

DOI: <http://doi.org/10.5281/zenodo.18859924>

MARCH 2026

Yapay Zekânın Sosyal Kabulü: CanikFest Örnek Olayı**Social Acceptance of Artificial Intelligence: The CanikFest Case Study****Berker KILIÇ**

Canik Mesleki ve Teknik Anadolu Lisesi

berker.kilic@gmail.com, ORCID: <https://orcid.org/0000-0001-8751-8192>**İrfan GÜMÜŞ**

Samsun Bilim ve Sanat Merkezi

mahmud.irfan@gmail.com, ORCID: <https://orcid.org/0000-0001-8757-8726>**Esra KILIÇ**

Canik Mesleki ve Teknik Anadolu Lisesi

esracaya_5281@hotmail.com, ORCID: <https://orcid.org/0000-0003-3009-9462>**Nagihan GÜMÜŞ**

Canik Özdemir Bayraktar Keşif Kampüsü

ylmznagihan55@gmail.com, ORCID: <https://orcid.org/0009-0004-9041-3278>**Elif DEMİRCAN**

Canik Özdemir Bayraktar Keşif Kampüsü

demircann55elif@gmail.com, ORCID: <https://orcid.org/0009-0003-7554-1493>**Aysima ŞENTÜRK**

Canik Özdemir Bayraktar Keşif Kampüsü

aysima.senturk@kesifkampusu.org, ORCID: <https://orcid.org/0009-0009-0831-6257>**Zehra KAYA**

Canik Özdemir Bayraktar Keşif Kampüsü

zhr.ky.1903@gmail.com, ORCID: <https://orcid.org/0009-0009-7294-0281>**Özge ÜNAL**

Canik Özdemir Bayraktar Keşif Kampüsü

ozgeunal9955@gmail.com, ORCID: <https://orcid.org/0009-0007-9765-0200>**Özge TURNA**

Canik Özdemir Bayraktar Keşif Kampüsü

ozgeturna@hotmail.com, ORCID: <https://orcid.org/0009-0000-0657-2825>**Cansu AYDIN**

Canik Özdemir Bayraktar Keşif Kampüsü

muzisyen55@icloud.com, ORCID: <https://orcid.org/0009-0006-9056-225>

Gökhan BÜYÜK

Canik Özdemir Bayraktar Keşif Kampüsü

gokhanbk55@gmail.com, ORCID: <https://orcid.org/0009-0000-1424-2551>

Fatma Beyza YILMAZ

Canik Özdemir Bayraktar Keşif Kampüsü

beyzakeless98@gmail.com, ORCID: <https://orcid.org/0009-0005-1045-254X>

Havvanur ORUÇ ÇOBAN

Canik Özdemir Bayraktar Keşif Kampüsü

oruchavvanur0@gmail.com, ORCID: <https://orcid.org/0009-0008-8118-3561>

Özet

Yapay zekâ (YZ) sistemlerinin toplumsal kabulü yalnızca teknolojik yeterliliğe değil; güven, şeffaflık ve etik yönetişim gibi sosyoteknik koşullara da bağlıdır. Çok uluslu kamuoyu araştırmaları, YZ'ye yönelik tutumların “fayda” ile “risk” algısı arasında dalgalandığını, güven düzeyinin ise çoğu ülkede temkinli bir çizgide seyrettiğini göstermektedir (KPMG, 2023; Poushter, Fagan, & Corichi, 2023; Vogels, 2023). Bu çalışma, YZ'ye yönelik güvenin artırılmasında *Bilgi Yolu* stratejisinin (YZ okuryazarlığını yükseltmeye dönük müdahaleler) rolünü ve bu rolün iki kritik mekanizma üzerinden işleyişini tartışır: (i) açıklanabilir yapay zekâ (XAI) ile belirsizliğin azaltılması (Berger ve Calabrese, 1975; Gunning vd., 2019; Lundberg & Lee, 2017; Ribeiro, Singh, & Guestrin, 2016) ve (ii) YZ mühendislerinin yetkinliğine duyulan bilişsel güvenin oluşumu (Mayer, Davis, & Schoorman, 1995). Çalışma, CanikFest Yapay Zekâ temalı etkinliğe katılan katılımcılardan toplanan anket verileri (N=713) üzerinden, uygulamalı atölye deneyiminin YZ okuryazarlığı boyutları (farkındalık, kullanım, değerlendirme ve etik) ile güven/kabul değişkenleri arasındaki ilişkilerini hipotezler aracılığıyla analiz etmektedir.

Anahtar Kelimeler: Yapay Zekâ, Güven, Açıklanabilirlik, Sosyal Kabul.

Abstract

The social acceptance of artificial intelligence (AI) systems depends not only on technological adequacy but also on socio-technical conditions such as trust, transparency, and ethical governance. Multinational public opinion surveys show that attitudes toward AI fluctuate between perceptions of “benefit” and “risk,” while confidence levels remain cautious in most countries (KPMG, 2023; Poushter, Fagan, & Corichi, 2023; Vogels, 2023). This study discusses the role of the Information Path strategy (interventions aimed at increasing AI literacy) in increasing trust in AI and how this role operates through two critical mechanisms: (i) reducing uncertainty through explainable artificial intelligence (XAI) (Berger and Calabrese, 1975; Gunning vd., 2019; Lundberg & Lee, 2017; Ribeiro, Singh, & Guestrin, 2016) and (ii) the formation of cognitive trust in the competence of AI engineers (Mayer, Davis, & Schoorman, 1995). The study analyzes the relationships between the dimensions of AI literacy (awareness, usage, evaluation, and ethics) and the variables of trust/acceptance through hypotheses based on survey data (N=713) collected from participants attending the CanikFest Artificial Intelligence themed event.

Keywords: Artificial Intelligence, Trust, Explainability, Social Acceptance.

1. Giriş

Yapay zekâ teknolojileri finans, sağlık, eğitim ve kamu hizmetleri gibi alanlarda verimlilik ve ölçeklenebilirlik vaat ederken; mahremiyet, önyargı, hesap verebilirlik, güvenlik ve istihdam etkileri gibi risk başlıklarını da aynı anda görünür kılmaktadır. Bu ikili yapı, YZ'nin toplumsal kabulünün salt performans tartışması olmaktan ziyade, güven temelli bir *sosyoteknik benimseme* problemi olarak ele alınmasını gerektirir.

Sosyoteknik yaklaşımda kabul (acceptance) yalnızca bireyin teknolojiye yönelik tutumuyla sınırlı değildir; kullanıcının teknolojiyle temas ettiği kurumların güvenilirliği, açıklama/hesap verebilirlik mekanizmaları, algılanan fayda ve risklerin dengesi, sistemin adaleti ve öngörülebilirliği gibi çoklu bileşenlerin bir araya gelmesiyle şekillenir. Dolayısıyla YZ'nin benimsenmesi, hem tasarım düzeyinde (model seçimi, veri, açıklama yöntemleri) hem de yönetim düzeyinde (denetim, şeffaflık, sorumluluk) eşzamanlı müdahale gerektiren bir alandır.

Küresel ölçekte yürütülen anketler, YZ'nin artan kullanımına ilişkin duyguların çoğu zaman “eşit derecede endişe ve heyecan” ekseninde kümелendiğini ve YZ'ye güvenin kırılgan olabildiğini göstermektedir (KPMG, 2023; Poushter vd., 2023). Bununla birlikte, YZ hakkında bilgi ve kavramsal farkındalık düzeyinin sınırlı oluşu, risk algısını pekiştirerek benimsemeyi yavaşlatabilmektedir (Vogels, 2023). Bu bağlamda literatürde öne çıkan yaklaşımlardan biri, güveni artırmaya dönük stratejileri “kurumsal yol” (denetim ve yönetim), “etik yol” (normlar ve değerler), “teknik yol” (doğruluk/güvenilirlik) ve “bilgi yolu” (okuryazarlık ve açıklama) olarak ayrıştırarak tartışmaktadır.

Bu makalenin amacı, Bilgi Yolu stratejisini CanikFest örnek olayı üzerinden kavramsallaştırmak; YZ okuryazarlığı, açıklanabilirlik (XAI) ve “mühendis yetkinliğine güven” arasında kurulabilecek nedensel bağlantıları hipotezler aracılığıyla sistematik hâle getirmektir. Çalışma ayrıca, CanikFest etkinliği sonrası uygulanan anketten elde edilen bulguların, YZ okuryazarlığı boyutları ile güven/kabul göstergeleri arasındaki ilişkileri nasıl görünür kıldığına odaklanır.

660

2. Literatür Taraması

2.1. YZ'ye Yönelik Güven: Sistem ve Aktör Ayrımı

YZ'ye güven, bir yandan “sisteme güven” (çıktının doğruluğu, tutarlılığı ve tahmin edilebilirliği), diğer yandan “aktöre güven” (sistemi geliştiren/uygulayan kişi ve kurumların niyeti, yetkinliği ve dürüstlüğü) üzerinden şekillenir. Güven literatüründe aktör boyutu; yetkinlik, yardımseverlik ve dürüstlük (integrity) alt boyutlarıyla ele alınmıştır (Mayer vd., 1995). YZ bağlamında bu ayrım özellikle önemlidir; çünkü model performansı yüksek olsa bile, verinin kökeni, modelin eğitim süreci ve uygulama amaçları hakkında şeffaflık eksikliği güveni zayıflatabilir.

2.2. Açıklanabilirlik (XAI) ve Kullanıcı Güveni

Açıklanabilir Yapay Zekâ (XAI), model çıktılarının nedenlerini kullanıcıya anlaşılır biçimde sunmayı amaçlayan yaklaşım ve yöntemler bütünüdür (Gunning vd., 2019). Özellikle LIME ve SHAP gibi yöntemler, yerel açıklamalar veya katkı analizleri yoluyla “neden bu çıktı?” sorusunu yanıtlamaya çalışır (Lundberg & Lee, 2017; Ribeiro vd., 2016). Bununla birlikte XAI'nin güven üzerindeki etkisi yalnızca “açıklama var/yok” düzeyinde değildir; açıklamanın doğruluğu, tutarlılığı, kullanıcıya uygunluğu ve karar bağlamına uyumu belirleyici olabilir.

2.3. Belirsizliği Azaltma Teorisi (URT) ve Bilgi Arayışı

Belirsizliği Azaltma Teorisi, belirsizlik koşullarında bireylerin bilgi arayışıyla durumu anlamlandırmaya çalıştığını ve belirsizlik azaldıkça olumlu tutumların güçlenebileceğini öne sürer (Berger ve Calabrese 1975). YZ sistemleri, karar mantığının opaklığı nedeniyle belirsizlik üretmeye yatkındır. Bu nedenle XAI ve YZ okuryazarlığı, belirsizliği azaltan tamamlayıcı iki mekanizma olarak değerlendirilebilir.

2.4. Makro Riskler: İstihdam ve Enformasyon Ekosistemi

YZ teknolojilerinin kabulü, yalnızca bireysel kullanıcı deneyimine bağlı değildir; otomasyonun işgücü piyasası üzerindeki etkileri ve dezenformasyon/derin sahtecilik gibi toplumsal riskler de genel güven iklimini belirler (Acemoglu & Restrepo, 2020; Autor, 2015; Chesney & Citron, 2019; Eloundou, Manning, Mishkin, & Rock, 2023; Frey & Osborne, 2017; Vaccari & Chadwick, 2020). Bu risk bağlamı, “bilgi yolu” müdahalelerinin neden önemli olduğunu güçlendiren bir arka plan sunar.

3. Kuramsal Çerçeve

3.1. Güven, Kabul ve Sosyoteknik Yaklaşım

Güven, hem teknik sisteme (YZ sisteminin doğruluğu, tutarlılığı, güvenilirliği) hem de sistemi tasarlayan/işleten aktörlere (kurum, geliştirici ekip, mühendisler) yönelen iki katmanlı bir beklenti yapısıdır. Bu nedenle YZ’nin kabulü, yalnızca “sistem ne kadar iyi çalışıyor?” sorusuna değil, aynı zamanda “kim, hangi değerlerle ve hangi denetim mekanizmalarıyla çalıştırıyor?” sorusuna da verilen yanıta bağlıdır.

3.2. YZ Okuryazarlığı (Bilgi Yolu)

YZ okuryazarlığı; bireylerin YZ’yi tanıma, araçları kullanma, çıktıları eleştirel değerlendirme ve etik boyutları tartışabilme kapasitesini ifade eder. Bilgi Yolu yaklaşımı, bu kapasiteyi artıran müdahalelerin, belirsizlik ve korku kaynaklı dirençleri azaltarak kabulü güçlendirebileceğini varsayar.

3.3. Açıklanabilir Yapay Zekâ (XAI) ve Belirsizliği Azaltma Teorisi

Karmaşık modellerin “kara kutu” niteliği, kullanıcıların karar gerekçelerini anlamasını zorlaştırarak güveni zedeleyebilir. XAI, model çıktılarının gerekçelendirilmesi ve kullanıcıya anlaşılır açıklamalar sunulması yoluyla bu sorunu hedefler (Gunning vd., 2019; Lundberg & Lee, 2017; Ribeiro vd., 2016).

Belirsizliği Azaltma Teorisi (URT), ilk etkileşim ve belirsizlik koşullarında bireylerin bilgi arayışıyla belirsizliği azaltmaya çalıştığını; belirsizliğin azalmasının olumlu tutumları ve işbirliğini desteklediğini öne sürer (Berger & Calabrese 1975). Bu çerçevede XAI, YZ’ye ilişkin belirsizliği azaltarak güven-kabul hattını güçlendiren bir mekanizma olarak yorumlanabilir.

3.4. İnsan Faktörü: YZ Mühendislerinin Yetkinliğine Güven

Güven literatüründe yetkinlik (ability), yardımseverlik (benevolence) ve dürüstlük (integrity) boyutları güvenin temel belirleyicileri arasında ele alınır (Mayer vd., 1995). YZ sistemlerinin önyargı ve hata riskleri, çoğu zaman veri seçimi, modelleme tercihleri ve değer tanımları gibi insan

kaynaklı aşamalarla ilişkili olduğundan, kullanıcıların “sistemi üreten ekip” hakkındaki algısı kabul üzerinde bağımsız bir etkiye sahip olabilir.

4. Hipotezler

CanikFest gibi uygulamalı etkinlikler, Bilgi Yolu yaklaşımının somutlaştığı ortamlardır: katılımcılar YZ araçlarını deneyimler, sınırlılıkları görür ve karar süreçlerini tartışır.

H1: (Bilgi Yolu etkisi) Uygulamalı, etkileşimli YZ atölyeleri katılımcıların YZ okuryazarlığı düzeyini ve YZ mühendislerinin yetkinliğine olan güvenini artırır.

H2: (Aracılık) YZ mühendislerinin yetkinliğine olan güven, YZ uygulamalarına yönelik olumlu tutumu artırır ve bu tutum, YZ araçlarını gelecekte kullanma niyeti üzerinde aracılık rolü oynar.

H3: (XAI etkisi) Algılanan açıklanabilirlik (XAI), yetkinliğe güveni pozitif yönde etkiler; bu etki, yardıseverlik ve dürüstlük boyutlarına kıyasla daha güçlüdür.

5. Yöntem

5.1. Örnek Olay Tasarımı: CanikFest

Çalışma, CanikFest kapsamındaki atölye deneyimini bir örnek olay olarak ele alır ve etkinlik sonrası uygulanan değerlendirme formu (anket) üzerinden hipotezlerin test edilmesini amaçlar.

5.2. Örneklem ve Veri Toplama

Veriler, CanikFest atölyelerine katılan katılımcılardan etkinlik sonrasında toplanmıştır (N=713). Ölçek maddeleri 7’li Likert tipinde puanlanmıştır.

5.3 Ölçüm Araçları ve Değişkenler

YZ okuryazarlığı; farkındalık, kullanım, değerlendirme ve etik alt boyutları ile ele alınır. Güven değişkenleri; mühendis yetkinliğine güven, algılanan açıklanabilirlik (XAI) ve genel tutum/kullanım niyeti olarak modellenmiştir.

6. Bulgular

Bu bölümde, CanikFest katılımcılarından toplanan (N=713) veri setinin betimsel istatistikleri ve güvenilirlik göstergeleri raporlanmıştır.

6.1. Tanımlayıcı İstatistikler

Alt boyutlar ve toplam puana ilişkin ortalama, standart sapma ve aralık değerleri Tablo 1’te sunulmuştur.

Tablo 1: Alt boyutlar ve toplam puana ilişkin tanımlayıcı istatistikler

Ölçüm	Ortalama	SS	Min	Maks
Farkındalık (3–21)	13,03	3,57	3	21
Kullanım (3–21)	14,20	3,96	3	21
Değerlendirme (3–21)	15,73	4,42	3	21
Etik (3–21)	14,22	4,02	3	21
Toplam (12–84)	57,18	13,23	12	84
Madde Ortalaması (1–7)	4,77	1,10	1,00	7

6.2. Güvenirlilik

Ölçeğin toplam iç tutarlılık katsayısı Cronbach $\alpha = 0,826$ olarak raporlanmıştır. Alt boyutlara ilişkin göstergeler Tablo 2’de verilmiştir.

Tablo 2: Güvenirlilik göstergeleri (α , MIC ve McDonald’s ω)

Ölçek	Madde sayısı	Cronbach α	MIC	McDonald’s ω (yakl.)
Farkındalık (3–21)	3	0,311	0,122	0,599
Kullanım (3–21)	3	0,457	0,234	0,730
Değerlendirme (3–21)	3	0,816	0,597	0,891
Etik (3–21)	3	0,332	0,171	0,669
Toplam (M1–M12)	12	0,826	0,300	0,874

6.3. Hipotezlerin Değerlendirilmesi

Bu alt bölümde hipotezler, anketten elde edilen betimsel bulgular ve kuramsal beklentiler birlikte değerlendirilmiştir. Bu çalışma kapsamında hipotezler SEM ile tam bir yapısal model olarak raporlanmadığından, aşağıdaki değerlendirme; (i) ölçek alt boyutlarındaki seviye farklılıkları, (ii) iç tutarlılık bulguları ve (iii) CanikFest bağlamında gözlenen genel eğilimler üzerinden *yorumlayıcı* bir çerçevede sunulmaktadır.

6.3.1. H1: Bilgi Yolu Etkisi (Uygulamalı Eğitimin Okuryazarlık ve Güvene Katkısı)

663

H1, CanikFest gibi uygulamalı atölyelerin katılımcıların YZ okuryazarlığını ve YZ mühendislerinin yetkinliğine duyulan güveni artıracaklarını öngörmektedir. Bulgular incelendiğinde, okuryazarlık alt boyutlarının genel olarak orta-üst düzeyde seyrettiği görülmektedir (Tablo 1). Özellikle *değerlendirme* alt boyutunun ortalamasının ($M=15,73$; 3–21 aralığı) diğer alt boyutlara kıyasla daha yüksek olması, katılımcıların atölye deneyimi sonrasında YZ araçlarının kapasite ve sınırlılıklarını ayırt etme becerisinin güçlenmiş olabileceğine işaret etmektedir.

Bununla birlikte *farkındalık* ve *kullanım* boyutlarındaki ortalama değerlerin de görece yüksek olması ($M=13,03$ ve $M=14,20$), CanikFest’in “bilgi yolu” müdahalesi olarak işlev gördüğünü destekler. Bu bulgular, etkinlikte araçların bizzat deneyimlenmesi ve uzmanlarla etkileşim kurulması sayesinde belirsizliğin azalması (URT) ve bilişsel açıklık kazanımı ile tutarlı biçimde yorumlanabilir.

H1 kapsamında ikinci vurgu “mühendis yetkinliğine güven” olsa da, mevcut tablolar bu değişkene ilişkin ayrı bir ölçek raporlamamaktadır. Buna rağmen, okuryazarlık boyutlarının özellikle değerlendirme ekseninde yükselmesi, katılımcıların teknik süreçlere ilişkin daha net bir kavrayış geliştirmiş olabileceğini ve bunun da dolaylı biçimde “yetkinlik” algısını besleyebileceğini düşündürmektedir. Bu nedenle H1, mevcut betimsel bulgularla *kısmen desteklenmektedir*; ancak doğrudan güven maddeleriyle ayrıca test edilmesi gerekir.

6.3.2. H2: Aracılık (Yetkinliğe Güven → Tutum → Kullanım Niyeti)

H2, mühendis yetkinliğine duyulan güvenin YZ'ye yönelik olumlu tutumu artıracığını ve bu tutumun da gelecekte kullanım niyeti üzerinde aracılık rolü oynayacağını öngörmektedir. Bu hipotez, kavramsal olarak teknoloji kabul modelleriyle uyumlu bir mantık izler: birey, hem sistemin hem de geliştirici aktörün yeterliliğine ikna olduğunda, “risk” algısı azalır ve “fayda” algısı güçlenir.

Mevcut bulgular H2'yi doğrudan test edecek korelasyon/regresyon sonuçları sunmadığından, değerlendirme daha sınırlı bir çerçevede yapılabilir. Özellikle *kullanım* alt boyutunun (M=14,20) görece yüksek olması, katılımcıların etkinlikten sonra YZ araçlarını kullanmaya daha yatkın olduklarını veya en azından kullanım yeterlik algısının yükseldiğini düşündürür. Bu eğilim, “tutum” ve “kullanım niyeti” değişkenlerinin olumlu yönde olabileceğine dair dolaylı bir işarettir.

Öte yandan, alt boyut güvenirlilikleri incelendiğinde farkındalık/kullanım/etik alt boyutlarında Cronbach α değerlerinin görece düşük olması (Tablo 2), bu alt boyutlardan türetilen çıkarımların dikkatle ele alınması gerektiğini gösterir. H2'nin güçlü biçimde değerlendirilebilmesi için tutum ve niyet maddelerinin (ve özellikle “mühendis yetkinliğine güven” ölçeğinin) ayrı olarak raporlanması, ardından aracılık analizinin gerçekleştirilmesi önerilir.

6.3.3. H3: XAI Etkisi (Açıklanabilirlik → Bilişsel Güvenin Üstünlüğü)

H3, algılanan açıklanabilirliğin (XAI) mühendislerin yetkinliğine duyulan güveni pozitif etkilediğini ve bu etkinin diğer güven boyutlarına kıyasla daha güçlü olabileceğini öne sürmektedir. Bu hipotez, URT perspektifiyle birlikte okunduğunda, açıklamanın “belirsizliği azaltma” işlevi görmesi nedeniyle özellikle yetkinlik algısını öne çıkaracağını varsayar: kullanıcı, sistemin neden böyle davrandığını daha iyi anladığında, bunu “uzmanın işi bildiği” şeklinde yorumlayabilir.

Bu çalışmada XAI algısı doğrudan ölçülmemiştir; ancak *değerlendirme* boyutunun yüksekliği, katılımcıların çıktılarının gerekçeleri ve sınırlılıkları hakkında daha refleksif bir bakış geliştirdiğini düşündürülebilir. Buradan hareketle H3, *kuramsal olarak tutarlı* ve CanikFest bağlamında makul olmakla birlikte, mevcut raporlanan göstergelerle doğrudan doğrulanamaz.

H3'ün güçlü bir biçimde test edilebilmesi için, (i) algılanan açıklanabilirlik maddeleri, (ii) güvenin üç boyutunu (yetkinlik, yardımseverlik, dürüstlük) ayrı ayrı ölçen maddeler ve (iii) bu değişkenler arasındaki ilişkileri test eden bir model gereklidir.

6.3.4. Genel Değerlendirme

Özetle, CanikFest anket bulguları okuryazarlık alt boyutlarında görece olumlu bir tabloya işaret etmektedir. Bununla birlikte hipotezlerin (özellikle H2 ve H3'ün) güçlü biçimde desteklenmesi; güven, açıklanabilirlik, tutum ve kullanım niyeti gibi değişkenlerin ayrı ölçeklerle ölçülmesi ve ilişkisel analizlerin raporlanmasıyla mümkün olacaktır.

7. Tartışma

Bu bölümde, CanikFest anketinden elde edilen bulgular (Tablo 1 ve Tablo 2) kuramsal çerçeve ve hipotezlerle ilişkilendirilerek yorumlanmaktadır. Tartışma, (i) Bilgi Yolu müdahalesinin çıktıları, (ii) hipotezlerin bulgular ışığında değerlendirilmesi ve (iii) daha geniş risk/yönetişim bağlamı olmak üzere üç düzeyde yapılandırılmıştır.

7.1. CanikFest Bulgularının Genel Deseni: Okuryazarlıkta “Değerlendirme” Ağırlığı

Tablo 1 incelendiğinde, dört alt boyut içinde *değerlendirme* puanının (M=15,73) en yüksek olması, CanikFest’in katılımcılarda özellikle “çıktıyı yorumlama” ve “sınırlılıkları görme” becerilerini güçlendirdiği yönünde bir sinyal üretmektedir. Bu durum, etkinlikte farklı araçların (çoklu stant/çoklu uygulama senaryosu) deneyimlenmesiyle katılımcıların karşılaştırma yapabilmesine ve “hangi araç hangi işte işe yarar?” sorusuna daha rasyonel cevaplar geliştirmesine bağlanabilir.

Kullanım (M=14,20) ve etik (M=14,22) alt boyutlarının da orta-üst düzeyde seyretmesi, CanikFest’in yalnızca “tanıtım” değil, “deneyim temelli öğrenme” alanı oluşturduğunu düşündürür. Bununla birlikte güvenirlilik sonuçları (Tablo 2) farkındalık/kullanım/etik alt boyutlarında düşük α değerlerine işaret etmektedir; bu nedenle bu alt boyutlara ilişkin çıkarımların, maddelerin homojenliğini ve ölçüm hassasiyetini dikkate alarak temkinli yorumlanması gerekir.

7.2. H1’in Tartışılması: Bilgi Yolu Müdahalesi Okuryazarlığı ve Bilişsel Güveni Nasıl Etkiler?

H1, uygulamalı atölyelerin YZ okuryazarlığını ve mühendis yetkinliğine güveni artıracaklarını öngörmektedir. Bulguların “değerlendirme” boyutunda görece yüksek olması, Bilgi Yolu yaklaşımının temel iddiasıyla uyumludur: kişi YZ’nin nasıl çalıştığını ve nerede yanılabilirliğini daha iyi anladığında, belirsizlik azalır ve teknolojiyle daha sağlıklı bir ilişki kurabilir (Berger & Calabrese 1975).

Bu noktada CanikFest’in pedagojik katkısı iki biçimde okunabilir. Birincisi, katılımcıların araçları doğrudan deneyimlemesi sayesinde “kör güven” yerine “koşullu güven” geliştirmesi (sisteme ne zaman güvenip ne zaman denetim isteyeceğini bilmesi) mümkün hâle gelmektedir. İkincisi, atölye ortamında uzmanlarla kurulan etkileşim, “aktöre güven” boyutunu güçlendirebilecek bir temas alanı yaratmaktadır. Ancak mevcut raporlanan tablolarda “mühendis yetkinliğine güven” doğrudan ölçülmediği için H1’in bu ikinci kısmı ampirik olarak gösterilememekte; yalnızca okuryazarlık göstergeleri üzerinden dolaylı bir çıkarım yapılabilmektedir.

7.3. H2’nin Tartışılması: Yetkinliğe Güvenin Tutum ve Kullanım Niyetiyle İlişkisi

H2’nin temel önerisi, “yetkinliğe güven”in olumlu tutumu beslediği ve bunun da kullanım niyetine taşındığıdır. CanikFest verileri içinde bu zinciri doğrudan test edecek korelasyon/ regresyon/SEM çıktıları raporlanmadığından, tartışma ilişkisel bir doğrulama yerine mekanizma temelli bir açıklama düzeyinde kalmaktadır.

Bununla birlikte kullanım alt boyutunun orta-üst düzeyde olması (M=14,20), katılımcıların YZ araçlarını kullanmaya dönük öz-yeterlik algısının ve pratik aşinalığın yükseldiğini göstermektedir. Bu tür bir aşinalık, kabul literatüründe tutumun güçlenmesi için kritik bir öncül olarak görülür: kullanıcı, teknolojiyi “yabancı ve riskli” olmaktan çıkarıp “deneyimlenmiş ve yönetilebilir” bir araç olarak çerçevelediğinde benimseme olasılığı artar.

Bu nedenle H2, CanikFest bağlamında *makul* ve bulgularla *uyumlu* bir açıklama önermektedir; ancak hipotezin “aracılık” kısmının doğrulanabilmesi için tutum ve kullanım niyeti maddelerinin (ve özellikle yetkinlik güveninin) ayrı ölçeklerle raporlanması ve aracılık analizinin gerçekleştirilmesi gereklidir.

7.4. H3'ün Tartışılması: XAI'nin Güven Boyutları Üzerindeki Olası Üstün Etkisi

H3, açıklanabilirliğin (XAI) özellikle “yetkinlik” algısı üzerinden güveni güçlendireceğini ve bunun diğer güven boyutlarından daha baskın olabileceğini savunur. Bu önerme, XAI literatüründe açıklamaların kullanıcı güvenliğini artırma potansiyeliyle uyumludur (Gunning vd., 2019; Lundberg & Lee, 2017; Ribeiro vd., 2016).

CanikFest verilerinde XAI algısı ölçülmediği için H3 doğrudan sınanamaz; ancak değerlendirme boyutunun yüksekliği, katılımcıların YZ çıktıları hakkında daha “açıklama talep eden” ve daha eleştirel bir bakış geliştirdiğine işaret eder. Bu da pratik bir sonuç doğurur: CanikFest gibi etkinliklerde yalnızca araç gösterimi yapmak yerine, araçların çıktılarının nedenlerini açıklayan mini XAI demoları veya “çıktıyı birlikte yorumlama” oturumları, H3'te öne sürülen mekanizmayı daha görünür hâle getirebilir.

7.5. Bulguların Yönetişim ve Risk Bağlamıyla Birlikte Okunması

CanikFest bulgularının okuryazarlık yönünde olumlu bir tablo sunması, daha geniş risk gündemini ortadan kaldırmaz; aksine, bu risklerle başa çıkmak için neden Bilgi Yolu yaklaşımına ihtiyaç duyulduğunu gösterir. Otomasyonun işgücü üzerindeki etkileri (Acemoglu & Restrepo, 2020; Autor, 2015; Eloundou vd., 2023; Frey & Osborne, 2017) ve dezenformasyon/derin sahtecilik riski (Chesney & Citron, 2019; Vaccari & Chadwick, 2020), kullanıcıların YZ'ye ilişkin genel güven iklimini etkileyen makro arka plandır.

Bu bağlamda CanikFest türü etkinlikler, yalnızca teknik bir “tanıtım” faaliyeti olarak değil; (i) eleştirel değerlendirme becerisi, (ii) etik farkındalık ve (iii) açıklama/denetim beklentisi oluşturan bir kamu eğitimi müdahalesi olarak konumlandırıldığında, sosyal kabul açısından daha güçlü bir rol üstlenebilir.

8. Sonuç ve Öneriler

Bu çalışma, CanikFest kapsamında yürütülen uygulamalı atölye deneyiminin, YZ okuryazarlığı alt boyutlarında (farkındalık, kullanım, değerlendirme ve etik) orta-üst düzeyde bir profil sunduğunu göstermektedir. Özellikle *değerlendirme* boyutunun görece yüksek seyretmesi (Tablo 1), katılımcıların YZ araçlarını yalnızca kullanmakla kalmayıp, çıktıları yorumlama ve sınırlılıkları fark etme yönünde daha eleştirel bir tutum geliştirebildiğine işaret etmektedir.

Bulguların kuramsal çerçeve ile birlikte okunması, Bilgi Yolu yaklaşımının toplumsal kabul açısından neden kritik olduğunu ortaya koymaktadır: kullanıcılar teknolojiye dair daha çok bilgiye ve deneyime sahip oldukça belirsizlik azalmakta; güven, daha koşullu ve denetim talep eden bir biçimde yeniden inşa edilebilmektedir. Bu bulgu, Belirsizliği Azaltma Teorisi'nin (URT) bilgi arayışı ve tutum ilişkisine dair öngörülerıyla de uyumludur.

Bununla birlikte, bu makalede hipotezlerin tamamı ilişki/etki düzeyinde (örn. aracılık, XAI'nin güven boyutları üzerindeki görece gücü) istatistiksel olarak test edilmemiştir. Dolayısıyla sonuçlar, CanikFest anket bulgularının ortaya koyduğu genel eğilimler üzerinden yorumlanmış; güven, açıklanabilirlik, tutum ve kullanım niyeti gibi değişkenlerin ayrı ölçeklerle ölçülmesi ve ilişkisel analizlerin raporlanması bir sonraki adım olarak değerlendirilmiştir.

8.1. Öneriler

Aşağıda, CanikFest deneyimi ve bulgular ışığında geliştirilebilecek politika ve uygulama önerileri özetlenmektedir:

1. Bağımsız denetim mekanizmalarını kurumsallaştırmak.
2. Yüksek riskli uygulamalarda açıklanabilirliği (XAI) asgari gereklilik hâline getirmek.
3. YZ okuryazarlığını ve yeniden becerilendirme programlarını yaygınlaştırmak.
4. Psikolojik etki ve dezenformasyon risk değerlendirmelerini yönetişime entegre etmek.
5. İnsan denetimi ve etik sorumluluk zincirini görünür kılmak.

8.2. Sınırlılıklar

Çalışmanın bulguları tek bir etkinlik bağlamında (CanikFest) ve öz-bildirim (anket) verileri üzerinden elde edilmiştir. Dolayısıyla (i) örneklemin temsiliyeti, (ii) sosyal beğenirlik yanlılığı, (iii) çapraz kesitsel tasarım nedeniyle nedensellik iddialarının sınırlılığı ve (iv) ölçümlerin etkinlik ortamından etkilenebilmesi gibi sınırlılıklar dikkate alınmalıdır.

Bilgilendirme: CanikFest etkinliği 2025 Eylül ayında Tubitak 4007 Bilim Şenliklerini Destekleme Programı kapsamında desteklenmiş, Samsun/Canik Belediyesinin koordinesinde Canik Özdemir Bayraktar Keşif Kampüsü tarafından yürütülmüştür.

Kaynakça

- Acemoglu, D., & Restrepo, P. (2020). Robots and jobs: Evidence from us labor markets. *Journal of Political Economy*, 128 (6), 2188–2244. doi: 10.1086/705716
- Autor, D. H. (2015). Why are there still so many jobs? the history and future of workplace automation. *Journal of Economic Perspectives*, 29 (3), 3–30. doi: 10.1257/jep.29.3.3
- Berger, C. R., & Calabrese, R. J. (1975). Some explorations in initial interaction and beyond: Toward a developmental theory of interpersonal communication. *Human Communication Research*, 1 (2), 99–112. doi: 10.1111/j.1468-2958.1975.tb00258.x
- Chesney, R., & Citron, D. K. (2019). Deep fakes: A looming challenge for privacy, democracy, and national security. *California Law Review*, 107 (6), 1753–1819.
- Eloundou, T., Manning, S., Mishkin, P., & Rock, D. (2023). Gpts are gpts: An early look at the labor market impact potential of large language models. *arXiv preprint*.
- Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254–280. doi: 10.1016/j.techfore.2016.08.019
- Gunning, D., Stefik, M., Choi, J., Miller, T., Stumpf, S., & Yang, G.-Z. (2019). Xai-explainable artificial intelligence. *Science Robotics*, 4 (37), eaay7120. doi: 10.1126/scirobotics.aay7120
- KPMG. (2023). *Trust in artificial intelligence: Global insights 2023*. Erişim: 2026-01-31, from <https://kpmg.com/xx/en/home/insights/2023/11/trust-in-artificial-intelligence.html>

- Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30 .
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *The Academy of Management Review*, 20 (3), 709–734. doi:10.2307/258792
- Poushter, J., Fagan, M., & Corichi, M. (2023). *How people around the world view artificial intelligence*. Erişim: 2026-01-31, from <https://www.pewresearch.org/global/2023/08/28/how-people-around-the-world-view-artificial-intelligence/>
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "why should i trust you?": Explaining the predictions of any classifier. In *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining (kdd '16)* (pp. 1135–1144). doi:10.1145/2939672.2939778
- Vaccari, C., & Chadwick, A. (2020). Deepfakes and disinformation: Exploring the impact of synthetic political video on deception, uncertainty, and trust in news. *Social Media+ Society*, 6 (1). doi: 10.1177/2056305120903408
- Vogels, E. A. (2023). *What the data says about americans' views of artificial intelligence*. Erişim: 2026-01-31, from <https://www.pewresearch.org/short-reads/2023/08/28/what-the-data-says-about-americans-views-of-artificial-intelligence/>