

CHATGPT OKURYAZARLIK ÖLÇEĞİNİN TÜRKÇEYE UYARLANMASI; GEÇERLİLİK VE GÜVENİRLİK ÇALIŞMASI

Adaptation of the ChatGPT Literacy Scale to Turkish; Validity and Reliability Study

Rahim ARSLAN¹ | Mesut KARAMAN^{2*}

¹ Doç. Dr., Sivas Cumhuriyet Üniversitesi, İktisadi ve İdari Bilimler Fakültesi, Sivas, Türkiye rahimarslan@cumhuriyet.edu.tr, ORCID: 0000-0003-4329-3651

² Arş. Gör., İstanbul Galata Üniversitesi, Sanat ve Sosyal Bilimler Fakültesi, İstanbul, Türkiye, mesut.karaman@galata.edu.tr, ORCID: 0000-0002-7584-0800

*Sorumlu yazar

ÖZ

Bu çalışmada Lee ve Park (2024) tarafından geliştirilen “ChatGPT Okuryazarlık Ölçeği’nin” Türkçe versiyonunun geçerlilik ve güvenirlik çalışmasını gerçekleştirmek amaçlanmıştır. Bu psikometrik çalışma bir devlet üniversitesinde uygulanmıştır. Çalışma 648 üniversitesinin katılımıyla gerçekleştirilmiştir. Veriler, Mayıs-Haziran 2025 tarihleri arasında çevrimiçi bir anket ve kolayda örnekleme yöntemi kullanılarak toplanmıştır. Verilerin analizinde SPSS 26 ve AMOS 24 programları kullanılmıştır. Ölçeğin psikometrik özelliklerini değerlendirmek amacıyla madde analizi, açıklayıcı faktör analizi, doğrulayıcı faktör analizi, yakınsak ve iraksak geçerlilik istatistikleri, Cronbach Alpha iç tutarlılık katsayısı, bağımlı örnekleme t-testi ve sınıf içi korelasyon analizleri uygulanmıştır. Ölçeğin Kapsam Geçerlilik İndeksi 0,92’dir. Ayrı örneklemlerde yapılan Açıklayıcı Faktör Analizi (AFA) ve Doğrulamalı Faktör Analizi DFA sonucunda uyarlama yapılan orijinal ölçekteki ile benzer şekilde toplam 25 madde ve beş alt faktör doğrulanmıştır. Ölçeğin Cronbach Alpha değeri her iki örnekleme için de 0,90’ın üzerindedir. Beş faktörlü yapıda toplam varyansın %72,275’ini açıklanmıştır. AFA’da elde edilen 25 maddelik ve beş faktörlü yapı DFA’da doğrulanmış, yakınsak ve iraksak geçerlilik istatistikleri sağlanmıştır. Test-tekrar test sonuçları ölçeğin yüksek güvenirliğe sahip olduğunu göstermiştir. ChatGPT okuryazarlık ölçeği, bireylerin ChatGPT okuryazarlık algılarının değerlendirilmesi için geçerli ve güvenilir bir ölçektir.

Anahtar Kelimeler: ChatGPT, Okuryazarlık, Geçerlilik, Güvenirlik

ABSTRACT

The aim of this study was to conduct a validity and reliability study of the Turkish version of the “ChatGPT Literacy Scale” developed by Lee and Park (2024). This psychometric study was conducted at a state university. The study involved 648 university students. Data were collected between May and June 2025 using an online survey and convenience sampling. SPSS 26 and AMOS 24 were used for data analysis. To evaluate the psychometric properties of the scale, item analysis, exploratory factor analysis, confirmatory factor analysis, convergent and divergent validity statistics, Cronbach’s Alpha internal consistency coefficient, paired-sample t-test, and intraclass correlation analyses were applied. The scale’s Content Validity Index is 0.92. As a result of the Exploratory Factor Analysis (EFA) and Confirmatory Factor Analysis (CFA) conducted on separate samples, a total of 25 items and five sub-factors were confirmed, similar to the original adapted scale. The Cronbach’s Alpha value of the scale is above 0.90 for both samples. The five-factor structure explains 72.275% of the total variance. The 25-item, five-factor structure obtained in EFA was confirmed in CFA, and convergent and discriminant validity statistics were established. Test-retest results indicate that the scale has high reliability. The ChatGPT Literacy Scale is valid and reliable for assessing individuals’ perceptions of ChatGPT literacy.

Keywords: ChatGPT, Literacy, Validity, Reliability.

Geliş Tarihi: 26 Eylül 2025

Kabul Tarihi: 8 Nisan 2026

Yayın Tarihi: 24 Mayıs 2026

Atfı: Arslan, R. ve Karaman, M. (2026). ChatGPT Okuryazarlık Ölçeğinin Türkçeye Uyarlanması; Geçerlilik ve Güvenirlik Çalışması. Selçuk Üniversitesi Sosyal Bilimler Meslek Yüksekokulu Dergisi, 29 (2026), 230-241.

1. GİRİŞ

Yapay zekâ, bilimsel disiplinlerin %98'inde kullanılmaktadır. Araştırmacılar araştırma sürecinin çeşitli yönlerini hızlandırmak ve desteklemek için yapay zekâ destekli araçları benimsemişlerdir (Hoffmann vd., 2024). Yapay zekâ, eğitimde öğrencilere yeni bilgileri öğrenme süreçlerinde rehberlik etmek, yardımcı ve destek olmak için kullanılmaktadır. Yapay zekâ; öğrenme süreçlerini iyileştirme, kişisel destek sağlama ve akademik performansı arttırmaktadır (Farhi vd., 2023). Yapay zekânın bireylerin kullanımına girmesiyle ilerleyen süreçlerde ortaya çıkan ve yapay zekâ tabanlı bir dil modeli olan ChatGPT (Chat Generative Pre-Trained Transformer), OpenAI tarafından geliştirilmiştir. Bir bireyin üreteceği benzer şekilde metinler oluşturma ve otomatik metin üretimi, soru-cevap, istenilen bölümlerle ilgili otomatik özetleme ve görüntü üretimi gibi karmaşık dil görevlerini yerine getirebilir (Yu vd., 2024). ChatGPT bilgi aktarımı ve yenilikçilik alanlarında devrim niteliğinde bir değişim ortaya çıkarmıştır (Gordijn & Have, 2023).

ChatGPT'nin ortaya çıkardığı bu devrim ile geleneksel yapay zekâ arasında da bir takım işlevsel farklılıklar kendisini göstermiştir. ChatGPT, "Üretken Yapay Zekâ"ya aittir. Geleneksel "Ayırt Edici Yapay Zekâ" ile karşılaştırıldığında, aşağıdaki farklılıklar vardır: (1) Etkileşim açısından, üretken yapay zekâ insan konuşmalarına doğal bir şekilde yanıt verebilirken, ayırt edici yapay zekâ daha zayıf etkileşim yeteneklerine sahiptir. Bu nedenle, üretken yapay zekâ daha fazla duygusal destek ve sosyal destek sağlayabilir. İnsanların ona daha fazla güvenmesini sağlar. (2) Yenilikçilik açısından, geleneksel yapay zekâ esas olarak mevcut verileri analiz etmek ve yorumlamak için kullanılır ve yenilikçilik yoksundur. Buna karşılık, üretken yapay zekâ; yeni metinler, görüntüler, müzik ve diğer içerikler oluşturabilir, bu da üretken yapay zekâyı sınırsız bir potansiyel kazandırır ve insanları ona daha fazla bağlamaktadır (Yu vd., 2024).

ChatGPT kullanıcı sayısı bakımında kısa sürede milyonlara ulaşmıştır. ChatGPT'nin ilk bir milyon kullanıcıya sadece beş gün içerisinde (Valova vd., 2024), 100 milyon kullanıcıya ulaşması ise iki ay sürmüştür. Bu da Instagram'ın 2,5 yılda ulaştığı kullanıcı sayısı ile aynıdır (Clay, 2023; Homolak, 2023; Reed, 2023; Lee & Park, 2024). Ağustos 2023'te 180 milyondan fazla kullanıcısı olan ChatGPT'nin Ocak 2023'te bir milyar kullanıcı etkileşimine ulaşmıştır. Bu istatistikler, ChatGPT'nin önemli bir etkisini ve çekiciliğini ortaya koymaktadır (Yu vd., 2024). Buna karşın bazı analisteler bazı meslekler için ChatGPT'yi bir

tehdit olarak görürken bazıları ise son derece başarılı ve yararlı olduğunu bildirmektedirler. Ayrıca bu eleştiriler noktasında bir önemli hususta ChatGPT ile ilgili etik hususlar, eğitimde öğrencilerin ileri düzey düşünme becerileri ve bilimsel dürüstlük üzerindeki olası olumsuz etkileri gündeme getirilmiştir (Qadir, 2022; Farhi vd., 2023).

Literatürde ChatGPT kullanımında araştırma önerileri, özet bölümler oluşturma (Altmäe ve ark., 2023; Gao ve ark., 2023), metinlerle ilgili dilbilgisi (Bouschery ve ark., 2023), ve genel düzenlemeyle ilgili iyileştirmelerde yardımcı olması (Macdonald ve ark., 2023; Seghier, 2023) bakımından bazı pratik yönleri ortaya konulmuştur. ChatGPT, hipotez oluşturma, büyük veri kümelerinden istenilen bilgileri elde etmesi, karmaşık araştırmaları basitleştirme ve çok daha fazlası olmak üzere ve çok daha fazlası dahil olmak üzere diğer araştırma görevleri için de kullanılabilir (Rossi ve ark., 2024; Susarla ve ark., 2023; venturebeat.com, 2023; Hoffmann vd., 2024).

ChatGPT ile ilgili literatürde teorik dayanağı olarak Planlı davranış teorisi gösterilebilir. Bu teori, bilgisel ve motivasyonel faktörlerin bireysel davranışı nasıl etkilediğini ortaya koymak üzere yaygın olarak kullanıldığı ifade edilebilir. Planlı davranış teorisi, mobil destekli öğrenme (Nie vd. 2020) ve çevrimiçi öğrenme (Pan vd. 2023) dâhil olmak üzere eğitim teknolojileriyle bütünleştirilmiş bireylerin hazır oluşunu gösteren güçlü bir teorik çerçevede uygulanmıştır (Luo & Zou, 2024). Literatürde ChatGPT'nin temel gücünün mevcut bilgiyi sıralamak, özetlemek olduğu ve uygun bir şekilde dikkatle kullanıldığında ve eğitimde yardımcı olmak için ideal olduğu bildirilmiştir (Cooper, 2023; Montenegro-Rueda ve ark., 2023). Boubker (2024) çalışmasında ChatGPT'nin algılanan faydasının ChatGPT kullanımını ve öğrencilerin memnuniyetini olumlu yönde etkilediğini ve bu aracın eğitimdeki pratik uygulamalarda sorumlu bir şekilde uygulanması gerektiği belirtilmiştir (Sedlbauer vd., 2024).

ChatGPT anında geri bildirim sağlayarak ve daha etkileşimli bir öğrenme ortamı oluşturarak öğrenci motivasyonunu ve öğrenme süreçlerine katılımını arttırabilmektedir (Kasneci vd., 2023). ChatGPT, dil becerilerini kullanarak ve içerikle etkileşim kurmak için düşük stresli bir öğrenme ortamı sunmaktadır. Bu doğrultuda güven ve akıcılık oluşturmaya yardımcı olabilir (Escalante vd., 2023; Pellas, 2023). Böylelikle destekleyici bir ortam ile öğrencileri öğrenmeye ve daha aktif katılmaya teşvik edebilir. Buna karşın öğrencileri öğrenmeye yönelik ilerleyen süreçlerde ChatGPT'nin üretmiş olduğu materyallere aşırı bağımlı hale de getirebilir. Bu da öğrencilerin

ilgili araştırma konularına yönelik bağlantı kurma, önemli analitik yetenekler edinme ve eleştirel düşünme noktasında engeller teşkil edebilir (Al Shloul vd., 2024; Elly Heung & Chiu, 2025).

ChatGPT ile ilgili her ne kadar olumlu yönler vurgulanmış olsa da görüldüğü üzere literatürde olumsuz yönler de değinilmiştir. Bu bakımdan yapay zeka tabanlı ChatGPT gibi oldukça önemli görülen bir yardımcı aracın kullanımı açısından teknolojik yönünden dijital okuryazarlık düzeyinin bireylerde olması aşikardır. Bu çalışmada ulusal literatürde her ne kadar ChatGPT ile ilgili çalışmalar yapılmış olsa da bireylerin ChatGPT okuryazarlık düzeylerini ölçmeye yönelik bir ölçeğin olmayışı çalışmanın motivasyon kaynağı olmakla birlikte çalışmanın önemini ortaya koymaktadır. ChatGPT ile ilgili ulusal literatürdeki ölçekler değerlendirildiğinde ChatGPT kullanım ölçeği (Taktak & Bafralı, 2025) ve problemli ChatGPT kullanım ölçeğinin (Maral vd., 2025) olduğu görülmektedir. Literatürde yapılan incelemelerde ulusal literatüre kazandırılan ölçeklerde yapay zeka okuryazarlığı ile ilgili ölçeklerin yer aldığı görülmektedir (Akyürek, 2025; Karaoğlu Yılmaz ve Yılmaz, 2023; Çelebi, 2023). Bu da genel olarak yapay zeka ile ilgili çalışmalarının olması yanında yapay zekanın bir alt tabanı olan ve de metin üzerinden insanla etkileşim kuran bireylerin ChatGPT okuryazarlık düzeylerini belirlemeye yönelik bir ölçeğin olmayışı çalışmanın bir diğer önemini ortaya koymaktadır.

ChatGPT okuryazarlığı teknik yeterlilik, eleştirel değerlendirme, iletişim yeterliliği, yenilikçi uygulama ve etik yeterlilik olmak üzere beş boyut altında kavramsallaştırılmıştır. Teknik yeterlilik, bireylerin farklı teknolojileri anlama ve kullanma becerisidir. Eleştirel değerlendirme ChatGPT'nin vermiş olduğu yanıtların doğruluğunu, güvenilirliğini, eksiksizliğini, ön yargısını ve hatalarını değerlendirme ve analiz etme yeteneğidir. İletişim yeterliliği bireylerin iletişim ve işbirliği becerilerini ChatGPT ile etkili bir şekilde iletişim kurma yeteneğidir. Yenilikçi uygulama bireylerin yenilikçi ve yenilikçilik becerilerini geliştirmek için yeni fikirler ya da çözümler üretme becerisidir. Etik yeterlilik, etik veya yasal sorunları belirleme ve ChatGPT'nin etik kullanımını ifade etmektedir (Lee & Park, 2024). Bu çalışmada aşağıdaki iki araştırma sorusuna cevap aranmıştır.

ChatGPT okuryazarlığı geçerli bir ölçme aracı mıdır?

ChatGPT okuryazarlığı güvenilir bir ölçme aracı mıdır?

2. YÖNTEM

Orijinal Lee ve Park (2024)'ın geliştirdiği ChatGPT Okuryazarlık Ölçeği Türkçe uyarlamasının geçerlik ve güvenilirliğin belirlemek amacıyla yürütülen bu çalışma nicel bir yaklaşım ile yürütülmüştür. Çalışma tarama ile gerçekleştirilmiş veriler nicel analiz yöntemleri ile test edilmiştir.

2.1. Amaç

Bu çalışmada Lee ve Park (2024) tarafından geliştirilen "ChatGPT Okuryazarlık Ölçeği'nin" Türkçe versiyonunun geçerlilik ve güvenilirlik çalışmasını gerçekleştirmek amaçlanmıştır. Bu çalışma ChatGPT Okuryazarlık Ölçeği'nin" Türkçeye kültürel uyarlanması, geçerlilik ve güvenilirlik olmak üzere iki aşamalıdır. Ölçeğin çeviri, uyarlama, geçerlilik ve güvenilirlik için uluslararası rehberlerden yararlanılmıştır (Hernández vd., 2020; Sousa ve Rojjanasrirat, 2011).

2.2. Uyarlama Süreci

Lee ve Park (2024)'ın geliştirdiği ChatGPT Okuryazarlık Ölçeği'ni İngilizce dil eğitimi veren ana dili Türkçe olan iki dil bilimci tarafından bağımsız olarak İngilizceden Türkçeye çevrilmiştir. Çevirmenler ve araştırmacılar daha sonrasında çeviri metinlerini ve orijinal ölçek ile birlikte değerlendirmiş, belirsizlik ve tutarsızlıklar incelenmiştir. Çevirmenler ve araştırmacılar, ölçek maddelerinin her birinde kelimelerin tam karşılığından çok, anlam bütünlüğü üzerinde uzlaşmıştır. Böylece ölçeğin ilk Türkçe versiyonu ortaya çıkmıştır. Bu ilk Türkçe versiyon bir sonraki aşamada ana dili İngilizce olan Türkçeyi ana dili düzeyinde bilen iki tercüman tarafından Türkçeden İngilizceye geri çevirme yapılmıştır. Ölçek ile ilgili fikir birliği sağlamak ve kavramları netleştirmek için ilk ve ikinci aşamada yapılan çeviriler orijinal verisyonla karşılaştırılarak uyarlanan ölçeğin son hali ortaya konulmuştur (Hernández vd., 2020; Sousa ve Rojjanasrirat, 2011).

Pilot test için 10-40 kişi alınmalıdır ilkesi gereği ilk veri toplamada ölçeği doldurma yönergesi ve yanıt formatının değerlendirilmesi istenmiştir (Hernández vd., 2020; Sousa ve Rojjanasrirat, 2011). Bu çalışmada pilot aşamasında 30 kişilik bir öğrenci grubundan yüz yüze görüşme tekniği ile görüş alınmıştır. Öğrenciler ölçek maddelerinin ve doldurma yönergesinin anlaşılır olduğunu belirtmiştir. Daha sonraki aşamada 10 kişilik bir uzman görüşü alınmıştır. Bu kapsamda ölçek konusuna yönelik çalışmaları olan, geçerlilik ve güvenilirlik alanında deneyimli 10 kişilik bir uzman görüşüne başvurulmuştur. Literatürde uzman görüşü sayısının 6-10 kişi olmasının yeterli olduğu belirtilmektedir (Hernández

vd., 2020, Sousa ve Rojjanasrirat, 2011). Uzman değerlendirme formunda her bir uzmandan ölçek değerlendirmesi ile ilgili “1-uygun değil, 2-maddenin uygun şekle getirilmesi gerekir ve 3- uygun, ancak küçük değişiklik gerekir ve 4- çok uygun” şeklinde puanlandırması istenmiştir. Uzmanlardan gelen yanıtlar doğrultusunda kapsam geçerlilik indeksi hesaplanmıştır. Bunun için Davis (1992) tekniğinden (3 ya da 4 puan veren uzman sayısı/tüm uzman sayısı) yararlanılmıştır (Grant ve Davis, 1997). Ölçek KGO ortalaması 0,88’dir. Her bir madde için KGİ değeri 0,80 üzeridir. Her bir madde ve ölçek geneli için KGO oranları 0,80 üzeri olduğu için bir sonraki aşama olan geçerlilik ve güvenirlik analizlerine geçilmiştir (Hernández vd., 2020; Sousa ve Rojjanasrirat, 2011).

2.3. Geçerlilik ve Güvenirlik Analizi

2.3.1. Örneklem

Bu çalışma bir devlet üniversitesinde öğrenim görmekte olan öğrenciler üzerinde uygulanmıştır. Araştırma öğrenci grubuna İktisadi ve İdari Bilimler alanında öğrenim görmekte olan öğrenciler dahil edilmiştir. Çalışma evreni üniversitenin İktisadi ve İdari Bilimler alanında öğrenim görmekte olan tüm öğrencilerdir. Araştırmanın örneklem seçimi ise kolayda örneklem yöntemi ile öğrencilerin gönüllü katılımı esas alınarak toplamda 648 kişilik bir öğrenci grubuna ulaşılmıştır. Örneklem sayısının belirlenmesinde her bir madde sayısının 5 ile 10 katı arasında alınması gerektiği önerilmektedir (Karaman, 2023). Bir diğer ifadeye göre de yapı geçerliliği analizlerinde örneklem ölçek madde sayısının 3-20 katı olmalıdır (Gunawan vd., 2021). Bu bakımdan ulaşılan örneklem yeterli olduğu söylenebilir.

2.3.2. Veri Toplama Araçları

Araştırmada veriler “Kişisel Bilgi Formu” ve “Chatgp Dijital Okuryazarlık Ölçeği” ile elde edilmiştir.

Kişisel Bilgi Formu; öğrencilerin cinsiyet, yaş, bölüm, sınıf düzeyleri, “Okul araştırmalarınızda ChatGPT’yi kullanır mısınız?” ve “Sosyal aktivitelerinizi yapay zeka ya da ChatGPT’de değerlendirir misiniz?” gibi ifadelerle yer verilmiştir.

Chatgp Dijital Okuryazarlık Ölçeği; Lee ve Park (2024) tarafından geliştirilmiş ölçek toplamda 25 madde ve 5 alt boyutludur. Ölçek puanlaması 5’li Likert tipinde olup “1-Kesinlikle Katılmıyorum ile 5-Kesinlikle Katılıyorum” arasındadır. Ölçekten alınan puanlar arttıkça Chatgp Dijital Okuryazarlık düzeyi artmaktadır. Ölçekte ters kodlama

bulunmamaktadır. Orijinal yapıda ölçek geneli Lee ve Park (2024) Cronbach Alpha değeri 0,88’dir.

2.3.3. Verilerin Toplanması

Çalışmada veriler bir devlet üniversitesinde öğrenim görmekte olan öğrencilerden yüz yüze görüşme tekniği ile elde edilmiştir. Anketin ilk sayfasında öğrencilere çalışma ile ilgili bilgilendirme yapılmıştır. Çalışmaya katılmayı kabul eden öğrencilerde anket uygulaması yapılmıştır. Test-tekrar test analizi için ise çalışma örnekleminde ayrı 30 kişilik bir öğrenci grubuna iki haftalık arayla anket yeniden uygulanmıştır.

2.3.4. Etik Hususlar

Bu çalışma için bir devlet üniversitesinden Sosyal Bilimler Bilimsel Araştırma Önerisi Etik Değerlendirme Kurulu’ndan 24.04.2025 tarihli 31 sayılı Karar No ile etik onay alınmıştır. Anketin uygulama aşaması öncesinde çalışmanın amacı açıklanmış ve öğrencilerden gönüllü olarak ankete katılmaları istenmiştir.

2.3.5. İstatistiksel Analiz

Araştırma kapsamında elde edilen verilerin analizinde SPSS 26 ve AMOS 24 paket programları kullanılmıştır. Araştırma da ölçek değişkenlerinin çarpıklık ve basıklık değerlerinin -1,5 ile +1,5 arasında ve normal dağılım aralığında olduğu belirlenmiştir. Öğrencilerin sosyo-demografik özelliklerinin belirlenmesinde tanımlayıcı istatistikler, ölçek faktör yapısının test edilmesinde Açıklayıcı faktör analizi (AFA) ve doğrulayıcı faktör analizi (DFA) uygulanmıştır.

İlk öncelikle AFA için 322 kişilik bir örneklem elde edilmiş daha sonrasında ölçek yapısından herhangi bir değişiklik olmadığı görülmüş DFA için ise 326 kişilik bir örneklemde uygulanmıştır. Ölçek güvenirliğinin değerlendirilmesinde madde toplam korelasyon puanları ve iç tutarlılık katsayısının belirlenmesinde Cronbach Alpha katsayısından yararlanılmıştır.

3. BULGULAR

3.1. Üniversite Öğrencilerinin Sosyo-Demografik Özellikleri

Üniversite öğrencilerinin sosyo-demografik özellikleri Tablo 1’de gösterilmiştir.

Tablo 1. Sosyo-demografik özellikler (N=648)

Değişkenler		N	%
Cinsiyet	Kadın	357	55,1
	Erkek	291	44,9
Yaş	18-21 yaş	498	76,9
	22 yaş ve üzeri	150	23,1
Sınıf düzeyiniz	1.Sınıf	170	26,2
	2.Sınıf	220	34
	3.Sınıf	195	30,1
	4.Sınıf	63	9,7
Bölümünüz	Bankacılık	53	8,2
	Çeko	30	4,6
	Ekonometri	36	5,6
	İktisat	58	9
	İşletme	98	15,1
	Kamu	66	10,2
	Maliye	146	22,5
	UTL	66	10,2
Okul araştırmalarınızda ChatGPT'yi kullanır mısınız?	Evet	491	75,8
	Hayır	157	24,2
Sosyal aktivitelerinizi yapay zeka ya da ChatGPT'de değerlendirir misiniz?	Evet	328	50,6
	Hayır	320	49,4

Öğrencilerin %55,1'i kadın, %76,9'u 18-21 yaş aralığında, %22,5'i maliye bölümünde, %75,8'i okul araştırmalarında ChatGPT'yi kullandığı ve %50,6'sının sosyal aktiviteleri yapay zeka ya da ChatGPT'de değerlendirdiği belirlenmiştir (Tablo 1).

3.2. Madde Analizi

Ölçek maddelerine yönelik 25 madde için düzeltilmiş madde toplam korelasyon katsayıları 0,30 üzeri olduğu belirlenmiştir (Tablo 2).

Tablo 2. Madde analizi ve AFA sonuçları (N=322)

Items	CITC	CAID	ED	TY	İY	EY	YU
ED_1	,587	,920	,770				
ED_2	,533	,921	,794				
ED_3	,505	,922	,799				
ED_4	,529	,921	,690				
ED_5	,533	,921	,722				
ED_6	,606	,920	,739				
TY_1	,624	,920		,716			
TY_2	,568	,921		,840			
TY_3	,638	,920		,702			
TY_4	,614	,920		,793			
TY_5	,599	,920		,664			
TY_6	,543	,921		,757			
İY_1	,620	,920			,747		
İY_2	,543	,921			,717		
İY_3	,555	,921			,736		
İY_4	,621	,920			,786		
İY_5	,500	,922			,693		
EY_1	,566	,921				,802	
EY_2	,492	,922				,796	
EY_3	,531	,921				,862	
EY_4	,459	,922				,752	
YU_1	,541	,921					,770
YU_2	,453	,923					,651
YU_3	,471	,922					,715
YU_4	,514	,922					,785
Öz değer katsayısı			8,929	3,048	2,494	2,004	1,593
Açıklanan varyans			35,717	12,191	9,978	8,018	6,372
Toplam açıklanan varyans			72,275				

CITC: Corrected Item-Total Correlation; CAID: Cronbach's Alpha if Item Deleted; ED: Eleştirel Değerlendirme; TY: Teknik Yeterlilik; İY: İletişim Yeterliliği; EY: Etik Yeterlilik; YU: Yenilikçi Uygulama

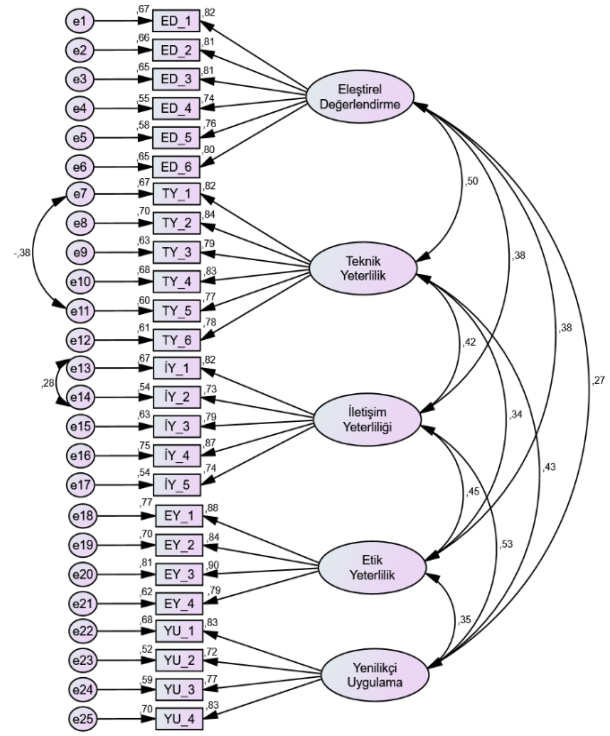
3.3. Yapı Geçerliliği

3.3.1. Açıklayıcı Faktör Analizi (AFA)

Yapı geçerliliği için ilk olarak AFA uygulanmıştır. AFA sonucunda madde faktör yük değerleri 0,50 üzeri olduğu için ölçekten madde çıkarımı yapılmamıştır. Ölçeğin KMO testi ve Bartlett küresellik testi yapılmış ve KMO değeri 0,906 ve küresellik testi sonucu Approx Chi-Square= 5402,772 (df = 300, p< ,001) olarak bulunmuştur. Bununla birlikte ölçek maddeleri için öz değer katsayısı 1'in üzerinde olan orijinal ölçekteki gibi 5 faktör altında toplandığı ve toplam varyansın %72,275'ini açıkladığı belirlenmiştir. Ölçek madde faktör yükleri faktör 1 için ,690 ile ,799; faktör 2 için ,664 ile ,840; faktör 3 ,693 ile ,786; faktör 4 için ,752 ile ,862 ve faktör 5 için ,651 ile ,785 arasındadır (Tablo 2).

3.3.2. Doğrulayıcı Faktör Analizi (DFA)

AFA sonrası keşfedilen yapının model uyumunu değerlendirmek için DFA yapılmıştır. DFA'da model uyum indeksleri χ^2/df (<5); NFI; IFI; CFI; RMSEA; iyi uyum, RMR; GFI; AGFI; TLI; değerlerinin kabul edilebilir uyum düzeyinde olduğu belirlenmiştir ($\chi^2 /sd = 1.520$; RMR= .065; GFI= .946; AGFI= .917; NFI= .961; IFI= .954; TLI= .949; CFI= .959; RMSEA= .043). Modifikasyon önerileri doğrultusunda Faktör 1 için e1 ile e5 ve e13 ile e14 arasında hata kovaryansı atanmıştır (Şekil 1).



Şekil 1. Ölçek DFA Görseli (N=326)

DFA sonrası alt ölçeklerde faktör yükleri Faktör 1> ,75; Faktör 2> ,78 Faktör 3> ,74 ve Faktör 4> ,84 ve Faktör 5> ,73 olduğu bulunmuştur (Şekil 1).

3.3.3. Yakınsak ve İraksak Geçerlilik Analizi

Taslak ölçeğin yakınsak ve ıraksak geçerlilik istatistikleri Tablo 3'te sunulmuştur.

Tablo 3. Yakınsak ve ıraksak geçerlilik istatistikleri (N=326)

	CR	AVE	MSV	ASV	YU	TY	İY	EY	ED
YU	0,873	0,633	0,307	0,190	0,795				
TY	0,921	0,660	0,281	0,208	0,463	0,812			
İY	0,896	0,633	0,307	0,223	0,554	0,445	0,796		
EY	0,916	0,732	0,225	0,169	0,376	0,372	0,474	0,855	
ED	0,915	0,643	0,281	0,179	0,315	0,530	0,405	0,413	0,802

ED: Eleştirel Değerlendirme; TY: Teknik Yeterlilik; İY: İletişim Yeterliliği; EY: Etik Yeterlilik; YU: Yenilikçi Uygulama

Taslak ölçeğin yakınsak geçerlilikleri için CR> ,70 ve CR>AVE; ıraksak geçerlilik için MSV; ASV<AVE olarak elde edilmiştir (Tablo 3).

3.4. Güvenirlik

Bu çalışma için güvenilirlik katsayısı iki farklı örneklem için ayrı ayrı hesaplanmıştır. AFA sonucunda (I. Örneklem N=326) ölçek geneli Cronbach Alpha katsayısı ,930 DFA sonucunda (II. örneklem N= 322) ,924 olarak belirlenmiştir. Ölçek alt boyutları

Cronbach Alpha değerleri her iki örneklem için de ,80 üzeridir.

3.5. Test-Tekrar Test Analizi

Ölçeğin test-tekrar test analizi ve bağımlı örneklem t testi ile kararlılığı değerlendirilmiştir ve ICC değerleri hesaplanmıştır. Ölçek geneli ICC ,970 olup güven aralığı (CI) ,938 ile ,986 arasında (F=33,849; p < 0.001) olduğu belirlenmiştir.

Tablo 4. Test-tekrar test (N=30)

Değişkenler	M ± SD birinci	M± SD ikinci	ICC	%95 (GA)		F	p	T	p
				Alt Sınır	Üst Sınır				
Ölçek Geneli	3,72±0,22	3,73±0,24	,970	,938	,986	33,849	,000	-0,678	0,503
Eleştirel Değerlendirme	3,81±0,23	3,80±0,27	,918	,828	,961	12,230	,000	0,225	0,824
Teknik Yeterlilik	3,70±0,33	3,75±0,34	,872	,730	,939	7,790	,000	-1,330	0,194
İletişim Yeterliliği	3,73±0,47	3,72±0,42	,917	,826	,961	12,060	,000	0,147	0,884
Etik Yeterlilik	3,72±0,22	3,73±0,24	,912	,816	,958	11,394	,000	-0,528	0,601
Yenilikçi Uygulama	3,60±0,23	3,59±0,26	,896	,781	,950	9,608	,000	0,297	0,769

ICC: Sınıf içi korelasyon; T:Bağımlı örneklem t testi

Ölçek alt faktörlerinde; eleştirel değerlendirme için ,918 (%95 CI = 0,828-0,961, F = 12,230, p < 0,001), Teknik yeterlilik için ,872 (%95 CI = 0,730-0,939, F = 7,790, p < 0,001), İletişim yeterliliği için 0,917 (%95 CI = 0,826-0,961, F = 12,060, p < 0,001), Etik yeterlilik için 0,912 (%95 CI = 0,816-0,958, F = 11,394, p < 0,001) ve Yenilikçi uygulama için 0,896 (%95 CI = 0,781-0,950, F = 9,608, p < 0,001), olarak hesaplanmıştır. İki ölçüm arasında anlamlı ve pozitif yönlü bir ilişkinin olduğu belirlenmiştir. Bununla birlikte bağımlı örneklem t testi sonucunda da anlamlı bir fark olmadığı saptanmıştır (Tablo 4).

4. TARTIŞMA

Günümüz dünyasında gelişen teknoloji ile birlikte birçok noktada hayatımıza giren ve birçok farklı alanda yaşamı kolaylaştıran ChatGPT kullanımının giderek yaygınlaştığı görülmektedir. Yapılan literatür incelemesinde ulusal yazında ChatGPT okuryazarlığını belirlemeye yönelik bir ölçeğin olmadığı bu sebeple ChatGPT'nin uluslararası literatürdeki yeri ve önemi dikkate alındığında bu çalışmanın motivasyon kaynağını ve de önemini ortaya koymaktadır. Bu çalışma ChatGPT okuryazarlık ölçeğinin Türkçe geçerlilik ve güvenilirliğini test etmek amaçlanmıştır. Ölçekte yer alan maddeler ile ilgili her bir madde için KGİ değeri 0,80 ile 1,00 arasındadır. KGİ hesaplamasında her bir madde için elde edilen değer 0,80 üzerinde olmalıdır (Çapık vd., 2018). Ölçek geneli KGİ değeri toplam uzman sayısına göre ≥0,56 üzeri olmalıdır (Yurdağül, 2005). Bu çalışma için ölçek geneli KGİ oranı 0,92'dir. KGİ sonuçları değerlendirildiğinde ölçmeyi amaçladığı özelliğin kavramsal alt yapısını tüm yönleriyle ölçtüğü ifade edilebilir (Kartal & Bardakçı, 2018).

Madde analizi için ölçek maddelerinde düzeltilmiş madde toplam korelasyon değerleri $r \geq 0,30$ olmalıdır (Leech, Barret and Morgan, 2015). ChatGPT okuryazarlık ölçeğinde düzeltilmiş madde toplam korelasyon değerleri 0,30 üzeridir. Analiz neticesinde ölçek maddelerinin birbirleri ile ilişkili olduğu söylenebilir (Karaman, 2023). Ölçeğin ölçmeye çalıştığı

yeterlilikleri ölçebilir anlamına geldiği ifade edilebilir. Taktak & Bafralı (2024) eğitimde ChatGPT kullanım ölçeğinde madde toplam korelasyon değerlerini 0,30 üzerinde olduğu görülmüştür. Benzer şekilde Batuk vd., (2025) tarafından yapılan ChatGPT kullanım ölçeğinde madde toplam korelasyon değerleri 0,30 üzerinde olduğu saptanmıştır.

Ölçek yapı geçerliliği AFA, DFA, yakınsak ve ıraksak geçerlilikleri olmak üzere tamamlanmıştır. AFA aşamalarında ilk kriter olarak KMO değerine dikkat edilmiştir. Bu çalışmada KMO değeri 0,906 olarak bulunmuştur. KMO değeri 0,80-1,00 arasında olursa örneklemin AFA için uygun olduğunu göstermektedir (Shrestha, 2021; Watkins, 2018), yine Bartlett testinin anlamlı olması da verilerin AFA iyi uygun olduğunu göstermektedir (Howard, 2016). Bir diğer önemli kriterde KMO değerinin 0,90 üzerinde olması bu değer mükemmel düzeyde olduğu yönünde ifade edilmektedir (Nikkhah vd., 2018; Marofi vd., 2020). Elde edilen değer literatür açısından uygun olduğu ifade edilebilir. Süleymanoğulları vd., (2024) çalışmasında yapay zeka okur yazarlığına yönelik geliştirmesinde KMO değerinin 0,89 olduğu görülmektedir Bu da elde edilen değer benzer ölçeklerle uyumlu ve literatürde belirtilen aralık seviyesinden yüksek düzeyde olduğunu göstermektedir.

Madde faktör yük değeri 0,50 üzerinde bir değer alması beklenmektedir. Faktör yük değerlerinin 0,50 üzeri olması modelin iyi tanımlanmış bir yapıda olduğunu göstermektedir (Hair et al., 2014). Bu çalışma için madde faktör yük değerleri ,651 ile ,862 arasında değiştiği, yüksek düzeyde olduğu ve ölçek modelinin iyi tanımlanmış bir yapıda olduğu söylenebilir. Süleymanoğulları ve diğerlerinin (2024) yapay zeka okuryazarlığına yönelik ölçek geliştirme çalışmasında madde faktör yükü 0,60 üzeri olduğu; Taktak & Bafralı'nın (2024) ChatGPT kullanım ölçeğinde ise madde faktör yükü 0,65 üzeri olduğu belirlenmiştir. Benzer şekilde elde edilen madde faktör yük değerlerinin literatür ile uyumlu olduğu söylenebilir.

AFA'da bir sonraki aşamada açıklanan varyans değerine dikkat edilmiştir. Bu çalışma için açıklanan varyans değeri %72,275'tir. Açıklanan varyansın literatürde %40-%60 arasında olması ilgili değer kabul edilebilir olduğunu göstermektedir (Çokluk vd., 2014). Çok faktörlü bir yapıda açıklanan varyans değeri en az 0,50 olmalıdır (Shrestha, 2021). Bu çalışmadaki elde edilen değer ile karşılaştırıldığında öğrencilerin ChatGPT okuryazarlık düzeyinin açıklanmasında oldukça iyi bir değer olduğu ifade edilebilir. AFA'da boyutlandırma neticesinde Lee ve Park (2024) çalışmasında olduğu gibi beş alt boyutta toplandığı ve bu çalışmada elde edilen açıklanan varyans değerinin orijinal ölçekteki değere göre yüksek düzeyde olduğu görülmüştür (Lee & Park, 2024).

AFA'dan sonra ölçek modelinin geçerliliğinin test edilmesinde DFA aşamasına geçilmiştir. DFA ayrı bir örnekleme (N=326) uygulanmıştır. DFA'da ölçek modeli oluşturulmuş ve model fit değerleri dikkate alınarak modelin geçerliliği test edilmiştir. Model fit değerlerinin literatürde hangilerinin alınması noktasında kesin bir görüş birliği olmasa da en sık kullanılan χ^2/df , RMR, GFI, AGFI, NFI, IFI, TLI, CFI ve RMSAE model fit değerleri kriter olarak esas alınmıştır (Kwon & Marzec, 2016; Karaman, 2023). Bu çalışmada DFA'da model uyum fit değerleri χ^2/df ; NFI; IFI; CFI; RMSEA; iyi uyum, RMR; GFI; AGFI; TLI değerlerinin kabul edilebilir uyum düzeyinde olduğu belirlenmiştir (Kline, 2011; Nikkhah vd., 2018).

DFA sonrasında ortaya çıkan yapının yakınsak ve ıraksak geçerlilik istatistikleri incelenmiştir. Yakınsak (uyum) geçerliliği aynı yapıyı temsil eden maddelerin (gözlenen değişkenlerin) birbiriyle ilgili olduğunu ve tek bir kavramsal yapıyı ölçtüğünü gösterir (Kartal & Bardakçı, 2018). Yakınsak geçerlilik kriteri için $CR > AVE$ ve $AVE > 0,50$ olmalıdır (Hair vd., 2011; Kartal & Bardakçı, 2018; Karaman, 2023). Bu çalışma için $AVE > 0,50$ ve $CR > AVE$ 'dir. Analiz sonuçları doğrultusunda yakınsak (uyum) geçerliliği sağlanmıştır. ıraksak (ayırma) geçerliliği çok faktörlü bir yapıda faktörlerin birbirinden bağımsız ve farklı yapıları ölçüp ölçmediğini inceler (Kartal & Bardakçı, 2018). ıraksak geçerlilik değerleri; $ASV < AVE$ ve $MSV < AVE$ olmalıdır (Nikkhah vd., 2018). Bu çalışmada ıraksak (ayırma) geçerlilik şartı ($ASV < AVE$ ve $MSV < AVE$) sağlanmıştır. Ayrıca $CR > 0,70$ olmalıdır. Bir diğer önemli kriterde eğer yakınsak geçerlilik için ilgili şart sağlanmazsa $AVE < 0,50$ fakat $CR > 0,60$ 'dan büyük olursa yakınsak geçerliliğin şartının sağlandığı yönünde yorumlanmaktadır (Fornell ve Larcker, 1981). Lee ve Park (2024) çalışmasında da benzer şekilde yakınsak geçerlilik istatistiklerinin sağlanmıştır. Bu çalışmada elde edilen

değerlerin orijinal ölçek ile uyumlu olduğu literatürde belirtilen yakınsak ve ıraksak geçerlilik istatistik değer aralıklarına göre yüksek düzeyde olduğu ifade edilebilir.

ChatGPT okuryazarlık ölçeği Likert tipinde olup güvenirlik katsayısı Cronbach Alpha ile incelenmiştir (Polit & Beck, 2012). Cronbach Alpha katsayısı 0,70-0,79 kabul edilebilir, 0,80-0,89 iyi, 0,90-1,00 mükemmel düzeyde değer aldığını göstermektedir (Sarmiento & Costa, 2019). Bu çalışmada ölçek geneli Cronbach Alpha katsayısının her iki örneklem için de 0,90 üzeri olduğu görülmüştür. Elde edilen değerlere göre uyarlama yapılan ölçeğin mükemmel düzeyde güvenilir olduğunu göstermektedir. Süleymanoğulları ve diğerlerinin (2024) yapay zeka ölçeklerinde Cronbach Alpha değeri 0,86; Karaoğlu Yılmaz ve Yılmaz (2023) çalışmasında Cronbach Alpha değeri 0,99; Akyürek (2025) çalışmasında yapay zeka okuryazarlığı ölçeği Cronbach Alpha değeri 0,83'tür. Bu çalışma da elde edilen Cronbach Alpha değerinin literatür ile uyumlu olduğu söylenebilir.

Ölçek test-tekrar test analizi ICC ve bağımlı örneklem t testi ile kararlılığı değerlendirilmiştir. Ölçek geneli ICC .947'dir. İki ölçüm arasında anlamlı ve pozitif yönlü bir ilişkinin olduğu bu doğrultuda ölçeğin zaman içinde değişmezliğini ortaya çıkmıştır (Polit, 2014; Koo and Li, 2016). Elde edilen korelasyon katsayısı değeri güvenirlik katsayısı olarak kabul edilmektedir. Diğer adı ile ölçeğin kararlılığını da belirttiği için kararlılık katsayısı olarak bilinmektedir. Kararlılık katsayısının yüksek olması (1'e yakın olması) ölçeğin ölçme sonuçlarının zaman içinde değişmediğini, dolayısıyla ölçeğin yüksek güvenirliğe sahip olduğunu göstermektedir (Kartal & Bardakçı, 2018). Bununla birlikte bağımlı örneklem t testi sonucunda da anlamlı bir fark olmadığı saptanmıştır. Yapılan analizler ve elde edilen bulgular dikkate alınarak ölçeğin kullanılabilecek bir ölçek olduğu söylenebilir.

ChatGPT okuryazarlığına yönelik Lee ve Park (2024) tarafından geliştirilen ölçeğin daha öncesinde ulusal literatürde uyarlamasının yapılamamış olması çalışmanın özgünlüğünü ortaya koymakla birlikte ele almış olduğu boyutlarla (eleştirel değerlendirme, teknik yeterlilik, iletişim yeterliliği, etik yeterlilik ve yenilikçi uygulama) yapay zeka alanında yapılan ölçeklere nazaran ChatGPT okuryazarlığını daha kapsamlı bir şekilde incelediği söylenebilir. Bu bakımdan benzer şekilde ulusal literatürde yapılan çalışmalarda örneğin Akyürek (2025) çalışmasında yapay zeka okuryazarlık ölçeği alt boyutlarının farkındalık, kullanım, değerlendirme ve etik olarak gruplandırıldığı görülmektedir. Yine Karaoğlu Yılmaz ve Yılmaz (2023) tarafından Türkçe uyarlaması yapılan

ölçek teknik anlama, eleştirel değerlendirme ve pratik uygulama olmak üzere üç alt boyutludur.

Çelebi vd., (2023) tarafından Türkçe uyarlaması yapılan ölçekte alt boyutları; farkındalık, kullanım, değerlendirme ve etik olmak üzere dört alt boyutludur. Mevcut yapılan çalışmanın alt boyutlarının ulusal literatür açısından uygunluk taşıdığı söylenebilir. Literatürde ChatGPT okuryazarlığına yönelik uyarlama bir ölçek çalışmasına rastlanılmamıştır. Literatürde yapay zeka okuryazarlığına yönelik çalışmaların olduğu görülmüştür (Karaoğlu Yılmaz ve Yılmaz 2023; Çelebi vd., 2023). Çelebi vd., (2023) tarafından Türkçe uyarlaması yapılan ölçekte bireyler üzerinde dahil edilmiş olup tüm öğrenim düzeylerinden kesimi kapsamaktadır. Bu çalışmada ise 18 yaş ve üzeri üniversite öğrencileri dahil edilmiş olup geleceğin bireyleri dahil edilmiştir. Uyarlama çalışması yapılan Lee ve Park (2024) tarafından geliştirilen ChatGPT okuryazarlık ölçeğinde üniversite öğrencileri kapsamında geliştirildiği görülmektedir. Bu bakımdan ölçek maddeleri dikkate alındığında ChatGPT okuryazarlığı ölçeği öğrenciler ve 18 yaş üzeri yetişkinler dahilinde uygulanabilir.

5. SONUÇ

ChatGPT okuryazarlık ölçeğinin geçerlilik ve güvenilirlik neticesinde 18 yaş üzeri bireylerde kullanılabılır beş faktörlü (eleştirel değerlendirme, teknik yeterlilik, iletişim yeterliliği, etik yeterlilik ve yenilikçi uygulama) bir ölçeğe ile sonuçlanmıştır. Bireylerin ChatGPT okuryazarlık düzeylerini değerlendirmede bu ölçek kabul edilebilir psikometrik özelliklere de sahip olduğunu göstermiştir. ChatGPT okuryazarlık ölçeği yetişkin bir kesime uygulanabileceği gibi belirli bir çalışan grubuna da uygulanabilir. Bununla birlikte bu ölçeğin farklı meslek gruplarında karşılaştırılması önerilebilir. Uyarlama çalışması yapılan bu ölçek 18 yaş üzeri üniversite öğrencileri dahilinde gerçekleştirilmiştir. Bu nedenle hizmet sektörünün farklı alanlarında çalışan (örneğin sağlık, eğitim, işletmeciler gibi) farklı örneklemelerde yeniden değerlendirilmesi önerilebilir. Bunun nedeni olarak farklı meslek gruplarında ChatGPT okuryazarlığının düzeylerini incelemek ve konunun algılanış biçimi bakımından meslekler arası ChatGPT farkındalığını ortaya koymak amaçlı olduğu söylenebilir.

ChatGPT okuryazarlık ölçeği mevcut teroillerin karşılaştırılması, literatürdeki teknoloji modellerine yönelik hipotez geliştirilmesinde, ChatGPT ile ilgili kavramların daha net ortaya konulması ve tanımlanmasında teorik olarak alana katkılar sağlayabilir. Ayrıca ek olarak ChatGPT ile ilgili yapılan

çalışmalarda bulguları karşılaştırmada bütüncül bir bakış açısı sağlayabilir. Pratik açıdan alana sağlayacağı katkılara bakıldığında uygulayıcılara günümüz teknolojilerine yönelik karar almada önemli düzeyde katkılar sağlayabilir. Eğitim alanında materyal geliştirme, örnek olay analizleri üretme ve farklı mesleklerde uygulama rehberler hazırlama da mesleki gelişim noktasında önemli düzeyde katkılar sunabilir. Günlük hayat veya iş akışlarında süreçlerin iyileştirilmesi, zamanı verimli kullanma ve de performansı arttırma noktasında önemli düzeyde faydalar sağlayabilir. Yapılan literatür incelemeğinde yerel örneklemde bireylerin ChatGPT okuryazarlık düzeylerini incelemeye yönelik bir ölçeğin olmayışı ve konunun derinlemesine daha öncesinde araştırılması gerektiği yönünde söylenebilir. Kavramsal ve psikometrik özellikleriyle bu çalışmanın literatürde önemli düzeyde katkı sağlayacağı düşünülmektedir.

ETİK BEYAN

Etik Kurul İzni: Bu çalışma için Sivas Cumhuriyet Üniversitesi Sosyal Bilimler Bilimsel Araştırma Önerisi Etik Değerlendirme Kurulu'ndan 24.04.2025 tarih ve 31 sayılı karar ile etik kurul izni alınmıştır.

Çıkar Çatışması: Yazarlar arasında çıkar çatışması bulunmamaktadır.

Finansman: Bu araştırma herhangi bir dış fon almamıştır.

Yazar Katkıları: Tüm yazarlar çalışmanın tasarımı, veri analizi ve makale yazımı süreçlerine eşit oranda katkı sağlamıştır.

KAYNAKLAR

- Akyürek, M. İ. (2025). Yapay zekâ okuryazarlığı ölçeğinin Türkçe uyarlaması: Geçerlik ve güvenilirlik çalışması (Eğitim paydaşlarına yönelik). *Gazi Üniversitesi Gazi Eğitim Fakültesi Dergisi*, 45(2), 649-663.
- Altın'ae, S., Sola-Leyva, A., Salumets, A., 2023. Artificial intelligence in scientific writing: a friend or a foe? Reproductive BioMedicine Online, in press. <https://doi.org/10.1016/j.rbmo.2023.04.009>
- Batuk, B., Türk, N., & Özmen, M. (2025). Psychometric Properties of the Turkish Version of the ChatGPT Usage Scale. *Siirt Sosyal Araştırmalar Dergisi*, 4(2), 47-60.
- Boubker, O. (2024). From chatting to self-educating: Can AI tools boost student learning outcomes? *Expert Systems with Applications*, 238, 121820. <https://doi.org/10.1016/j.eswa.2023.121820>
- Bouschery, S.G., Blazeovic, V., Piller, F.T., 2023. Augmenting human innovation teams with artificial intelligence:

- exploring transformer-based language models. *Journal of Product Innovation Management* 40 (2), 139–153.
- Clay, I. (2023). Fact of the week: OpenAI's ChatGPT user base has grown faster than TikTok's or Instagram's. Information Technology & Innovation Foundation. Retrieved February 9, 2024, from <https://itif.org/publications/2023/02/13/openai-chatgpt-user-base-has-grown-faster-than-tiktok-or-instagram/>
- Cooper, G. (2023). Examining science education in ChatGPT: An exploratory study of generative artificial intelligence. *Journal of Science Education and Technology*, 32, 444–452.
- Çapık, C., Gözümlü, S., & Aksayan, S. (2018). Intercultural Scale Adaptation Stages, Language and Culture Adaptation: Updated Guideline. *Florence Nightingale Journal of Nursing*. <https://doi.org/10.26650/fnjin397481>
- Çelebi, C., Yılmaz, F., Demir, U., & Karakuş, F. (2023). Artificial intelligence literacy: An adaptation study. *Instructional Technology and Lifelong Learning*, 4(2), 291–306.
- Çokluk Ö., Şekercioğlu, G. ve Büyüköztürk, Ş. (2014) *Sosyal bilimler için çok değişkenli istatistik: Spss ve lisrel uygulamaları*, Ankara: Pegem Akademi Yayıncılık, pp. 211–275.
- Escalante, J., Pack, A., & Barrett, A. (2023). AI-Generated feedback on writing: Insights into efficacy and ENL student preference. *International Journal of Educational Technology in Higher Education*, 20(1), 57. <https://doi.org/10.1186/s41239-023-00425-2>
- Farhi, F., Jeljeli, R., Aburezeq, I., Dweikat, F. F., Al-shami, S. A., & Slamene, R. (2023). Analyzing the students' views, concerns, and perceived ethics about chat GPT usage. *Computers and Education: Artificial Intelligence*, 5, 1–8.
- Fornell, C. ve Larcker, D. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50.
- Gao, C.A., Howard, F.M., Markov, N.S., Dyer, E.C., Ramesh, S., Luo, Y., Pearson, A.T., 2023. Comparing scientific abstracts generated by ChatGPT to real abstracts with detectors and blinded human reviewers. *NPJ Digital Medicine* 6 (1), 75.
- Gordijn, B., & Have, H. T. (2023). ChatGPT: Evolution or revolution? *Medicine Health Care and Philosophy*, 1–2.
- Grant, J. S. ve Davis, L. L. (1997). Selection and use of content experts for instrument development. *Research in Nursing & Health*, 20, 269–274.
- Gunawan, G., Jufri, A. W., Nisrina, N., Al-Idrus, A., Ramdani, A., ve Harjono, A. (2021). Guided inquiry blended learning tools (GIBL) for school magnetic matter in junior high school to improve students' scientific literacy. In *Journal of Physics: Conference Series*, 1747, 1–9
- Hair, J. F., Black, W. C., Babin, B.J. ve Anderson, R.E. (2014). *Exploratory factor analysis. Multivariate data analysis. (7th ed)*. Prentice Hall.
- Hair, JR, J. F., Ringle, C. M. ve Sarstedt, M. (2011). PLS-SEM: Indeed a Silver Bullet. *Journal of Marketing Theory and Practice*, 19(2), 139–151.
- Hernández, A., Hidalgo, M. D., Hambleton, R. K. ve Gomez-Benito, J. (2020). International test commission guidelines for test adaptation: A criterion checklist. *Psicothema*, 32(2), 390–398.
- Heung, Y. M. E., & Chiu, T. K. (2025). How ChatGPT impacts student engagement from a systematic review and meta-analysis study. *Computers and Education: Artificial Intelligence*, 8, 1–13.
- Hoffmann, S., Lasarov, W., & Dwivedi, Y. K. (2024). AI-empowered scale development: Testing the potential of ChatGPT. *Technological Forecasting and Social Change*, 205, 1–20.
- Homolák, J. (2023). Opportunities and risks of ChatGPT in medicine, science, and academic publishing: A modern Promethean dilemma. *Croatian Medical Journal*, 64(1), 1–3.
- Howard MC (2016) A review of exploratory factor analysis decisions and overview of current practices: What we are doing and how can we improve? *International Journal of Human-Computer Interaction* 32(1): 51–62.
- Karaman, M. (2023). Keşfedici ve Doğrulamalı Faktör Analizi: Kavramsal Bir Çalışma. *Uluslararası İktisadi ve İdari Bilimler Dergisi*, 9(1), 47–63. <https://doi.org/10.29131/uiibd.1279602>
- Karaoğlu Yılmaz, F. G., & Yılmaz, R. (2023). Adaptation of artificial intelligence literacy scale into Turkish. *Journal of Information and Communication Technologies*, 5(2), 172–190.
- Kartal, M., & Bardakçı, S. (2018). *SPSS ve AMOS uygulamalı örneklerle güvenirlik ve geçerlilik analizleri*. Ankara: Akademisyen Kitabevi.
- 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, Article 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Kline R.B. (2011). *Principles and Practice of Structural Equation Modeling*. New York: The Guilford Press.
- Koo, T. K., & Li, M. Y. (2016). A Guideline of Selecting and Reporting Intraclass Correlation Coefficients for Reliability Research. *Journal of Chiropractic Medicine*, 15(2). <https://doi.org/10.1016/j.jcm.2016.02.012>
- Kwon, Y. and Marzec, M. L. (2016). Does worksite culture of health (CoH) matter to employees? Empirical evidence using job-related metrics. *Journal of Occupational and Environmental Medicine*, 58(5), 448–454.

- Lee, S., & Park, G. (2024). Development and validation of ChatGPT literacy scale. *Current Psychology*, 43(21), 18992-19004.
- Leech, N.L., Barret, K.C., and Morgan, G.A. (2015). *IBMSPSS for Intermediate Statistics: Use and Interpretation. Fifth Edition*. New Jersey: Lawrence Erlbaum Associates, Inc
- Luo, S., & Zou, D. (2025). University Learners' Readiness for ChatGPT-Assisted English Learning: Scale Development and Validation. *European Journal of Education*, 60(1), 1-13.
- Macdonald, C., Adeloye, D., Sheikh, A., Rudan, I., 2023. Can ChatGPT draft a research article? An example of population-level vaccine effectiveness analysis. *Journal of Global Health* 13 (01003). <https://doi.org/10.7189/jogh.13.01003>
- Maral, S., Naycı, N., Bilmez, H., Erdemir, E. İ., & Satıcı, S. A. (2025). Problematic ChatGPT Use Scale: AI-Human Collaboration or Unraveling the Dark Side of ChatGPT. *International Journal of Mental Health and Addiction*, 1-27.
- Marofi, Z., Bandari, R., Heravi-Karimooi, M., Rejeh, N. ve Montazeri, A. (2020). Cultural adoption, and validation of the Persian version of the coronary artery disease education questionnaire (CADE-Q): a second-order confirmatory factor analysis. *BMC Cardiovascular Disorders*, 20, 1-9.
- Montenegro-Rueda, M., Fernández-Cerero, J., Fernández-Batanero, J. M., & Lopez-Meneses, E. (2023). Impact of the implementation of ChatGPT in education: A systematic review. *Computers*, 12(8), 153. <https://doi.org/10.3390/computers12080153>
- Nie, J., C. Zheng, P. Zeng, B. Zhou, L. Lei, and P. Wang. 2020. "Using the Theory of Planned Behavior and the Role of Social Image to Understand Mobile English Learning Check- In Behavior." *Computers & Education* 156: 103942.
- Nikkhah, M., Heravi-Karimooi, M., Montazeri, A., Rejeh, N. and Sharif Nia, H. (2018). Psychometric properties the Iranian version of older People's quality of life questionnaire (OPQOL). *Health and Quality of Life Outcomes*, 16, 1-10. <https://doi.org/10.1186/s12955-018-1002-z>
- OpenAI. 2024. "ChatGPT (GPT- 4, September version) [Multimodal Large Language Model]." <https://openai.com/ChatGPT>
- Pan, Y., Y. Huang, H. Kim, and X. Zheng. 2023. "Factors Influencing Students' Intention to Adopt Online Interactive Behaviors: Merging the Theory of Planned Behavior With Cognitive and Motivational Factors." *Asia- Pacific Education Researcher* 32: 27–36.
- Pellas, N. (2023). The effects of generative AI platforms on undergraduates' narrative intelligence and writing self-efficacy. *Education Sciences*, 13(11), 1155. <https://doi.org/10.3390/educsci13111155>
- Polit, D. F. (2014). Getting serious about test-retest reliability: A critique of retest research and some recommendations. *Quality of Life Research*, 23(6), 1713-1720. <https://doi.org/10.1007/s11136-014-0632-9>
- Polit, D.F., & Beck, C.T. (2012). *Nursing Research Principles and Method, 6th ed.* Philadelphia, PA: Lippincott Williams & Wilkins.
- Qadir, J. (2022). Engineering education in the era of ChatGPT: Promise and pitfalls of generative AI for education. *TechRxiv*. <https://doi.org/10.36227/techrxiv.21789434.v1>
- Reed, B. (2023). ChatGPT reaches 100 million users two months after launch. The Guardian. Retrieved February 9, 2024, from <https://www.theguardian.com/technology/2023/feb/02/chatgpt-100-million-users-open-ai-fast-test-growing-app>
- Rossi, S., Rossi, M., Mukkamala, R. R., Thatcher, J. B., Dwivedi, Y. K. (2024). Augmenting research methods with foundation models and generative AI. *International Journal of Information Management*, 102749, doi: <https://doi.org/https://doi.org/10.1016/j.ijinfomgt.2024.102749>
- Sarmiento, R. P., & Costa, V. (2019). Confirmatory factor analysis--a case study. <https://www.semanticscholar.org/reader/a4e19cbb6cc766ad197935d2fc4974bb425e2698>
- Šedlbauer, J., Činčera, J., Slavík, M., & Hartlová, A. (2024). Students' reflections on their experience with ChatGPT. *Journal of Computer Assisted Learning*, 40(4), 1526-1534.
- Seghier, M.L. (2023). Using ChatGPT and other AI-assisted tools to improve manuscripts readability and language. *International Journal of Imaging Systems and Technology*. <https://doi.org/10.1002/ima.22902>
- Shrestha, N. (2021). Factor analysis as a tool for survey analysis. *American Journal of Applied Mathematics and Statistics*, 9(1), 4-11.
- Sousa, V. ve Rojjanasrirat, W. (2011). Translation, adaptation and validation of instruments or scales for use in cross-cultural health care research: A clear and user-friendly guideline. *Journal of Evaluation in Clinical Practice*, 17(2), 268-274.
- Susarla, A., Gopal, R., Thatcher, J.B., & Sarker, S. (2023). The Janus effect of generative AI: charting the path for responsible conduct of scholarly activities in information systems. *Information Systems Research* 34 (2), 399-408.
- Süleymanoğulları, M., Özdemir, A., & Tekin, A. (2024). Yapay zekâ ölçeği: Geçerlik ve güvenirlik çalışması. *Education Science and Sports*, 6(1), 13-27.
- Taktak, M., & Bafralı, G. (2025). ChatGPT Usage Scale in education: Validity and reliability study. *International Journal of Technology in Education*, 8(1), 193-207.
- Valova, I., Mladenova, T., & Kanev, G. (2024). Students' Perception of ChatGPT Usage in Education. *International Journal of Advanced Computer Science & Applications*, 15(1), 466-473.
- Venturebeat, 2023. Harnessing the power of GPT-3 in scientific research, Venturebeat, <https://venturebeat.com/ai/harnessing-the-power-of-gpt-3-in-scientific-research/>, accessed May 21th, 2023.

Watkins, M. W. (2018). Exploratory factor analysis: A guide to best practice. *Journal of Black Psychology*, 44(3), 219–246

Yu, S. C., Chen, H. R., & Yang, Y. W. (2024). Development and validation the Problematic ChatGPT Use Scale: a preliminary report. *Current Psychology*, 43(31), 26080-26092.

Yurdugül, H. (2005). Ölçek geliştirme çalışmalarında kapsam geçerliği için kapsam geçerlik indekslerinin kullanılması, XIV. Ulusal Eğitim Bilimleri Kongresi, Pamukkale Üniversitesi Eğitim Fakültesi, 20-22 Ocak 2025 tarihinde <http://yunus-hacettepe.edu.tr/~yurdugul/3/indir/PamukkaleBildiri.pdf> adresinden alınmıştır.