



BANDIRMA ONYEDİ EYLÜL ÜNİVERSİTESİ SAĞLIK BİLİMLERİ VE ARAŞTIRMALARI DERGİSİ BANU Journal of Health Science and Research

DOI: 10.46413/boneyusbad.1472077

Özgün Araştırma / Original Research

CHATGPT 3.5, CHATGPT 4.0 ve Hemşirelik Öğrencilerinin Çocuk Acillerde Hemşirelik Yaklaşımı Dersi Sınavındaki Performans Karşılaştırmaları

Comparison of CHATGPT 3.5, CHATGPT 4.0 and Nursing Students' Performance in the Nursing Approach in Child Emergencies Course Examination

Ayla İrem AYDIN¹ Doğukan REİS²

¹ Dr. Öğr. Üyesi, Uludağ Üniversitesi, Sağlık Bilimleri Fakültesi, Hemşirelik Bölümü, Bursa

² Hemşirelik Lisans Öğrencisi, Uludağ Üniversitesi, Sağlık Bilimleri Fakültesi, Hemşirelik Bölümü, Bursa

Sorumlu yazar /
Corresponding author

Ayla İrem AYDIN

iremaydin@uludag.edu.tr

Geliş tarihi / Date of receipt:
22.04.2024

Kabul tarihi / Date of
acceptance: 08.02.2025

Atıf / Citation: Aydın, A. İ., Reis, D. (2025). ChatGPT 3.5, ChatGPT 4.0 ve hemşirelik öğrencilerinin çocuk acillerde hemşirelik yaklaşımı dersi sınavındaki performans karşılaştırmaları. *BANU Sağlık Bilimleri ve Araştırmaları Dergisi*, 7(1), 73-79. doi: 10.46413/boneyusbad.1472077

*Bu çalışma sözel bildiri olarak 20. Ulusal Uludağ Pediatri Kış Kongresinde sunulmuş ve özet metni kongre kitapçığında yayınlanmıştır.

ÖZET

Amaç: Bu çalışmanın amacı, ChatGPT modelleri ve öğrenci hemşirelerin çocuk acillerde hemşirelik yaklaşımı dersindeki başarı performanslarını karşılaştırmaktır.

Gereç ve Yöntem: Bu çalışma retrospektif ve karşılaştırmalı bir analiz olarak tasarlanmıştır. Çalışma bir devlet üniversitesinin hemşirelik bölümünde çocuk acillerde hemşirelik yaklaşımı dersinin yarıyıl sonu sınav soruları kullanılarak yapılmıştır. Sınavda yer alan beş seçenekli çoktan seçmeli 40 soru, ChatGPT 3.5 ve ChatGPT 4.0 modellerine yöneltilerek yanıtlanmıştır. Sorulara verilen cevaplar yapay zeka robotlarının verdikleri cevaplarla ve dersi alan öğrenci hemşirelerin verdikleri cevapların ortalamalarıyla karşılaştırılmıştır.

Bulgular: ChatGPT 4.0 sorulara %90 oranında doğru yanıt verirken, öğrenciler %82.32 oranında doğru yanıt vermiştir. En düşük doğru yanıt oranı %55 ile ChatGPT 3.5 modelindedir. Bu sınavda en yüksek başarıyı sırasıyla ChatGPT 4.0, öğrenciler ve ChatGPT 3.5 elde etmiştir ($p < 0.05$).

Sonuç: ChatGPT 4.0, çocuk acillerde hemşirelik yaklaşımı konusunda sağladığı doğru önerilerle iyi bir performans sergilemiştir. Yapay zeka temelli sistemler hemşirelik eğitiminde dersin öğretim elemanına ve öğrencilerine yardımcı bir referans aracı olarak kullanılabilir. Bu teknolojilerin hemşirelik eğitiminde nasıl daha etkin kullanılabileceğine yönelik stratejiler geliştirmek ve çalışmalar yapmak önemlidir.

Anahtar Kelimeler: Büyük dil modelleri, ChatGPT, Çocuk acil hemşireliği, Yapay zeka

ABSTRACT

Aim: The aim of this study was to compare the success performances of ChatGPT models and student nurses in the course of nursing approach in pediatric emergencies.

Material and Method: This study was designed as a retrospective and comparative analysis. The study was conducted by using the end-of-semester exam questions of the nursing approach in pediatric emergencies course in the nursing department of a state university. In the exam, 40 multiple-choice questions with five options were directed to ChatGPT 3.5 and ChatGPT 4.0 models. The answers given to the questions were compared with the answers given by the artificial intelligence robots and the averages of the answers given by the student nurses taking the course.

Results: ChatGPT 4.0 gave 90% correct answers to the questions, while the students gave 82.32% correct answers. The lowest correct response rate was 55% for ChatGPT 3.5. In this exam, ChatGPT 4.0, students and ChatGPT 3.5 achieved the highest success respectively ($p < 0.05$).

Conclusion: ChatGPT 4.0 performed well with the correct suggestions it provided on nursing approach in pediatric emergencies. Artificial intelligence-based systems can be used in nursing education as a reference tool to assist the instructor and students of the course. It is important to develop strategies and conduct studies on how these technologies can be used more effectively in nursing education.

Keywords: Large language models, ChatGPT, Pediatric emergency nursing, Artificial intelligence



This work is licensed under a Creative Commons Attribution-Non Commercial 4.0 International License.

GİRİŞ

ChatGPT (Generative Pre-trained Transformer), metin tabanlı sohbetlerde doğal dil anlama ve üretme yeteneklerini kullanarak etkileşimli bir deneyim sunan bir dil modeli uygulamasıdır. OpenAI firması tarafından geliştirilmiştir. Firma geliştirdiği ChatGPT 3.5 modelini 2021 yılında, ChatGPT 4.0 ise Mart 2023'de piyasaya sürmüştür. Kullanıcılar, herhangi bir konuda bilgi almak, yaratıcı yazılar oluşturmak, metin tabanlı oyunlar oynamak veya sadece sohbet etmek için ChatGPT'yi kullanabilir (OpenAI, 2023; Sanchez-Ramos, Lin ve Romero, 2023). ChatGPT 4.0, ChatGPT 3.5'e göre daha çok fonksiyonel yeteneklere sahiptir. ChatGPT 4.0 bu yönüyle tıp dünyasında şimdiden birçok çalışmanın ana noktası olmayı başarmıştır (Handa, Chhabra, Goel ve Krishnan, 2023; Huang ve ark., 2023a; Huang ve ark., 2023b).

Çocuk acil hemşireliği, pediatrik hastaların acil durumlarıyla ilgilenen özel bir sağlık dalıdır (Hockenberry, Wilson ve Rodgers, 2021). Bu alanda çalışan hemşireler, pediatrik triyaj, pediatrik travma, resüsitasyon, tanılama ve hemşirelik girişimleri gibi konulara hakimdirler (Moore ve ark., 2020). Hızlı karar alma yeteneği, çocuk acil hemşiresinin acil durumların dinamik doğasına uyum sağlamasını sağlar. Zira çocuklar, ani sağlık sorunlarına karşı daha hassas bir tepki gösterebilirler (Conlon ve ark., 2021). Hemşireler acil durumlarda doğru hemşirelik tanısı koymak, tedavi planlarını uygulamak ve gerektiğinde ekiplerle etkili iletişim kurmak konusunda hızlı ve kararlı olmalıdır. Bunun yanı sıra çocuk acil servislerinde çalışan hemşireler hem çocuklarla hem de aileleriyle etkili bir şekilde iletişim kurmayı öğrenmelidir (Barata ve ark., 2020; Sahyoun ve ark., 2021; Walker ve Tolentino, 2020).

Hemşirelikte yapay zeka destekli teknolojilerin kullanımı artarak devam etmektedir (Aydın ve Özyazıcıoğlu, 2023; von Gerich ve ark., 2022). ChatGPT gibi yapay zeka destekli araçlar öğrenme deneyimlerini geliştirmek ve öğrencilerin akademik performanslarını iyileştirmelerine yardımcı olmak için kullanılmaktadır (Alkhaqani, 2023). Yapay zeka destekli dil modeli olan ChatGPT çeşitli alanlarda olduğu gibi sağlık bilimleri alanında da kullanılmaya başlanmıştır. Japonya ulusal hemşirelik sınavında ChatGPT başarılı olup sınavı geçmiştir (Taira, Itaya ve Hanada, 2023). Japonya'da yapılan bir başka çalışmada ChatGPT

4.0'ün tıp asistanlarının girdiği tıp eğitimi sınavında tıp asistanlarından yüksek puan aldıkları belirlenmiştir (Watari ve ark., 2023). Bir başka çalışmada ülkemizdeki tıp fakültesi öğrencileriyle yapılmıştır. Bu sınavda en yüksek başarı puanını sırasıyla ChatGPT 4.0, ChatGPT 3.5 ve öğrenciler elde etmiştir (Asker, Özgür, Eriç ve Bekiroğlu, 2023). Tayvan'da yapılan hemşire lisans sınavında ChatGPT'nin sınavda başarılı olduğu ve yüksek düzeyde uyum gösterdiği belirtilmiştir (Huang, 2023). Hemşirelik eğitiminde yapılan bir diğer çalışmada Çin'in ulusal hemşirelik lisanslama sınavındaki çeşitli dil modellerinin performansı karşılaştırılmış ve ChatGPT'nin diğer dil modellerinden daha yüksek oranda doğruluk oranı gösterdiği bulunmuştur (Ni, Peng, Zheng ve Xie, 2024). Türkiye'deki bir devlet üniversitesinin Sağlık Bilimleri Fakültesi öğrencileriyle yapılan 40 soruluk anatomi testinde ChatGPT'nin daha iyi performans gösterdiği tespit edilmiştir (Talan ve Kalinkara, 2023). Çalışmalarda da görüldüğü gibi ChatGPT'nin kullanımı giderek yaygınlaşmaktadır. Hemşirelik eğitiminin geleceğini öngörmek ve ona hazırlanmak için, ChatGPT gibi büyük dil modelleriyle ilişkili çalışmalar yapmak ve bu dil modellerinin potansiyel faydaları ve riskleri kapsamlı bir şekilde değerlendirmek esastır (Huang, 2023). Bu çalışma, hemşirelik eğitiminde yapay zeka destekli dil modellerinin etkinliğini değerlendirmektedir. ChatGPT modelleri çeşitli tıbbi ortamlarda test edilmiş olsa da çocuk acil hemşireliğine ilişkin çalışmalarda henüz test edilmemiştir. Dolayısıyla bu çalışma alanında özgün bir çalışmadır. Bununla birlikte bu çalışma dil modellerinin çocuk hemşireliği alanında getirdiği avantajlar ve sınırlamalara ilişkin veriler sunmasıyla literatüre yeni bir bakış açısı kazandırabilir. Bu kapsamda, araştırma sonuçlarının büyük dil modellerinin eğitim ve klinik uygulamalardaki potansiyelini daha iyi anlamamıza yardımcı olması beklenmektedir. Öğrenci ile performans karşılaştırmalı çalışma olması modellerin öğrencilere göre hangi alanlarda üstün veya eksik olduğunu belirlemede de yardımcı olabilir. Bu çalışmanın amacı ChatGPT modelleri ve öğrenci hemşirelerin çocuk acillerde hemşirelik yaklaşımı dersindeki başarı performanslarını karşılaştırmaktır.

GEREÇ VE YÖNTEM

Araştırmanın Türü (Tasarımı)

Bu çalışma, retrospektif ve karşılaştırmalı bir araştırma niteliğindedir.

Araştırmanın Evreni ve Örneklemi

Bu çalışma, Şubat 2024 tarihinde bir devlet üniversitesinin Sağlık Bilimleri Fakültesi Hemşirelik Bölümünde yürütülmüştür. Çalışmada 2022-2023 yılı eğitim-öğretim döneminde çocuk acillerde hemşirelik yaklaşımı dersini alan öğrencilerin (28 öğrenci) yarıyıl sonu sınavlarının soruları çalışma örnekleme alınmıştır.

Veri Toplama Araçları

Bu çalışmada çocuk acillerde hemşirelik yaklaşımı dersine ait yarıyıl sonu sınav soruları kullanılmıştır. Çocuk acillerde hemşirelik yaklaşımı dersi yarıyıl sonu sınavının soruları beş şıklı seçmeli sorulardan oluşan toplam 40 soruyu içermiştir. Sorular ilgili dersin müfredatı kapsamında ders yürütücüsü tarafından hazırlanmıştır. Sorular aynı prompt kullanılarak Türkçe (Tr) ve İngilizce (Ing) dillerinde ayrı ayrı ChatGPT 3.5 ve ChatGPT 4.0 modellerine sorulmuştur. Yapay zekanın pekiştirerek öğrenmesini engellemek için her soruda yeni bir internet sayfası (<https://chat.openai.com/>) kullanılmıştır. Soruların tamamı bittikten sonra ChatGPT 4.0 ve ChatGPT 3.5'un cevapları puanlandırılıp hesaplanmıştır. Sonuçlar hem yapay zeka robotlarının kendi içinde hem de çocuk acillerde hemşirelik yaklaşımı dersini alan öğrenci hemşirelerin başarı ortalamalarıyla karşılaştırılmıştır.

Araştırmanın Etik Yönü

Bu çalışmanın yürütülmesi için bir devlet üniversitesinin Sağlık Bilimleri Araştırma ve Yayın Etik Kurulu'ndan onay (Tarih: 24.01.2024 ve Karar No: 2024-01) ve çalışmanın yapılacağı kurumundan izin alınmıştır.

Verilerin Değerlendirilmesi

Bu çalışmadan elde edilen verilerin istatistiksel analizi IBM SPSS Statistics 28.0 yazılımı kullanılarak gerçekleştirilmiştir. Çalışmada sınav sorularına verilen yanıtların dağılımını incelemek için yanıtların sayısı (n), minimum ve maksimum değerleri (min-max), ortalama ve standart sapma (ortalama $\pm$ standart sapma) hesaplanmıştır. Örneklem dağılımının normal olduğunu doğrulamak için Shapiro-Wilk testi kullanılmış ve

verilerin normal dağılıma uygun olduğu tespit edilmiştir. Öğrencilerin ve ChatGPT modellerinin verdiği yanıtların ortalamaları arasındaki farklılıkların belirlenmesi amacıyla örneklemin normal dağılıma uygunluğu ve verilerin bağımsızlığı sağlanarak one-sample t-testi uygulanmıştır.

BULGULAR

Çalışmada sınav sorularına verilen cevapların tanımlayıcı verileri Tablo 1'de gösterildi. Öğrencilerin sınav sorularına verdikleri doğru yanıtların ortalaması 82.32 ± 16.30 idi. ChatGPT 4.0, Tr ve Ing sorulan sorulara %90 oranında doğru yanıt verdi. ChatGPT 3.5 ise Ing sorulan sorulara %57.5 oranında doğru yanıt verdi.

Tablo 1. Sınav Sorularına Verilen Cevapların Tanımlayıcı Verileri

	Sayı (n)	Min-Max	Ortalama $\pm$ Standart Sapma
Öğrenci	28	30-100	82.32 $\pm$ 16.30
ChatGPT 4.0 (Tr)	1	90	90
ChatGPT 4.0 (Ing)	1	90	90
ChatGPT 3.5 (Tr)	1	55	55
ChatGPT 3.5 (Ing)	1	57.5	57.5

Öğrencilerle ChatGPT modellerinin yanıtları arasındaki karşılaştırma sonuçları Tablo 2'de verildi. ChatGPT 4.0 modelinin sorulara verdiği doğru cevapların ortalaması, öğrencilerin sorulara verdikleri doğru cevapların ortalamasından anlamlı derecede daha yüksekti ($p<0.05$). Öğrencilerin verdiği doğru cevap ortalamalarıyla ChatGPT 3.5 modelinin doğru cevap ortalaması karşılaştırıldığında, öğrencilerin anlamlı derecede daha yüksek oranda doğru cevap verdiği tespit edildi ($p<0.01$). Bu sonuçlara göre başarı sıralaması ChatGPT 4.0, öğrenciler ve ChatGPT 3.5 olarak bulundu.

Şekil 1'de ChatGPT 4.0 modelinin Türkçe sorulan sınav sorusuna verdiği doğru yanıtın ekran görüntüsü verilmiştir. Burada ChatGPT 4.0, travma gelişen çocukta Glaskow koma skalası skorlama sorusunu doğru yanıtlamıştır.

ChatGPT 4.0 modelinin İngilizce sorulan sınav sorusuna verdiği yanlış yanıt Şekil 2'de verilmiştir. Burada Hemşirelik Yönetmeliğinde yer alan hemşirenin rol ve sorumluluğuna ilişkin

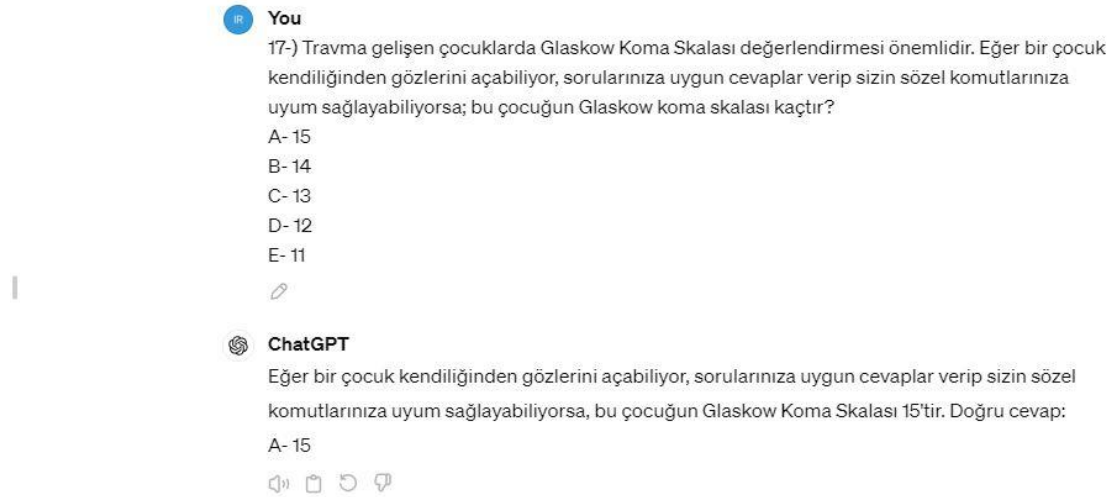
soruya ChatGPT 4.0 modelinin yanlış yanıt verdiği görülmektedir.

Tablo 2. Öğrencilerle ChatGPT Yanıtları Arasındaki Karşılaştırma Sonuçları

Karşılaştırma	t†	df	P değeri	Ortalama Fark
Öğrenci vs ChatGPT-4.0 (Tr)	-2.493	27	0.019	-7.67
Öğrenci vs ChatGPT-4.0 (Ing)	-2.493	27	0.019	-7.67
Öğrenci vs ChatGPT-3.5 (Tr)	8.869	27	<0.01‡	27.32
Öğrenci vs ChatGPT-3.5 (Ing)	8.058	27	<0.01‡	24.82

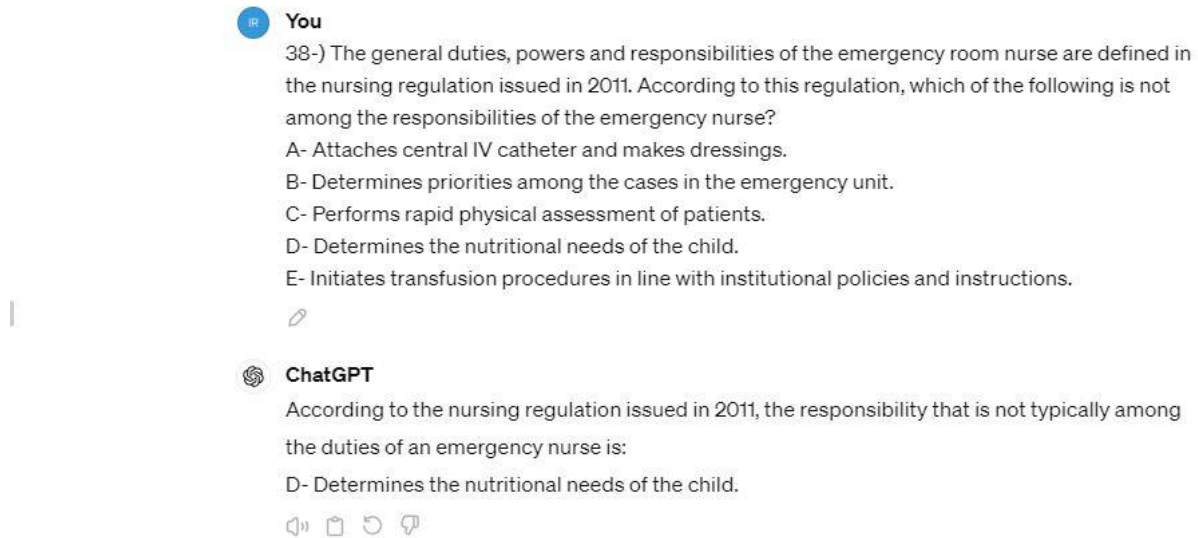
t†: one-sample t-testi; ‡ P <0,01

ChatGPT 4 ▾



Şekil 1: ChatGPT 4.0 (Tr) Doğru Cevaplanan Yanıt Örneği

ChatGPT 4 ▾



Şekil 2: ChatGPT 4.0 (Ing) Yanlış Cevaplanan Yanıt Örneği

TARTIŞMA

Tıp literatüründe oldukça yeni bir araştırma konusu olan ChatGPT gibi dil modellerinin sağladıkları tıbbi bilgilerin doğruluğu ve

güvenilirliği konusundaki araştırmalar artarak devam etmektedir (Chow ve ark., 2023). Bu çalışmada ChatGPT modelleri ve öğrenci hemşirelerin çocuk acil hemşireliği yaklaşımı dersindeki başarı performansları

karşılaştırılmıştır. Bu sınav soruları kapsamında çocuk acil hemşireliğinde hastaya yaklaşım, pediatrik ilaç uygulamaları, temel ve ileri yaşam desteği gibi konulara ilişkin sorulara yer verilmiştir. Bu çalışmada ChatGPT 4.0 modeli sorulara %90 oranında doğru yanıt vermiştir. ChatGPT 4.0 'ün başarı performansı ChatGPT 3.5 ve öğrencilere göre anlamlı düzeyde daha yüksek bulunmuştur ($p<0.05$). Bizim çalışmamıza benzer bir çalışma ülkemizdeki tıp fakültesi öğrencileriyle yapılmıştır. 245 tıp fakültesi öğrencisinin biyoistatistik sınavındaki başarıları ChatGPT 4.0 ve ChatGPT 3.5 ile karşılaştırılmıştır. Bu sınavda en yüksek başarı puanını sırasıyla ChatGPT 4.0, ChatGPT 3.5 ve öğrenciler elde etmiştir (Asker, Özgür, Eriç ve Bekiroğlu, 2023). Japonya'da yapılan bir çalışmada ChatGPT 4.0'ün tıp asistanlarının girdiği tıp eğitimi sınavındaki başarı oranı incelenmiştir. Yapılan sınavda asistanlar %55.8, ChatGPT 4.0 ise %70.1 doğruluk yüzdesine ulaşmıştır. Çalışmada ChatGPT 4.0'ün puanlarının asistanlara göre anlamlı derecede yüksek olduğu bulunmuştur (Watari ve ark., 2023). Bir başka çalışmada Tayvan'da yapılmış ve ChatGPT 4.0, Tayvan Ulusal Hemşirelik Sınavında %80.75 oranında genel doğruluk performansına ulaşmıştır (Su, Lin, Lin ve Chen, 2024). Çalışma sonuçları bizim çalışma sonuçlarımızla benzerdir. ChatGPT 4.0 sağlık bilimleri alanında yapılan çalışmalarda başarılı olmuştur. Bununla birlikte yapay zeka robotları ve öğrencilerle karşılaştırıldığında yüksek performans göstermiştir.

Genel olarak, dil modellerinde yapay zekanın İngilizce'de diğer dillere göre daha iyi performans gösterdiği bilinmektedir. Bu çalışmada ChatGPT 4.0 modeli hem İngilizce hemde Türkçe sorulan sorularda aynı oranda doğru yanıt verirken, ChatGPT 3.5 İngilizce sorulan sorularda daha yüksek oranda doğru yanıt vermiştir. Ancak bu fark anlamlı bulunmamıştır ($p>0.05$). Çalışmamıza benzer olarak ChatGPT 3.5 ve ChatGPT 4.0 modellerinin dermatoloji uzmanlık sınavında başarı oranlarının karşılaştırıldığı çalışmada, sorular İngilizce ve Lehçe dillerinde ChatGPT modellerine sorulmuştur. ChatGPT İngilizce sorulan sorulardan daha yüksek puan almasına rağmen bu fark anlamlı bulunmamıştır (Lewandowski, Łukowicz, Świetlik ve Barańska-Rybak, 2023). Literatürde yer alan bir çalışmada Japoncanın kullanıldığı Japon Ulusal Hemşirelik sınavındaki sorulara yüksek oranda doğru yanıtlar verilmiştir (Watari ve ark., 2023). Bu sonuçlar

çalışma bulgularımızla benzerdir ve ChatGPT'nin İngilizce dışındaki dillerde de iyiye yakın performans gösterdiğini bildirmektedir.

Dil modellerinin çalışma sorularına verdikleri yanıtlar incelendiğinde; ChatGPT 4.0'ün yanıt veremediği sorulardan birisi Hemşirelik Yönetmeliği içerisinde yer alan hemşirelik rol ve sorumluluklarına ilişkin soru olmuştur. ChatGPT modellerinin her ikisi de hem İngilizce hemde Türkçe olarak sorulan bu soruya yanlış cevap vermiştir. ChatGPT modellerinin bu soruya yanlış cevap verme nedenleri dil çeviri hataları veya İngilizce dışındaki dillerde karşılaştığı zorluklar olabilir. Çünkü bu modellerin eğitim verilerinin çoğunluğu İngilizce metinlerden oluşmaktadır. Bu durum ChatGPT modelleri için sınırlılıkları oluşturabilir. Bu sınırlamalar bu dil modellerinin yanlış yanıt vermesine yol açabilir.

Araştırmanın Sınırlılıkları

Bu çalışmada dil modellerinin başarı performansını değerlendirmede sınav sorularının hazırlığı bakımından sınırlılıklar mevcuttur. Bu çalışmada kullanılan soruların sayısının azlığı ve sadece ders yürütücüsü tarafından hazırlanması gibi durumlar çalışmanın sınırlılıklarındandır. Çalışmada seçenekli soruların kullanılması ChatGPT'nin açık uçlu sorulara verebileceği yanıtları incelemeye ve performansını değerlendirmede kısıtlılık oluşturmaktadır.

Bu çalışma sadece ChatGPT dil modelinin performansını incelemeye odaklanmıştır. Diğer dil modellerine yer verilmemesi çalışmanın sınırlılıkları arasındadır. Bu çalışma yakın zamanda geliştirilen alternatif dil modellerinin sağlık bakımı alanındaki potansiyellerini gözden kaçırabilir. Bu sebeple gelecekteki çalışmalarda kullanılacak soruların kapsamlı bir şekilde hazırlanması ve bu çalışmalarda alternatif büyük dil modellerinin kullanılması önerilir.

SONUÇ

Bu çalışma literatürde büyük dil modellerinin çocuk acil hemşireliği yaklaşımı dersindeki başarı performanslarını değerlendiren ilk çalışmadır. ChatGPT araştırmacılar tarafından sağlanan seçmeli sorulara insan benzeri yanıtlar sağlamıştır. Bu çalışma sonucunda ChatGPT 4.0 dil modeli, ChatGPT 3.5 ve hemşirelik öğrencilerinden daha yüksek başarı puanı elde etmiştir. Bu sonuca göre yapay zeka temelli sistemler hemşirelik eğitiminde dersin öğretim elemanına ve öğrencilerine yardımcı bir referans aracı olarak

kullanılabilir. Bu teknolojilerin hemşirelik eğitiminde nasıl daha etkin kullanılabileceğine yönelik stratejiler geliştirmek ve çalışmalar yapmak önemlidir. Gelecekteki çalışmalar için; ChatGPT gibi dil modellerini hemşirelik derslerine entegre etmek, çeşitli araçlar ve yöntemler aracılığıyla sınırlamalarını, etkinliğini değerlendirmek, farklı klinik durumlarda, pediatri hemşireliğindeki uygulamalarda kullanım olanaklarını değerlendirmek önerilir.

Araştırmanın Etik Yönü / Ethics Committee Approval

Araştırma için Uludağ Üniversitesi Sağlık Bilimleri Araştırma ve Yayın Etik Kurulu'ndan etik onayı alınmıştır (Tarih: 24.01.2024 ve Karar No: 2024-01).

Yazar Katkısı / Author Contributions

Fikir/Kavram: A.İ.A., D.R., Tasarım: A.İ.A., D.R., Denetleme/Danışmanlık: A.İ.A., Analiz ve/veya Yorum: A.İ.A., D.R., Kaynak Taraması: A.İ.A., D.R., Makalenin Yazımı: A.İ.A., D.R., Eleştirel İnceleme: A.İ.A.

Hakem Değerlendirmesi / Peer-review

Dış bağımsız

Çıkar Çatışması / Conflict of Interest

Yazarlar araştırmanın yürütülmesinde herhangi bir çıkar çatışması olmadığını beyan etmişlerdir.

Finansal Destek / Financial Disclosure

Yazarlar araştırmanın yürütülmesi sürecinde bir finansal destek almadıklarını beyan etmiştir.

KAYNAKLAR

- Alkhaqani, A. L. (2023). ChatGPT and academic integrity in nursing and health sciences education. *Journal of Medical Research and Reviews*, 1(1), 5-8. doi: 10.5455/JMRR.20230624044947
- Asker, Ö. F., Özgür, E. G., Eriç, A., Bekiroğlu, N. (2023). Comparing the performance of medical students, chatgpt-3.5 and chatgpt-4 in biostatistics exam: pros and cons as an education assistant: a cross-sectional research. *International Journal of Management Information Systems and Computer Science*, 7(2), 85-94. doi: 10.33461/uybisbbd.1329650
- Aydın, A. İ., Özyazıcıoğlu, N. (2023). Assessment of postoperative pain in children with computer assisted facial expression analysis. *Journal of Pediatric Nursing*, 71, 60-65. doi:10.1016/j.pedn.2023.03.008
- Barata, I. A., Stadnyck, J. M., Akerman, M., O'Neill, K., Castaneda, J., Subramony, A., ... D'Angelo, J. (2020). Novel approach to emergency departments' pediatric readiness across a health system. *Pediatric Emergency Care*, 36(6), 274-276. doi:10.1097/PEC.0000000000001385.
- Chow, J. C. L., Sanders, L., Li, K. (2023). Impact of

ChatGPT on medical chatbots as a disruptive technology. *Frontiers in Artificial Intelligence*, 6, 1166014. doi: 10.3389/frai.2023.1166014

- Conlon, C., McDonnell, T., Barrett, M., Cummins, F., Deasy, C., Hensey, C., ...Nicholson, E. (2021). The impact of the COVID-19 pandemic on child health and the provision of Care in Paediatric Emergency Departments: a qualitative study of frontline emergency care staff. *BMC Health Services Research*, 21(1), 279. doi:10.1186/s12913-021-06284-9
- Handa, P., Chhabra, D., Goel, N., Krishnan, S. (2023). Exploring the role of ChatGPT in medical image analysis. *Biomedical Signal Processing and Control*, 86, 105292. doi: 10.1016/j.bspc.2023.105292
- Hockenberry, M. J., Wilson, D., Rodgers, C. C. (2021). *Wong's essentials of pediatric nursing-e-book. Elsevier Health Sciences*.
- Huang, H., Zheng, O., Wang, D., Yin, J., Wang, Z., Ding, S., ...Shi, B. (2023b). ChatGPT for shaping the future of dentistry: the potential of multi-modal large language model. *International Journal of Oral Science*, 15(1), 29. doi: 10.1038/s41368-023-00239-y
- Huang, Y., Gomaa, A., Semrau, S., Haderlein, M., Lettmaier, S., Weissmann, T., ...Putz, F. (2023a). Benchmarking ChatGPT-4 on a radiation oncology in-training exam and Red Journal Gray Zone cases: potentials and challenges for ai-assisted medical education and decision making in radiation oncology. *Frontiers in Oncology*, 13, 1265024. doi: 10.3389/fonc.2023.1265024
- Huang, H. (2023). Performance of ChatGPT on registered nurse license exam in Taiwan: A descriptive study. *Healthcare*, 11(21), 2855. doi: 10.3390/healthcare11212855
- Lewandowski, M., Łukowicz, P., Świetlik, D., Barańska-Rybak, W. (2023). An original study of ChatGPT-3.5 and ChatGPT-4 dermatological knowledge level based on the dermatology specialty certificate examinations. *Clinical and Experimental Dermatology*, 11ad255. doi: 10.1093/ced/llad255
- Moore, B., Shah, M. I., Owusu-Ansah, S., Gross, T., Brown, K., Gausche-Hill, M., ...American Academy of Pediatrics Committee on Pediatric Emergency Medicine. (2020). Pediatric readiness in emergency medical services systems. *Pediatrics*, 145(1). doi: 10.1542/peds.2019-3307
- Ni, Z., Peng, R., Zheng, X., Xie, P. (2024). Embracing the future: Integrating ChatGPT into China's nursing education system. *International Journal of Nursing Sciences*, 11(2), 295-299. doi: 10.1016/j.ijnss.2024.03.006

- OpenAI. (2023). ChatGPT. Erişim adresi (01.04.2024): <https://openai.com/chatgpt>
- Sahyoun, C., Cantais, A., Gervaix, A., Bressan, S., Löllgen, R., Krauss, B., ...Pediatric Emergency Medicine Comfort and Analgesia Research in Europe (PemCARE) group of the Research in European Pediatric Emergency Medicine (2021). Pediatric procedural sedation and analgesia in the emergency department: surveying the current European practice. *European Journal of Pediatrics*, 180(6), 1799–1813. doi: 10.1007/s00431-021-03930-6
- Sanchez-Ramos, L., Lin, L., Romero, R. (2023). Beware of references when using ChatGPT as a source of information to write scientific articles. *American Journal of Obstetrics and Gynecology*, 229(3), 356–357. doi: 10.1016/j.ajog.2023.04.004
- Su, M. C., Lin, L. E., Lin, L. H., Chen, Y. C. (2024). Assessing question characteristic influences on ChatGPT's performance and response-explanation consistency: Insights from Taiwan's nursing licensing exam. *International Journal of Nursing Studies*, 153, 104717. Advance online publication. doi: 10.1016/j.ijnurstu.2024.104717
- Taira, K., Itaya, T., Hanada, A. (2023). Performance of the large language model ChatGPT on the national nurse examinations in Japan: Evaluation Study. *JMIR Nursing*, 6, e47305. doi: 10.2196/47305
- Talan, T., Kalinkara, Y. (2023). The role of artificial intelligence in higher education: ChatGPT assessment for anatomy course. *International Journal of Management Information Systems and Computer Science*, 7(1), 33-40. <https://doi.org/10.33461/uybisbbd.1244777>
- von Gerich, H., Moen, H., Block, L. J., Chu, C. H., DeForest, H., Hobensack, M., ...Peltonen, L. M. (2022). Artificial Intelligence -based technologies in nursing: A scoping literature review of the evidence. *International Journal of Nursing Studies*, 127, 104153. doi: 10.1016/j.ijnurstu.2021.104153
- Walker, D. M., Tolentino, V. R. (2020). COVID-19: The impact on pediatric emergency care. *Pediatric Emergency Medicine Practice*, 17(Suppl 6-1), 1–27.
- Watari, T., Takagi, S., Sakaguchi, K., Nishizaki, Y., Shimizu, T., Yamamoto, Y., ...Tokuda, Y. (2023). Performance comparison of Chatgpt-4 and Japanese medical residents in the general medicine in-training examination: Comparison Study. *JMIR Medical Education*, 9, e52202. doi: 10.2196/52202