna hem de 35 yaş ve üzeri gruba kıyasla anlamlı düzeyde daha düşük puan almıştır.

Ahlaki değerler alt boyutunda da yaş grupları arasında anlamlı farklılık olduğu belirlenmiştir [$F_{(2,555)}=4.688$; $p=.010$]. Tukey HSD sonuçlarına göre anlamlı farklılık yalnızca 18–24 yaş grubu ile 25–34 yaş grubu arasında ortaya çıkmıştır. 35 yaş ve üzeri grup ile diğer gruplar arasında anlamlı bir fark bulunmamıştır.

Sorumluluk ve hesap verilebilirlik alt boyutunda gruplar arasında anlamlı farklılık tespit edilmiştir [$F_{(2,555)}= 3.754$; $p=.024$]. Tukey HSD sonuçlarına göre 18–24 yaş grubu hem 25–34 yaş grubuna hem de 35 yaş ve üzeri gruba göre daha düşük puanlara sahiptir.

Toplumsal etki alt boyutunda da anlamlı farklılık bulunmuştur [$F_{(2,555)} = 4.729$; $p = .009$]. Tukey HSD sonuçlarına göre 18–24 yaş grubu, 25–34 yaş grubuna göre anlamlı düzeyde daha düşük puan almıştır. Diğer karşılaştırmalarda anlamlı farklılık gözlenmemiştir.

Olumlu beklentiler alt boyutunda ise en güçlü farklılık tespit edilmiştir [$F_{(2,555)} = 8.447$; $p < .001$]. Tukey HSD sonuçlarına göre 18–24 yaş grubu, hem 25–34 yaş hem de 35 yaş ve üzeri gruba kıyasla anlamlı düzeyde daha düşük puan almıştır.

Tüm alt boyutlarda 18–24 yaş grubunun diğer yaş gruplarına göre daha düşük puanlara sahip olduğu görülmektedir. Özellikle olumlu beklentiler ve sorumluluk/hesap verilebilirlik alt boyutlarında bu fark daha belirgin düzeydedir. Bu bulgular, yaş arttıkça yapay zekâya yönelik güvenlik algısı ve etik değerlendirmelerin güçlendiğini göstermektedir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının eğitim durumlarına göre dağılımını belirlemek amacıyla tanımlayıcı istatistikler hesaplanmıştır. Eğitim düzeylerine ilişkin ortalama puanlar ve standart sapma değerleri Tablo 41'de sunulmuştur.

Tablo 41. Eğitim Durumlarına göre katılımcıların YZGAAÖ puanlarının tanımlayıcı istatistikleri

Eğitim Durumu	N	$\bar{X}$	Ss
1. Lisans	431	144.73	32.934
2. Ön lisans	65	130.86	46.129
3. Lisansüstü	62	149.37	35.561
Toplam	558	143.63	35.270

Tablo 41 incelendiğinde katılımcıların YZGAAÖ puan ortalamalarının eğitim durumuna göre farklılaştığı görülmektedir. Lisans öğrencilerinin puan ortalaması 144.73 ($Ss = 32.934$), ön lisans öğrencilerinin puan ortalaması 130.86 ($Ss = 46.129$) ve lisansüstü öğrencilerin puan ortalaması ise 149.37 ($Ss = 35.561$) olarak belirlenmiştir. Toplam puan ortalaması 143.63 ($Ss = 35.270$)'tür.

Ortalamalar karşılaştırıldığında en yüksek puan ortalamasının lisansüstü eğitim düzeyindeki katılımcılara, en düşük puan ortalamasının ise ön lisans düzeyindeki katılımcılara ait olduğu görülmektedir. Bu durum, eğitim düzeyi yükseldikçe katılımcıların YZGAAÖ kapsamında ölçülen özelliklere ilişkin puanlarının artma eğiliminde olduğunu göstermektedir. Ayrıca ön

lisans grubunun standart sapma değerinin diğer gruplara göre daha yüksek olması ($S_s = 46.129$), bu gruptaki bireylerin puanlarının daha heterojen bir dağılım gösterdiğine işaret etmektedir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının eğitim durumuna göre farklılaşıp farklılaşmadığını belirlemek amacıyla tek yönlü varyans analizi (ANOVA) uygulanmıştır. Eğitim düzeyleri arasında anlamlı farklılık bulunması durumunda, farklılığın hangi gruplardan kaynaklandığını belirlemek amacıyla Tukey çoklu karşılaştırma testi gerçekleştirilmiştir. Analiz sonuçları Tablo 42'de sunulmuştur.

Tablo 42. Eğitim Durumlarına göre YZGAAÖ puanları için tek yönlü ANOVA sonuçları

Kaynak	Kareler Toplamı	df	Kare Ortalaması	F	p	Tukey
Grup içi	13165.149	2	6582.575	5.375	.005*	1-2, 2-3
Gruplar arası	679718.537	555	1224.718			
Toplam	692883.686	557				

*p < .05

Analiz sonucunda eğitim durumuna göre YZGAAÖ puanları arasında istatistiksel olarak anlamlı bir fark olduğu belirlenmiştir ($F_{(2,555)} = 5.375$, $p = .005$). Farkın hangi gruplardan kaynaklandığını belirlemek amacıyla gerçekleştirilen Tukey çoklu karşılaştırma testi sonuçlarına göre anlamlı farklılıkların lisans ile ön lisans grupları arasında (1-2) ve ön lisans ile lisansüstü grupları arasında (2-3) olduğu görülmüştür. Buna göre lisans öğrencilerinin ve lisansüstü öğrencilerin YZGAAÖ puanları, ön lisans öğrencilerinin puanlarından anlamlı düzeyde daha yüksektir. Bunun yanı sıra lisans ve lisansüstü grupları arasında anlamlı bir farklılık bulunmamıştır.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) alt boyut puanlarının eğitim durumlarına göre dağılımını belirlemek amacıyla tanımlayıcı istatistikler hesaplanmıştır. Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verilebilirlik, Toplumsal Etki ile Olumlu Beklentiler alt boyutlarına ilişkin ortalama ve standart sapma değerleri Tablo 43'te sunulmuştur.

Tablo 43. Eğitim Durumlarına göre katılımcıların YZGAAÖ alt boyut puanlarının tanımlayıcı istatistikleri

Alt Boyutlar	Eğitim Durumu	N	$\bar{X}$	S_s
Güvenlik Farkındalığı	1. Lisans	431	24.84	6.858

Alt Boyutlar	Eğitim Durumu	N	$\bar{X}$	Ss
Ahlaki Değerler	2. Ön lisans	65	22.65	8.207
	3. Lisansüstü	62	26.02	7.487
	Toplam	558	24.72	7.134
	1. Lisans	431	37.89	10.109
Sorumluluk ve Hesap Verilebilirlik	2. Ön lisans	65	34.38	13.188
	3. Lisansüstü	62	39.16	10.179
	Toplam	558	37.62	10.572
	1. Lisans	431	29.43	7.883
Toplumsal Etki	2. Ön lisans	65	27.25	10.034
	3. Lisansüstü	62	30.84	8.147
	Toplam	558	29.33	8.220
	1. Lisans	431	29.68	8.193
Olumlu Beklentiler	2. Ön lisans	65	26.51	11.122
	3. Lisansüstü	62	29.48	9.344
	Toplam	558	29.29	8.753
	1. Lisans	431	22.88	5.996
	2. Ön lisans	65	20.08	7.942
	3. Lisansüstü	62	23.87	6.266
	Toplam	558	22.66	6.347

Tablo incelendiğinde katılımcıların YZGAAÖ alt boyut puanlarının eğitim durumuna göre farklılaştığı görülmektedir. YZGAAÖ'nün tüm alt boyutlarında ön lisans düzeyindeki katılımcıların en düşük ortalama puanlara, lisansüstü düzeyindeki katılımcıların ise en yüksek ortalama puanlara sahip olduğu görülmektedir. Lisans grubunun puan ortalamaları ise çoğu alt boyutta bu iki grup arasında yer almaktadır.

Güvenlik farkındalığı alt boyutunda lisans öğrencilerinin ortalama puanı 24.84 (Ss = 6.858), ön lisans öğrencilerinin 22.65 (Ss = 8.207) ve lisansüstü öğrencilerin 26.02 (Ss = 7.487) olarak belirlenmiştir. En yüksek ortalama lisansüstü öğrencilerde, en düşük ortalama ise ön lisans öğrencilerinde görülmektedir. Bu sonuç, eğitim düzeyi arttıkça yapay zekâ sistemlerinin güvenlik riskleri ve veri koruma süreçlerine ilişkin farkındalığın yükseldiğine işaret etmektedir.

Ahlaki değerler alt boyutunda lisans öğrencilerinin ortalaması 37.89 (Ss = 10.109), ön lisans öğrencilerinin 34.38 (Ss = 13.188) ve lisansüstü öğrencilerin 39.16 (Ss = 10.179) olarak

hesaplanmıştır. Her ne kadar lisansüstü öğrenciler en yüksek ortalamaya sahip olsa da gruplar arasındaki farkların oldukça sınırlı olduğu görülmektedir.

Sorumluluk ve hesap verilebilirlik alt boyutunda lisans öğrencilerinin ortalaması 29.43 ($S_s = 7.883$), ön lisans öğrencilerinin 27.25 ($S_s = 10.034$) ve lisansüstü öğrencilerin 30.84 ($S_s = 8.147$) olarak bulunmuştur. Bu boyutta da en yüksek puan lisansüstü öğrencilerde, en düşük puan ise ön lisans öğrencilerde görülmektedir.

Toplumsal etki alt boyutunda lisans öğrencilerinin ortalama puanı 29.68 ($S_s = 8.193$), ön lisans öğrencilerinin 26.51 ($S_s = 11.122$) ve lisansüstü öğrencilerin 29.48 ($S_s = 9.344$) olarak belirlenmiştir. Lisans ve lisansüstü öğrencilerin puanlarının birbirine oldukça yakın olduğu görülmektedir.

Olumlu beklentiler alt boyutunda lisans öğrencilerinin ortalama puanı 22.88 ($S_s = 5.996$), ön lisans öğrencilerinin 20.08 ($S_s = 7.942$) ve lisansüstü öğrencilerin 23.87 ($S_s = 6.266$) olarak hesaplanmıştır. Bu alt boyutta en yüksek ortalama lisansüstü öğrencilerde, en düşük ortalama ise ön lisans öğrencilerde elde edilmiştir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) alt boyut puanlarının eğitim durumuna göre farklılaşıp farklılaşmadığını belirlemek amacıyla tek yönlü varyans analizi (ANOVA) uygulanmıştır. Eğitim düzeyleri arasında anlamlı farklılık bulunan alt boyutlarda, farklılığın hangi gruplardan kaynaklandığını belirlemek amacıyla Tukey çoklu karşılaştırma testi yapılmıştır. Analiz sonuçları Tablo 44'te sunulmuştur.

Tablo 44. Eğitim Durumlarına göre YZGAAÖ alt boyut puanları için tek yönlü ANOVA sonuçları

Alt Boyutlar	Kaynak	Kareler Toplamı	df	Kare Ortalaması	F	p	Tukey
Güvenlik Farkındalığı	Grup içi	390.396	2	195.198	3.875	.021*	2-3
	Gruplar arası	27954.430	555	50.368			
	Toplam	28344.826	557				
Ahlaki Değerler	Grup içi	859.321	2	429.661	3.884	.021*	1-2, 2-3
	Gruplar arası	61395.646	555	110.623			
	Toplam	62254.968	557				
Sorumluluk ve Hesap Verilebilirlik	Grup içi	427.821	2	213.910	3.190	.042*	2-3
	Gruplar arası	37212.179	555	67.049			

Alt Boyutlar	Kaynak	Kareler Toplamı	df	Kare Ortalaması	F	p	Tukey
Toplumsal Etki	Toplam	37640.000	557				
	Grup içi	572.569	2	286.285	3.774	.024*	1-2
	Gruplar arası	42104.816	555	75.865			
Olumlu Beklentiler	Toplam	42677.385	557				
	Grup içi	545.783	2	272.892	6.919	.001*	1-2, 2-3
	Gruplar arası	21890.548	555	39.442			
	Toplam	22436.332	557				

*p < .05

ANOVA sonuçları, Güvenlik Farkındalığı alt boyutunda eğitim durumuna göre anlamlı bir farklılık bulunmuştur ($F_{(2,555)}=3.875$, $p=.021<.05$). Bu sonuç, ön lisans, lisans ve lisansüstü gruplarının Güvenlik Farkındalığı puanlarının en az bir grup açısından birbirinden anlamlı biçimde farklılaştığını göstermektedir. Yapılan Tukey testi sonucunda anlamlı farklılığın 2-3, yani lisans ile lisansüstü grupları arasında olduğu belirlenmiştir. Buna göre eğitim düzeyi, katılımcıların yapay zekâya ilişkin güvenlik farkındalıklarında anlamlı bir farklılaşmaya yol açmaktadır. Ancak ANOVA tablosu tek başına farkın hangi grubun lehine olduğunu göstermediğinden, bu yönün kesin olarak ifade edilebilmesi için eğitim gruplarına ait ortalamaların incelenmesi gerekir.

Ahlaki Değerler alt boyutunda da eğitim durumuna göre istatistiksel olarak anlamlı bir farklılık tespit edilmiştir ($F_{(2,555)}=3.884$, $p=.021<.05$). Tukey testi sonucunda anlamlı farklılığın 1-2 ve 2-3 grupları arasında olduğu görülmektedir. Başka bir ifadeyle, ön lisans-lisans ve lisans-lisansüstü grupları arasında anlamlı farklılık bulunurken, ön lisans-lisansüstü grupları arasındaki fark anlamlı değildir. Bu bulgu, eğitim düzeyinin yapay zekânın ahlaki ve etik boyutlarına ilişkin algılarda farklılaşmaya neden olduğunu ve farklılaşmanın özellikle lisans grubunun diğer iki eğitim grubundan ayrılmasıyla ortaya çıktığını göstermektedir.

Sorumluluk ve Hesap Verilebilirlik alt boyutunda eğitim durumuna göre anlamlı bir farklılık bulunmuştur ($F_{(2,555)}=3.190$, $p=.042<.05$). Tukey testi sonucunda farklılığın 2-3, yani lisans ve lisansüstü grupları arasında olduğu belirlenmiştir. Buna göre, katılımcıların yapay zekâ kullanımında sorumluluk ve hesap verebilirliğe ilişkin algıları eğitim düzeyine göre farklılaşmaktadır. Bununla birlikte, bu farklılığın yönünü belirlemek için ilgili grupların ortalama puanlarının incelenmesi gerekmektedir.

Toplumsal Etki alt boyutunda da eğitim durumuna göre anlamlı bir farklılık saptanmıştır ($F_{(2,555)}=3.774$, $p=.024<.05$). Tukey testi sonucunda anlamlı farklılığın 1-2, yani ön lisans ve lisans grupları arasında olduğu görülmektedir. Buna göre, ön lisans ve lisans düzeyindeki katılımcıların yapay zekânın toplumsal etkilerine ilişkin algıları birbirinden anlamlı düzeyde farklılaşmaktadır. Lisansüstü grubunun ise diğer gruplardan istatistiksel olarak anlamlı biçimde ayrılmadığı anlaşılmaktadır.

Olumlu Beklentiler alt boyutu, eğitim durumuna göre en belirgin farklılaşmanın görüldüğü boyuttur. ANOVA sonucu $F_{(2,555)}=6.919$, $p=.001<.05$ olarak bulunmuştur. Bu sonuç, eğitim gruplarının Olumlu Beklentiler puanları arasında istatistiksel olarak anlamlı bir farklılık olduğunu göstermektedir. Tukey testi sonucunda farklılığın 1-2 ve 2-3 grupları arasında olduğu belirlenmiştir. Dolayısıyla ön lisans-lisans ve lisans-lisansüstü grupları arasında anlamlı farklılık bulunurken, ön lisans-lisansüstü grupları arasında anlamlı bir fark bulunmamıştır. Bu bulgu, eğitim düzeyinin katılımcıların yapay zekânın gelecekteki yararlarına ve olumlu sonuçlarına yönelik beklentileri üzerinde anlamlı bir farklılaşma oluşturduğunu göstermektedir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının akademik alanlarına göre farklılaşıp farklılaşmadığını belirlemek amacıyla bağımsız örneklem t-testi uygulanmıştır. Analiz sonuçları Tablo 45'te sunulmuştur.

Tablo 45. Katılımcıların Akademik alanlarına göre YZGAAÖ toplam puanları için bağımsız örneklem sonuçları

Akademik Alanlar	N	$\bar{X}$	Ss	df	t	p
Sayısal	282	143.26	34.418	556	-.253	.801
Sözel	276	144.01	36.178			

$p > .05$

Tablo 45 incelendiğinde, Sayısal alanda 282 katılımcının YZGAAÖ toplam puan ortalaması 143.26 ($Ss=34.418$), sözel alanda 276 katılımcının ortalaması ise 144.01 ($Ss=36.178$) olarak bulunmuştur. Sözel alandaki katılımcıların ortalaması sayısal alandaki katılımcılardan yaklaşık 0.75 puan daha yüksek olmasına rağmen, iki grup arasındaki bu fark istatistiksel olarak anlamlı değildir.

Yapılan bağımsız örneklem t-testi sonucunda $t(556)=-.253$, $p=.801>.05$ değeri elde edilmiştir. Buna göre, katılımcıların YZGAAÖ toplam puanları akademik alanlarına göre

anlamalı bir farklılık göstermemektedir. Başka bir ifadeyle, sayısal ve sözel alanda öğrenim gören katılımcıların yapay zekâ güvenlik ve ahlak algıları birbirine oldukça yakın düzeydedir.

Standart sapma değerleri de iki grubun puan dağılımlarının benzer olduğunu göstermektedir. Sayısal alandaki katılımcıların standart sapması 34.418, sözel alandaki katılımcıların standart sapması ise 36.178'dir. Sözel grubun değişkenliği bir miktar daha yüksek olmakla birlikte, bu durum gruplar arasında anlamlı bir farklılık bulunduğunu göstermemektedir.

Tablo 46. Katılımcıların Akademik alanlarına göre YZGAAÖ alt boyut puanları için bağımsız örneklem sonuçları

Alt Boyut	A.Alanlar	N	$\bar{X}$	Ss	df	t	p
Güvenlik Farkındalığı	Sayısal	282	24.45	7.062	556	-.898	.370
	Sözel	276	24.99	7.208			
Ahlaki Değerler	Sayısal	282	37.30	10.602	556	-.719	.472
	Sözel	276	37.95	10.551			
Sorumluluk ve Hesap Verilebilirlik	Sayısal	282	29.30	8.306	556	-.082	.934
	Sözel	276	29.36	8.147			
Toplumsal Etki	Sayısal	282	29.40	8.258	556	.296	.767
	Sözel	276	29.18	9.246			
Olumlu Beklentiler	Sayısal	282	22.80	6.150	556	.501	.617
	Sözel	276	22.53	6.550			

$p > .05$

Tablo 46'da Güvenlik Farkındalığı alt boyutunda sayısal alan katılımcılarının puan ortalaması 24.45 ($Ss=7.062$), sözel alan katılımcılarının puan ortalaması ise 24.99 ($Ss=7.208$) olarak bulunmuştur. Sözel alan grubunun ortalaması sayısal alan grubundan 0.54 puan daha yüksek olmasına rağmen, bu fark istatistiksel olarak anlamlı değildir ($t_{(556)}=-.898$, $p=.370>.05$). Buna göre, katılımcıların yapay zekâ kullanımına ilişkin güvenlik farkındalıkları akademik alanlarına göre anlamlı biçimde farklılaşmamaktadır.

Ahlaki Değerler alt boyutunda sayısal alan grubunun ortalaması 37.30 ($Ss=10.602$), sözel alan grubunun ortalaması ise 37.95 ($Ss=10.551$) olarak belirlenmiştir. Sözel alandaki katılımcıların ortalaması yaklaşık 0.65 puan daha yüksek olmakla birlikte, bu fark anlamlı değildir ($t_{(556)}=-.719$, $p=.472>.05$). Dolayısıyla yapay zekânın etik ve ahlaki boyutlarına ilişkin algıların akademik alanlara göre farklılaşmadığı sonucuna ulaşılmıştır.

Sorumluluk ve Hesap Verilebilirlik alt boyutunda sayısal alan katılımcılarının ortalaması 29.30 ($S_s=8.306$), sözel alan katılımcılarının ortalaması ise 29.36 ($S_s=8.147$) olarak hesaplanmıştır. İki grubun ortalamaları arasındaki fark yalnızca 0.06 puandır. Test sonucunun $t_{(556)}=-.082$, $p=.934>.05$ olması, grupların bu alt boyuttaki puanlarının birbirinden anlamlı biçimde farklılaşmadığını göstermektedir. Bu sonuç, akademik alanın yapay zekâ kullanımında sorumluluk ve hesap verebilirlik algısı üzerinde anlamlı bir farklılaştırıcı olmadığını ortaya koymaktadır.

Toplumsal Etki alt boyutunda sayısal alan grubunun puan ortalaması 29.40 ($S_s=8.258$), sözel alan grubunun ortalaması ise 29.18 ($S_s=9.246$) olarak bulunmuştur. Bu kez sayısal alan grubunun ortalaması sözel alandan 0.22 puan daha yüksektir. Ancak bu fark da istatistiksel olarak anlamlı değildir ($t_{(556)}=.296$, $p=.767>.05$). Dolayısıyla yapay zekânın toplumsal etkilerine ilişkin algıların sayısal ve sözel akademik alanlara göre anlamlı bir farklılık göstermediği söylenebilir.

Olumlu Beklentiler alt boyutunda sayısal alan grubunun ortalaması 22.80 ($S_s=6.150$), sözel alan grubunun ortalaması ise 22.53 ($S_s=6.550$) olarak belirlenmiştir. Sayısal alan grubunun ortalaması 0.27 puan daha yüksek olmasına rağmen, bu fark istatistiksel olarak anlamlı değildir ($t_{(556)}=.501$, $p=.617>.05$). Buna göre, katılımcıların yapay zekânın gelecekteki olumlu etkilerine ve sağlayabileceği yararlarla ilişkin beklentileri de akademik alanlarına göre anlamlı biçimde farklılaşmamaktadır.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının gelir durum algılarına göre dağılımını belirlemek amacıyla tanımlayıcı istatistikler hesaplanmıştır. Gelir durum algısına ilişkin ortalama puanlar ve standart sapma değerleri Tablo 47'de sunulmuştur.

Tablo 47. Gelir durum algılarına göre katılımcıların YZGAAÖ puanlarının tanımlayıcı istatistikleri

Gelir Durum Algısı	N	$\bar{X}$	S_s
1. Kötü	192	146.52	36.5324
2. Orta	195	139.31	37.2557
3. İyi	171	145.32	30.9627
Toplam	558	143.32	30.963

Tablo 47 incelendiğinde katılımcıların YZGAAÖ puan ortalamalarının gelir durumu algılarına göre kısmen farklılaştığı görülmektedir. Gelir durumunu kötü olarak değerlendiren 192 katılımcının YZGAAÖ toplam puan ortalaması 146.52 ($S_s=36.5324$) olarak bulunmuştur. Bu grup, üç gelir algısı grubu içerisinde en yüksek YZGAAÖ puan ortalamasına sahiptir. Bu sonuç, gelir durumunu kötü olarak algılayan katılımcıların yapay zekâ güvenlik ve ahlak algılarının betimsel olarak diğer gruplara kıyasla daha yüksek olduğunu göstermektedir.

Gelir durumunu orta olarak değerlendiren 195 katılımcının YZGAAÖ toplam puan ortalaması 139.31 ($S_s=37.2557$)'dir. Bu grup, üç grup içerisinde en düşük ortalamaya sahiptir. Kötü gelir algısına sahip grupla karşılaştırıldığında yaklaşık 7.21 puanlık bir fark bulunmaktadır. Bununla birlikte, söz konusu farkın istatistiksel olarak anlamlı olup olmadığı yalnızca tanımlayıcı istatistiklere bakılarak söylenemez.

Gelir durumunu iyi olarak değerlendiren 171 katılımcının YZGAAÖ toplam puan ortalaması ise 145.32 ($S_s=30.9627$) olarak belirlenmiştir. Bu grubun ortalaması, gelir durumunu kötü olarak algılayan gruptan yaklaşık 1.20 puan daha düşük, orta olarak algılayan gruptan ise yaklaşık 6.01 puan daha yüksektir..

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının gelir durum algısına göre farklılaşıp farklılaşmadığını belirlemek amacıyla tek yönlü varyans analizi (ANOVA) uygulanmıştır. Analiz sonuçları Tablo 48'de sunulmuştur.

Tablo 48. Gelir durum algılarına göre YZGAAÖ puanları için tek yönlü ANOVA sonuçları

Kaynak	Kareler Toplamı	df	Kare Ortalaması	F	p
Grup içi	5724.904	2	2862.452	2.312	.100
Gruplar arası	687158.782	555	1238.124		
Toplam	692883.686	557			

*p < .05

Analiz sonucunda gelir durumu algısına göre YZGAAÖ puanları arasında istatistiksel olarak anlamlı bir farklılık bulunmamıştır ($F(2,555) = 2.312$, $p = .100$). Başka bir ifadeyle, katılımcıların gelir durumlarını kötü, orta veya iyi olarak algılamaları yapay zekâ güvenliği ve etik algılarını anlamlı düzeyde etkilememektedir.

Tanımlayıcı istatistiklerde gruplar arasında belirli düzeyde puan farklılıkları gözlenmesine rağmen, bu farklılıklar istatistiksel olarak anlamlı değildir. Dolayısıyla gelir durumu algısının YZGAAÖ puanlarını açıklayan önemli bir değişken olmadığı söylenebilir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) alt boyut puanlarının gelir durum algılarına göre dağılımını belirlemek amacıyla tanımlayıcı istatistikler hesaplanmıştır. Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verilebilirlik, Toplumsal Etki ile Olumlu Beklentiler alt boyutlarına ilişkin ortalama ve standart sapma değerleri Tablo 49'da sunulmuştur.

Tablo 49. Gelir durum algısına göre katılımcıların YZGAAÖ alt boyut puanlarının tanımlayıcı istatistikleri

Alt Boyutlar	Gelir Algısı	Durum	N	$\bar{X}$	Ss
Güvenlik Farkındalığı	1. Kötü		192	25.07	7.526
	2. Orta		195	24.25	7.338
	3. İyi		171	24.85	6.422
	Toplam		558	24.72	7.134
Ahlaki Değerler	1. Kötü		192	38.26	10.647
	2. Orta		195	36.12	11.251
	3. İyi		171	38.62	9.505
	Toplam		558	37.62	10.572
Sorumluluk ve Hesap Verilebilirlik	1. Kötü		192	30.25	8.251
	2. Orta		195	28.39	8.785
	3. İyi		171	29.37	7.410
	Toplam		558	29.33	8.220
Toplumsal Etki	1. Kötü		192	30.08	8.857
	2. Orta		195	28.39	9.335
	3. İyi		171	29.43	7.862
	Toplam		558	29.29	8.753
Olumlu Beklentiler	1. Kötü		192	22.86	6.552
	2. Orta		195	22.15	6.566
	3. İyi		171	23.04	5.837
	Toplam		558	22.66	6.347

Tablo 49 incelendiğinde, katılımcıların gelir durum algılarına göre YZGAAÖ'nün Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verilebilirlik, Toplumsal Etki ve Olumlu Beklentiler alt boyutlarından aldıkları puanların farklılaştığı görülmektedir. Katılımcıların gelir durum algıları “kötü”, “orta” ve “iyi” olmak üzere üç grupta değerlendirilmiştir. Katılımcıların 192'si gelir durumunu kötü, 195'i orta ve 171'i ise iyi olarak algılamaktadır.

Güvenlik Farkındalığı alt boyutu açısından değerlendirildiğinde, gelir durumunu kötü olarak algılayan katılımcıların ortalamasının 25.07 (SS=7.526), orta olarak algılayanların 24.25 (SS=7.338) ve iyi olarak algılayanların 24.85 (SS=6.422) olduğu görülmektedir. Buna göre en yüksek Güvenlik Farkındalığı puanına gelir durumunu kötü olarak algılayan katılımcılar sahipken, en düşük ortalama orta gelir algısına sahip katılımcılarda görülmektedir. Kötü ve iyi gelir algısı gruplarının ortalamaları birbirine oldukça yakınken, orta gelir algısı grubunun ortalaması diğer iki gruba göre bir miktar daha düşüktür. Standart sapmalar incelendiğinde ise kötü gelir algısı grubundaki puanların daha fazla değişkenlik gösterdiği, iyi gelir algısı grubundaki puanların ise görece daha homojen olduğu söylenebilir.

Ahlaki Değerler alt boyutunda kötü gelir algısına sahip katılımcıların ortalaması 38.26 (SS=10.647), orta gelir algısına sahip katılımcıların ortalaması 36.12 (SS=11.251) ve iyi gelir algısına sahip katılımcıların ortalaması 38.62 (SS=9.505) olarak belirlenmiştir. Bu sonuçlara göre en yüksek Ahlaki Değerler puanı iyi gelir algısı grubunda, en düşük puan ise orta gelir algısı grubunda bulunmaktadır. İyi ve kötü gelir algısı gruplarının ortalamalarının birbirine oldukça yakın olması dikkat çekmektedir. Buna karşılık orta gelir algısı grubunun ortalaması bu iki gruptan daha düşüktür. Standart sapmalar açısından değerlendirildiğinde, orta gelir algısı grubunda puanların diğer gruplara göre daha fazla değişkenlik gösterdiği görülmektedir.

Sorumluluk ve Hesap Verilebilirlik alt boyutunda gelir durumunu kötü olarak algılayan katılımcıların ortalaması 30.25 (SS=8.251), orta olarak algılayanların 28.39 (SS=8.785) ve iyi olarak algılayanların 29.37 (SS=7.410) olarak hesaplanmıştır. Bu alt boyutta da en yüksek ortalamanın kötü gelir algısı grubunda, en düşük ortalamanın ise orta gelir algısı grubunda olduğu görülmektedir. Kötü ve iyi gelir algısı grupları arasındaki ortalama farkı yaklaşık 0.88 puan iken, kötü ve orta gelir algısı grupları arasındaki fark 1.86 puandır. Orta ve iyi gelir algısı grupları arasındaki fark ise yaklaşık 0.98 puandır. Bu bulgu, betimsel açıdan kötü gelir algısına sahip katılımcıların yapay zekâ kullanımında sorumluluk ve hesap verilebilirlik konusunda daha yüksek puanlara sahip olduklarını göstermektedir. Ancak söz konusu farkların istatistiksel olarak anlamlı olup olmadığı ANOVA sonuçlarıyla belirlenmelidir.

Toplumsal Etki alt boyutu incelendiğinde, kötü gelir algısı grubunun ortalamasının 30.08 (SS=8.857), orta gelir algısı grubunun 28.39 (SS=9.335) ve iyi gelir algısı grubunun 29.43 (SS=7.862) olduğu görülmektedir. Bu alt boyutta da en yüksek ortalama kötü gelir algısı grubunda, en düşük ortalama ise orta gelir algısı grubunda ortaya çıkmıştır. Kötü gelir algısı ile orta gelir algısı arasındaki ortalama farkı 1.69 puan, kötü gelir algısı ile iyi gelir algısı arasındaki fark 0.65 puan, iyi gelir algısı ile orta gelir algısı arasındaki fark ise 1.04 puandır. Dolayısıyla betimsel olarak gelir durumunu kötü olarak değerlendiren katılımcıların yapay zekânın toplumsal etkilerine ilişkin puanlarının diğer gruplardan daha yüksek olduğu görülmektedir.

Olumlu Beklentiler alt boyutunda ise kötü gelir algısına sahip katılımcıların ortalaması 22.86 (SS=6.552), orta gelir algısına sahip katılımcıların 22.15 (SS=6.566) ve iyi gelir algısına sahip katılımcıların 23.04 (SS=5.837) olarak bulunmuştur. Bu alt boyutta en yüksek ortalama iyi gelir algısı grubunda, en düşük ortalama ise orta gelir algısı grubundadır. İyi gelir algısı grubunun ortalaması kötü gelir algısı grubundan yalnızca 0.18 puan daha yüksektir. Buna karşılık iyi ve orta gelir algısı grupları arasındaki fark 0.89 puandır. Standart sapmalar incelendiğinde, orta gelir algısı grubunda puanların değişkenliğinin kötü gelir algısı grubuna oldukça yakın olduğu, iyi gelir algısı grubunda ise daha düşük bir dağılım gösterdiği görülmektedir.

Tüm alt boyutlarda orta gelir grubunun puan ortalamalarının diğer gruplara göre bir miktar daha düşük olduğu görülse de gruplar arasındaki farkların oldukça sınırlı olduğu dikkat çekmektedir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) alt boyut puanlarının gelir durum algısına göre farklılaşıp farklılaşmadığını belirlemek amacıyla tek yönlü varyans analizi (ANOVA) uygulanmıştır. Analiz sonuçları Tablo 50'de sunulmuştur.

Tablo 50. Gelir durum algısı göre YZGAAÖ alt boyut puanları için tek yönlü ANOVA sonuçları

Alt Boyutlar	Kaynak	Kareler Toplamı	df	Kare Ortalaması	F	p
Güvenlik Farkındalığı	Grup içi	69.815	2	34.907	.685	.504
	Gruplar arası	28275.011	555	50.946		
	Toplam	28344.826	557			
Ahlaki Değerler	Grup içi	686.650	2	343.325	3.095	.046*
	Gruplar arası	61568.318	555	110.934		
	Toplam	62254.968	557			
	Grup içi	333.358	2	166.679	2.480	.085

Alt Boyutlar	Kaynak	Kareler Toplamı	df	Kare Ortalaması	F	p
Sorumluluk ve Hesap Verilebilirlik	Gruplar arası	37306.642	555	67.219		
	Toplam	37640.000	557			
Toplumsal Etki	Grup içi	278.986	2	139.493	1.826	.162
	Gruplar arası	42398.400	555	76.394		
	Toplam	42677.385	557			
Olumlu Beklentiler	Grup içi	82.652	2	41.326	1.026	.359
	Gruplar arası	22353.680	555	40.277		
	Toplam	22436.332	557			

*p < .05

Analiz sonuçlarına göre Güvenlik farkındalığı puanları gelir durumu algısına göre anlamlı farklılık göstermemektedir ($F_{(2,555)}=0.685$, $p=.504$). Bu bulgu, katılımcıların gelir düzeyini nasıl algıladıklarının yapay zekâ güvenliği konusundaki farkındalıklarını etkilemediğini göstermektedir.

Ahlaki Değerler alt boyutunda gelir durum algısına göre gruplar arasında istatistiksel olarak anlamlı bir farklılık bulunmuştur ($F_{(2,555)}=3.095$, $p=.046$). $p=.046$ değeri .05 anlamlılık düzeyinden küçük olduğundan, gelir durum algısı gruplarının Ahlaki Değerler puanlarının birbirinden tamamen aynı olduğu varsayımı reddedilmektedir. Bununla birlikte, farklılığın hangi gruplar arasından kaynaklandığını belirlemek amacıyla gerçekleştirilen Tukey HSD çoklu karşılaştırma testi sonuçları incelendiğinde, gelir durum algısı gruplarının tamamının aynı homojen alt kümede yer aldığı görülmüştür. Buna göre orta gelir algısı grubunun ortalaması 36.12, kötü gelir algısı grubunun ortalaması 38.26 ve iyi gelir algısı grubunun ortalaması 38.62'dir. Betimsel olarak iyi ve kötü gelir algısı gruplarının Ahlaki Değerler puanlarının orta gelir algısı grubundan daha yüksek olduğu görülmekle birlikte, Tukey HSD testi sonucunda gruplar arasındaki ikili karşılaştırmaların hiçbirinde .05 anlamlılık düzeyinde istatistiksel olarak anlamlı bir fark belirlenmemiştir (Sig.=.059). Bu durum, omnibus ANOVA sonucunun sınırda anlamlı olmasına karşın, çoklu karşılaştırmalarda anlamlı bir ikili grup farklılığının ortaya çıkmadığını göstermektedir. Dolayısıyla gelir durum algısının Ahlaki Değerler alt boyutundaki farklılaşmasına ilişkin bulgunun temkinli biçimde değerlendirilmesi gerekmektedir.

Sorumluluk ve hesap verebilirlik boyutunda gelir durumu algısına göre anlamlı farklılık bulunmamıştır ($F_{(2,555)} = 2.480$, $p = .085$). Bu sonuç, yapay zekâ sistemlerinde sorumluluk ve hesap verebilirlik konusundaki algının gelir durumundan bağımsız olduğunu göstermektedir.

Toplumsal etki alt boyutunda da gelir durumu algısına göre anlamlı farklılık tespit edilmemiştir ($F_{(2,555)} = 1.826$, $p = .162$). Katılımcılar, yapay zekânın toplumsal etkilerini gelir durumlarından bağımsız olarak benzer biçimde değerlendirmektedir.

Olumlu beklentiler alt boyutunda da gelir durumu algısına göre anlamlı farklılık bulunmamıştır ($F_{(2,555)} = 1.026$, $p = .359$). Bu bulgu, yapay zekânın gelecekte sağlayacağı yararlarla ilişkin beklentilerin sosyoekonomik algıya göre değişmediğini göstermektedir.

Katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının yapay zekâ uygulamalarını kullanma durumuna göre farklılaşıp farklılaşmadığını belirlemek amacıyla bağımsız örneklem t-testi uygulanmıştır. Analiz sonuçları Tablo 51'de sunulmuştur.

Tablo 51. Katılımcıların YZ uygulamalarını kullanımına göre YZGAAÖ toplam puanları için bağımsız örneklem sonuçları

Akademik Alanlar	N	$\bar{X}$	Ss	df	t	p
Evet	310	142.41	32.309	556	-.918	.359
Hayır	248	145.17	38.666			

$p > .05$

Bağımsız örneklem t-testi sonuçlarına göre katılımcıların YZGAAÖ toplam puanları, yapay zekâ uygulamalarını kullanma durumuna göre anlamlı bir farklılık göstermemektedir ($t_{(556)} = -.918$, $p = .359$).

YZ uygulamalarını kullanan katılımcıların ortalama puanı 142.41 ($Ss = 32.309$), kullanmayan katılımcıların ortalama puanı ise 145.17 ($Ss = 38.666$) olarak belirlenmiştir. Gruplar arasındaki fark istatistiksel olarak anlamlı değildir ($p > .05$).

Tablo 51. Katılımcıların YZ uygulamalarını kullanımına göre YZGAAÖ alt boyut puanları için bağımsız örneklem sonuçları

Alt Boyut	YZ Uyg. Kullanımı	N	$\bar{X}$	Ss	df	t	p
Güvenlik Farkındalığı	Evet	310	24.67	6.452	556	-.164	.869
	Hayır	248	24.77	7.917			

Alt Boyut	YZ Uyg. Kullanımı	N	$\bar{X}$	Ss	df	t	p
Ahlaki Değerler	Evet	310	37.03	9.933	556	-1.487	.138
	Hayır	248	38.37	11.297			
Sorumluluk ve Hesap Verilebilirlik	Evet	310	28.80	7.714	556	-1.727	.085
	Hayır	248	30.00	8.783			
Toplumsal Etki	Evet	310	29.49	8.116	556	.588	.557
	Hayır	248	29.05	9.501			
Olumlu Beklentiler	Evet	310	22.42	5.940	556	-1.022	.307
	Hayır	248	22.97	6.821			

$p > .05$

YZ uygulamalarını kullanan katılımcıların güvenlik farkındalığı puanı 24.67 ($Ss = 6.452$), kullanmayanların ise 24.77 ($Ss = 7.917$) olarak bulunmuştur. Gruplar arasında anlamlı fark yoktur ($t_{(556)} = -0.164$, $p = .869$). Ahlaki değerler boyutunda kullanıcıların ortalaması 37.03 ($Ss = 9.933$), kullanmayanların ortalaması 38.37 ($Ss = 11.297$)'dir. Fark istatistiksel olarak anlamlı değildir ($t_{(556)} = -1.487$, $p = .138$).

Sorumluluk ve Hesap Verilebilirlik boyutunda YZ kullanan katılımcıların ortalaması 28.80 ($Ss = 7.714$), kullanmayanların 30.00 ($Ss = 8.783$) olarak belirlenmiştir. Fark anlamlı değildir ($t_{(556)} = -1.727$, $p = .085$). Toplumsal etki boyutunda YZ kullananların ortalaması 29.49 ($Ss = 8.116$), kullanmayanların 29.05 ($Ss = 9.501$)'dir. Gruplar arasında anlamlı fark yoktur ($t_{(556)} = .588$, $p = .557$). Olumlu beklentiler boyutunda kullanıcıların ortalaması 22.42 ($Ss = 5.940$), kullanmayanların ortalaması 22.97 ($Ss = 6.821$)'dir. Fark istatistiksel olarak anlamlı değildir ($t_{(556)} = -1.022$, $p = .307$).

5. TARTIŞMA VE SONUÇ

5.1. Sonuç

Bu araştırmada, bireylerin yapay zekâya ilişkin güvenlik ve ahlaki algılarını çok boyutlu olarak ölçebilecek geçerli ve güvenilir bir ölçme aracı geliştirilmesi amaçlanmıştır. Bu doğrultuda geliştirilen Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ), kapsam geçerliği, yapı geçerliği ve güvenirlik analizleri ile psikometrik açıdan ayrıntılı olarak değerlendirilmiştir. Ölçek geliştirme sürecinde uluslararası literatür, yapay zekâ etiği, güvenlik, sorumluluk ve toplumsal etki alanındaki kuramsal yaklaşımlar ile uzman görüşleri dikkate alınarak oluşturulan madde havuzu, aşamalı istatistiksel analizler sonucunda yeniden düzenlenmiş ve son hâline ulaştırılmıştır.

Ölçeğin geliştirilme süreci, klasik ölçek geliştirme çalışmalarında önerilen aşamalara uygun olarak yürütülmüştür. İlk aşamada kapsamlı literatür taraması yapılarak yapay zekâ güvenliği, etik ilkeler, hesap verebilirlik, toplumsal etkiler ve yapay zekâyâ yönelik olumlu beklentiler kapsamında toplam 60 maddelik bir madde havuzu oluşturulmuştur. Oluşturulan maddeler alan uzmanlarının görüşleri doğrultusunda kapsam geçerliği açısından değerlendirilmiş, dil ve içerik bakımından gerekli düzenlemeler yapıldıktan sonra ölçek uygulamaya hazır hâle getirilmiştir. Ölçeğin psikometrik özelliklerinin belirlenmesi amacıyla gerçekleştirilen analizler 558 katılımcıdan elde edilen veriler üzerinden yürütülmüştür.

Araştırmada öncelikle yapı geçerliğinin incelenmesi amacıyla her bir alt boyut için ayrı ayrı Açımlayıcı Faktör Analizi (AFA) gerçekleştirilmiştir. Analizler sonucunda faktör yükü düşük olan, kuramsal yapıyla yeterince örtüşmeyen veya faktör yapısını olumsuz etkileyen maddeler ölçekten çıkarılmıştır. Güvenlik Farkındalığı alt boyutunda beş, Ahlaki Değerler alt boyutunda iki, Sorumluluk ve Hesap Verebilirlik alt boyutunda dört, Toplumsal Etki alt boyutunda dört ve Olumlu Beklentiler alt boyutunda beş madde analizlerden çıkarılmıştır. Böylece başlangıçta 60 maddeden oluşan ölçek, toplam 40 maddelik son yapıya ulaşmıştır.

AFA sonuçları, her bir alt boyutun tek faktörlü bir yapı sergilediğini ortaya koymuştur. Faktör analizlerinin uygulanabilirliğini değerlendirmek amacıyla hesaplanan Kaiser-Meyer-Olkin (KMO) örneklem yeterlilik katsayılarının tüm alt boyutlarda .94'ün üzerinde olduğu görülmüştür. Güvenlik Farkındalığı (.947), Ahlaki Değerler (.947), Sorumluluk ve Hesap Verebilirlik (.956), Toplumsal Etki (.954) ve Olumlu Beklentiler (.949) alt boyutlarına ait KMO değerleri, örneklem büyüklüğünün faktör analizi için "mükemmel" düzeyde yeterli olduğunu göstermektedir. Benzer şekilde Bartlett Küresellik Testi sonuçlarının tüm alt boyutlarda anlamlı bulunması ($p < .001$), maddeler arasında faktör analizinin yapılmasını sağlayacak düzeyde korelasyon bulunduğunu göstermektedir.

Alt boyutların açıklanan varyans oranları da oldukça yüksek bulunmuştur. Güvenlik Farkındalığı alt boyutu toplam varyansın %92.92'sini, Ahlaki Değerler %82.69'unu, Sorumluluk ve Hesap Verebilirlik %83.82'sini, Toplumsal Etki %93.69'unu ve Olumlu Beklentiler ise %86.92'sini açıklamaktadır. Sosyal bilimlerde tek faktörlü ölçeklerde %40-60 düzeyindeki açıklanan varyansın yeterli kabul edildiği dikkate alındığında, elde edilen varyans oranlarının oldukça yüksek olduğu ve her bir alt boyutun kendi yapısını güçlü biçimde temsil ettiği söylenebilir.

AFA sonrasında elde edilen faktör yapısının doğrulanması amacıyla her alt boyut için Doğrulayıcı Faktör Analizi (DFA) uygulanmıştır. DFA sonuçları incelendiğinde, tüm alt boyutlarda model uyum indekslerinin kabul edilebilir hatta çoğunlukla iyi uyum düzeyinde olduğu görülmektedir. Özellikle CFI, NFI, NNFI (TLI), IFI, GFI ve AGFI değerlerinin büyük çoğunluğunun .95'in üzerinde olması, geliştirilen modellerin verilerle yüksek düzeyde uyum gösterdiğini ortaya koymaktadır. RMSEA değerlerinin .041 ile .079 arasında, RMR değerlerinin .015 ile .026 arasında ve SRMR değerlerinin ise .011 ile .018 arasında değişmesi, ölçme modellerinin kabul edilebilir hatta iyi uyum sergilediğini göstermektedir. Böylece AFA ile elde edilen tek faktörlü yapıların DFA ile de doğrulandığı söylenebilir.

Ölçeğin güvenilirliğini belirlemek amacıyla gerçekleştirilen analizlerde oldukça güçlü sonuçlar elde edilmiştir. Alt boyutlara ilişkin Cronbach Alfa katsayıları Güvenlik Farkındalığı için .942, Ahlaki Değerler için .953, Sorumluluk ve Hesap Verebilirlik için .952, Toplumsal Etki için .966 ve Olumlu Beklentiler için .966 olarak hesaplanmıştır. Bu değerler, ölçeğin tüm alt boyutlarının çok yüksek düzeyde iç tutarlılığa sahip olduğunu göstermektedir. Ayrıca madde-toplam korelasyonlarının tüm alt boyutlarda .70'in üzerinde olması ve alt-üst %27'lik gruplar arasında gerçekleştirilen bağımsız örneklem t-testlerinin bütün maddelerde istatistiksel olarak anlamlı bulunması ($p < .001$), maddelerin ayırt edicilik gücünün oldukça yüksek olduğunu ortaya koymaktadır. Başka bir ifadeyle, ölçeği oluşturan maddelerin tamamı ölçülmek istenen yapıyı başarılı biçimde temsil etmekte ve bireyleri ilgili özellik bakımından ayırt edebilmektedir.

Araştırmada geliştirilen ölçek, yapay zekâya yönelik algıyı yalnızca etik ya da yalnızca güvenlik boyutunda ele almak yerine, birbirini tamamlayan beş temel boyut altında değerlendirmektedir. Bu yönüyle ölçek, yapay zekâya ilişkin çok boyutlu algıyı bütüncül bir bakış açısıyla ölçebilmektedir.

İlk alt boyut olan Güvenlik Farkındalığı, bireylerin yapay zekâ sistemlerinin oluşturabileceği güvenlik tehditleri, veri güvenliği, siber saldırılar ve kişisel mahremiyet konularındaki farkındalık düzeylerini ölçmektedir. Bu alt boyut, yapay zekâ sistemlerinin yalnızca teknolojik avantajlarına değil, aynı zamanda oluşturabileceği risklere ilişkin bilinç düzeyini değerlendirmektedir.

İkinci alt boyut olan Ahlaki Değerler, bireylerin yapay zekânın etik ilkelere uygun geliştirilmesi, adalet, tarafsızlık, etik denetim, insan vicdanı ve toplumsal değerlerle uyumlu

kullanımı konularındaki görüşlerini ortaya koymaktadır. Bu boyut, yapay zekâ sistemlerinin etik yönüne ilişkin bireysel değerlendirmelerin ölçülmesine katkı sağlamaktadır.

Üçüncü alt boyut olan Sorumluluk ve Hesap Verebilirlik, yapay zekâ sistemlerinde ortaya çıkabilecek hatalardan kimlerin sorumlu olacağı, şeffaflık, hukuki sorumluluk, geliştirici sorumluluğu ve hesap verebilirlik mekanizmalarına ilişkin algıları değerlendirmektedir. Günümüzde yapay zekâ yönetişimi açısından öne çıkan hesap verebilirlik kavramını ölçmesi bakımından bu boyut önemli bir katkı sunmaktadır.

Dördüncü alt boyut olan Toplumsal Etki, yapay zekânın işsizlik, demokrasi, sosyal ilişkiler, insan hakları, özgürlükler ve toplumsal güven üzerindeki olası etkilerine ilişkin algıları kapsamaktadır. Böylece bireylerin yapay zekânın toplumsal dönüşüm oluşturabilecek yönlerine ilişkin değerlendirmeleri ölçülebilmektedir.

Son alt boyut olan Olumlu Beklentiler ise yapay zekânın sağlık, eğitim, sürdürülebilir kalkınma, yaşam kalitesi, iş yükünün azaltılması ve karar verme süreçlerinin iyileştirilmesi gibi alanlarda sağlayabileceği faydalara yönelik algıları belirlemektedir. Bu boyut, bireylerin yapay zekâyâ ilişkin olumlu tutumlarını ve teknolojiye yönelik beklentilerini değerlendirmesi bakımından ölçeğin dengeleyici yönünü oluşturmaktadır.

Araştırma sonucunda geliştirilen YZGAAÖ, yalnızca yapay zekâ risklerini ya da yalnızca olumlu yönlerini ölçen bir araç olmaktan öte, güvenlik, etik, sorumluluk, toplumsal etki ve olumlu beklentileri birlikte değerlendiren bütüncül bir ölçme aracı niteliği taşımaktadır. Bu yönüyle ölçek, yapay zekâ okuryazarlığı, yapay zekâ etiği, dijital vatandaşlık, teknoloji kabulü ve yapay zekâ yönetişimi alanlarında yürütülecek araştırmalarda kullanılabilecek güçlü bir veri toplama aracı sunmaktadır.

Bu doğrultuda, gerçekleştirilen kapsam geçerliği, açımlayıcı ve doğrulayıcı faktör analizleri ile güvenilirlik analizleri birlikte değerlendirildiğinde, Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği'nin (YZGAAÖ) 40 maddeden ve beş alt boyuttan oluşan yapısının geçerli ve güvenilir olduğu belirlenmiştir. Ölçek, yüksek iç tutarlılık katsayıları, güçlü faktör yapısı, yüksek açıklanan varyans oranları ve kabul edilebilir düzeyin üzerinde model uyum indeksleri ile desteklenmiştir. Dolayısıyla geliştirilen ölçeğin, farklı örneklem gruplarında yapay zekâyâ yönelik güvenlik ve ahlak algısını ölçmek amacıyla güvenle kullanılabilecek psikometrik açıdan güçlü bir ölçme aracı olduğu söylenebilir. Ayrıca ölçeğin, Türkiye'de bu alanda

geliştirilen kapsamlı ve çok boyutlu ölçme araçlarından biri olması nedeniyle, gelecekte gerçekleştirilecek akademik çalışmalara ve yapay zekâ politikalarının geliştirilmesine önemli katkılar sağlayacağı değerlendirilmektedir.

5.2. Bulguların Tartışılması

5.2.1. Ölçek Gelişim Sürecinin literatür ile tartışılması

Bu araştırmada geliştirilen Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ), bireylerin yapay zekâ teknolojilerine yönelik güvenlik ve etik algılarını bütüncül bir yaklaşımla değerlendirmek amacıyla geliştirilmiştir. Ölçek geliştirme sürecinde klasik ölçek geliştirme ilkeleri izlenmiş; kapsam geçerliği, yapı geçerliği ve güvenirlik analizleri sistematik olarak yürütülmüştür. Elde edilen bulgular, geliştirilen ölçeğin hem kuramsal altyapı hem de psikometrik özellikler bakımından güçlü bir ölçme aracı olduğunu göstermektedir.

Ölçeğin geliştirilmesinde öncelikle yapay zekâ güvenliği, etik ilkeler, sorumluluk ve hesap verebilirlik, toplumsal etkiler ile yapay zekâdan beklenen yararlar konularına ilişkin güncel uluslararası literatür incelenmiştir. Son yıllarda yapay zekâ alanında yapılan çalışmalar, etik ve güvenlik konularının teknik performans kadar önemli hâle geldiğini göstermektedir (Díaz-Rodríguez et al., 2023; Liu et al., 2022; Floridi & Cows, 2019; Jobin, Ienca, & Vayena, 2019). Özellikle şeffaflık, adalet, zarar vermeme, hesap verebilirlik ve insan özerkliğinin korunması gibi ilkeler, günümüzde yapay zekâ yönetişiminin temel bileşenleri olarak kabul edilmektedir (European Commission, 2019; OECD, 2019; UNESCO, 2021). Bu doğrultuda geliştirilen YZGAAÖ'nün alt boyutlarının, uluslararası kuruluşlar tarafından önerilen etik yapay zekâ ilkeleriyle büyük ölçüde örtüştüğü görülmektedir.

Ölçeğin kapsam geçerliği sürecinde uzman görüşlerinden yararlanılması ve madde havuzunun kuramsal çerçeve doğrultusunda oluşturulması, ölçek geliştirme literatüründe önerilen yöntemlerle uyumludur (DeVellis, 2017; Boateng et al., 2018). Uzman değerlendirmeleri sonucunda maddelerin kapsam, ifade açıklığı ve ölçülmek istenen yapıyı temsil etme düzeyi gözden geçirilmiş, gerekli düzenlemeler yapıldıktan sonra uygulama aşamasına geçilmiştir. Bu süreç, geliştirilen ölçeğin içerik geçerliğini güçlendirmiştir.

Araştırmada yapı geçerliğinin belirlenmesi amacıyla gerçekleştirilen Açıklayıcı Faktör Analizi (AFA) sonuçları, her alt boyutun tek faktörlü bir yapı oluşturduğunu göstermiştir. Analizler

sonucunda faktör yükü düşük veya kuramsal yapıyla uyum göstermeyen maddelerin ölçekten çıkarılması, ölçek geliştirme çalışmalarında önerilen psikometrik yaklaşımlarla uyumludur (Hair et al., 2019). Başlangıçta 60 maddeden oluşan ölçeğin analizler sonucunda 40 maddeye indirgenmesi, ölçeğin ölçme gücünü azaltmamış; aksine her bir alt boyutun daha homojen ve güçlü bir yapıya kavuşmasını sağlamıştır. Benzer şekilde birçok ölçek geliştirme çalışmasında, ilk madde havuzunun yaklaşık %20–40'ının analizler sonucunda çıkarıldığı görülmektedir (Worthington & Whittaker, 2006). Bu açıdan değerlendirildiğinde, araştırmada gerçekleştirilen madde eleme süreci ölçek geliştirme literatürüyle tutarlıdır.

Araştırmada elde edilen KMO katsayılarının (.947-.956) oldukça yüksek olması, örneklem büyüklüğünün faktör analizi için mükemmel düzeyde yeterli olduğunu göstermektedir. Kaiser (1974), KMO değerinin .90'ın üzerinde olmasını "marvelous (mükemmel)" olarak değerlendirmektedir. Aynı şekilde Bartlett Küresellik Testi sonuçlarının tüm alt boyutlarda anlamlı bulunması, maddeler arasında faktör analizi yapılabilecek düzeyde ilişki bulunduğunu göstermektedir. Bu bulgular, ölçeğin yapı geçerliğinin güçlü biçimde desteklendiğini ortaya koymaktadır.

Alt boyutlarda açıklanan varyans oranlarının %82 ile %94 arasında değişmesi dikkat çekici bir diğer bulgudur. Sosyal bilimlerde tek faktörlü ölçeklerde %40'ın üzerindeki açıklanan varyans yeterli kabul edilirken (Hair et al., 2019), geliştirilen ölçeğin tüm alt boyutlarında bu değerin oldukça üzerinde sonuçlar elde edilmiştir. Özellikle Güvenlik Farkındalığı (%92.92), Toplumsal Etki (%93.69) ve Olumlu Beklentiler (%86.92) alt boyutlarında ulaşılan varyans oranları, maddelerin ilgili yapıyı oldukça güçlü biçimde temsil ettiğini göstermektedir. Bu durum, her alt boyutun kavramsal bütünlüğünün yüksek olduğuna işaret etmektedir.

Araştırmada Açıklayıcı Faktör Analizi sonrasında gerçekleştirilen Doğrulayıcı Faktör Analizi (DFA) sonuçları da geliştirilen modelin güçlü yapısını desteklemektedir. Özellikle CFI, IFI, NFI ve NNFI değerlerinin büyük bölümünün .98'in üzerinde olması, model-veri uyumunun oldukça yüksek olduğunu göstermektedir. RMSEA değerlerinin .041 ile .079 arasında değişmesi ise modellerin iyi ya da kabul edilebilir uyum sergilediğini ortaya koymaktadır. Hu ve Bentler (1999), CFI ve TLI değerlerinin .95'in üzerinde, RMSEA değerinin ise .08'in altında olmasının iyi model uyumuna işaret ettiğini belirtmektedir. Dolayısıyla araştırmada elde edilen DFA bulguları, geliştirilen ölçeğin faktör yapısının istatistiksel açıdan doğrulandığını göstermektedir.

Ölçeğin güvenilirlik analizleri de oldukça güçlü sonuçlar ortaya koymuştur. Cronbach Alfa katsayılarının .942 ile .966 arasında değişmesi, tüm alt boyutların yüksek düzeyde iç tutarlılığa sahip olduğunu göstermektedir. Nunnally ve Bernstein (1994), .70'in üzerindeki alfa katsayılarının kabul edilebilir, .80'in üzerindeki değerlerin yüksek, .90'ın üzerindeki değerlerin ise çok yüksek güvenilirliği ifade ettiğini belirtmektedir. Bu doğrultuda geliştirilen ölçeğin tüm alt boyutlarının oldukça yüksek güvenilirlik düzeyine sahip olduğu söylenebilir. Ayrıca madde-toplam korelasyonlarının .70'in üzerinde olması ve alt-üst %27 gruplar arasında gerçekleştirilen t-testlerinde tüm maddelerin anlamlı biçimde ayırt edici bulunması, ölçeğin madde kalitesinin yüksek olduğunu göstermektedir.

Araştırmada geliştirilen ölçeğin en önemli özelliklerinden biri, yapay zekâya ilişkin algıyı yalnızca etik veya yalnızca teknoloji kabulü açısından ele almamasıdır. Günümüzde geliştirilen birçok ölçek belirli bir boyuta odaklanmaktadır. Örneğin, Wang ve arkadaşları (2022) tarafından geliştirilen Yapay Zekâ Okuryazarlığı Ölçeği daha çok bireylerin yapay zekâ bilgisi ve kullanım becerilerini değerlendirmektedir. Benzer şekilde Schepman ve Rodway (2020) tarafından geliştirilen General Attitudes toward Artificial Intelligence Scale (GAAIS), bireylerin yapay zekâya yönelik genel olumlu ve olumsuz tutumlarını ölçmektedir. Buna karşın YZGAAÖ; güvenlik farkındalığı, ahlaki değerler, sorumluluk ve hesap verebilirlik, toplumsal etki ve olumlu beklentiler olmak üzere yapay zekâya ilişkin güncel tartışma alanlarının tamamını kapsayan çok boyutlu bir yapı sunmaktadır. Bu yönüyle ölçek, mevcut ölçme araçlarına göre daha kapsamlı bir değerlendirme yapabilme potansiyeline sahiptir.

Özellikle son yıllarda yayımlanan yapay zekâ etik rehberleri incelendiğinde, güvenlik, şeffaflık, hesap verebilirlik, insan denetimi ve toplumsal fayda ilkelerinin ortak payda olarak öne çıktığı görülmektedir (Camilleri, 2023; Jobin et al., 2019; OECD, 2019; UNESCO, 2021). Geliştirilen ölçeğin alt boyutlarının bu temel ilkeleri kapsaması, ölçeğin yalnızca psikometrik açıdan değil, aynı zamanda kavramsal açıdan da güçlü bir temele dayandığını göstermektedir.

Bu araştırmada geliştirilen Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği, kapsam geçerliği, yapı geçerliği ve güvenilirlik analizlerinden elde edilen bulgular doğrultusunda psikometrik özellikleri güçlü bir ölçme aracı olarak değerlendirilmektedir. Ölçeğin yüksek faktör yükleri, yüksek açıklanan varyans oranları, doğrulayıcı faktör analizinde elde edilen iyi uyum indeksleri, yüksek iç tutarlılık katsayıları ve güçlü madde ayırt edicilik değerleri, ölçeğin bilimsel açıdan güvenilir ve geçerli olduğunu göstermektedir. Ayrıca yapay zekâya ilişkin güvenlik, etik, sorumluluk, toplumsal etki ve olumlu beklentileri tek bir ölçme aracı altında bir

araya getirmesi, YZGAAÖ'yü literatürdeki mevcut ölçeklerden ayıran önemli bir özellik olarak değerlendirilebilir. Bu yönüyle ölçeğin, gelecekte eğitim, psikoloji, bilişim, yapay zekâ etiği ve teknoloji politikaları alanlarında yürütülecek çalışmalarda kullanılabilecek önemli bir ölçme aracı olduğu düşünülmektedir.

5.2.2. Cinsiyet bulgularının tartışılması

Bu araştırmada katılımcıların Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının cinsiyete göre anlamlı bir farklılık göstermediği belirlenmiştir ($t=-1.446$, $p>.05$). Benzer şekilde, ölçeğin Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verebilirlik ve Toplumsal Etki alt boyutlarında da kadın ve erkek katılımcılar arasında istatistiksel olarak anlamlı bir farklılık bulunmamıştır ($p>.05$). Erkek katılımcıların ortalama puanları kadın katılımcılardan sayısal olarak biraz daha yüksek olmasına rağmen, bu fark istatistiksel anlamlılık düzeyine ulaşmamıştır. Bu bulgu, yapay zekâya ilişkin güvenlik ve etik algılarının biyolojik cinsiyetten çok bireylerin bilgi düzeyi, eğitim deneyimleri ve teknolojiye maruz kalma süreçleri gibi değişkenlerden etkilendiğini düşündürmektedir. Fakat bu durum Olumlu Beklentiler alt boyutunda öyle değildir bu alt boyutta anlamlı bir farklılık tespit edilmiştir.

Son yıllarda yapay zekâya yönelik tutumları inceleyen çalışmalar da cinsiyetin tek başına belirleyici bir değişken olmadığını göstermektedir (Russo et al., 2025; Rahman et al., 2024; Kim & Lee, 2023). Örneğin, Schepman ve Rodway (2020) tarafından geliştirilen General Attitudes toward Artificial Intelligence Scale (GAAIS) ile yürütülen araştırmada, kadın ve erkek katılımcılar arasında yapay zekâya yönelik genel tutumlarda yalnızca sınırlı farklılıklar belirlenmiş; bu farklılıkların teknoloji kullanım deneyimi ve bireysel özellikler kontrol edildiğinde büyük ölçüde ortadan kalktığı belirtilmiştir. Benzer şekilde Araujo, Helberger, Kruikemeier ve De Vreese (2020), algoritmik karar sistemlerine duyulan güvenin cinsiyetten ziyade algoritmaların nasıl çalıştığını bilme düzeyi ve algoritmik okuryazarlıkla ilişkili olduğunu ortaya koymuştur.

Araştırmanın bulguları, yapay zekâ etiği alanında geliştirilen uluslararası politika belgeleriyle de dolaylı olarak örtüşmektedir. OECD (2019), UNESCO (2021) ve Avrupa Komisyonu (2019) tarafından yayımlanan yapay zekâ etik ilkelerinde güvenlik, şeffaflık, adalet ve hesap verebilirlik ilkelerinin tüm bireyler için ortak değerler olduğu vurgulanmaktadır. Bu ilkeler

cinsiyete özgü değil, evrensel etik ilkeler olarak ele alınmaktadır. Dolayısıyla katılımcıların benzer etik değerlendirmeler yapmaları beklenen bir durumdur.

Elde edilen sonuçlar, yapay zekâ okuryazarlığı alanındaki bazı araştırmalarla da paralellik göstermektedir. Ng, Leung, Chu ve Qiao (2021) tarafından gerçekleştirilen sistematik derleme çalışmasında, yapay zekâ okuryazarlığını etkileyen temel değişkenlerin eğitim fırsatları, dijital yeterlik, teknoloji deneyimi ve problem çözme becerileri olduğu; cinsiyet değişkeninin ise tutarlı bir etkisinin bulunmadığı belirtilmiştir. Benzer şekilde Long ve Magerko (2020), yapay zekâ okuryazarlığının bireylerin teknolojiyle etkileşim sıklığı ve eğitim geçmişiyle daha yakından ilişkili olduğunu ifade etmektedir.

Araştırmanın bulguları, daha önce yürütülen bazı çalışmalarla farklılık göstermektedir (Borwein et al., 2026; Naiseh et al., 2025). Örneğin, Chaudhry ve Kazim (2022) kadın katılımcıların yapay zekânın etik risklerine karşı daha duyarlı olduklarını bildirirken, bazı çalışmalarda erkek katılımcıların yapay zekâ teknolojilerine daha olumlu yaklaştıkları belirtilmiştir (Schepman & Rodway, 2020). Ancak bu araştırmaların büyük bölümünde incelenen değişkenler yapay zekâyâ yönelik genel tutum, teknoloji kabulü veya kullanım isteği olup, doğrudan güvenlik ve ahlaki algıyı ölçmemektedir. Bu nedenle bulgular arasındaki farklılıkların kullanılan ölçme araçlarından ve araştırma kapsamından kaynaklanmış olması olasıdır.

Bu çalışmada kullanılan YZGAAÖ, güvenlik farkındalığı, etik değerler, sorumluluk, toplumsal etki ve olumlu beklentileri birlikte değerlendiren çok boyutlu bir yapıya sahiptir. Bu nedenle ölçek, bireylerin yalnızca teknolojiye yönelik olumlu veya olumsuz tutumlarını değil, aynı zamanda etik değerlendirme süreçlerini de ölçmektedir. Etik değerlendirmelerin ise bireysel cinsiyetten ziyade toplumsal normlar, eğitim ve mesleki deneyimlerden daha fazla etkilenmesi beklenmektedir. Nitekim çalışmada hem kadın hem de erkek katılımcıların tüm alt boyutlarda birbirine oldukça yakın ortalamalara sahip olması bu yorumu desteklemektedir.

Öte yandan, örneklem dağılımı incelendiğinde kadın (%53,4) ve erkek (%46,6) katılımcı sayılarının birbirine yakın olduğu görülmektedir. Dengeli örneklem yapısı, cinsiyet karşılaştırmalarının daha güvenilir biçimde yapılmasına olanak sağlamıştır. Dolayısıyla anlamlı farklılığın bulunmaması, örneklem büyüklüğünün yetersizliğinden değil, gerçekten cinsiyet değişkeninin bu araştırma grubunda belirleyici olmamasından kaynaklanmış olabilir.

Araştırmanın bu bulgusu, daha önce aynı araştırmacı tarafından geliştirilen Yapay Zekâ Kaygısı Ölçeği ile elde edilen sonuçlardan farklılık göstermektedir. Önceki çalışmada kadın katılımcıların yapay zekâ kaygısı düzeylerinin erkeklerden anlamlı düzeyde daha yüksek olduğu belirlenmiştir. Buna karşın mevcut araştırmada güvenlik ve ahlaki algılar açısından cinsiyet farklılığı ortaya çıkmamıştır. Bu durum, kaygı ile etik farkındalık kavramlarının birbirinden farklı psikolojik yapılar olduğunu göstermektedir. Yapay zekâyâ ilişkin kaygılar bireysel özelliklerden ve duygusal süreçlerden daha fazla etkilenirken, güvenlik ve etik değerlendirmeler bilgi, deneyim ve toplumsal değerlerle daha yakından ilişkilidir.

Araştırmada cinsiyet değişkeninin yapay zekâ güvenlik ve ahlak algısı üzerinde anlamlı bir etkisinin bulunmaması, yapay zekâyâ ilişkin etik farkındalık ve güvenlik bilincinin toplumun her kesiminde benzer biçimde gelişmeye başladığını göstermektedir. Bu bulgu, yapay zekâ etiği eğitimlerinin veya farkındalık çalışmalarının kadın ve erkek bireyler için farklılaştırılmasına gerek olmadığını; bunun yerine tüm bireyleri kapsayan ortak eğitim programlarının daha uygun olacağını düşündürmektedir. Bununla birlikte, gelecekte yapılacak araştırmalarda cinsiyet değişkeninin teknoloji kullanım deneyimi, yapay zekâ okuryazarlığı, dijital yeterlik ve mesleki alan gibi değişkenlerle birlikte incelenmesi, yapay zekâ güvenliği ve etik algısını etkileyen faktörlerin daha kapsamlı biçimde açıklanmasına katkı sağlayacaktır.

5.2.3 Yaş bulgularının tartışılması

Bu araştırmada katılımcıların yaş gruplarına göre Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarında istatistiksel olarak anlamlı farklılık olduğu belirlenmiştir ($F=6.518$, $p<.01$). Tukey çoklu karşılaştırma testi sonuçlarına göre anlamlı farklılık, 18–24 yaş grubu ile 25–34 yaş grubu ve 18–24 yaş grubu ile 35 yaş ve üzeri grup arasında ortaya çıkmıştır. Ortalama puanlar incelendiğinde, 18–24 yaş grubunun ($\bar{X}=144,86$) en düşük puana sahip olduğu, 25–34 yaş grubunun ($\bar{X}=128,73$) ve özellikle 35 yaş ve üzeri grubun ($\bar{X}=150,08$) daha yüksek puanlar aldığı görülmektedir. Bu sonuç, yaş ilerledikçe bireylerin yapay zekâyâ ilişkin güvenlik ve etik konularında daha yüksek farkındalık geliştirdiklerini göstermektedir.

Alt boyutlar incelendiğinde de benzer bir eğilim gözlenmiştir. Güvenlik Farkındalığı, Sorumluluk ve Hesap Verebilirlik ile Olumlu Beklentiler alt boyutlarında 18–24 yaş grubunun hem 25–34 yaş hem de 35 yaş ve üzeri gruplardan anlamlı düzeyde daha düşük puan aldığı belirlenmiştir. Ahlaki Değerler ve Toplumsal Etki alt boyutlarında ise anlamlı farklılık yalnızca 18–24 yaş ile 25–34 yaş grupları arasında ortaya çıkmıştır. Bu bulgular, özellikle genç

yetiřkinlerin yapay zekânın etik, güvenlik ve toplumsal sonuçlarına ilişkin deęerlendirmelerinin dięer yař gruplarına gre daha sınırlı olduęunu gstermektedir. (Gazulla et al., 2024; Quince & Nikolic, 2025; Ghotbi & Ho, 2021).

Bu sonular, yapay zekâ okuryazarlıęı ve dijital etik alanındaki uluslararası alıřmalarla byk lde rtřmektedir. Long ve Magerko (2020), yapay zekâ okuryazarlıęının yalnızca teknolojiyi kullanabilme becerisini deęil, aynı zamanda algoritmaların alıřma mantıęını, etik sorunlarını, risklerini ve toplumsal etkilerini anlayabilmeyi kapsadıęını belirtmektedir. Bu bilgi ve farkındalık dzeyinin ise zaman ierisinde eęitim, mesleki deneyim ve teknolojiyle etkileřim sonucu geliřtięi ifade edilmektedir. Dolayısıyla yařın ilerlemesiyle birlikte bireylerin yapay zekâ sistemlerini daha eleřtirel deęerlendirmeleri beklenen bir durumdur.

Benzer řekilde Ng, Leung, Chu ve Qiao (2021) tarafından gerekleřtirilen sistematik derleme alıřmasında, yapay zekâ okuryazarlıęının geliřiminde yařam boyu ęrenme, dijital deneyim ve disiplinler arası bilgi birikiminin nemli rol oynadıęı vurgulanmıřtır. Arařtırmacılar, gen bireylerin yapay zekâyı gnlk yařamlarında yoęun olarak kullanmalarına raęmen, etik ilkeler, veri gvenlięi, algoritmik nyargılar ve hesap verebilirlik konularında yeterli bilgiye her zaman sahip olmadıklarını belirtmektedir. Mevcut arařtırmada 18–24 yař grubunun daha dřk puan alması bu deęerlendirmeyi desteklemektedir.

Yař arttıķça güvenlik farkındalıęının ykselmesi, bireylerin teknoloji kullanım srecinde karřılařtıķları deneyimlerle de aıklanabilir. (Haunschild et al., 2023; Torrijos-Fincias et al., 2021). European Commission (2019) ve UNESCO (2021) tarafından yayımlanan gvenilir yapay zekâ ilkelerinde veri gizlilięi, siber gvenlik, řeffaflık ve insan denetimi temel bileřenler arasında gsterilmektedir. İř yařamında veya profesyonel ortamlarda daha fazla deneyime sahip olan bireylerin bu ilkelerle daha fazla karřılařmaları, güvenlik farkındalıklarının artmasına katkı saęlamaktadır. zellikle kamu kurumları ve zel sektrde veri koruma, bilgi gvenlięi ve etik kullanıma ilişkin dzenlemelerin yaygınlařması, ileri yař gruplarının bu konularda daha bilinli olmalarına neden olabilmektedir.

Arařtırmada yař arttıķça Sorumluluk ve Hesap Verebilirlik puanlarının ykselmesi de dikkat ekmektedir. Bu bulgu, yapay zekâ uygulamalarının hukuki ve etik sonularına ilişkin farkındalıęın mesleki deneyimle birlikte geliřtięini dřndrmektedir. Floridi ve Cowls (2019), yapay zekâ etięinin temel ilkelerinden birinin hesap verebilirlik olduęunu ve bireylerin bu ilkeyi anlayabilmeleri iin yalnızca teknoloji bilgisine deęil, aynı zamanda etik ve hukuki bakıř

açısına da ihtiyaç duyduklarını belirtmektedir. Benzer şekilde Jobin, Ienca ve Vayena (2019) tarafından incelenen 84 uluslararası etik rehberinin büyük bölümünde hesap verebilirlik ve insan sorumluluğu en sık vurgulanan ilkeler arasında yer almaktadır. Dolayısıyla yaş ilerledikçe bireylerin bu konulara ilişkin farkındalıklarının artması beklenen bir sonuçtur.

Araştırmada elde edilen önemli bulgulardan biri de Olumlu Beklentiler alt boyutunda gözlenmiştir. Özellikle 35 yaş ve üzerindeki katılımcıların yapay zekânın gelecekte sağlayacağı katkılara ilişkin daha olumlu değerlendirmelerde bulundukları görülmektedir. İlk bakışta genç bireylerin teknolojiye daha yakın olmaları nedeniyle daha olumlu beklentilere sahip olmaları beklenebilir. Ancak mevcut bulgular bunun tersini göstermektedir. Bunun nedeni, genç bireylerin yapay zekâyı çoğunlukla günlük uygulamalar (sohbet robotları, sosyal medya algoritmaları veya üretken yapay zekâ araçları) üzerinden deneyimlemeleri, buna karşılık daha ileri yaş gruplarının yapay zekânın sağlık, üretim, mühendislik, kamu hizmetleri ve karar destek sistemleri gibi profesyonel kullanım alanlarına daha fazla aşina olmaları olabilir. Bu durum yapay zekânın toplumsal yararlarına ilişkin algıyı güçlendirebilmektedir.

Araştırmanın bulguları, teknoloji kabulü literatüründeki bazı çalışmalarla da paralellik göstermektedir. Venkatesh, Morris, Davis ve Davis (2003) tarafından geliştirilen Birleştirilmiş Teknoloji Kabul ve Kullanım Kuramı (UTAUT), yaştan teknoloji kullanım davranışını etkileyen önemli düzenleyici değişkenlerden biri olduğunu ileri sürmektedir. Bununla birlikte model, yaştan etkisinin doğrudan değil; deneyim, kullanım amacı ve performans beklentisi gibi değişkenler aracılığıyla ortaya çıktığını belirtmektedir. Mevcut araştırmada yaş ilerledikçe güvenlik ve etik algılarının yükselmesi, bu kuramsal açıklamayla uyumludur.

Öte yandan bazı araştırmalar genç bireylerin yapay zekâ teknolojilerine daha olumlu yaklaştıklarını göstermektedir (Schepman & Rodway, 2020). Ancak söz konusu çalışmaların büyük çoğunluğu yapay zekâyı yönelik genel tutum veya kullanım isteğini ölçmektedir. Mevcut araştırma ise güvenlik, etik, hesap verebilirlik ve toplumsal etkiler gibi daha karmaşık bilişsel değerlendirmeleri incelemektedir. Bu nedenle araştırmalar arasındaki farklılıkların kullanılan ölçme araçlarından kaynaklandığı söylenebilir. Günlük yaşamda yapay zekâyı sık kullanan genç bireyler teknolojiyi daha fazla benimseyebilir; ancak bu durum etik ve güvenlik konularında daha yüksek farkındalığa sahip oldukları anlamına gelmemektedir.

Araştırmada özellikle 18–24 yaş grubunun bütün alt boyutlarda daha düşük ortalamalara sahip olması, yükseköğretim kurumlarında yapay zekâ etiği ve güvenliği konularının daha sistematik

biçimde ele alınması gerektiğini göstermektedir. Günümüzde öğrenciler üretken yapay zekâ araçlarını yoğun biçimde kullanmalarına rağmen, veri gizliliği, algoritmik önyargılar, telif hakları, şeffaflık, hesap verebilirlik ve etik kullanım ilkeleri konusunda yeterli eğitimi her zaman alamamaktadırlar (UNESCO, 2021; OECD, 2019). Bu nedenle üniversitelerde yalnızca yapay zekâ kullanım becerisini geliştirmeye yönelik değil, aynı zamanda güvenli ve etik kullanım kültürünü kazandırmaya yönelik ders ve eğitim programlarının yaygınlaştırılması önem taşımaktadır.

Bu araştırmada yaş değişkeninin yapay zekâ güvenlik ve ahlak algısını etkileyen önemli bir demografik değişken olduğu belirlenmiştir. Özellikle 18–24 yaş grubunun tüm alt boyutlarda daha düşük puanlara sahip olması, yapay zekâ teknolojilerinin etik ve güvenlik boyutlarına ilişkin farkındalığın yaş, deneyim ve eğitimle birlikte geliştiğini göstermektedir. Bu bulgu, yapay zekâ okuryazarlığına yönelik eğitim programlarının özellikle genç yetişkinleri hedeflemesi gerektiğini ortaya koymaktadır. (Chee et al., 2024; Laupichler et al., 2022). Bununla birlikte gelecekte gerçekleştirilecek boylamsal çalışmalarla, bireylerin yaş ilerledikçe yapay zekâyâ ilişkin güvenlik ve etik algılarında meydana gelen değişimlerin daha ayrıntılı biçimde incelenmesi, alan yazına önemli katkılar sağlayacaktır.

5.2.4 Eğitim düzeyi bulgularının tartışılması

Bu araştırmada katılımcıların eğitim düzeylerine göre Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanlarının anlamlı düzeyde farklılaştığı belirlenmiştir ($F=5.375$, $p<.05$). Tukey çoklu karşılaştırma testi sonuçlarına göre farklılık, lisans ile ön lisans ve ön lisans ile lisansüstü grupları arasında ortaya çıkmıştır. Ortalama puanlar incelendiğinde, ön lisans öğrencilerinin ($\bar{X}=130,86$) en düşük puana sahip olduğu, lisans ($\bar{X}=144,73$) ve özellikle lisansüstü ($\bar{X}=149,37$) katılımcıların ise daha yüksek puanlar aldığı görülmektedir. Bu bulgu, eğitim düzeyi yükseldikçe bireylerin yapay zekâ güvenliği ve etik konularına ilişkin farkındalıklarının arttığını göstermektedir.

Alt boyutlar incelendiğinde de benzer bir eğilim gözlenmiştir. Güvenlik Farkındalığı alt boyutunda ön lisans ile lisansüstü grupları arasında, Sorumluluk ve Hesap Verebilirlik alt boyutunda yine ön lisans ile lisansüstü grupları arasında, Toplumsal Etki alt boyutunda lisans ile ön lisans grupları arasında ve Olumlu Beklentiler alt boyutunda hem lisans ile ön lisans hem de ön lisans ile lisansüstü grupları arasında anlamlı farklılıklar belirlenmiştir. Buna karşın Ahlaki Değerler alt boyutunda eğitim düzeyine göre anlamlı bir farklılık bulunmamıştır.

($p > .05$). Bu sonuç, yapay zekânın etik boyutuna ilişkin temel değerlerin eğitim düzeyinden bağımsız olarak benzer şekilde benimsendiğini, buna karşılık güvenlik, sorumluluk ve yapay zekânın toplumsal etkilerine ilişkin değerlendirmelerin eğitim süreciyle birlikte geliştiğini düşündürmektedir. (Biagini, 2025; Kong et al., 2024; Walter, 2024).

Araştırmanın bu bulgusu, yapay zekâ okuryazarlığı alanındaki uluslararası çalışmalarla büyük ölçüde uyumludur. Long ve Magerko (2020), yapay zekâ okuryazarlığını yalnızca teknolojiyi kullanabilme becerisi olarak değil; algoritmaların çalışma mantığını, etik ilkeleri, toplumsal etkileri ve olası riskleri anlayabilme yeterliği olarak tanımlamaktadır. Bu yeterliklerin gelişmesinde formal eğitimin önemli bir rol oynadığı belirtilmektedir. Özellikle lisans ve lisansüstü eğitim sürecinde bireylerin araştırma yapma, eleştirel düşünme, bilimsel bilgiye ulaşma ve teknoloji politikalarını değerlendirme becerilerinin gelişmesi, yapay zekâyâ ilişkin daha bilinçli değerlendirmeler yapmalarını sağlayabilmektedir.

Benzer şekilde Ng, Leung, Chu ve Qiao (2021) tarafından gerçekleştirilen sistematik derleme çalışmasında, yapay zekâ okuryazarlığının gelişiminde eğitim düzeyinin belirleyici değişkenlerden biri olduğu vurgulanmıştır. Araştırmacılar, özellikle yükseköğretim sürecinde yapay zekâ teknolojilerinin teknik, etik ve toplumsal yönlerinin birlikte ele alınmasının bireylerin farkındalık düzeyini artırdığını ifade etmektedir. Mevcut araştırmada lisans ve lisansüstü katılımcıların ön lisans grubuna göre daha yüksek puan alması bu görüşü desteklemektedir.

Özellikle Güvenlik Farkındalığı alt boyutunda lisansüstü grubunun en yüksek puana sahip olması, ileri düzey akademik eğitim alan bireylerin veri güvenliği, kişisel verilerin korunması, algoritmik riskler ve siber güvenlik gibi konulara daha fazla önem verdiklerini göstermektedir. Bu durum, lisansüstü eğitim sürecinde bilimsel araştırma etiği, araştırma yöntemleri ve veri yönetimi konularında edinilen bilgi ve deneyimlerle açıklanabilir. European Commission (2019) tarafından yayımlanan Güvenilir Yapay Zekâ Etik İlkeleri'nde veri yönetişimi, teknik güvenlik ve insan denetimi temel ilkeler arasında yer almaktadır. Bu ilkeler hakkında daha fazla bilgi sahibi olan bireylerin güvenlik farkındalığı puanlarının daha yüksek olması beklenen bir sonuçtur.

Araştırmada Sorumluluk ve Hesap Verebilirlik alt boyutunda da benzer şekilde lisansüstü katılımcılar ön lisans grubundan anlamlı düzeyde daha yüksek puan almıştır. Bu bulgu, eğitim düzeyi arttıkça bireylerin yapay zekâ uygulamalarının hukuki ve etik sorumluluklarına ilişkin

farkındalıklarının geliştiğini göstermektedir. Floridi ve Cows (2019) ile Jobin, Ienca ve Vayena (2019), hesap verebilirlik ilkesini güvenilir yapay zekânın temel bileşenlerinden biri olarak değerlendirmektedir. Özellikle bilimsel araştırma deneyimi bulunan bireylerin etik sorumluluk, şeffaflık ve denetlenebilirlik gibi kavramlara daha aşina olmaları, bu alt boyutta daha yüksek puan almalarını açıklayabilir.

Araştırmada Toplumsal Etki alt boyutunda lisans grubunun ön lisans grubundan daha yüksek puan aldığı belirlenmiştir. Yapay zekânın istihdam, demokrasi, sosyal eşitsizlik, insan hakları ve toplumsal ilişkiler üzerindeki etkilerinin değerlendirilmesi, yalnızca teknoloji bilgisi değil aynı zamanda sosyal, ekonomik ve etik bakış açısını da gerektirmektedir. Lisans programlarında disiplinler arası derslerin daha yaygın olması, öğrencilerin bu konuları daha geniş bir perspektiften değerlendirebilmelerine katkı sağlamaktadır. UNESCO (2021), yapay zekâ eğitiminin yalnızca teknik yeterlik kazandırmaması gerektiğini; insan hakları, toplumsal adalet ve etik sorumluluk gibi konuları da kapsamaması gerektiğini vurgulamaktadır. Mevcut araştırma bulguları da bu yaklaşımı desteklemektedir.

Araştırmada dikkat çeken bir diğer bulgu ise Olumlu Beklentiler alt boyutunda elde edilmiştir. Ön lisans grubunun hem lisans hem de lisansüstü gruplarına göre daha düşük puan alması, eğitim düzeyi yükseldikçe bireylerin yapay zekânın sağlayabileceği fırsatlara ilişkin daha olumlu beklentilere sahip olduklarını göstermektedir. Bu durum, ileri eğitim düzeyindeki bireylerin yapay zekânın sağlık, eğitim, bilimsel araştırma, üretim, sürdürülebilir kalkınma ve kamu hizmetleri gibi alanlardaki uygulamalarına daha fazla aşina olmalarından kaynaklanabilir. OECD (2019), yapay zekânın ekonomik büyüme ve toplumsal refah üzerindeki katkılarının ancak etik ilkeler doğrultusunda geliştirildiğinde sürdürülebilir olacağını belirtmektedir. Daha yüksek eğitim düzeyine sahip bireylerin bu potansiyeli daha iyi değerlendirebildikleri söylenebilir.

Buna karşılık Ahlaki Değerler alt boyutunda eğitim düzeyine göre anlamlı bir farklılık bulunmaması dikkat çekmektedir. Bu bulgu, yapay zekâ uygulamalarında adalet, tarafsızlık, insan haklarına saygı ve etik ilkeler gibi temel değerlerin bireyler tarafından eğitim düzeyinden bağımsız olarak benimsendiğini göstermektedir. Başka bir ifadeyle, etik değerlere ilişkin genel farkındalık toplumun farklı eğitim düzeylerinde benzer şekilde oluşmuş görünmektedir. Bu durum, son yıllarda yapay zekâ etiği konusunun medya, kamuoyu ve eğitim ortamlarında sıkça tartışılmasıyla da açıklanabilir.

Elde edilen bulgular, teknoloji kabul modelleri açısından da değerlendirilebilir. Venkatesh ve arkadaşları (2003) tarafından geliştirilen Birleştirilmiş Teknoloji Kabul ve Kullanım Kuramı (UTAUT), bireylerin yeni teknolojilere yönelik algılarında bilgi ve deneyimin önemli belirleyiciler olduğunu ifade etmektedir. Eğitim düzeyi yükseldikçe bireylerin bilgiye erişim, eleştirel değerlendirme ve teknoloji kullanım deneyimlerinin artması, yapay zekâya ilişkin daha bilinçli ve dengeli değerlendirmeler yapmalarını sağlamaktadır. (Usher & Barak, 2024; Abuadas et al., 2025; Choi et al., 2024). Mevcut araştırma sonuçları da bu kuramsal çerçeveyi desteklemektedir.

Araştırmanın bulguları uygulama açısından da önemli sonuçlar ortaya koymaktadır. Özellikle ön lisans öğrencilerinin hemen hemen tüm alt boyutlarda daha düşük puan alması, bu gruba yönelik yapay zekâ güvenliği ve etik eğitimlerinin güçlendirilmesi gerektiğini göstermektedir. Ön lisans programlarında yapay zekâ teknolojilerinin yalnızca kullanımına değil; veri güvenliği, algoritmik önyargılar, etik sorumluluk, şeffaflık ve toplumsal etkiler gibi konulara da yer verilmesi, öğrencilerin yapay zekâ okuryazarlık düzeylerinin artırılmasına katkı sağlayacaktır.

Araştırmada eğitim düzeyinin Yapay Zekâ Güvenlik ve Ahlak Algısı üzerinde etkili bir değişken olduğu belirlenmiştir. Özellikle ön lisans grubunun toplam puan ve alt boyut puanlarının lisans ve lisansüstü gruplara göre daha düşük olması, yapay zekâ güvenliği ve etik farkındalığının eğitim süreciyle birlikte geliştiğini göstermektedir. (Usher & Barak, 2024; Buna karşın ahlaki değerler alt boyutunda anlamlı farklılık bulunmaması, etik ilkelerin toplumun farklı eğitim düzeylerinde ortak kabul gören evrensel değerler hâline geldiğini düşündürmektedir. Bu bulgular doğrultusunda, yapay zekâ güvenliği ve etik eğitimlerinin özellikle ön lisans düzeyinde yaygınlaştırılması ve yükseköğretim programlarında disiplinler arası bir yaklaşımla ele alınması önerilmektedir.

5.2.5 Akademik alan bulgularının tartışılması

Bu araştırmada katılımcıların öğrenim gördükleri akademik alanın (sayısal–sözel) Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanları üzerinde anlamlı bir farklılık oluşturmadığı belirlenmiştir ($t=-0.195$, $p>.05$). Sayısal alanda öğrenim gören katılımcıların ortalama puanı 219.19, sözel alanda öğrenim gören katılımcıların ortalama puanı ise 220.07 olarak hesaplanmıştır. Ortalama puanlar arasındaki farkın oldukça düşük olması, yapay zekâ

güvenliği ve etik konularına ilişkin farkındalığın akademik alanlardan bağımsız olarak benzer düzeyde geliştiğini göstermektedir.

Alt boyut analizleri de bu sonucu desteklemektedir. Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verebilirlik, Toplumsal Etki ve Olumlu Beklentiler alt boyutlarının tamamında sayısal ve sözel alanlarda öğrenim gören katılımcılar arasında istatistiksel olarak anlamlı bir farklılık bulunmamıştır ($p>.05$). Bu bulgu, yapay zekâ güvenliği ve etik algısının yalnızca teknik bilgiye bağlı olmadığını; etik değerler, toplumsal farkındalık ve dijital okuryazarlık gibi disiplinler üstü özelliklerden etkilendiğini göstermektedir. (Biagini, 2025; Dianova, & Schultz, 2023).

Geçmiş yıllarda yapay zekâ ve bilgi teknolojilerine ilişkin çalışmaların önemli bir kısmı, teknik alanlarda öğrenim gören bireylerin yapay zekâya ilişkin bilgi düzeylerinin daha yüksek olduğunu ortaya koymuştur. Ancak son yıllarda yapay zekânın eğitim, sağlık, hukuk, iletişim, sosyal bilimler ve sanat gibi çok farklı disiplinlerde kullanılmaya başlanmasıyla birlikte bu farklılığın giderek azaldığı görülmektedir. Long ve Magerko (2020), yapay zekâ okuryazarlığının yalnızca bilgisayar bilimi öğrencilerine yönelik bir yeterlik olmadığını, bütün bireylerin sahip olması gereken temel bir dijital okuryazarlık alanı hâline geldiğini belirtmektedir. Bu nedenle farklı akademik alanlarda öğrenim gören bireylerin benzer farkındalık düzeylerine ulaşması beklenen bir durumdur.

Benzer şekilde Ng, Leung, Chu ve Qiao (2021) tarafından gerçekleştirilen sistematik derleme çalışmasında, yapay zekâ okuryazarlığının disiplinler arası bir kavram olduğu vurgulanmaktadır. Araştırmacılar, yapay zekâ eğitiminin yalnızca mühendislik ve bilgisayar bilimleri programlarında değil, sosyal bilimler, eğitim bilimleri ve sağlık bilimleri gibi farklı alanlarda da yaygınlaştırılması gerektiğini ifade etmektedir. Bu yaklaşım, mevcut araştırmada akademik alanlar arasında anlamlı farklılık bulunmamasını açıklayan önemli etkenlerden biri olarak değerlendirilebilir.

Araştırmada sayısal ve sözel alanlar arasında Güvenlik Farkındalığı açısından anlamlı bir fark bulunmaması dikkat çekicidir. Geleneksel olarak teknik alanlarda öğrenim gören bireylerin siber güvenlik ve veri güvenliği konularında daha yüksek farkındalığa sahip olmaları beklenebilir. Ancak günümüzde kişisel verilerin korunması, dijital güvenlik, siber tehditler ve yapay zekâ destekli uygulamalar toplumun tüm kesimlerini ilgilendiren konular hâline gelmiştir. European Commission (2019) tarafından yayımlanan Güvenilir Yapay Zekâ Etik

İlkeleri'nde güvenlik ve veri yönetişiminin tüm kullanıcılar açısından temel bir gereklilik olduğu belirtilmektedir. Bu doğrultuda güvenlik farkındalığının yalnızca teknik alanlara özgü bir özellik olmaktan çıktığı söylenebilir.

Benzer şekilde Ahlaki Değerler ve Sorumluluk ve Hesap Verebilirlik alt boyutlarında da akademik alanlar arasında anlamlı farklılık bulunmaması, etik ilkelerin evrensel niteliğini ortaya koymaktadır. Floridi ve Cows (2019) ile Jobin, Ienca ve Vayena (2019), güvenilir yapay zekânın temel ilkeleri arasında adalet, şeffaflık, hesap verebilirlik ve insan haklarına saygıyı ön plana çıkarmaktadır. Bu ilkeler yalnızca teknik uzmanların değil, yapay zekâyı kullanan tüm bireylerin benimsemesi gereken ortak etik değerlerdir. Araştırmada hem sayısal hem de sözel alan öğrencilerinin bu konularda benzer puanlar almaları, etik farkındalığın disiplinler arasında ortak bir anlayış hâline gelmeye başladığını göstermektedir.

Toplumsal Etki alt boyutunda da anlamlı farklılık bulunmaması, katılımcıların yapay zekânın toplum üzerindeki etkilerini benzer biçimde değerlendirdiklerini göstermektedir. Günümüzde üretken yapay zekâ uygulamalarının eğitimden medyaya, sağlıktan kamu hizmetlerine kadar çok geniş alanlarda kullanılması, yapay zekânın toplumsal etkilerine ilişkin tartışmaların bütün disiplinlerde yer almasına neden olmuştur. UNESCO (2021), yapay zekânın etik kullanımına ilişkin tavsiye kararlarında toplumun tüm bireylerinin yapay zekânın sosyal, ekonomik ve kültürel etkileri konusunda bilinçlendirilmesi gerektiğini vurgulamaktadır. Bu durum, farklı akademik alanlarda öğrenim gören katılımcıların benzer değerlendirmeler yapmalarını açıklayabilir.

Araştırmada Olumlu Beklentiler alt boyutunda da anlamlı farklılık bulunmaması, sayısal ve sözel alan öğrencilerinin yapay zekânın gelecekte sağlayacağı katkılara ilişkin benzer düzeyde iyimser olduklarını göstermektedir. Günümüzde yapay zekânın eğitim, sağlık, bilimsel araştırmalar, kamu hizmetleri ve iş dünyasında sağladığı katkılar medya ve akademik çevrelerde yoğun biçimde tartışılmaktadır. Bu nedenle yapay zekânın potansiyel faydalarına ilişkin farkındalığın belirli disiplinlerle sınırlı kalmadığı söylenebilir. OECD (2019) de yapay zekânın ekonomik ve toplumsal gelişime katkı sağlayabilmesi için toplumun bütün kesimlerinin bu teknolojiyi anlayabilmesi gerektiğini ifade etmektedir.

Araştırma bulguları, yapay zekâ teknolojilerinin giderek disiplinler arası bir yapıya dönüştüğünü göstermektedir. Geçmişte yapay zekâ yalnızca bilgisayar bilimleri ve mühendislik alanlarının çalışma konusu olarak görülürken, günümüzde eğitim bilimleri, psikoloji, hukuk,

iletiřim, iřletme, saęlık bilimleri ve sosyal bilimler gibi pek ok alanda aktif olarak kullanılmaktadır. (Dwivedi et al., 2019; Sadiku et al., 2021). Bu geliřme, farklı akademik alanlarda ğrenim gren bireylerin yapay zekâya iliřkin benzer dzeyde bilgi edinmelerine ve benzer etik deęerlendirmeler yapmalarına katkı saęlamaktadır.

Bunun yanında arařtırmanın rneklemine oluřturan katılımcıların byk blmnn niversite ęrencisi veya niversite mezunu bireylerden oluřması da bu sonucu etkilemiř olabilir. niversite ortamında farklı disiplinlerden ęrencilerin ortak semeli dersler alması, bilimsel etkinliklere katılması, dijital teknolojileri yoęun biimde kullanması ve yapay zekâ uygulamalarıyla gnlk yařamlarında sık karřılařmaları, akademik alanlar arasındaki farkın azalmasına neden olmuř olabilir.

Bu arařtırmanın bulguları, yapay zekâ gvenlięi ve etik eęitimlerinin yalnızca teknik blmlerde verilmesinin yeterli olmayacaęını gstermektedir. Akademik alanlar arasında anlamlı farklılık bulunmaması, yapay zekâ gvenlięi ve etik konularının tm yksekęretim programlarında ortak bir dijital yetkinlik olarak ele alınması gerektięini ortaya koymaktadır. zellikle veri gvenlięi, algoritmik nyargılar, etik karar verme, řeffaflık ve hesap verebilirlik gibi konuların disiplinler arası ders ieriklerine dâhil edilmesi, tm ęrencilerin gvenilir yapay zekâ anlayıřını geliřtirmelerine katkı saęlayacaktır.

Bu arařtırmada akademik alan deęiřkeninin Yapay Zekâ Gvenlik ve Ahlak Algısı zerinde anlamlı bir etkisinin bulunmadıęı belirlenmiřtir. Sayısal ve szel alanlarda ğrenim gren katılımcılar hem toplam puanlarda hem de lęin beř alt boyutunda benzer dzeyde puan almıřlardır. Bu bulgu, yapay zekâ gvenlięi ve etik farkındalıęının gnmzde belirli bir disipline zg olmaktan ıkarak btn akademik alanları kapsayan ortak bir dijital yeterlik hâline geldięini gstermektedir. Dolayısıyla yksekęretim kurumlarında yapay zekâ gvenlięi ve etik eęitimlerinin yalnızca teknoloji odaklı programlarla sınırlı kalmaması, tm disiplinlerde yaygınlařtırılması nemli grlmektedir.

5.2.6 Gelir algısı bulgularının tartıřılması

Bu arařtırmada katılımcıların gelir durum algılarına gre Yapay Zekâ Gvenlik ve Ahlak Algısı lęi (YZGAA) toplam puanları arasında istatistiksel olarak anlamlı bir farklılık bulunmamıřtır ($F=2.125$, $p>.05$). Ortalama puanlar incelendięinde, gelir durumunu kt olarak deęerlendiren katılımcıların ortalama puanının 223.82, orta olarak deęerlendirenlerin 213.38 ve iyi olarak deęerlendirenlerin ise 222.05 olduęu grlmektedir. Ortalama puanlar arasında

küçük farklılıklar bulunmasına rağmen bu farklılıkların istatistiksel açıdan anlamlı olmaması, bireylerin yapay zekâ güvenliği ve etik algılarının ekonomik durumlarından bağımsız olarak benzer düzeyde geliştiğini göstermektedir.

Benzer şekilde, ölçeğin Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verebilirlik, Toplumsal Etki ve Olumlu Beklentiler alt boyutlarının hiçbirinde gelir durum algısına göre anlamlı bir farklılık belirlenmemiştir ($p > .05$). Bu sonuç, yapay zekâ güvenliği ve etik konularına ilişkin farkındalığın bireylerin ekonomik imkânlarından çok eğitim, dijital deneyim ve bilgiye erişim gibi faktörlerle ilişkili olduğunu düşündürmektedir. (Biagini, 2025; Yang, 2024; G et al., 2025; Faisal, 2025).

Literatürde dijital teknolojilere erişim ile sosyoekonomik düzey arasında belirli ilişkiler olduğu bilinmektedir. Özellikle erken dönem çalışmalarda yüksek gelir düzeyine sahip bireylerin teknolojiye erişim, dijital cihaz kullanımı ve internet olanakları bakımından daha avantajlı olduğu belirtilmiştir. Ancak son yıllarda internet erişiminin yaygınlaşması, mobil teknolojilerin ucuzlaması ve üretken yapay zekâ uygulamalarının büyük bölümünün ücretsiz veya düşük maliyetli olarak kullanılabilmesi, gelir düzeyinin teknoloji kullanımındaki belirleyici etkisini önemli ölçüde azaltmıştır. Van Dijk (2020), dijital eşitsizliğin günümüzde yalnızca ekonomik kaynaklarla açıklanamayacağını; eğitim, dijital beceriler ve teknoloji kullanım amaçlarının daha belirleyici hâle geldiğini ifade etmektedir. Bu araştırmada gelir durumuna göre anlamlı farklılık bulunmaması da bu görüşü desteklemektedir.

Araştırmada Güvenlik Farkındalığı alt boyutunda gelir grupları arasında anlamlı farklılık bulunmaması, veri güvenliği ve kişisel verilerin korunmasına ilişkin farkındalığın toplumun tüm kesimlerinde benzer düzeyde oluştuğunu göstermektedir. Son yıllarda kişisel verilerin korunması, siber saldırılar, kimlik hırsızlığı ve yapay zekâ destekli dolandırıcılık gibi konuların medya ve kamuoyu tarafından yoğun biçimde gündeme getirilmesi, farklı gelir gruplarındaki bireylerin bu riskler konusunda benzer bilinç düzeyine ulaşmalarını sağlamış olabilir. European Commission (2019) tarafından yayımlanan Güvenilir Yapay Zekâ Etik İlkeleri'nde güvenlik ve veri yönetişimi tüm bireylerin ortak sorumluluğu olarak değerlendirilmektedir. Bu nedenle güvenlik farkındalığının gelir düzeyinden bağımsız olması beklenen bir sonuç olarak değerlendirilebilir.

Ahlaki Değerler alt boyutunda anlamlı farklılık bulunmaması ise etik ilkelerin ekonomik durumdan bağımsız evrensel değerler olduğunu göstermektedir. Yapay zekânın adalet,

tarafsızlık, insan haklarına saygı ve etik ilkeler doğrultusunda geliştirilmesi gerektiği yönündeki görüşler farklı gelir gruplarında benzer biçimde benimsenmiştir. Floridi ve Cowls (2019) ile Jobin, Ienca ve Vayena (2019), güvenilir yapay zekânın temelini oluşturan etik ilkelerin toplumun tüm kesimleri tarafından paylaşılması gerektiğini vurgulamaktadır. Bu araştırmanın bulguları da söz konusu yaklaşımı desteklemektedir.

Araştırmada Sorumluluk ve Hesap Verebilirlik alt boyutunda da gelir durumuna göre anlamlı farklılık bulunmamıştır. Yapay zekâ uygulamalarında geliştiricilerin, kullanıcıların ve kurumların sorumluluklarının açık biçimde belirlenmesi gerektiğine ilişkin görüşlerin tüm gelir gruplarında benzer düzeyde olduğu görülmektedir. Bu durum, yapay zekâyâ ilişkin etik tartışmaların yalnızca belirli sosyoekonomik grupların değil, toplumun genelinin gündeminde yer aldığını göstermektedir. (Huang et al., 2023; Hagerty & Rubinov, 2019; Schiff et al., 2021; Shukla, 2024).

Toplumsal Etki alt boyutunda da gelir grupları arasında anlamlı farklılık belirlenmemiştir. Katılımcılar yapay zekânın iş gücü piyasası, sosyal eşitsizlik, insan hakları ve toplumsal yaşam üzerindeki olası etkilerini benzer şekilde değerlendirmişlerdir. UNESCO (2021), yapay zekânın toplumsal etkilerinin tüm bireyleri ilgilendiren küresel bir konu olduğunu ve bu nedenle etik farkındalığın toplumun her kesiminde geliştirilmesi gerektiğini belirtmektedir. Araştırmanın bulguları, yapay zekânın toplumsal etkilerine ilişkin algının ekonomik durumdan bağımsız biçimde şekillendiğini göstermektedir.

Benzer şekilde Olumlu Beklentiler alt boyutunda da gelir durum algısına göre anlamlı farklılık bulunmamıştır. Katılımcılar gelir düzeylerinden bağımsız olarak yapay zekânın sağlık, eğitim, bilim, kamu hizmetleri ve ekonomik gelişme gibi alanlarda sağlayabileceği katkıları benzer düzeyde değerlendirmişlerdir. OECD (2019), yapay zekânın ekonomik ve sosyal faydalarının toplumun tüm kesimlerine eşit biçimde ulaştırılması gerektiğini belirtmektedir. Bu bağlamda elde edilen bulgu, katılımcıların yapay zekânın gelecekte sağlayabileceği fırsatlara ilişkin ortak bir bakış açısına sahip olduklarını göstermektedir.

Araştırmanın bu sonucu, yapay zekâ teknolojilerinin son yıllarda toplumun hemen her kesimine ulaşmasıyla da açıklanabilir. ChatGPT, Gemini, Copilot ve benzeri üretken yapay zekâ uygulamalarının ücretsiz sürümlerinin yaygınlaşması, yapay zekâ teknolojilerine erişimde ekonomik engelleri önemli ölçüde azaltmıştır. Dolayısıyla bireylerin yapay zekâyâ ilişkin bilgi edinmeleri ve bu teknolojileri deneyimlemeleri için yüksek gelir düzeyine sahip olmaları artık

zorunlu değildir. Bu durum, gelir durumuna bağlı farklılıkların azalmasına katkı sağlamış olabilir.

Bunun yanında araştırmada kullanılan değişkenin objektif gelir düzeyi değil, bireylerin gelir durum algısı olması da sonuçların yorumlanmasında dikkate alınmalıdır. Gelir algısı bireylerin öznel değerlendirmelerine dayandığından, aynı ekonomik koşullara sahip bireyler kendilerini farklı gelir gruplarında değerlendirebilmektedir. Bu nedenle gelir algısının yapay zekâ güvenliği ve etik algısı üzerindeki etkisi sınırlı kalmış olabilir. Gelecek araştırmalarda gelir düzeyinin objektif ekonomik göstergelerle birlikte incelenmesi, bu ilişkinin daha ayrıntılı değerlendirilmesine katkı sağlayabilir.

Araştırmanın uygulama açısından önemli bir sonucu da yapay zekâ güvenliği ve etik eğitimlerinin yalnızca belirli sosyoekonomik gruplara yönelik planlanmasının gerekli olmadığını göstermesidir. Gelir düzeyinden bağımsız olarak tüm bireylerin benzer farkındalık düzeyine sahip olması, yapay zekâ güvenliği eğitimlerinin toplumun tamamını kapsayacak biçimde tasarlanmasının daha uygun olacağını ortaya koymaktadır. Özellikle veri güvenliği, algoritmik önyargılar, etik karar verme ve yapay zekânın toplumsal etkileri gibi konuların herkese yönelik dijital okuryazarlık programlarına entegre edilmesi önem taşımaktadır. (Biagini, 2025; Wazir et al., 2025).

Bu araştırmada gelir durum algısının Yapay Zekâ Güvenlik ve Ahlak Algısı üzerinde anlamlı bir etkisinin olmadığı belirlenmiştir. Hem toplam puanlarda hem de ölçeğin beş alt boyutunda gelir grupları arasında istatistiksel olarak anlamlı farklılık bulunmamıştır. Bu bulgu, yapay zekâ güvenliği ve etik farkındalığının ekonomik imkânlardan ziyade eğitim, bilgiye erişim, dijital deneyim ve toplumsal farkındalık gibi değişkenlerden etkilendiğini göstermektedir. Günümüzde yapay zekâ teknolojilerinin yaygınlaşması ve geniş kitleler tarafından erişilebilir hâle gelmesi, gelir düzeyinin bu alandaki belirleyici rolünü azaltmış görünmektedir.

5.2.7 Yapay zekâ kullanımı bulgularının tartışılması

Bu araştırmada katılımcıların yapay zekâ uygulamalarını kullanma durumlarına göre Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) toplam puanları arasında istatistiksel olarak anlamlı bir farklılık bulunmamıştır ($t=-0.823$, $p>.05$). Yapay zekâ kullanan katılımcıların ortalama puanı 217.97, kullanmayan katılımcıların ortalama puanı ise 221.70 olarak belirlenmiştir. Her ne kadar yapay zekâ kullanmayan katılımcıların ortalama puanı daha yüksek görünse de bu fark istatistiksel olarak anlamlı değildir. Bu bulgu, yapay zekâ uygulamalarını

kullanmanın tek başına bireylerin güvenlik ve etik algılarını belirleyen bir değişken olmadığını göstermektedir.

Araştırmanın bu sonucu, yapay zekâ kullanım sıklığı ile yapay zekâ okuryazarlığı ya da etik farkındalık kavramlarının birbirinden farklı yapılar olduğunu ortaya koymaktadır. (Wang et al., 2025). Günümüzde bireyler ChatGPT, Gemini, Copilot veya benzeri üretken yapay zekâ uygulamalarını yoğun biçimde kullanabilmektedir. Ancak bu kullanım, yapay zekânın çalışma prensiplerini, etik ilkelerini, güvenlik risklerini veya toplumsal etkilerini yeterince bildikleri anlamına gelmemektedir. Long ve Magerko (2020), yapay zekâ okuryazarlığını yalnızca yapay zekâ araçlarını kullanabilme becerisi olarak değil; bu sistemlerin nasıl çalıştığını, hangi sınırlılıklara sahip olduğunu, etik sorunlarını ve toplumsal etkilerini anlayabilme yeterliği olarak tanımlamaktadır. Dolayısıyla kullanım deneyimi ile etik farkındalık arasında doğrudan bir ilişki bulunmaması beklenen bir durumdur.

Benzer şekilde Ng, Leung, Chu ve Qiao (2021), yapay zekâ okuryazarlığının bilgi, beceri ve etik farkındalığın birlikte gelişmesiyle oluştuğunu belirtmektedir. Araştırmacılar, yapay zekâ uygulamalarını kullanmanın bireylerin yapay zekâ okuryazarlığını otomatik olarak geliştirmediklerini, bu süreçte planlı eğitim ve eleştirel düşünme becerilerinin önemli rol oynadığını ifade etmektedir. Bu açıdan değerlendirildiğinde mevcut araştırmada yapay zekâ kullanan ve kullanmayan bireylerin benzer puanlara sahip olmaları literatürle uyumludur.

Araştırmada anlamlı farklılık bulunmamasının bir diğer nedeni, günümüzde yapay zekâ uygulamalarının oldukça yaygın hâle gelmiş olması olabilir. Özellikle üretken yapay zekâ sistemlerinin eğitim, iş yaşamı ve günlük yaşamda hızla yaygınlaşması sonucunda yapay zekâ kullanmayan bireyler dahi medya, sosyal ağlar ve eğitim ortamları aracılığıyla yapay zekâyâ ilişkin etik ve güvenlik tartışmalarına maruz kalmaktadır. (Huang, 2023; Mao & Shi-Kupfer, 2021; Ouchchy et al., 2020; Vilchez et al., 2026). Bu durum, yapay zekâyı aktif olarak kullanmayan bireylerin de güvenlik ve etik konularında belirli bir farkındalık geliştirmelerini sağlamış olabilir.

Diğer taraftan araştırmada yapay zekâ kullanım durumunun yalnızca "evet-hayır" biçiminde değerlendirilmiş olması da bulguların yorumlanmasında dikkate alınmalıdır. Yapay zekâyı kullanan bireyler arasında kullanım amacı, kullanım sıklığı, kullanılan uygulamalar ve kullanım düzeyi bakımından önemli farklılıklar bulunabilir. Örneğin yalnızca metin düzenleme amacıyla zaman zaman ChatGPT kullanan bir birey ile yapay zekâyı akademik araştırmalarında, yazılım

geliřtirmede veya veri analizinde yoğun olarak kullanan bir bireyin yapay zekâ deneyimi aynı deęildir. Ancak bu arařtırmada kullanımın derinlięi veya yoğunluęu ölçölmedięi için bu farklılıkların sonuçlara yansması mümkün olmamıřtır.

Literatürde yapay zekâ kullanım deneyiminin bazı çalıřmalarda etik farkındalıkla iliřkili olduęu, bazı çalıřmalarda ise anlamlı bir iliřki göstermedięi görölmektedir. Schepman ve Rodway (2020), yapay zekâ ile daha fazla etkileřim kuran bireylerin teknolojiye yönelik daha olumlu tutum geliřtirdiklerini belirtmektedir. Ancak arařtırmacılar bu olumlu tutumun etik farkındalık veya güvenlik bilinciyle aynı kavram olmadıęını özellikle vurgulamaktadır. Benzer řekilde Chiu, Ismailov ve arkadaşları (2024) tarafından gerçekteřtirilen güncel çalıřmalarda, üretken yapay zekâ araçlarının yoğun kullanımının bireylerin kullanım becerilerini geliřtirdięi; ancak etik karar verme ve güvenli kullanım davranıřlarının ayrıca eęitim gerektirdięi ifade edilmektedir.

Arařtırmanın bulguları, güvenilir yapay zekâ yaklařımını benimseyen uluslararası politika belgeleriyle de örtüřmektedir. European Commission (2019) ve UNESCO (2021), yapay zekâ teknolojilerinin güvenli ve etik biçimde kullanılabilmesi için yalnızca teknolojinin yaygınlařmasının yeterli olmadıęını; bireylerin etik ilkeler, veri gizlilięi, algoritmik önyargılar, insan denetimi ve hesap verebilirlik konularında bilinçlendirilmesi gerektięini vurgulamaktadır. Dolayısıyla yapay zekâ kullanımının artması, etik farkındalıęın da aynı hızda artacaęı anlamına gelmemektedir.

Bu arařtırmada yapay zekâ kullanan bireylerin ortalama puanlarının kullanmayan bireylerden biraz daha düşük olmasına raęmen farkın anlamlı olmaması da dikkat çekmektedir. Bu durum, yapay zekâyı aktif kullanan bireylerin teknolojinin sınırlarını ve mevcut eksikliklerini daha yakından gözlemlemeleri nedeniyle daha eleřtirel deęerlendirmeler yapmalarıyla açıklanabilir. Yapay zekâ uygulamalarını doğrudan deneyimleyen kullanıcılar, halüsinasyon (yanlıř bilgi üretme), veri gizlilięi sorunları, telif hakları, algoritmik önyargılar ve karar verme hataları gibi problemleri daha fazla fark edebilmekte ve bu nedenle deęerlendirmelerinde daha temkinli davranabilmektedirler. Ancak mevcut arařtırmada bu farklılık istatistiksel olarak anlamlı düzeye ulařmamıřtır.

Elde edilen bulgular, yapay zekâ güvenlięi ve etik eęitimlerinin yalnızca yapay zekâyı hiç kullanmamıř bireylere deęil, aktif kullanıcıları da kapsayacak biçimde planlanması gerektięini göstermektedir. Günümüzde birçok kullanıcı üretken yapay zekâ araçlarını etkin biçimde

kullanmasına rağmen veri güvenliği, kişisel verilerin korunması, telif hakları, akademik dürüstlük, algoritmik önyargılar ve şeffaflık gibi konularda yeterli bilgiye sahip olmayabilmektedir. (Deric et al., 2025; Yusuf et al., 2024). Bu nedenle yükseköğretim kurumlarında ve hizmet içi eğitim programlarında yapay zekâ kullanım becerisinin yanında güvenli ve etik kullanım kültürünün de geliştirilmesi büyük önem taşımaktadır.

Ayrıca gelecekte gerçekleştirilecek araştırmalarda yalnızca "kullanıyor-kullanmıyor" değişkeninin değil; kullanım sıklığı, kullanım amacı, kullanılan yapay zekâ uygulaması, günlük kullanım süresi, yapay zekâ eğitimi alma durumu ve yapay zekâ okuryazarlık düzeyi gibi değişkenlerin de birlikte incelenmesi önerilmektedir. Böylece yapay zekâ deneyiminin güvenlik ve etik algısı üzerindeki etkisi daha ayrıntılı biçimde ortaya konulabilecektir.

Bu araştırmada yapay zekâ uygulamalarını kullanma durumunun Yapay Zekâ Güvenlik ve Ahlak Algısı üzerinde anlamlı bir etkisinin olmadığı belirlenmiştir. Bu bulgu, yapay zekâ araçlarını kullanmanın tek başına güvenlik ve etik farkındalığını artırmadığını; bu farkındalığın planlı eğitim, dijital okuryazarlık, eleştirel düşünme becerileri ve etik bilinçle birlikte geliştiğini göstermektedir. Dolayısıyla yapay zekânın toplumda güvenli ve sorumlu biçimde kullanılabilmesi için yalnızca kullanımın yaygınlaştırılması değil, etik ve güvenlik odaklı eğitimlerin de eş zamanlı olarak geliştirilmesi gerekmektedir.

5.2.8 Genel değerlendirme

Bu araştırmada geliştirilen Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ), bireylerin yapay zekâ teknolojilerine yönelik güvenlik, etik ve toplumsal bakış açılarını çok boyutlu olarak değerlendirebilen geçerli ve güvenilir bir ölçme aracı olarak alanyazına kazandırılmıştır. Ölçek geliştirme sürecinde gerçekleştirilen kapsam geçerliği çalışmaları, Açımlayıcı Faktör Analizi (AFA), Doğrulayıcı Faktör Analizi (DFA), madde analizleri ve güvenirlik analizleri sonucunda ölçeğin güçlü psikometrik özelliklere sahip olduğu belirlenmiştir. Elde edilen bulgular, ölçeğin yalnızca kuramsal açıdan değil, aynı zamanda uygulama açısından da kullanılabilecek nitelikte olduğunu göstermektedir.

Araştırmada elde edilen demografik bulgular birlikte değerlendirildiğinde, yaş ve eğitim düzeyi değişkenlerinin yapay zekâ güvenlik ve ahlak algısını etkileyen önemli değişkenler olduğu, buna karşın cinsiyet, akademik alan, gelir durum algısı ve yapay zekâ uygulamalarını kullanma durumu değişkenlerinin anlamlı farklılık oluşturmadığı belirlenmiştir. Bu sonuç, yapay zekâ güvenliği ve etik farkındalığının daha çok bireylerin bilgi birikimi, eğitim deneyimi ve yaşam

deneyimleriyle ilişkili olduğunu; buna karşılık biyolojik ya da sosyoekonomik özelliklerden daha az etkilendiğini göstermektedir.

Araştırmanın en dikkat çekici bulgularından biri, 18–24 yaş grubunun ölçek toplam puanında ve alt boyutların büyük bölümünde diğer yaş gruplarından anlamlı düzeyde daha düşük puan almasıdır. Bu durum, genç bireylerin yapay zekâ teknolojilerini yoğun biçimde kullanmalarına rağmen güvenlik, etik, hesap verebilirlik ve toplumsal etkiler konusunda aynı düzeyde farkındalığa sahip olmadıklarını göstermektedir. Günümüzde üretken yapay zekâ uygulamalarının özellikle genç kullanıcılar arasında hızla yaygınlaşması dikkate alındığında, bu bulgu yükseköğretim kurumlarında yapay zekâ etiği ve güvenliği eğitimlerinin güçlendirilmesi gerektiğini ortaya koymaktadır. Teknolojiyi kullanabilmek ile teknolojinin etik sonuçlarını değerlendirebilmek birbirinden farklı yeterliklerdir ve bu iki becerinin birlikte geliştirilmesi gerekmektedir (Long & Magerko, 2020; Ng et al., 2021).

Benzer şekilde eğitim düzeyine ilişkin bulgular da yapay zekâ güvenliği ve etik farkındalığının eğitim süreciyle birlikte geliştiğini göstermektedir. Özellikle ön lisans öğrencilerinin lisans ve lisansüstü katılımcılara göre daha düşük puan almaları, yapay zekâ güvenliği ve etik konularının yükseköğretimin tüm kademelerinde daha sistematik biçimde ele alınmasının önemini ortaya koymaktadır. Bu sonuç, UNESCO (2021) ve European Commission (2019) tarafından önerilen "insan merkezli ve güvenilir yapay zekâ" yaklaşımıyla da uyumludur.

Araştırmada cinsiyet değişkenine göre anlamlı farklılık bulunmaması, yapay zekâ güvenliği ve etik algısının kadın ve erkek katılımcılar arasında benzer düzeyde geliştiğini göstermektedir. Son yıllarda dijital teknolojilere erişim olanaklarının artması, yapay zekâ uygulamalarının günlük yaşamın doğal bir parçası hâline gelmesi ve etik tartışmaların toplumun tüm kesimlerine yayılması, bu sonucun ortaya çıkmasında etkili olmuş olabilir. Aynı şekilde akademik alanlar arasında anlamlı farklılık bulunmaması da yapay zekânın artık yalnızca teknik disiplinlerin konusu olmaktan çıkarak tüm bilim alanlarını ilgilendiren disiplinler arası bir teknoloji hâline geldiğini göstermektedir.

Gelir durum algısına göre anlamlı farklılık bulunmaması ise yapay zekâ güvenliği ve etik farkındalığının ekonomik koşullardan çok bilgiye erişim, eğitim ve dijital deneyimle ilişkili olduğunu düşündürmektedir. Özellikle ücretsiz üretken yapay zekâ araçlarının yaygınlaşması ve internet erişiminin toplumun geniş kesimlerine ulaşması, gelir düzeyinin yapay zekâ farkındalığı üzerindeki etkisini azaltmış görünmektedir. Benzer biçimde yapay zekâ

uygulamalarını kullanma durumuna göre de anlamlı farklılık bulunmaması, yapay zekâ kullanımının tek başına etik ve güvenlik farkındalığını geliştirmediğini göstermektedir. Bu bulgu, yapay zekâ okuryazarlığının yalnızca kullanım becerilerinden değil; etik bilgi, eleştirel düşünme, dijital güvenlik farkındalığı ve toplumsal sorumluluk gibi çok boyutlu yeterliklerden oluştuğunu ortaya koymaktadır.

Araştırmanın bulguları, uluslararası yapay zekâ etik ilkeleriyle de büyük ölçüde örtüşmektedir. European Commission (2019), OECD (2019), UNESCO (2021), Floridi ve Cowls (2019) ile Jobin, Ienca ve Vayena (2019) tarafından ortaya konulan temel ilkeler; insan denetimi, güvenlik, şeffaflık, hesap verebilirlik, adalet ve toplumsal yarar kavramları etrafında şekillenmektedir. Bu çalışmada geliştirilen ölçeğin alt boyutları da söz konusu uluslararası etik çerçevelerle önemli ölçüde uyum göstermektedir. Bu yönüyle YZGAAÖ, yalnızca Türkiye'de değil, farklı kültürel bağlamlarda gerçekleştirilecek çalışmalara da uyarlanabilecek güçlü bir kuramsal altyapıya sahiptir.

Araştırmanın önemli katkılarından biri de yapay zekâ güvenliği ve etik konularını tek boyutlu değil, Güvenlik Farkındalığı, Ahlaki Değerler, Sorumluluk ve Hesap Verebilirlik, Toplumsal Etki ve Olumlu Beklentiler olmak üzere beş farklı boyutta ele almasıdır. Bu yapı, bireylerin yapay zekâyâ ilişkin algılarının yalnızca risk veya yalnızca fayda odaklı değerlendirilmesinin yetersiz olduğunu göstermektedir. Yapay zekâ teknolojileri hem önemli fırsatlar hem de önemli etik ve güvenlik riskleri barındırmaktadır. Dolayısıyla bireylerin bu teknolojilere ilişkin algılarının çok boyutlu olarak değerlendirilmesi daha kapsamlı sonuçlar elde edilmesini sağlamaktadır.

Araştırma sonuçları uygulama açısından da önemli çıkarımlar sunmaktadır. Öncelikle yükseköğretim kurumlarında yapay zekâyâ ilişkin eğitimlerin yalnızca teknik kullanım becerilerine odaklanmaması; etik ilkeler, veri güvenliği, algoritmik önyargılar, kişisel verilerin korunması, şeffaflık, hesap verebilirlik ve insan hakları gibi konuları da içermesi gerekmektedir. Özellikle öğretmen yetiştiren programlarda, eğitim fakültelerinde ve hizmet içi eğitimlerde yapay zekâ okuryazarlığı ile güvenilir yapay zekâ ilkelerinin birlikte ele alınması, gelecekte bu teknolojilerin daha bilinçli kullanılmasına katkı sağlayacaktır.

Bunun yanında geliştirilen ölçek, gelecekte yapılacak araştırmalar için önemli bir ölçme aracı niteliği taşımaktadır. Ölçek kullanılarak farklı yaş grupları, meslek grupları, kamu çalışanları, öğretmenler, akademisyenler, lise öğrencileri veya özel sektör çalışanları üzerinde

karşılaştırmalı çalışmalar gerçekleştirilebilir. Ayrıca farklı ülkelerde yapılacak kültürlerarası uyarlama çalışmaları ile ölçeğin uluslararası geçerliği test edilebilir. Bunun yanı sıra yapay zekâ okuryazarlığı, dijital vatandaşlık, siber güvenlik farkındalığı, teknoloji kabulü, dijital etik, algoritmik güven ve yapay zekâ kaygısı gibi değişkenlerle ilişkisini inceleyen yapısal modeller oluşturulması da alanyazına önemli katkılar sağlayacaktır.

Sonuç olarak bu araştırma, yapay zekâ güvenliği ve etik alanında Türkiye'de geliştirilen kapsamlı ölçeklerden biri olarak önemli bir boşluğu doldurmaktadır. Elde edilen bulgular, yapay zekâ teknolojilerinin güvenli, etik ve insan merkezli biçimde kullanılabilmesi için teknik bilgi kadar etik bilinç, dijital okuryazarlık ve toplumsal sorumluluk anlayışının da geliştirilmesi gerektiğini göstermektedir. Yapay zekânın gelecekte eğitimden sağlığa, hukuktan kamu yönetimine kadar yaşamın her alanında daha fazla yer alacağı düşünüldüğünde, bireylerin bu teknolojilere ilişkin güvenlik ve etik farkındalıklarını ölçebilecek bilimsel araçların geliştirilmesi büyük önem taşımaktadır. Bu bağlamda YZGAAÖ'nün, hem akademik araştırmalarda hem de eğitim ve politika geliştirme süreçlerinde kullanılabilecek güçlü, geçerli ve güvenilir bir ölçme aracı olduğu değerlendirilmektedir.

5.3 Araştırmanın Sınırlılıkları

Bu araştırmanın ilk sınırlılığı, çalışma grubunun belirli bir örneklem ile sınırlandırılmış olmasıdır. Araştırma, toplam 558 katılımcıdan elde edilen veriler doğrultusunda yürütülmüş olup, örneklem ağırlıklı olarak 25–34 yaş grubundaki bireylerden oluşmaktadır. Bu nedenle elde edilen sonuçların farklı yaş gruplarına veya tüm topluma doğrudan genellenmesi uygun olmayabilir.

İkinci olarak, araştırmanın katılımcıları ön lisans, lisans ve lisansüstü düzeyindeki bireylerden oluşmaktadır. Bu nedenle geliştirilen ölçeğin ortaöğretim öğrencileri, öğretmenler, akademisyenler, kamu çalışanları, özel sektör çalışanları veya farklı meslek gruplarında uygulanması durumunda farklı psikometrik özellikler gösterebilmesi mümkündür.

Araştırmanın bir diğer sınırlılığı, verilerin öz-bildirim (self-report) yöntemiyle toplanmış olmasıdır. Katılımcılar ölçek maddelerini kendi algı ve değerlendirmelerine göre yanıtlamışlardır. Bu durum sosyal beğenirlik eğilimi, bireysel yorum farklılıkları ve cevaplayıcı yanlılığı gibi etkenlerin sonuçlar üzerinde etkili olmasına neden olabilir.

Araştırma kesitsel (cross-sectional) bir araştırma desenine sahiptir. Bu nedenle katılımcıların yapay zekâ güvenlik ve ahlak algılarının zaman içerisindeki değişimi değerlendirilememiştir. Yapay zekâ teknolojilerinin çok hızlı geliştiği dikkate alındığında, bireylerin algılarının da zamanla değişebileceği göz önünde bulundurulmalıdır.

Araştırmada yapay zekâ kullanım durumu yalnızca "kullanıyorum" ve "kullanmıyorum" biçiminde değerlendirilmiştir. Katılımcıların yapay zekâyı hangi amaçla kullandıkları, kullanım sıklıkları, kullandıkları yapay zekâ uygulamaları, kullanım deneyim süreleri veya yapay zekâ konusunda eğitim alıp almadıkları gibi değişkenler incelenmemiştir. Bu nedenle yapay zekâ deneyiminin güvenlik ve ahlak algısı üzerindeki ayrıntılı etkisi değerlendirilememiştir.

Araştırmada incelenen demografik değişkenler cinsiyet, yaş, eğitim düzeyi, akademik alan, gelir durum algısı ve yapay zekâ kullanım durumu ile sınırlıdır. Dijital okuryazarlık düzeyi, teknoloji kabulü, yapay zekâ okuryazarlığı, bilgisayar deneyimi, mesleki deneyim, internet kullanım süresi ve dijital etik eğitimi gibi değişkenler araştırma kapsamına dâhil edilmemiştir. Bu değişkenlerin yapay zekâ güvenlik ve ahlak algısı üzerinde etkili olabileceği düşünülmektedir.

Ölçeğin yapı geçerliği Açıklayıcı Faktör Analizi (AFA) ve Doğrulayıcı Faktör Analizi (DFA) ile test edilmiş olmakla birlikte, analizler aynı veri seti üzerinde gerçekleştirilmiştir. Psikometrik açıdan daha güçlü kanıtlar elde edebilmek için farklı örneklemeler üzerinde yeniden doğrulama çalışmalarının yapılması önerilmektedir.

Araştırmada ölçüt bağıntılı geçerlik (criterion-related validity), test-tekrar test güvenilirliği ve farklı kültürlerde ölçme değişmezliği (measurement invariance) analizleri gerçekleştirilememiştir. Bu nedenle ölçeğin zamana karşı kararlılığı ve farklı kültürlerdeki kullanımına ilişkin ek araştırmalara ihtiyaç bulunmaktadır.

Bunun yanında araştırma yalnızca Türkiye'deki katılımcılar üzerinde gerçekleştirilmiştir. Yapay zekâyı ilişkin etik anlayış, güvenlik algısı ve teknolojik deneyimler kültürler arasında farklılık gösterebileceğinden, geliştirilen ölçeğin farklı ülkelerde uyarlanması ve kültürlerarası geçerlik çalışmalarının yapılması önem taşımaktadır.

Son olarak, yapay zekâ teknolojileri oldukça hızlı gelişen ve sürekli değişen bir alandır. Üretken yapay zekâ sistemlerinin yaygınlaşmasıyla birlikte güvenlik riskleri, etik ilkeler ve toplumsal beklentiler de sürekli değişmektedir. Bu nedenle geliştirilen ölçeğin ilerleyen yıllarda yeni

teknolojik gelişmeler doğrultusunda güncellenmesi ve yeniden psikometrik açıdan değerlendirilmesi gerekebilir.

Bu sınırlılıklara rağmen araştırma, güçlü örneklem büyüklüğü, kapsamlı psikometrik analizleri, yüksek geçerlik ve güvenilirlik değerleri ile yapay zekâ güvenlik ve ahlak algısını ölçmeye yönelik bilimsel açıdan geçerli ve güvenilir bir ölçek geliştirilmesine önemli katkı sağlamaktadır. Dolayısıyla araştırmanın bulgularının, ileride gerçekleştirilecek çalışmalara ve yapay zekâ etiği alanındaki uygulamalara önemli bir temel oluşturacağı değerlendirilmektedir.

5.4 Öneriler

5.4.1. Araştırmacılar İçin Öneriler

Bu araştırmadan elde edilen bulgular ve çalışmanın sınırlılıkları doğrultusunda gelecekte gerçekleştirilecek araştırmalar için aşağıdaki öneriler sunulmaktadır:

1. Geliştirilen Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ), farklı örneklem gruplarında yeniden uygulanarak psikometrik özellikleri test edilmelidir. Özellikle ortaöğretim öğrencileri, öğretmenler, akademisyenler, kamu çalışanları, sağlık çalışanları, hukukçular ve özel sektör çalışanları gibi farklı hedef gruplar üzerinde geçerlik ve güvenilirlik çalışmaları yapılabilir.
2. Ölçeğin farklı kültürlerde uyarlanmasına yönelik kültürlerarası geçerlik çalışmaları gerçekleştirilmelidir. Yapay zekâyâ ilişkin etik anlayış ve güvenlik algısı kültürel farklılıklardan etkilenebileceği için ölçeğin farklı ülkelerde uygulanması alanyazına önemli katkılar sağlayacaktır.
3. Araştırmada yalnızca kesitsel veriler kullanılmıştır. Gelecekte boylamsal (longitudinal) araştırmalar yapılarak bireylerin yapay zekâ güvenlik ve ahlak algılarının zaman içerisindeki değişimi incelenebilir. Özellikle yapay zekâ teknolojilerinin hızlı gelişimi dikkate alındığında bu tür çalışmalar önemli bilgiler sağlayacaktır.
4. Yapay zekâ kullanımının yalnızca varlığı değil, kullanım düzeyi de araştırılmalıdır. Gelecek çalışmalarda kullanım sıklığı, kullanım amacı, kullanılan yapay zekâ uygulamaları, kullanım süresi, yapay zekâ eğitimi alma durumu ve kullanım deneyimi gibi değişkenler ayrıntılı olarak incelenebilir.
5. YZGAAÖ ile yapay zekâ okuryazarlığı, dijital okuryazarlık, siber güvenlik farkındalığı, teknoloji kabulü, dijital vatandaşlık, algoritmik güven, teknoloji kaygısı ve yapay zekâ kaygısı gibi değişkenler arasındaki ilişkiler yapısal eşitlik modeli (YEM) kullanılarak

araştırılabilir. Böylece yapay zekâ güvenlik ve ahlak algısını etkileyen faktörler daha kapsamlı biçimde ortaya konulabilir.

6. Ölçeğin ölçüt bağıntılı geçerliği ve test-tekrar test güvenilirliği farklı örneklemeler üzerinde incelenmelidir. Böylece ölçeğin zamana karşı kararlılığı ve farklı ölçütlerle ilişkisi hakkında daha güçlü kanıtlar elde edilebilir.
7. Araştırmada nicel yöntem kullanılmıştır. Gelecekte nitel ve karma yöntem araştırmaları yapılarak bireylerin yapay zekâ güvenliği ve etik konularına ilişkin görüşleri, deneyimleri ve gerekçeleri daha ayrıntılı biçimde incelenebilir.
8. Farklı demografik değişkenlerin etkisi araştırılmalıdır. Özellikle meslek, çalışma deneyimi, dijital yeterlik düzeyi, internet kullanım süresi, teknolojiye yönelik tutum, yapay zekâ eğitimi alma durumu ve dijital etik eğitimi gibi değişkenlerin yapay zekâ güvenlik ve ahlak algısı üzerindeki etkisi incelenebilir.
9. Geliştirilen ölçek kullanılarak deneysel araştırmalar gerçekleştirilebilir. Yapay zekâ etiği, güvenliği veya yapay zekâ okuryazarlığı eğitimlerinin ön test-son test desenleriyle değerlendirilmesi, bu eğitimlerin bireylerin güvenlik ve etik algıları üzerindeki etkisini ortaya koyacaktır.
10. Uluslararası etik ilkeler doğrultusunda yeni alt boyutlar geliştirilebilir. Yapay zekâ teknolojilerindeki hızlı gelişmeler dikkate alınarak açıklanabilir yapay zekâ (Explainable AI), algoritmik önyargı, veri yönetimi, sürdürülebilir yapay zekâ, insan merkezli yapay zekâ ve güvenilir yapay zekâ gibi güncel konuların ölçeğe eklenmesi önerilmektedir.
11. Farklı disiplinlerde karşılaştırmalı çalışmalar yapılmalıdır. Eğitim, sağlık, mühendislik, hukuk, işletme ve sosyal bilimler gibi alanlarda öğrenim gören veya çalışan bireylerin yapay zekâ güvenlik ve ahlak algıları karşılaştırılarak disiplinler arası farklılıklar incelenebilir.
12. Ölçeğin farklı istatistiksel yaklaşımlarla yeniden değerlendirilmesi önerilmektedir. Özellikle Madde Tepki Kuramı (Item Response Theory-IRT), Rasch analizi ve ölçme değişmezliği (Measurement Invariance) analizleri gerçekleştirilerek ölçeğin psikometrik özellikleri daha ayrıntılı biçimde incelenebilir.

Bu önerilerin hayata geçirilmesi, yapay zekâ güvenliği ve etik alanında yapılacak çalışmaların kuramsal ve uygulamalı yönlerini güçlendirecek; geliştirilen ölçeğin farklı örneklem gruplarında kullanımını yaygınlaştırarak alanyazına önemli katkılar sağlayacaktır. Özellikle yapay zekâ teknolojilerinin hızla geliştiği günümüzde, güvenlik ve etik algılarının düzenli

olarak izlenmesi ve bilimsel ölçme araçlarının güncellenmesi, gelecekte yapılacak araştırmalar açısından büyük önem taşımaktadır.

5.4.2. Eğitim Kurumları İçin Öneriler

Bu araştırmadan elde edilen bulgular doğrultusunda eğitim kurumlarına yönelik aşağıdaki öneriler geliştirilmiştir:

1. Yapay zekâ güvenliği ve etiği konuları tüm eğitim kademelerinde öğretim programlarına dâhil edilmelidir. Yapay zekânın yalnızca teknik yönü değil; veri güvenliği, etik ilkeler, hesap verebilirlik, mahremiyet, algoritmik önyargı ve insan hakları gibi boyutları da sistematik olarak öğrencilere kazandırılmalıdır.
2. Özellikle üniversitelerde "Yapay Zekâ Okuryazarlığı", "Yapay Zekâ Etiği" ve "Güvenilir Yapay Zekâ" içerikli seçmeli veya zorunlu dersler açılmalıdır. Bu derslerde öğrencilerin yalnızca yapay zekâ araçlarını kullanmaları değil, bu teknolojileri eleştirel ve etik bir bakış açısıyla değerlendirmeleri de sağlanmalıdır.
3. Öğretmen yetiştiren programlarda yapay zekâ güvenliği ve etik farkındalığını geliştirmeye yönelik uygulamalı eğitimlere yer verilmelidir. Geleceğin öğretmenlerinin yapay zekâ teknolojilerini güvenli, sorumlu ve bilinçli şekilde kullanabilmeleri, öğrencilerine doğru rehberlik edebilmeleri açısından büyük önem taşımaktadır.
4. Öğretim elemanları ve öğretmenlere yönelik hizmet içi eğitim programları düzenlenmelidir. Bu programlarda üretken yapay zekâ araçlarının eğitimde kullanımı, akademik dürüstlük, telif hakları, kişisel verilerin korunması, güvenlik riskleri ve etik kullanım ilkeleri ele alınmalıdır.
5. Üniversitelerde yapay zekâ kullanımına ilişkin kurumsal etik ilkeler ve kullanım rehberleri hazırlanmalıdır. Öğrenciler ve akademik personelin yapay zekâ uygulamalarını hangi amaçlarla ve hangi sınırlar içerisinde kullanabileceklerini açıklayan açık politika belgeleri oluşturulmalıdır.
6. Akademik dürüstlüğü destekleyen yapay zekâ kullanım politikaları geliştirilmelidir. Yapay zekâ destekli ödev, proje, tez ve bilimsel yayın süreçlerinde etik kullanım esasları belirlenmeli; öğrencilerin ve araştırmacıların bu kurallara ilişkin farkındalıkları artırılmalıdır.
7. Derslerde yalnızca yapay zekâ uygulamalarının avantajlarına değil, olası risklerine de yer verilmelidir. Özellikle yanlış bilgi üretimi (hallucination), algoritmik önyargılar,

- veri güvenliği, mahremiyet ihlalleri, telif hakları ve siber güvenlik riskleri uygulamalı örneklerle öğrencilere aktarılmalıdır.
8. Yapay zekâ destekli öğrenme ortamlarında eleştirel düşünme becerileri geliştirilmelidir. Öğrencilerin yapay zekâ tarafından üretilen bilgileri sorgulayabilmeleri, doğrulayabilmeleri ve güvenilir kaynaklarla karşılaştırabilmeleri teşvik edilmelidir.
 9. Siber güvenlik ve kişisel veri koruma eğitimleri yapay zekâ eğitimleriyle bütünleştirilmelidir. Özellikle üretken yapay zekâ sistemlerine kişisel veya kurumsal verilerin girilmesi sırasında oluşabilecek riskler konusunda öğrenciler bilinçlendirilmelidir.
 10. Yapay zekâ teknolojilerinin eğitimde güvenli kullanımını destekleyecek dijital altyapılar oluşturulmalıdır. Kurumsal yapay zekâ platformları, lisanslı uygulamalar ve güvenli erişim sistemleri kullanılarak öğrencilerin güvenilir ortamlarda yapay zekâ deneyimi kazanmaları sağlanmalıdır.
 11. Üniversitelerde disiplinler arası yapay zekâ etik komisyonları veya çalışma grupları oluşturulmalıdır. Eğitim, mühendislik, hukuk, sağlık ve sosyal bilimler alanlarından uzmanların ortak çalışmalarıyla yapay zekâ kullanımına yönelik etik standartlar geliştirilebilir.
 12. Eğitim kurumları, öğrencilerin yapay zekâ güvenlik ve ahlak algılarını belirli aralıklarla değerlendirmelidir. Bu çalışmada geliştirilen Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ), öğrencilerin etik ve güvenlik farkındalıklarının izlenmesi, ihtiyaç analizlerinin yapılması ve eğitim programlarının etkililiğinin değerlendirilmesinde kullanılabilir.
 13. Öğrencilere yapay zekânın toplumsal etkileri konusunda farkındalık kazandıracak etkinlikler düzenlenmelidir. Seminerler, çalıştaylar, panel ve konferanslar aracılığıyla yapay zekânın istihdam, hukuk, sağlık, eğitim, medya ve demokrasi üzerindeki etkileri tartışılmalı; öğrencilerin çok yönlü bakış açısı geliştirmeleri desteklenmelidir.
 14. Yapay zekâ eğitimi yaşam boyu öğrenme anlayışıyla ele alınmalıdır. Yapay zekâ teknolojilerinin sürekli gelişmesi nedeniyle eğitim kurumları, öğrencilerin ve çalışanların güncel gelişmeleri takip edebilecekleri sürekli eğitim programları ve sertifika uygulamaları geliştirmelidir.

Sonuç olarak eğitim kurumlarının temel amacı, bireylere yalnızca yapay zekâ teknolojilerini kullanma becerisi kazandırmak değil; aynı zamanda bu teknolojileri güvenli, etik, şeffaf, hesap verebilir ve toplumsal sorumluluk bilinciyle kullanabilecek bireyler yetiştirmek olmalıdır. Bu

araştırmada geliştirilen YZGAAÖ, eğitim kurumlarının bu hedefe yönelik uygulamalarını değerlendirmede ve geliştirmede kullanılabilecek önemli bir ölçme aracı niteliğindedir.

5.4.3. Politika Yapıcılar İçin Öneriler

Bu araştırmadan elde edilen bulgular doğrultusunda politika yapıcılara yönelik aşağıdaki öneriler sunulmaktadır:

1. Ulusal düzeyde yapay zekâ güvenliği ve etik kullanımına ilişkin kapsamlı politika ve strateji belgeleri hazırlanmalı ve düzenli olarak güncellenmelidir. Bu belgeler, teknolojik gelişmelere uyum sağlayacak esnek bir yapıda olmalı ve uluslararası standartlarla uyumlu şekilde geliştirilmelidir.
2. Yapay zekâya ilişkin yasal düzenlemeler etik ilkeler temelinde güçlendirilmelidir. Özellikle veri gizliliği, kişisel verilerin korunması, algoritmik şeffaflık, hesap verebilirlik, ayrımcılığın önlenmesi ve insan haklarının korunmasına yönelik hukuki altyapı sürekli geliştirilmelidir.
3. Milli Eğitim Bakanlığı ve Yükseköğretim Kurulu tarafından yapay zekâ okuryazarlığı, güvenliği ve etik kullanımına yönelik ulusal eğitim standartları oluşturulmalıdır. Bu standartlar okul öncesinden yükseköğretime kadar tüm eğitim kademelerinde uygulanabilecek ortak yeterlikleri içermelidir.
4. Öğretmenler, akademisyenler ve kamu çalışanlarına yönelik sürekli mesleki gelişim programları desteklenmelidir. Yapay zekâ teknolojilerinin güvenli ve etik kullanımına ilişkin hizmet içi eğitimler yaygınlaştırılmalı ve düzenli aralıklarla güncellenmelidir.
5. Kamu kurumlarında yapay zekâ kullanımına ilişkin etik rehberler ve kurumsal kullanım politikaları hazırlanmalıdır. Kamu hizmetlerinde kullanılan yapay zekâ sistemlerinin şeffaf, denetlenebilir, adil ve hesap verebilir olması sağlanmalıdır.
6. Ulusal yapay zekâ etik kurulları veya danışma komisyonları oluşturulmalıdır. Bu kurullarda eğitim, hukuk, mühendislik, sağlık, sosyal bilimler ve bilişim alanlarından uzmanlar birlikte görev alarak yapay zekâ uygulamalarına ilişkin etik değerlendirmeler yapmalıdır.
7. Yapay zekâ teknolojilerinin geliştirilmesi kadar güvenli kullanımına yönelik araştırmalar da desteklenmelidir. Özellikle etik yapay zekâ, açıklanabilir yapay zekâ (Explainable AI), güvenilir yapay zekâ (Trustworthy AI), algoritmik adalet ve veri yönetişimi alanlarında yürütülen bilimsel projelere öncelik verilmelidir.

8. Yapay zekâ sistemlerinin kamu ve özel sektörde kullanımına yönelik denetim mekanizmaları oluşturulmalıdır. Yapay zekâ tabanlı karar destek sistemlerinin düzenli olarak etik, güvenlik ve ayrımcılık açısından değerlendirilmesi sağlanmalıdır.
9. Toplumun yapay zekâ konusunda bilinçlendirilmesine yönelik ulusal farkındalık çalışmaları yürütülmelidir. Kamu spotları, çevrim içi eğitimler, medya kampanyaları ve bilgilendirme faaliyetleriyle bireylerin yapay zekâ teknolojilerinin fırsatları ve riskleri hakkında bilinç düzeyleri artırılmalıdır.
10. Ulusal yapay zekâ politikalarında dijital eşitlik ve kapsayıcılık ilkeleri gözetilmelidir. Yapay zekâ teknolojilerine erişimde fırsat eşitliğini artıracak uygulamalar geliştirilmeli; dezavantajlı grupların bu teknolojilerden güvenli biçimde yararlanmaları desteklenmelidir.
11. Yapay zekâ sistemlerinin eğitim, sağlık, hukuk ve kamu yönetimi gibi kritik alanlarda kullanımına yönelik etik etki değerlendirmeleri zorunlu hâle getirilmelidir. Böylece olası riskler uygulama öncesinde belirlenerek gerekli önlemler alınabilir.
12. Uluslararası kuruluşlarla iş birliği güçlendirilmelidir. Avrupa Birliği, UNESCO, OECD ve diğer uluslararası kuruluşların yapay zekâ etik ilkeleri ve güvenlik standartları dikkate alınarak ulusal politikalar güncellenmeli ve uluslararası iyi uygulama örneklerinden yararlanılmalıdır.
13. Yapay zekâ güvenliği ve etik farkındalığının düzenli olarak izlenmesi için ulusal ölçme ve değerlendirme çalışmaları yapılmalıdır. Bu kapsamda geliştirilen Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ) farklı yaş ve meslek gruplarında uygulanarak toplumun yapay zekâ güvenliği ve etik algısı belirli aralıklarla değerlendirilebilir. Böylece geliştirilen politikaların etkililiği bilimsel verilerle izlenebilir.
14. Politika geliştirme süreçlerinde çok paydaşlı bir yaklaşım benimsenmelidir. Kamu kurumları, üniversiteler, özel sektör, sivil toplum kuruluşları ve meslek örgütlerinin iş birliği içerisinde hareket etmesi, yapay zekâ politikalarının daha kapsayıcı ve sürdürülebilir olmasını sağlayacaktır.

Sonuç olarak yapay zekâ teknolojilerinin güvenli, etik ve insan odaklı biçimde geliştirilmesi ve kullanılabilmesi, yalnızca teknik düzenlemelerle değil; eğitim, hukuk, kamu yönetimi ve toplumsal farkındalık alanlarında bütüncül politikaların oluşturulmasıyla mümkündür. Bu araştırmada geliştirilen YZGAAÖ, politika yapıcıların toplumun yapay zekâ güvenliği ve etik farkındalık düzeyini bilimsel olarak izlemelerine, ihtiyaç alanlarını belirlemelerine ve

geliştirilen politikaların etkisini değerlendirmelerine katkı sağlayabilecek geçerli ve güvenilir bir ölçme aracı olarak kullanılabilir.

5.4.4. Yapay Zekâ Geliştiricileri İçin Öneriler

Bu araştırmanın bulguları doğrultusunda yapay zekâ sistemlerini geliştiren kurum ve geliştiricilere yönelik aşağıdaki öneriler sunulmaktadır:

1. Yapay zekâ sistemlerinin tasarım sürecinde etik ilkeler temel alınmalıdır. Geliştirilen sistemlerde adalet, şeffaflık, hesap verebilirlik, insan haklarına saygı ve zarar vermeme ilkeleri yazılım geliştirme yaşam döngüsünün ayrılmaz bir parçası hâline getirilmelidir.
2. Geliştirilen yapay zekâ sistemleri "güvenlik tasarımıyla birlikte gelir (security by design)" yaklaşımıyla geliştirilmelidir. Veri güvenliği, siber güvenlik, kimlik doğrulama, yetkilendirme ve saldırılara karşı dayanıklılık gibi unsurlar geliştirme sürecinin başlangıcından itibaren dikkate alınmalıdır.
3. Yapay zekâ sistemlerinin karar verme süreçleri mümkün olduğunca açıklanabilir (Explainable AI) hâle getirilmelidir. Kullanıcıların sistemin hangi verilerden yararlanarak nasıl sonuç ürettiğini anlayabilmesi, güven ve hesap verebilirliğin artırılmasına katkı sağlayacaktır.
4. Algoritmik önyargıları azaltmaya yönelik düzenli testler yapılmalıdır. Eğitim verileri farklı demografik grupları temsil edecek şekilde oluşturulmalı ve geliştirilen modeller ayrımcılık, adalet ve tarafsızlık açısından sürekli değerlendirilmelidir.
5. Kişisel verilerin korunmasına yönelik ulusal ve uluslararası veri koruma standartlarına uygun hareket edilmelidir. Veri minimizasyonu, anonimleştirme, şifreleme ve kullanıcı onayı gibi uygulamalar geliştirme sürecinin temel bileşenleri olmalıdır.
6. Yapay zekâ sistemlerinde insan denetimi (human oversight) korunmalıdır. Özellikle eğitim, sağlık, hukuk, finans ve kamu hizmetleri gibi kritik alanlarda yapay zekâ sistemlerinin tamamen bağımsız karar vermesi yerine, nihai kararın insan tarafından gözden geçirilebildiği mekanizmalar oluşturulmalıdır.
7. Kullanıcıların yapay zekâ sistemlerini güvenli kullanabilmeleri için bilgilendirici uyarılar ve kullanım rehberleri sunulmalıdır. Sistemlerin güçlü yönleri kadar sınırlılıkları, hata yapabileceği durumlar ve güvenli kullanım ilkeleri de açık biçimde kullanıcılarla paylaşılmalıdır.

8. Yapay zekâ modelleri düzenli olarak etik ve güvenlik denetimlerinden geçirilmelidir. Yeni veri kümeleri ve güncellemeler sonrasında performans, güvenlik, doğruluk ve etik riskler yeniden değerlendirilmeli; gerekli iyileştirmeler yapılmalıdır.
9. Yapay zekâ sistemlerinde kullanıcı geri bildirim mekanizmaları etkin biçimde kullanılmalıdır. Kullanıcıların yanlış, önyargılı veya zararlı çıktıları kolaylıkla bildirebilecekleri sistemler oluşturularak modeller sürekli iyileştirilebilir.
10. Yapay zekâ geliştiricileri disiplinler arası ekiplerle çalışmalıdır. Yazılım geliştiricilerinin yanı sıra etik uzmanları, hukukçular, eğitim bilimciler, psikologlar, sosyologlar ve alan uzmanlarının geliştirme süreçlerine katılımı daha güvenilir ve toplumsal açıdan kabul edilebilir sistemlerin oluşturulmasına katkı sağlayacaktır.
11. Geliştirilen sistemlerin toplumsal etkileri ürün geliştirme sürecinde değerlendirilmelidir. Özellikle istihdam, eğitim, mahremiyet, demokratik süreçler, çocukların korunması ve sosyal eşitlik üzerindeki olası etkiler uygulamaya geçmeden önce analiz edilmelidir.
12. Yapay zekâ uygulamalarında etik kullanım ilkeleri ve kullanım koşulları kullanıcıların kolay anlayabileceği biçimde sunulmalıdır. Karmaşık teknik açıklamalar yerine açık, anlaşılır ve erişilebilir bilgilendirme metinleri hazırlanmalıdır.
13. Yapay zekâ sistemlerinin uluslararası etik standartlarla uyumu düzenli olarak gözden geçirilmelidir. Özellikle Avrupa Birliği Yapay Zekâ Tüzüğü (EU AI Act), UNESCO Yapay Zekâ Etiği Tavsiye Kararı ve OECD Yapay Zekâ İlkeleri gibi uluslararası çerçeveler dikkate alınarak sistemler güncellenmelidir.
14. Yapay zekâ geliştiricileri güven oluşturmaya teknik performans kadar önemli bir hedef olarak görmelidir. Kullanıcı güvenini artıran şeffaflık, hesap verebilirlik, veri güvenliği ve etik tasarım ilkeleri, sistemlerin uzun vadeli kabulü ve sürdürülebilirliği açısından temel unsurlar olarak değerlendirilmelidir.
15. Geliştirilen sistemlerin kullanıcılar üzerindeki etkileri bilimsel yöntemlerle düzenli olarak değerlendirilmelidir. Bu kapsamda geliştirilen Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği (YZGAAÖ), kullanıcıların yapay zekâ sistemlerine ilişkin güvenlik ve etik algılarındaki değişimi ölçmek, yeni geliştirilen uygulamaların toplumsal kabulünü değerlendirmek ve kullanıcı odaklı iyileştirmeler yapmak amacıyla kullanılabilir.

Sonuç olarak yapay zekâ sistemlerinin başarısı yalnızca yüksek doğruluk oranları veya teknik performansla değil; güvenilir, şeffaf, adil, hesap verebilir ve insan merkezli bir yaklaşımla geliştirilmelerine bağlıdır. Bu araştırmada elde edilen bulgular, kullanıcıların yapay zekâ

teknolojilerine yönelik güvenlik ve etik konularına önem verdiklerini göstermektedir. Bu nedenle geliştiricilerin teknik yenilikleri etik sorumlulukla birlikte ele almaları, yapay zekâ teknolojilerinin toplumsal kabulünü ve sürdürülebilir kullanımını önemli ölçüde artıracaktır.

KAYNAKÇA

- Abdel-Jaber, H., Devassy, D., Salam, A., Hidaytallah, L., & El-Amir, M. (2022). A Review of Deep Learning Algorithms and Their Applications in Healthcare. *Algorithms*, 15, 71. <https://doi.org/10.3390/a15020071>.
- Abuadas, M., & Albikawi, Z. (2025). AI ethical awareness and academic integrity in higher education: development and validation of a new scale. *Ethics & Behavior*, 36, 125 - 142. <https://doi.org/10.1080/10508422.2025.2511336>
- Abuadas, M., Albikawi, Z., & Rayani, AM (2025).Yapay zekâ odaklı bir etik eğitim programının hemşirelik öğrencilerinin etik farkındalığı, ahlaki duyarlılığı, tutumları ve üretken yapay zekâ benimseme niyeti üzerindeki etkisi: yarı deneysel bir çalışma. BMC Nursing, 24. <https://doi.org/10.1186/s12912-025-03458-2>
- Abuelezz, I., Barhamgi, M., Alshakhsi, S., Yankouskaya, A., Nhlabatsi, A., Khan, K. M., & Ali, R. (2024). How do gender and age similarities with a potential social engineer influence one's trust and willingness to take security risks? *International Journal of Information Security*, 24(1). <https://doi.org/10.1007/s10207-024-00954-5>
- Adadi, A., & Berrada, M. (2018). Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI). *IEEE Access*, 6, 52138-52160. <https://doi.org/10.1109/access.2018.2870052>
- Advincula, A. N. C., Meza,, J. D. E., Reyes Pacheco, L. P., & Guevara Ruiz, A. E. (2026). Ethics And Morality: The Philosophical Dilemma of Human Actions. *International Journal of Multidisciplinary and Innovative Research*. <https://doi.org/10.58806/ijmir.2026.v3i3n02>
- Afroogh, S., Akbari, A., Malone, E., Kargar, M., & Alambeigi, H. (2024). Trust in AI: progress, challenges, and future directions. *Humanities and Social Sciences Communications*, 11. <https://doi.org/10.1057/s41599-024-04044-8>.
- Akinrinola, O., Okoye, C., Ofodile, O., & Ugochukwu, C. (2024). Navigating and reviewing ethical dilemmas in AI development: Strategies for transparency, fairness, and accountability. *GSC Advanced Research and Reviews*. <https://doi.org/10.30574/gscarr.2024.18.3.0088>.
- Alasmari, A. A., Alruwaili, R. F., Alotaibi, R. F., Youssef, I. K., & Asklany, S. A. (2025). Demographic influences on trust in artificial intelligence across cognitive domains: A statistical perspective. *PLoS ONE*, 20(11), e0331003. <https://doi.org/10.1371/journal.pone.0331003>

- Alzubaidi, L., Al-Sabaawi, A., Bai, J., Dukhan, A., Alkenani, A. H., Al-Asadi, A., Alwzawzy, H. A., Manoufali, M., Fadhel, M. A., Albahri, A. S., Moreira, C., Ouyang, C., Zhang, J., Santamaría, J., Salhi, A., Hollman, F., Gupta, A., Duan, Y., Rabczuk, T., . . . Gu, Y. (2023). Towards Risk-Free Trustworthy Artificial intelligence: Significance and requirements. *International Journal of Intelligent Systems*, 2023(1).
<https://doi.org/10.1155/2023/4459198>
- Amanullah, M., Habeeb, R., Nasaruddin, F., Gani, A., Ahmed, E., Nainar, A., Akim, N., & Imran, M. (2020). Deep learning and big data technologies for IoT security. *Comput. Commun.*, 151, 495-517. <https://doi.org/10.1016/j.comcom.2020.01.016>.
- Amaro, I., Barra, P., Della Greca, A., Francese, R., & Tucci, C. (2023). Believe in artificial intelligence? A user study on the ChatGPT's fake information impact. *IEEE Transactions on Computational Social Systems*, 11(4), 5168–5177.
<https://doi.org/10.1109/tcss.2023.3291539>
- Ambani, D., & Rathod, D. (2025). The evolution of artificial intelligence: from rule-based systems to deep learning. *Interantional Journal of Scientific Research in Engineering and Management*. <https://doi.org/10.55041/ijsrem30072>
- Araujo, T., Helberger, N., Kruikemeier, S., & de Vreese, C. H. (2020). In AI we trust? Perceptions about automated decision-making by artificial intelligence. *AI & Society*, 35(3), 611–623.
- Author, F. (2026). The Importance of Ethics and Morals: Their Foundations, their Differences, their Values and Applications within all Human Actions as a Necessity and Ideal, Despite Human Adversities and Limitations, in the Search for Hope for a Better Future. *Global Journal of Human-Social Science*. <https://doi.org/10.34257/gjhssh257765>
- Aunugu, D., & Vathsavai, V. (2022). Privacy-First Artificial Intelligence: Toward Fair, Transparent, and Accountable Systems. *International Journal Of Scientific Advances*.
<https://doi.org/10.51542/ijscia.v3i6.26>.
- Ayanwale, M., Adelana, O., Molefi, R., Adeeko, O., & Ishola, A. (2024). Examining artificial intelligence literacy among pre-service teachers for future classrooms. *Computers and Education Open*. <https://doi.org/10.1016/j.caeo.2024.100179>.
- Bach, T. A., Khan, A., Hallock, H. P., Beltrao, G., & Sousa, S. (2022). A Systematic Literature Review of User Trust in AI-Enabled Systems: An HCI Perspective. *International Journal of Human-Computer Interaction*, 40, 1251 - 1266.
<https://doi.org/10.1080/10447318.2022.2138826>

- Bahangulu, J., & Owusu-Berko, L. (2025). Algorithmic bias, data ethics, and governance: Ensuring fairness, transparency and compliance in AI-powered business analytics applications. *World Journal of Advanced Research and Reviews*. <https://doi.org/10.30574/wjarr.2025.25.2.0571>.
- Bansal, R., & Sangwan, A. (2024). A Comprehensive Review of Artificial Intelligence. *International Journal For Multidisciplinary Research*. <https://doi.org/10.36948/ijfmr.2024.v06i05.28932>
- Bartlett, M. S. (1954). A Note on the Multiplying Factors for Various χ^2 Approximations. *Journal of the Royal Statistical Society Series B (Statistical Methodology)*, 16(2), 296–298. <https://doi.org/10.1111/j.2517-6161.1954.tb00174.x>
- Basch, C., Hillyer, G., Gold, B., & Yousaf, H. (2025). Artificial Intelligence in Higher Education: Student Knowledge, Attitudes, and Ethical Perceptions in the United States. *SDGs Studies Review*. <https://doi.org/10.37497/sdgs.v6istudies.34>.
- Bayram, V. (2025). Yapay Zekâ Etiği: Bir Ölçek Geliştirme Çalışması. *Uluslararası İktisadi ve İdari Akademik Araştırmalar Dergisi*, 5(1), 123-148. <https://izlik.org/JA99TS89BN>
- Biagini, G. (2025). Towards an AI-Literate Future: A Systematic Literature Review Exploring Education, Ethics, and Applications. *International Journal of Artificial Intelligence in Education*, 35, 2616 - 2666. <https://doi.org/10.1007/s40593-025-00466-w>
- Biggio, B., & Roli, F. (2018). Wild patterns: Ten years after the rise of adversarial machine learning. *Pattern Recognition*, 84, 317–331. <https://doi.org/10.1016/j.patcog.2018.07.023>
- Bilbao-Eraña, A., & Arroyo-Sagasta, A. (2025). Fostering AI literacy in pre-service teachers: impact of a training intervention on awareness, attitude and trust in AI. *Frontiers in Education*. <https://doi.org/10.3389/feduc.2025.1668078>.
- Blagoj, D., Chrysi, T., & Uros, K. (2020). Historical Evolution of Artificial Intelligence: Analysis of the three main paradigm shifts in AI. <https://doi.org/10.2760/801580>
- Boateng, G. O., Neilands, T. B., Frongillo, E. A., Melgar-Quinonez, H. R., & Young, S. L. (2018). Best practices for developing and validating scales for health, social, and behavioral research: A primer. *Frontiers in Public Health*, 6, 149.
- Bonem, E. M., Ellsworth, P. C., & Gonzalez, R. (2015). Age differences in risk: perceptions, intentions and domains. *Journal of Behavioral Decision Making*, 28(4), 317–330. <https://doi.org/10.1002/bdm.1848>
- Bote-Curiel, L., Muñoz-Romero, S., Gerrero-Curiseses, A., & Rojo-Álvarez, J. (2019). Deep Learning and Big Data in Healthcare: A Double Review for Critical Beginners. *Applied Sciences*. <https://doi.org/10.3390/app9112331>.

- Bozdoğanoglu, B., Haspolat, İ., & Yücel, A. (2024). Kamu İdarelerinde Yapay Zekâ Kullanımının Ülke Uygulamaları ve Temel Kamusal İlkeler Çerçevesinde Değerlendirilmesi. *Ankara Hacı Bayram Veli Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi*. <https://doi.org/10.26745/ahbvuibfd.1424290>
- Branley-Bell, D., Coventry, L., Dixon, M., Joinson, A., & Briggs, P. (2022). Exploring age and gender differences in ICT cybersecurity behaviour. *Human Behavior and Emerging Technologies*, 2022, 1–10. <https://doi.org/10.1155/2022/2693080>
- Brauner, P., Hick, A., Philipsen, R., & Ziefle, M. (2023). What does the public think about artificial intelligence?—A criticality map to understand bias in the public perception of AI. ,5. <https://doi.org/10.3389/fcomp.2023.1113903>.
- Borwein, S., Magistro, B., Alvarez, R. M., Bonikowski, B., & Loewen, P. J. (2026). Explaining women’s skepticism toward artificial intelligence: The role of risk orientation and risk exposure. *PNAS Nexus*, 5(1), pgaf399. <https://doi.org/10.1093/pnasnexus/pgaf399>
- Brynjolfsson, E., & McAfee, A. P. (2017). *Machine, platform, crowd: Harnessing our digital future*. <https://www.amazon.com/Machine-Platform-Crowd-Harnessing-Digital/dp/039335606X>
- Budak, V. Ö., & Erol, Ç. (2020). Web Kullanıcılarının Bilgi Erişim ve Ziyaret Desenlerinin Web Madenciliği ile Keşfi: Kırklareli Üniversitesi Örneği. *AJIT-e: Academic Journal of Information Technology*, 11(43), 37-55.
- Bukovec, M., & Antoliš, K. (2024). The impact of age and education on cyber security in digital banking. *Medijska Istraživanja*, 30(2), 129–152. <https://doi.org/10.22572/mi.30.2.6>
- Bullock, J., Pauketat, J., Huang, H., Wang, Y., & Anthis, J. (2025). Public Opinion and the Rise of Digital Minds: Perceived Risk, Trust, and Regulation Support. *Public Performance & Management Review*, 48, 1357 - 1388. <https://doi.org/10.1080/15309576.2025.2495094>.
- Büyüköztürk, Ş. (2020). *Sosyal Bilimler İçin Veri Analizi El Kitabı* (28. Baskı). Pegem Akademi.
- Camilleri, M. A. (2023). Artificial intelligence governance: Ethical considerations and implications for social responsibility. *Expert Systems*, 41(7), e13406. <https://doi.org/10.1111/exsy.13406>
- Čančer, V., Rožman, M., & Tominc, P. (2025). Beyond differences: a generational convergence in technology use among business students. *Systems*, 13(12), 1095. <https://doi.org/10.3390/systems13121095>

- Cao, G., Duan, Y., Edwards, J., & Dwivedi, Y. (2021). Understanding managers' attitudes and behavioral intentions towards using artificial intelligence for organizational decision-making. *Technovation*. <https://doi.org/10.1016/j.technovation.2021.102312>.
- Čartolovni, A., Tomićić, A., & Mosler, E. (2022). Ethical, legal, and social considerations of AI-based medical decision-support tools: A scoping review. *International journal of medical informatics*, 161, 104738 . <https://doi.org/10.1016/j.ijmedinf.2022.104738>.
- Carvalho, D., Corrêa, F., Faria, V., Lima, L., De Souza, A., De Moura Gromato, M., & Ferro, M. (2025). Expanding a Study on Users' Perception and Attitudes Toward Artificial Intelligence in Computational Systems. *Journal on Interactive Systems*. <https://doi.org/10.5753/jis.2025.5347>.
- Cattell, R. B. (1966). The scree test for the number of factors. *Multivariate Behavioral Research*, 1(2), 245–276. https://doi.org/10.1207/s15327906mbr0102_10
- Chaddha, R., & Agrawal, G. (2023). Ethics and Morality. *Indian Journal of Orthopaedics*, 57, 1707 - 1713. <https://doi.org/10.1007/s43465-023-01004-3>
- Chairul, M., Umanailo, B., Retrofilia, A., Umanailo, A., Sophia, D., & U. (2020). Functions of Values, Morals, Justice, Order and Community Welfare. <https://doi.org/10.22541/au.158680346.60332439>
- Challen, R., Denny, J., Pitt, M., Gompels, L., Edwards, T., & Tsaneva-Atanasova, K. (2019). Artificial intelligence, bias and clinical safety. *BMJ Quality & Safety*, 28, 231 - 237. <https://doi.org/10.1136/bmjqs-2018-008370>.
- Chan, C. K. Y. (2023). A comprehensive AI policy education framework for university teaching and learning. *International Journal of Educational Technology in Higher Education*, 20(1). <https://doi.org/10.1186/s41239-023-00408-3>
- Chaudhry, M. A., & Kazim, E. (2022). Artificial intelligence in education (AIEd): A high-level academic and industry note. *AI and Ethics*, 2(1), 157–165.
- Chavarria, A., Palau, R., & Santiago, R. (2025). Navigating stakeholders perspectives on artificial intelligence in higher education. *Algorithms*, 18(6), 336. <https://doi.org/10.3390/a18060336>
- Chee, H., Ahn, S., & Lee, J. (2024). A Competency Framework for AI Literacy: Variations by Different Learner Groups and an Implied Learning Pathway. *Br. J. Educ. Technol.*, 56, 2146-2182. <https://doi.org/10.1111/bjet.13556>
- Cheong, B. (2024). Transparency and accountability in AI systems: safeguarding wellbeing in the age of algorithmic decision-making. *Frontiers in Human Dynamics*. <https://doi.org/10.3389/fhumd.2024.1421273>.

- Choi, J.-I., Yang, E., & Goo, E.-H. (2024). Ortaokul Öğrencileri Arasında Yapay Zekaya Yönelik Etik Eğitim Programının Etkileri: Algı ve Tutum Değişikliklerinin Analizi. *Uygulamalı Bilimler*. <https://doi.org/10.3390/app14041588>
- Choung, H., David, P., & Ross, A. (2022). Trust and ethics in AI. *AI & SOCIETY*, 38, 733-745. <https://doi.org/10.1007/s00146-022-01473-4>.
- Choung, H., David, P., & Ross, A. (2022). Trust in AI and Its Role in the Acceptance of AI Technologies. *International Journal of Human–Computer Interaction*, 39, 1727 - 1739. <https://doi.org/10.1080/10447318.2022.2050543>
- Daneshjou, R., Smith, M., Sun, M., Rotemberg, V., & Zou, J. (2021). Lack of Transparency and Potential Bias in Artificial Intelligence Data Sets and Algorithms: A Scoping Review.. *JAMA dermatology*. <https://doi.org/10.1001/jamadermatol.2021.3129>.
- Davis, F. (1989). Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology. *MIS Q.*, 13, 319-340. <https://doi.org/10.2307/249008>.
- Deng, L. (2018). Artificial Intelligence in the Rising Wave of Deep Learning: The Historical Path and Future Outlook [Perspectives]. *IEEE Signal Processing Magazine*, 35, 180-177. <https://doi.org/10.1109/msp.2017.2762725>
- Deric, E., Frank, D., & Vuković, D. (2025). Yüksek Öğretimde Üretken Yapay Zeka Araçlarının Kullanımının Etik Sonuçlarını Keşfetmek. *Bilişim*, 12 , 36. <https://doi.org/10.3390/informatics12020036>
- DeVellis, R. F. (2017). *Scale Development: Theory and Applications* (4th ed.). Sage.
- DeVellis, R. F., & Thorpe, C. T. (2021). *Scale development: Theory and applications*. Sage publications.
- Dianova, V. G., & Schultz, M. D. (2023). Discussing ChatGPT's implications for industry and higher education: The case for transdisciplinarity and digital humanities. *Industry and Higher Education*, 37, 593 - 600. <https://doi.org/10.1177/09504222231199989>
- Díaz-Rodríguez, N., Javier, D. S., Coeckelbergh, M., López, D. P. M., Herrera-Viedma, E., & Herrera, F. (2023). Connecting the dots in trustworthy artificial intelligence: from AI principles, ethics, and key requirements to responsible AI systems and regulation. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2305.02231>
- Dignum, V. (2018). Ethics in artificial intelligence: introduction to the special issue. *Ethics and Information Technology*, 20, 1 - 3. <https://doi.org/10.1007/s10676-018-9450-z>.
- Ding, W., Wu, Y., & Wang, J. (2026). Beyond Risk Reduction: Vigilant Trust in Artificial Intelligence Based on Evidence from China. *Behavioral Sciences*, 16. <https://doi.org/10.3390/bs16010095>.

- Dwivedi, YK, Hughes, L., Ismagilova, ER, Aarts, G., Coombs, CR, Crick, T., Duan, Y., Dwivedi, R., Edwards, J., Eirug, A., Galanos, V., Ilavarasan, P., Janssen, M., Jones, P., Kar, A., Kizgin, H., Kronemann, B., Lal, B., Lucini, B., . . . Williams, MD (2019). Yapay Zeka (YZ): Araştırma, uygulama ve politika için ortaya çıkan zorluklar, fırsatlar ve gündem üzerine çok disiplinli bakış açıları. *Uluslararası Bilgi Yönetimi Dergisi* . <https://doi.org/10.1016/j.jjinfomgt.2019.08.002>
- Elamin, M. O. I. (2024). AI Through the Ages: Unlocking Key Opportunities and Navigating Challenges in the History and Future of Artificial Intelligence. *International Journal of Religion*. <https://doi.org/10.61707/tarzcg89>
- Ellemers, N. (2017). Morality and Social Identity. <https://doi.org/10.1093/oxfordhb/9780190247577.013.5>
- European Commission. (2019). *Ethics Guidelines for Trustworthy AI*.
- Faisal, R. (2025). Ethical AI Integration in Education: Policies, Challenges, and Strategies for Bridging Digital Divides in Emerging Economies. *Crossroads of Social Inquiry*. <https://doi.org/10.64851/csi.v1i2.28>
- Fjeld, J., Achten, N., Hilligoss, H., Nagy, Á., & Srikumar, M. (2020). Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-Based Approaches to Principles for AI. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3518482>.
- Fjelland, R. (2020). Why general artificial intelligence will not be realized. *Humanities and Social Sciences Communications*, 7. <https://doi.org/10.1057/s41599-020-0494-4>
- Floridi, L., & Cows, J. (2019). A Unified Framework of Five Principles for AI in Society. *Harvard Data Science Review*, 1(1). <https://doi.org/10.1162/99608f92.8cd550d1>
- Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. *Minds and Machines*, 28, 689 - 707. <https://doi.org/10.1007/s11023-018-9482-5>.
- G, R. V., Nair, S. G., & J. (2025). Exploring Ethical Concerns Among Students: The Impact of AI Usage in Education. *International Scientific Journal of Engineering and Management*. <https://doi.org/10.55041/isjem04845>
- Garnelo, M., & Shanahan, M. (2019). Reconciling deep learning with symbolic artificial intelligence: representing objects and relations. *Current Opinion in Behavioral Sciences*. <https://doi.org/10.1016/j.cobeha.2018.12.010>

- Gazulla, E. D., Hirvonen, N., Sharma, S., Hartikainen, H., Jylhä, V., Iivari, N., Kinnula, M., & Baizhanova, A. (2024). Youth perspectives on technology ethics: analysis of teens' ethical reflections on AI in learning activities. *Behaviour & Information Technology*, 44, 888 - 911. <https://doi.org/10.1080/0144929x.2024.2350666>
- Gerlich, M. (2023). Perceptions and Acceptance of Artificial Intelligence: A Multi-Dimensional Study. *Social Sciences*. <https://doi.org/10.3390/socsci12090502>.
- Ghotbi, N., & Ho, M.-T. (2021). Moral Awareness of College Students Regarding Artificial Intelligence. *Asian Bioethics Review*, 13, 421 - 433. <https://doi.org/10.1007/s41649-021-00182-2>
- Gignac, G. E., & Szodorai, E. T. (2024). Defining intelligence: Bridging the gap between human and artificial perspectives. *Intelligence*. <https://doi.org/10.1016/j.intell.2024.101832>
- Gillespie, N., Lockey, S., Curtis, C., Pool, J., & Akbari, A. (2023). Trust in Artificial Intelligence: A global study. In *The University of Queensland*. <https://doi.org/10.14264/00d3c94>
- Glikson, E., & Woolley, A. W. (2020). Human Trust in Artificial Intelligence: Review of Empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Göktepe, Y. (2019). Yapay zekâ tekniklerinin kullanımıyla protein etkileşimlerinin sekans bilgisine dayalı tahmini (Prediction of protein interactions by using artificial intelligence techniques based on protein sequence data).
- Granić, A., & Marangunic, N. (2019). Technology acceptance model in educational context: A systematic literature review. *Br. J. Educ. Technol.*, 50, 2572-2593. <https://doi.org/10.1111/bjet.12864>.
- Grassini, S. (2023). Development and validation of the AI attitude scale (AIAS-4): a brief measure of general attitude toward artificial intelligence. *Frontiers in Psychology*, 14. <https://doi.org/10.3389/fpsyg.2023.1191628>.
- Guan, L., Zhang, Y., & Gu, M. (2024). Pre-service teachers preparedness for AI-integrated education: An investigation from perceptions, capabilities, and teachers' identity changes. *Comput. Educ. Artif. Intell.*, 8, 100341. <https://doi.org/10.1016/j.caeai.2024.100341>.
- Gümüş, E., Medetoğlu, B., & Tutar, S. (2020). Finans ve Bankacılık Sisteminde Yapay Zekâ Kullanımı: Kullanıcılar Üzerine Bir Uygulama. , 3, 28-53. <https://doi.org/10.38057/bifd.690982>.

- Gümüş, M., Çakır, R., & Korkmaz, Ö. (2023). Investigation of pre-service teachers' sensitivity to cyberbullying, perceptions of digital ethics and awareness of digital data security. *Education and Information Technologies*, 28, 14399-14421. <https://doi.org/10.1007/s10639-023-11785-7>.
- Hagendorff, T. (2024). Mapping the Ethics of Generative AI: A Comprehensive Scoping Review. *Minds and Machines*, 34. <https://doi.org/10.1007/s11023-024-09694-w>.
- Haidt, J. (2007). The New Synthesis in Moral Psychology. *Science*, 316, 1002 - 998. <https://doi.org/10.1126/science.1137651>
- Hair, J., Black, W. C., Babin, B. J., & Anderson, R. E. Multivariate data analysis 8th ed., 2019. *Boston, USA: Cengage Learning*.
- Hamamcı, H., & Uyar, S. (2024). AI bias in auditing: A systematic review of sources, consequences, and solutions. *Scientific African*, 24, e02256. <https://doi.org/10.1016/j.sciaf.2024.e02256>
- Hartwig, T., Ikkatai, Y., Takanashi, N., & Yokoyama, H. (2022). Artificial intelligence ELSI score for science and technology: a comparison between Japan and the US. *AI & SOCIETY*, 38, 1609-1626. <https://doi.org/10.1007/s00146-021-01323-9>.
- Hassija, V., Chamola, V., Mahapatra, A., Singal, A., Goel, D., Huang, K., Scardapane, S., Spinelli, I., Mahmud, M., & Hussain, A. (2023). Interpreting Black-Box Models: A Review on Explainable Artificial Intelligence. *Cognitive Computation*, 16, 45-74. <https://doi.org/10.1007/s12559-023-10179-8>
- Haunschild, I. M. S., & Leipold, B. (2023). The Use of Digital Media and Security Precautions in Adulthood. *Human Behavior and Emerging Technologies*. <https://doi.org/10.1155/2023/1436710>
- Henkoğlu, T. (2019). Bilgi Saklama Süreçlerinde Yapay Zekâ Sistemlerinin Kullanımına Yönelik Risk Değerlendirmesi. , 6, 134-147.
- Hagerty, A., & Rubinov, I. (2019). Küresel Yapay Zeka Etiği: Yapay Zekanın Sosyal Etkileri ve Etik Sonuçlarının Gözden Geçirilmesi. *ArXiv*, abs/1907.07892 . <https://doi.org/10.48550/arxiv.1907.07892>
- HiddenLayer. (2025). *AI threat landscape report 2025*. <https://www.hiddenlayer.com/ai-threat-landscape-report-2025>
- Hornberger, M., Bewersdorff, A., & Nerdel, C. (2023). What do university students know about Artificial Intelligence? Development and validation of an AI literacy test. *Comput. Educ. Artif. Intell.*, 5, 100165. <https://doi.org/10.1016/j.caeai.2023.100165>.

- Horner, J. (2003). Morality, Ethics, and Law: Introductory Concepts. *SEMINARS IN SPEECH AND LANGUAGE*, 24, 263 - 274. <https://doi.org/10.1055/s-2004-815580>
- Hu, L., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis. *Structural Equation Modeling*, 6(1), 1–55.
- Huang, L. (2023). Ethics of Artificial Intelligence in Education: Student Privacy and Data Protection. *Science Insights Education Frontiers*. <https://doi.org/10.15354/sief.23.re202>
- Huang, C., Zhang, Z., Mao, B., & Yao, X. (2023). An Overview of Artificial Intelligence Ethics. *IEEE Transactions on Artificial Intelligence*, 4, 799-819. <https://doi.org/10.1109/tai.2022.3194503>.
- Huang, L., Zhang, Y., Zhang, Y., Chen, Y., & Li, Y. (2024). Composite backdoor attacks on large language models. *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (ACL 2024)*.
- Ibrahim, F., Münscher, J., Daseking, M., & Telle, N. (2025). The technology acceptance model and adopter type analysis in the context of artificial intelligence. *Frontiers in Artificial Intelligence*, 7. <https://doi.org/10.3389/frai.2024.1496518>.
- Ikkatai, Y., Hartwig, T., Takanashi, N., & Yokoyama, H. (2022). Octagon Measurement: Public Attitudes toward AI Ethics. *International Journal of Human–Computer Interaction*, 38, 1589 - 1606. <https://doi.org/10.1080/10447318.2021.2009669>.
- ISACA. (2025). *Adversarial machine learning: The next frontier in cybersecurity*. <https://www.isaca.org/resources/news-and-trends/industry-news/2025/adversarial-machine-learning-the-next-frontier-in-cybersecurity>
- Ismatullaev, U., & Kim, S.-H. (2022). Review of the Factors Affecting Acceptance of AI-Infused Systems. *Human Factors*, 66, 126-144. <https://doi.org/10.1177/00187208211064707>
- Jang, Y., Choi, S., & Kim, H. (2022). Development and validation of an instrument to measure undergraduate students' attitudes toward the ethics of artificial intelligence (AT-EAI) and analysis of its difference by gender and experience of AI education. *Education and Information Technologies*, 27, 11635-11667. <https://doi.org/10.1007/s10639-022-11086-5>
- Jedličková, A. (2024). Ethical approaches in designing autonomous and intelligent systems: a comprehensive survey towards responsible development. *AI & SOCIETY*, 40, 2703 - 2716. <https://doi.org/10.1007/s00146-024-02040-9>.
- Jiang, Y., Li, X., Luo, H., Yin, S., & Kaynak, O. (2022). Quo vadis artificial intelligence?. *Discover Artificial Intelligence*, 2. <https://doi.org/10.1007/s44163-022-00022-8>

- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389- 399. <https://doi.org/10.1038/s42256-019-0088-2>.
- Kaiser, H. F. (1974). An index of factorial simplicity. *Psychometrika*, 39(1), 31–36. <https://doi.org/10.1007/bf02291575>
- Kamila, M., & Jasrotia, S. (2023). Ethical issues in the development of artificial intelligence: recognizing the risks. *International Journal of Ethics and Systems*. <https://doi.org/10.1108/ijoes-05-2023-0107>.
- Karg, J., Jäger, J., & Asprion, P. M. (2025). Trusting the machine – the role of gender and personality in shaping the propensity to trust artificial intelligence. *AHFE International*, 5. <https://doi.org/10.54941/ahfe1006734>
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., . . . Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Kaur, D., Uslu, S., Rittichier, K., & Duresi, A. (2022). Trustworthy Artificial Intelligence: A Review. *ACM Computing Surveys (CSUR)*, 55, 1- 38. <https://doi.org/10.1145/3491209>.
- Kazim, E., & Koshiyama, A. (2020). A high-level overview of AI ethics. *Patterns*, 2. <https://doi.org/10.1016/j.patter.2021.100314>.
- Kelly, S., Kaye, S., & Oviedo-Trespalacios, O. (2022). What factors contribute to the acceptance of artificial intelligence? A systematic review. *Telematics Informatics*, 77, 101925. <https://doi.org/10.1016/j.tele.2022.101925>.
- Kerr, A., Barry, M., & Kelleher, J. (2020). Expectations of artificial intelligence and the performativity of ethics: Implications for communication governance. *Big Data & Society*, 7. <https://doi.org/10.1177/2053951720915939>.
- Kim, S., & Lee, Y. (2023). Investigation into the Influence of Socio-Cultural Factors on Attitudes toward Artificial Intelligence. *Education and Information Technologies*, 29(8), 9907–9935. <https://doi.org/10.1007/s10639-023-12172-y>
- Kong, E., Dilinika, J., Nie, X., Gautam, A., & Huang, K. (2025). Measuring Socio-Ethical Engagement with Generative AI: Scale Development via Exploratory Factor Analysis. *Proceedings of the Association for Information Science and Technology*, 62. <https://doi.org/10.1002/ptra2.1325>.

- Kong, S.-C., Korte, S.-M., Burton, S., Keskitalo, P., Turunen, T., Smith, D. P., Wang, L., Lee, J. C., & Beaton, M. (2024). Artificial Intelligence (AI) literacy – an argument for AI literacy in education. *Innovations in Education and Teaching International*, 62, 477 - 483. <https://doi.org/10.1080/14703297.2024.2332744>
- Kortak, İ. (2022). YAPAY ZEKÂ VE HABER İLİŞKİSİNE KULLANICI GÖZÜNDEN BAKMAK: SOSYAL MEDYADA ROBOT HABER SPİKERLERİNE GELEN YORUMLARIN İNCELENMESİ. *Fırat Üniversitesi Sosyal Bilimler Dergisi*. <https://doi.org/10.18069/firatsbed.1060917>.
- Köse, U. (2020). Yapay Zeka Etiği Çerçevesinde Geleceğin İşletmeleri: Dönüşüm ve Paradigma Değişiklikleri. *Mühendislik Bilimleri ve Tasarım Dergisi*, 8(5), 290-305. <https://doi.org/10.21923/jesd.833224>
- Kulaklıoğlu, D. (2024). Ethical AI in Autonomous Systems and Decision-making. *Human Computer Interaction*. <https://doi.org/10.62802/jb2ptt89>
- Kumari, J., Kumar, E., & Kumar, D. (2023). A Structured Analysis to study the Role of Machine Learning and Deep Learning in The Healthcare Sector with Big Data Analytics. *Archives of Computational Methods in Engineering*, 1- 29. <https://doi.org/10.1007/s11831-023-09915-y>.
- Lappin, S. (2025). The Deep Learning Revolution in AI. *European Review*, 33, 416- 425. <https://doi.org/10.1017/s1062798725100379>
- Laupichler, M. C., Aster, A., Schirch, J., & Raupach, T. (2022). Artificial intelligence literacy in higher and adult education: A scoping literature review. *Computers and Education Artificial Intelligence*, 3, 100101. <https://doi.org/10.1016/j.caeai.2022.100101>
- Lepri, B., Oliver, N., & Pentland, A. (2021). Ethical machines: The human-centric use of artificial intelligence. *iScience*, 24. <https://doi.org/10.1016/j.isci.2021.102249>.
- Leschanowsky, A., Rech, S., Popp, B., & Bäckström, T. (2024). Evaluating Privacy, Security, and Trust Perceptions in Conversational AI: A Systematic Review. *Comput. Hum. Behav.*, 159, 108344. <https://doi.org/10.1016/j.chb.2024.108344>
- Leslie, D. (2019). Understanding artificial intelligence ethics and safety. *ArXiv*, abs/1906.05684. <https://doi.org/10.5281/zenodo.3240529>.
- Liang, B., Wang, Y., & Tong, C. (2025). AI Reasoning in Deep Learning Era: From Symbolic AI to Neural-Symbolic AI. *Mathematics*. <https://doi.org/10.3390/math13111707>
- Liehner, G., Hick, A., Biermann, H., Brauner, P., & Ziefle, M. (2023). Perceptions, attitudes and trust toward artificial intelligence — An assessment of the public opinion. *Artificial Intelligence and Social Computing*. <https://doi.org/10.54941/ahfe1003271>.

- Lim, E. (2023). The effects of pre-service early childhood teachers' digital literacy and self-efficacy on their perception of AI education for young children. *Education and Information Technologies*, 28, 12969-12995.
<https://doi.org/10.1007/s10639-023-11724-6>
- Lindner, T., Puck, J., & Puhr, H. (2025). Artificial intelligence in international business: IB theory under augmented decision-making. *Journal of World Business*.
<https://doi.org/10.1016/j.jwb.2025.101676>
- Liu, H., Wang, Y., Fan, W., Liu, X., Li, Y., Jain, S., Liu, Y., Jain, A., & Tang, J. (2022). Trustworthy AI: a computational perspective. *ACM Transactions on Intelligent Systems and Technology*, 14(1), 1–59. <https://doi.org/10.1145/3546872>
- Long, D., & Magerko, B. (2020). What is AI Literacy? Competencies and Design Considerations. *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3313831.3376727>.
- Ma, J., Wang, P., Li, B., Wang, T., Pang, XS ve Wang, D. (2025). ChatGPT Kullanıcı Benimsemesini Keşfetmek: Bir Teknoloji Kabul Modeli Perspektifi. *Uluslararası İnsan-Bilgisayar Etkileşimi Dergisi* , 41 (2), 1431–1445.
<https://doi.org/10.1080/10447318.2024.2314358>.
- Mao, Y., & Shi-Kupfer, K. (2021). Online public discourse on artificial intelligence and ethics in China: context, content, and implications. *Ai & Society*, 38, 373 - 389.
<https://doi.org/10.1007/s00146-021-01309-7>
- Meissner, F., Wilke, A. J., & Puikytė, M. (2025). Can young people be persuaded to adopt secure behavior in cyberspace? an experimental study based on protection motivation theory. *International Crisis and Risk Communication Association Reports*, 13(1).
<https://doi.org/10.69931/001c.144515>
- Mensah, G. (2024). Artificial Intelligence and Ethics: A Comprehensive Reviews of Bias Mitigation, Transparency, and Accountability in AI Systems. *Africa Journal For Regulatory Affairs*. <https://doi.org/10.62839/ajfra/2024.v1.i1.32-45>.
- Mittelstadt, B. (2019). Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*, 1, 501 - 507. <https://doi.org/10.1038/s42256-019-0114-4>.
- Møgelvang, A., & Grassini, S. (2025). Validating the AI attitude scale (AIAS-4) and exploring attitudinal differences in a large sample of Norwegian university students. *Discover Education*, 4. <https://doi.org/10.1007/s44217-025-00657-6>.
- Mohammed, S. (2025). Navigating Algorithmic Accountability and Ethical Governance in Autonomous Data Analytics Systems: Toward Transparent, Bias-Resistant, and Human-

- Centric AI Frameworks for Critical Decision-Making. *International Journal of Innovative Science and Research Technology*. <https://doi.org/10.38124/ijisrt/25oct919>
- Morales-García, W. C., Sairitupa-Sanchez, L. Z., Morales-García, S. B., & Morales-García, M. (2024). Development and validation of a scale for dependence on artificial intelligence in university students. *Frontiers in Education*. <https://doi.org/10.3389/feduc.2024.1323898>
- Morandín-Ahuerma, F. (2022). What is Artificial Intelligence?. *International Journal of Research Publication and Reviews*. <https://doi.org/10.55248/gengpi.2022.31261>
- Mundlamuri, R., Gunnam, G., Mysari, N. K., & Pujuri, J. (2025). The Evolution of AI: From Classical Machine Learning to Modern Large Language Models. *IEEE Access*, 13, 178302-178341. <https://doi.org/10.1109/access.2025.3621344>
- Murikah, W., Nthenge, J., & Musyoka, F. (2024). Bias and Ethics of AI Systems Applied in Auditing - A Systematic Review. *Scientific African*. <https://doi.org/10.1016/j.sciaf.2024.e02281>.
- Naiseh, M., Babiker, A., Al-Shakhsi, S., Cemiloglu, D., Al-Thani, D., Montag, C., & Ali, R. (2025). Attitudes towards AI: the interplay of Self-Efficacy, Well-Being, and Competency. *Journal of Technology in Behavioral Science*, 10(4), 705–718. <https://doi.org/10.1007/s41347-025-00486-2>
- Najafabadi, M., Villanustre, F., Khoshgoftaar, T., Seliya, N., Wald, R., & Muharemagic, E. (2015). Deep learning applications and challenges in big data analytics. *Journal of Big Data*, 2. <https://doi.org/10.1186/s40537-014-0007-7>.
- Nathanson, V. (2015). What Is Ethics?. <https://doi.org/10.4324/9780203771006-1>
- Neri, H., & Cozman, F. (2019). The role of experts in the public perception of risk of artificial intelligence. *AI & SOCIETY*, 35, 663–673. <https://doi.org/10.1007/s00146-019-00924-9>.
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). AI literacy: Definition, teaching, evaluation and ethical issues. *Computers and Education: Artificial Intelligence*, 2, 100041.
- Ng, D., Wu, W., Leung, J., Chiu, T., & Chu, S. (2023). Design and validation of the AI literacy questionnaire: The affective, behavioural, cognitive and ethical approach. *Br. J. Educ. Technol.*, 55, 1082-1104. <https://doi.org/10.1111/bjet.13411>.
- Nguyen, B. N., Trinh, M. Q., Luong, T. H., & Nguyen, H. M. (2025). Learning with algorithms: Artificial Intelligence (AI) usage, trust, and literacy among university students. *Ho Chi Minh City Open University Journal of Science - Social Sciences*, 16(12). <https://doi.org/10.46223/HCMCOUJS.soci.en.16.12.4505.2026>

- NIST. (2024). *Artificial intelligence risk management framework (AI RMF 1.0)*. U.S. Department of Commerce. <https://www.nist.gov/itl/ai-risk-management-framework>
- Nouis, S., Uren, V., & Jariwala, S. (2025). Evaluating accountability, transparency, and bias in AI-assisted healthcare decision- making: a qualitative study of healthcare professionals' perspectives in the UK. *BMC Medical Ethics*, 26. <https://doi.org/10.1186/s12910-025-01243-z>.
- Novotny, M., Weber, W., Kern, C., & Kreuter, F. (2025). Measuring public opinion towards artificial intelligence: development and validation of a general AI attitude short scale. *AI & SOCIETY*. <https://doi.org/10.1007/s00146-025-02478-5>.
- Novozhilova, E., Mays, K. K., Paik, S., & Katz, J. E. (2024). More Capable, Less Benevolent: Trust Perceptions of AI Systems across Societal Contexts. *Mach. Learn. Knowl. Extr.*, 6, 342-366. <https://doi.org/10.3390/make6010017>
- Nunnally, J. C., & Bernstein, I. H. (1994). *Psychometric Theory* (3rd ed.). McGraw-Hill.
- OECD. (2019). *OECD Principles on Artificial Intelligence*.
- OECD. (2023). OECD Digital Education Outlook 2023 : Towards an Effective Digital Education Ecosystem. In *Digital education outlook*. <https://doi.org/10.1787/c74f03de-en>
- Omrani, N., Riviuccio, G., Fiore, U., Schiavone, F., & Agreda, S. (2022). To trust or not to trust? An assessment of trust in AI-based systems: Concerns, ethics and contexts. *Technological Forecasting and Social Change*. <https://doi.org/10.1016/j.techfore.2022.121763>.
- Onan, G., & Dalmış, A. (2025). Yapay zekâ kaygısının gölgesinde: Psikolojik güçlendirmenin işe adanmışlık ve işte kendini yetiştirme üzerindeki etkisi. *Business & Management Studies: An International Journal*. <https://doi.org/10.15295/bmij.v13i1.2508>.
- Osasona, F., Amoo, O., Atadoga, A., Abrahams, T., Farayola, O., & Ayinla, B. (2024). REVIEWING THE ETHICAL IMPLICATIONS OF AI IN DECISION MAKING PROCESSES. *International Journal of Management & Entrepreneurship Research*. <https://doi.org/10.51594/ijmer.v6i2.773>.
- Ouchchy, L., Coin, A., & Dubljević, V. (2020). AI in the headlines: the portrayal of the ethical issues of artificial intelligence in the media. *AI & SOCIETY*, 35, 927 - 936. <https://doi.org/10.1007/s00146-020-00965-5>
- Öztrak, M. (2023). Yapay Zekâ Kaygısının Çalışanların Yenilik Odaklı Davranışlarına Etkisine Yönelik Bir Araştırma. *Optimum Ekonomi ve Yönetim Bilimleri Dergisi*. <https://doi.org/10.17541/optimum.1255576>.

- Özüdoğru, G., & Durak, H. (2025). Conceptualizing pre-service teachers' artificial intelligence readiness and examining its relationship with various variables: The role of artificial intelligence literacy, digital citizenship, artificial intelligence-enhanced innovation and perceived threats from artificial intelligence. *Information Development*, 41, 916 - 932. <https://doi.org/10.1177/02666669251335657>.
- Papernot, N., McDaniel, P., Goodfellow, I., Jha, S., Celik, Z. B., & Swami, A. (2017). Practical black-box attacks against machine learning. *Proceedings of the 2017 ACM on Asia Conference on Computer and Communications Security (ASIA CCS '17)*, 506–519. <https://doi.org/10.1145/3052973.3053009>
- Park, J., Woo, S., & Kim, J. (2024). Attitudes towards artificial intelligence at work: Scale development and validation. *Journal of Occupational and Organizational Psychology*. <https://doi.org/10.1111/joop.12502>.
- Pham, T. (2025). Ethical and legal considerations in healthcare AI: innovation and policy for safe and fair use. *Royal Society Open Science*, 12. <https://doi.org/10.1098/rsos.241873>.
- Qi, X., Zeng, Y., Xie, T., Chen, P.-Y., Jia, R., Mittal, P., & Henderson, P. (2024). Fine-tuning aligned language models compromises safety, even when users do not intend to! *The Twelfth International Conference on Learning Representations (ICLR 2024)*.
- Quince, Z., & Nikolic, S. (2025). Student identification of the social, economic and environmental implications of using Generative Artificial Intelligence (GenAI): identifying student ethical awareness of ChatGPT from a scaffolded multi-stage assessment. *European Journal of Engineering Education*, 51, 43 - 62. <https://doi.org/10.1080/03043797.2025.2482830>
- Radanliev, P. (2025). AI Ethics: Integrating Transparency, Fairness, and Privacy in AI Development. *Applied Artificial Intelligence*, 39. <https://doi.org/10.1080/08839514.2025.2463722>.
- Radanliev, P. (2025). Privacy, ethics, transparency, and accountability in AI systems for wearable devices. *Frontiers in Digital Health*, 7. <https://doi.org/10.3389/fdgth.2025.1431246>.
- Rahman, M. M., Babiker, A., & Ali, R. (2024). Motivation, concerns, and attitudes towards AI: differences by gender, age, and culture. In *Lecture notes in computer science* (pp. 375–391). https://doi.org/10.1007/978-981-96-0573-6_28
- Ramos-Soler, I., López-Sánchez, C., & Torrecillas-Lacave, T. (2018). Online risk perception in young people and its effects on digital behaviour. *Comunicar*, 26(56), 71–79. <https://doi.org/10.3916/c56-2018-07>

- Rao, Y., Yang, R., He, S., Lin, W., & Cai, R. (2026). Gender Differences in Human Trust and Contextual Perception of Robots. *EAI Endorsed Transactions on Pervasive Health and Technology*, 11. <https://doi.org/10.4108/eetpht.11.11058>
- Ravi, D., Wong, C., Deligianni, F., Berthelot, M., Andreu-Perez, J., Lo, B., & Yang, G. (2016). Deep Learning for Health Informatics. *IEEE Journal of Biomedical and Health Informatics*, 21, 4-21. <https://doi.org/10.1109/jbhi.2016.2636665>.
- Refonte Learning. (2025). *Adversarial machine learning: Attacks and defenses*. <https://www.refontelearning.com/blog/adversarial-machine-learning-attacks-and-defenses>
- Riedl, R. (2022). Is trust in artificial intelligence systems related to user personality? Review of empirical evidence and future research directions. *Electronic Markets*, 32, 2021 - 2051. <https://doi.org/10.1007/s12525-022-00594-4>
- Rodríguez, N., Ser, J., Coeckelbergh, M., De Prado, M., Herrera-Viedma, E., & Herrera, F. (2023). Connecting the Dots in Trustworthy Artificial Intelligence: From AI Principles, Ethics, and Key Requirements to Responsible AI Systems and Regulation. *Inf. Fusion*, 99, 101896. <https://doi.org/10.48550/arxiv.2305.02231>.
- Russo, C., Romano, L., Clemente, D., Iacovone, L., Gladwin, T. E., & Panno, A. (2025). Gender differences in artificial intelligence: the role of artificial intelligence anxiety. *Frontiers in Psychology*, 16, 1559457. <https://doi.org/10.3389/fpsyg.2025.1559457>
- Saatci, E. (2025). AI and Ethics: Scale Development for Measuring Ethical Perceptions of Artificial Intelligence Across Sectors and Countries. *International Journal of Economic Behavior and Organization*. <https://doi.org/10.11648/j.ijebo.20251301.14>.
- Sadiku, M., Ashaolu, TJ, Ajayi-Majebi, A., & Musa, S. (2021). Eğitimde Yapay Zeka. Uluslararası Bilimsel Gelişmeler Dergisi. <https://doi.org/10.51542/ijscia.v2i1.2>
- Samsul, S. (2026). Morality/Ethics and Politics: An Inter-Relationship. *NBPA Journal for Arts, Humanities & Social Sciences*. <https://doi.org/10.65842/nbpa.v2.i1.006>
- Satıcı, S., Okur, S., Yilmaz, F., & Grassini, S. (2025). Psychometric properties and Turkish adaptation of the artificial intelligence attitude scale (AIAS-4): evidence for construct validity. *BMC Psychology*, 13. <https://doi.org/10.1186/s40359-025-02505-6>.
- Schepman, A., & Rodway, P. (2020). Initial validation of the general attitudes towards Artificial Intelligence Scale. *Computers in Human Behavior Reports*, 1, 100014 - 100014. <https://doi.org/10.1016/j.chbr.2020.100014>.

- Schiff, DS, Rakova, B., Ayesh, A., Fanti, A., & Lennon, M. (2021). Yapay Zekada Prensiplerden Uygulamalara Geçiş Açığını Açıklamak. *IEEE Teknoloji ve Toplum Dergisi*, 40 , 81-94. <https://doi.org/10.1109/mts.2021.3056286>
- Shakor, M., & Khaleel, M. (2024). Recent Advances in Big Medical Image Data Analysis Through Deep Learning and Cloud Computing. *Electronics*. <https://doi.org/10.3390/electronics13244860>.
- Shick, M., Johnson, N., & Fan, Y. (2023). Artificial intelligence and the end of bounded rationality: a new era in organizational decision making. *Development and Learning in Organizations: An International Journal*. <https://doi.org/10.1108/dlo-02-2023-0048>
- Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI. *Int. J. Hum. Comput. Stud.*, 146, 102551. <https://doi.org/10.1016/j.ijhcs.2020.102551>
- Shrivastava, A. (2024). Artificial Intelligence (AI): Evolution, Methodologies, and Applications. *International Journal for Research in Applied Science and Engineering Technology*. <https://doi.org/10.22214/ijraset.2024.61241>
- Shukla, S. (2024). Principles Governing Ethical Development and Deployment of AI. *International Journal of Engineering, Business and Management*. <https://doi.org/10.22161/ijebm.8.2.5>.
- Stein, J., Messingschlager, T., Gnambs, T., Hutmacher, F., & Appel, M. (2024). Attitudes towards AI: measurement and associations with personality. *Scientific Reports*, 14. <https://doi.org/10.1038/s41598-024-53335-2>.
- Suo, J., Li, M., Guo, J., & Sun, Y. (2024). Engineering Safety and Ethical Challenges in 2045 Artificial Intelligence Singularity. *Sustainability*. <https://doi.org/10.3390/su162310337>.
- Supti, T. I., Yankouskaya, A., Barhmagi, M., Nhlabatsi, A., Khan, K. M., Domingo, J. M., Erbad, A., & Ali, R. (2025). Digital Dependency and Security risk: Investigating the connections between fear of missing out, problematic social media use, and vulnerability perceptions. *Journal of Technology in Behavioral Science*. <https://doi.org/10.1007/s41347-025-00515-0>
- Süer, I., & Müftüoğlu, Z. (2024). Döngüde Etik: Yapay Zekâ Yas,am Döngüsü Yönetişimine Kapsamlı ve Multidisipliner Metodolojik Yaklaşım. *2024 8th International Artificial Intelligence and Data Processing Symposium (IDAP)*, 1-13. <https://doi.org/10.1109/idap64064.2024.10710652>.
- Swapna, A., Vemula, H., Amarnath, B., & Sasidhar, B. (2024). Leveraging the Power of Deep Learning, Machine Learning, Artificial Intelligence, and Big Data in Finance and

- Education. 2024 4th International Conference on Technological Advancements in Computational Sciences (ICTACS), 322-327. <https://doi.org/10.1109/ictacs62700.2024.10841229>.
- Timmons, M. (2020). Normative Ethics. International Encyclopedia of Ethics. <https://doi.org/10.1002/9781444367072.wbiee907>
- Torrijos-Fincias, P., Serrate-González, S., Martín-Lucas, J., & Muñoz-Rodríguez, J. (2021). Perception of Risk in the Use of Technologies and Social Media. Implications for Identity Building during Adolescence. Education Sciences. <https://doi.org/10.3390/educsci11090523>
- Towards Risk-Free Trustworthy Artificial Intelligence: Significance and Requirements. *Int. J. Intell. Syst.*, 2023, 1-41. <https://doi.org/10.1155/2023/4459198>.
- UNESCO. (2021). *Recommendation on the Ethics of Artificial Intelligence*. <https://www.unesco.org/en/articles/recommendation-ethics-artificial-intelligence>
- Upase, P. B. (2026). Reframing Organizational Decision-Making in the Age of Artificial Intelligence: A Conceptual Review of Human–AI Augmentation. *International Journal of Scientific Research in Engineering and Management*. <https://doi.org/10.55041/ijssrem.ibfe071>
- Usher, M., & Barak, M. (2024). Unpacking the role of AI ethics online education for science and engineering students. *International Journal of STEM Education*, 11. <https://doi.org/10.1186/s40594-024-00493-4>.
- Uzun, E. (2025). Investigation of Artificial Intelligence Awareness Among Pre-Service Teachers With Cluster Analysis. *Mersin Üniversitesi Eğitim Fakültesi Dergisi*. <https://doi.org/10.17860/mersinefd.1583282>.
- Valerio, A. (2024). Anticipating the Impact of Artificial Intelligence in Higher Education: Student Awareness and Ethical Concerns in Zamboanga City, Philippines. *Cognizance Journal of Multidisciplinary Studies*. <https://doi.org/10.47760/cognizance.2024.v04i06.024>.
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478.
- Vesnić-Alujević, L., Nascimento, S., & Pólvara, A. (2020). Societal and ethical impacts of artificial intelligence: Critical notes on European policy frameworks. *Telecommunications Policy*, 44, 101961. <https://doi.org/10.1016/j.telpol.2020.101961>.

- Vilchez, A. M., Bera, S., Lewenstein, B., & Hilgartner, S. (2026). Examining media coverage of ethical dimensions of advanced algorithmic technology.. Public understanding of science, 9636625261419755 . <https://doi.org/10.1177/09636625261419755>
- Villiers, D. D., & Villiers, E. D. (2023). What is morality? A historical exploration. Verbum et Ecclesia. <https://doi.org/10.4102/ve.v44i1.2935>
- Vincent, V. U. (2021). Integrating intuition and artificial intelligence in organizational decision-making. *Business Horizons*. <https://doi.org/10.1016/j.bushor.2021.02.008>
- Vo, V., Chen, G., Aquino, Y., Carter, S., N., Q., & Woode, M. (2023). Multi-stakeholder preferences for the use of artificial intelligence in healthcare: A systematic review and thematic analysis.. *Social science & medicine*, 338, 116357 . <https://doi.org/10.1016/j.socscimed.2023.116357>.
- Von Eschenbach, W. J. (2021). Transparency and the Black Box Problem: Why We Do Not Trust AI. *Philosophy & Technology*, 34, 1607 - 1622. <https://doi.org/10.1007/s13347-021-00477-0>
- Walter, Y. (2024). Embracing the future of Artificial Intelligence in the classroom: the relevance of AI literacy, prompt engineering, and critical thinking in modern education. *International Journal of Educational Technology in Higher Education*, 21, 1-29. <https://doi.org/10.1186/s41239-024-00448-3>
- Wan, A., Wallace, E., Shen, S., & Klein, D. (2023). Poisoning language models during instruction tuning. *International Conference on Machine Learning (ICML 2023)*, 35413–35425. PMLR.
- Wang, X., Myers, M. D., & Sundaram, D. (2022). Artificial Intelligence Literacy: Conceptualization and Scale Development. (*İlgili ölçek geliştirme çalışması*).
- Wang, Z., Chai, C., Li, J., & Lee, V. (2025). Assessment of AI ethical reflection: the development and validation of the AI ethical reflection scale (AIERS) for university students. *International Journal of Educational Technology in Higher Education*, 22. <https://doi.org/10.1186/s41239-025-00519-z>.
- Wazir, SM, Iqbal, S., & Shafia, B. (2025). YAPAY ZEKA ÇAĞINDA DİJİTAL OKURYAZARLIK: GELECEĞE HAZIR ÖĞRENCİLER YETİŞTİRMEK İÇİN SONUÇLARI. *Çağdaş Sosyal Bilimler Dergisi*. <https://doi.org/10.63878/cjssr.v3i4.1776>
- Weiner, E., Dankwa-Mullan, I., Nelson, W., & Hassanpour, S. (2024). Ethical challenges and evolving strategies in the integration of artificial intelligence into clinical practice. *PLOS Digital Health*, 4. <https://doi.org/10.1371/journal.pdig.0000810>.

- Winfield, A., & Jirotko, M. (2018). Ethical governance is essential to building trust in robotics and artificial intelligence systems. *Philosophical transactions. Series A, Mathematical, physical, and engineering sciences*, 376. <https://doi.org/10.1098/rsta.2018.0085>.
- WitnessAI. (2025). *AI red teaming: A comprehensive guide for security teams*. <https://www.witness.ai/blog/ai-red-teaming>
- Wolf, G., & Knoepffler, N. (2022). Philosophische Begründung einer Medizinethik. *Die Nephrologie*, 18, 65-71. <https://doi.org/10.1007/s11560-022-00610-w>
- Worthington, R. L., & Whittaker, T. A. (2006). Scale development research. *The Counseling Psychologist*, 34(6), 806–838.
- Xie, N., & Zhang, W. (2024). Empowering big data analytics through machine learning: Applications, challenges, and future directions. *Applied and Computational Engineering*. <https://doi.org/10.54254/2755-2721/93/20240917>.
- Yang, Y. (2024). Influences of Digital Literacy and Moral Sensitivity on Artificial Intelligence Ethics Awareness Among Nursing Students. *Healthcare*, 12. <https://doi.org/10.3390/healthcare12212172>
- Yanli, X., & Danni, L. (2021). Prospect of Vocational Education under the Background of Digital Age: Analysis of European Union’s “Digital Education Action Plan (2021-2027).” *2021 International Conference on Internet, Education and Information Technology (IEIT)*, 164–167. <https://doi.org/10.1109/ieit53597.2021.00042>
- Yılmaz, N., & Korkmaz, Ö. (2025). Eğitimde Yapay Zekâ Kullanımına İlişkin Geliştirilen ve Uyarlanan Ölçeklerin İncelenmesi. *Abant İzzet Baysal Üniversitesi Eğitim Fakültesi Dergisi*. <https://doi.org/10.17240/aibuefd.2025..-1563316>
- Yılmaz, G. (2020). YAPAY ZEKÂNIN YARGI SİSTEMLERİNDE KULLANILMASINA İLİŞKİN AVRUPA ETİK ŞARTI. *Marmara Üniversitesi Avrupa Topluluğu Enstitüsü Avrupa Araştırmaları Dergisi*. <https://doi.org/10.29228/mjes.8>.
- Yusuf, A., Pervin, N., & Román-González, M. (2024). Üretken Yapay Zeka ve Yüksek Öğretimin Geleceği: Akademik Bütünlüğe Bir Tehdit mi Yoksa Bir Reform mu? Çok Kültürlü Perspektiflerden Kanıtlar. *Uluslararası Yüksek Öğretimde Eğitim Teknolojisi Dergisi*, 21 , 1-29. <https://doi.org/10.1186/s41239-024-00453-6>
- Zaidi, N., Singh, B., & Yadav, S. (2018). The evolution of machine learning algorithms: A comprehensive historical review. *International Journal of Applied Research*. <https://doi.org/10.22271/allresearch.2018.v4.i9a.11451>

- Zaripova, R., Kosulin, V., Shkinderov, M., & Rakhmatullin, I. (2023). Unlocking the potential of artificial intelligence for big data analytics. *E3S Web of Conferences*. <https://doi.org/10.1051/e3sconf/202346004011>.
- Zhang, B., & Dafoe, A. (2019). Artificial intelligence: American attitudes and trends. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3312874>
- Zhang, B., Anderljung, M., Kahn, L., Dreksler, N., Horowitz, M., & Dafoe, A. (2021). Ethics and Governance of Artificial Intelligence: Evidence from a Survey of Machine Learning Researchers. *ArXiv, abs/2105.02117*. <https://doi.org/10.1613/jair.1.12895>.
- Zhang, J., & Zhang, Z. (2023). Ethics and governance of trustworthy medical artificial intelligence. *BMC Medical Informatics and Decision Making*, 23. <https://doi.org/10.1186/s12911-023-02103-9>.
- Zhang, X., Gu, Y., Yin, J., Zhang, Y., Jin, C., Wang, W., Li, A., Wang, Y., Su, L., Xu, H., Ge, X., Ye, C., Tang, L., Shen, B., Fang, J., Wan, D., & Feng, R. (2022). Development, Reliability, and Structural Validity of the Scale for Knowledge, Attitude, and Practice in Ethics Implementation Among AI Researchers: Cross-Sectional Study. *JMIR Formative Research*, 7. <https://doi.org/10.2196/42202>.
- Zinyemba, L., Ndoro, C. J., & Zinyemba, M. A. (2024). Public perception of AI: Awareness and trust in Artificial intelligence (AI). *Jocap - Journal of Contemporary African Philosophy*, 5(1). <https://doi.org/10.65846/uysv2685>
- Zubair, N., Hashmat, S., & Aimen, U. (2025). Artificial Intelligence and the Generational Divide: A Study on Trust and Acceptance. *Annual Methodological Archive Research Review*, 3(6), 19–44. <https://doi.org/10.63075/s0a9ec84>

EKLER

Ek A. Etik Kurul Onayı



T.C.
ERZİNCAN BİNALİ YILDIRIM ÜNİVERSİTESİ REKTÖRLÜĞÜ
Eğitim Bilimleri Etik Kurulu



Sayı : E-88012460-050.04-481127
Konu : Etik Kurul Kararı (Soner DEMİR)

12.09.2025

DAĞITIM YERLERİNE

Üniversitemiz İnsan Araştırmaları Eğitim Bilimleri Etik Kurulunun **11 Eylül 2025** tarihli ve **10** sayılı oturumunda alınan 10/18 sayılı kararı yazımız ekinde gönderilmiştir.
Bilgilerini rica ederim.

Prof.Dr. Güldem DÖNEL AKGÜL
Eğitim Bilimleri Etik Kurulu Başkanı

Ek:Karar 18 (1 Sayfa)

Dağıtım:
Gereği:
Soner DEMİR

Bilgi:
Doç.Dr. Recep ÖZ

Bu belge, güvenli elektronik imza ile imzalanmıştır.

Belge Doğrulama Kodu :BS9K27J74U

Belge Takip Adresi : <https://www.turkiye.gov.tr/eby-ebys>

Adres:Erzincan Binali Yıldırım Üniversitesi Rektörlüğü Yalnızbağ yerleşkesi Erzincan Sivas
karayolu 12. km 24002 Erzincan
Telefon:444 8 024 - (0446) 226 66 66 Faks:(0446) 226 66 65
e-Posta:rektörluk@erzincan.edu.tr Web:<https://ebyu.edu.tr/tr/>
Kep Adresi:erzincanmv@hs02.kep.tr

Bilgi için: Şehriban GAZİ
Unvanı: Birim Evrak Sorumlusu
Tel No: (0446) 226 6666 - 10058



Bu belge, güvenli elektronik imza ile imzalanmıştır.



T.C
ERZİNCAN BİNALİ YILDIRIM ÜNİVERSİTESİ
İNSAN ARAŞTIRMALARI EĞİTİM BİLİMLERİ
ETİK KURULU KARARI

Etik Kurul Toplantı Tarihi	11/09/2025
Protokol No	10/18
Araştırma Başlığı	Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği: Geçerlik ve Güvenirlilik Çalışması
Araştırma Türü	Karma yöntem araştırma
Araştırmacılar	Soner DEMİR (Sorumlu Araştırmacı) Doç. Dr. Recep ÖZ (Danışman)
Karar	Başvuru dosyanıza ait araştırmanız etik açıdan uygun bulunmuştur.
Açıklama: 1. Etik Kurul Onayı, uygulama ve/veya veri toplama için araştırmacının ilgili kurum veya kuruluşlardan izin alma sorumluluğunu ortadan kaldırmaz. 2. Kurul üyelerine ait araştırma önerileri görüşülürken, ilgili yönerge gereğince, öneri sahibi üye görüşmelere katılmamış ve oy kullanmamıştır.	

e-imzalıdır

Prof. Dr. Güldem DÖNEL AKGÜL
İnsan Araştırmaları Eğitim Bilimleri
Etik Kurul Başkanı



Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği

Sayın Katılımcı;

Yapay zekâ (YZ) teknolojilerinin hızlı gelişimi, bireylerin güvenlik ve ahlaki değerlere ilişkin algılarını önemli ölçüde etkilemektedir. Günümüzde YZ'nin sağlık, eğitim, ulaşım ve güvenlik gibi farklı alanlarda yoğun biçimde kullanılması, bu teknolojilerin etik ve güvenlik boyutlarının da araştırılmasını zorunlu hale getirmiştir (Floridi & Cowls, 2019). YZ uygulamalarının kişisel verilerin gizliliği, karar süreçlerinde tarafsızlık, adalet, hesap verebilirlik ve toplumsal etkiler gibi konular üzerinde önemli tartışmalar yarattığı bilinmektedir (Jobin, Ienca & Vayena, 2019). Bu bağlamda, bireylerin YZ'ye yönelik güvenlik ve ahlak algılarını ölçmeye imkân veren geçerli ve güvenilir bir ölçeğin geliştirilmesi literatüre önemli katkı sağlayacaktır.

Bu

araştırmanın amacı, bireylerin YZ güvenlik ve ahlak algılarını ölçmek üzere bir ölçek geliştirmektir. Özellikle eğitim, bilişim, hukuk ve sosyal bilimler alanlarında yapılacak araştırmalara katkı sağlaması beklenmektedir.

Ölçek, beş alt boyuttan oluşmaktadır: güvenlik farkındalığı, ahlaki değerler, sorumluluk ve hesap verebilirlik, toplumsal etki ve olumlu beklentiler.

Anket toplam 60 sorudan oluşmaktadır. Cevaplamak tamamen gönüllülük esasına dayanmaktadır.

Katılımınız için şimdiden teşekkür ederiz.

Sonraki

Formu temizle



Yapay Zekâ Güvenlik ve Ahlak Algısı Ölçeği

* Zorunlu soruyu belirtir

Katılımcı Bilgileri Maddeleri

Cinsiyetiniz *

- ☐ Kadın
- ☐ Erkek

Yaşınız *

- ☐ 18-24 Yaş
- ☐ 25-34 Yaş
- ☐ 35 Yaş ve Üstü

Ek C. YAPAY ZEKÂ GÜVENLİK VE AHLAK ALGISI ÖLÇEĞİ

No	Sorular	Kesinlikle Katılmıyorum	Katılmıyorum	Kararsızım	Katılıyorum	Kesinlikle katılıyorum
A. Güvenlik Farkındalığı						
1	YZ teknolojileri kişisel bilgilerimin çalınmasına yol açabilir.	1	2	3	4	5
2	YZ güvenlik açıkları nedeniyle toplumsal riskler oluşturabilir.	1	2	3	4	5
3	YZ kullanan kurumlara güven duymakta zorlanıyorum.	1	2	3	4	5
4	YZ verilerimi benim iznim olmadan paylaşabilir.	1	2	3	4	5
5	YZ uygulamaları insan güvenliğini tehdit edebilir.	1	2	3	4	5
6	YZ güvenlik açısından tam olarak güvenemem.	1	2	3	4	5
7	YZ askeri alanlarda kullanıldığında küresel güvenliği tehlikeye atabilir.	1	2	3	4	5
B. Ahlaki Değerler						
8	YZ etik ilkelere göre geliştirilmelidir.	1	2	3	4	5
9	YZ tarafsız kararlar alabilmelidir.	1	2	3	4	5
10	YZ etik denetim mekanizmalarına tabi olmalıdır.	1	2	3	4	5
11	YZ toplumun ahlaki değerleriyle uyumlu kullanılmalıdır.	1	2	3	4	5
12	YZ yanlış kullanıldığında etik sorunlara yol açabilir.	1	2	3	4	5
13	YZ insan vicdanının yerine geçemez.	1	2	3	4	5
14	YZ kararlarında eşitlik ilkesi gözetilmelidir.	1	2	3	4	5
15	YZ ya sınırsız özgürlük tanımak doğru değildir.	1	2	3	4	5
16	YZ toplumun değerlerini dönüştürebilir.	1	2	3	4	5
17	YZ'nın sınırlarını belirlemek ahlaki bir zorunluluktur.	1	2	3	4	5
C. Sorumluluk ve Hesap Verilebilirlik						
18	YZ sistemlerinin hatalarından geliştiriciler sorumludur.	1	2	3	4	5
19	YZ kararları şeffaf biçimde açıklanmalıdır.	1	2	3	4	5
20	YZ denetimsiz bırakıldığında hesap verebilirlik azalır.	1	2	3	4	5
21	YZ kaynaklı hatalarda kullanıcılar suçlanmamalıdır.	1	2	3	4	5
22	YZ kullanımında hukuki sorumluluk net tanımlanmalıdır.	1	2	3	4	5
23	YZ yanlış karar verdiğinde, sorumluluk kime ait olduğu açık olmalıdır.	1	2	3	4	5
24	YZ uygulamalarının kimler tarafından geliştirildiği şeffaf olmalıdır.	1	2	3	4	5
25	YZ uygulamalarında hesap verebilirlik yetersizdir.	1	2	3	4	5
D. Toplumsal Etki						
26	YZ işsizliği artırabilir.	1	2	3	4	5
27	YZ toplumda adaletsizliklere yol açabilir.	1	2	3	4	5
28	YZ bireylerin özgür iradesini sınırlayabilir.	1	2	3	4	5
29	YZ insan ilişkilerinde yabancılaşmaya sebep olabilir.	1	2	3	4	5
30	YZ toplumsal güveni zedeleyebilir.	1	2	3	4	5
31	YZ bireylerin davranışlarını yönlendirebilir.	1	2	3	4	5
32	YZ toplumda eşitsizlikleri derinleştirebilir.	1	2	3	4	5

33	YZ çocukların gelişimini olumsuz etkileyebilir.	1	2	3	4	5
E. Olumlu Beklentiler						
34	YZ yaşam kalitesini artırabilir.	1	2	3	4	5
35	YZ etik çerçevede kullanıldığında topluma katkı sağlar.	1	2	3	4	5
36	YZ eğitim süreçlerini geliştirebilir.	1	2	3	4	5
37	YZ insanların iş yükünü azaltabilir.	1	2	3	4	5
38	YZ doğru yönlendirildiğinde güvenli bir teknoloji olabilir.	1	2	3	4	5
39	YZ insanlara yeni iş imkânları yaratabilir.	1	2	3	4	5
40	YZ karar alma süreçlerini hızlandırabilir.	1	2	3	4	5