

Osmangazi Journal of Medicine
e-ISSN: 2587-1579

Acil Tıpta Yapay Zekâ Kullanımı: ChatGPT'nin Türkiye Sağlık Bakanlığı Kuduz Profilaksisi Rehberi ile Uyumunun Değerlendirilmesi

Use of Artificial Intelligence in Emergency Medicine: Evaluation of the Compliance of ChatGPT with the Rabies Prophylaxis Guide of the Ministry of Health of Turkey

¹Esat Barut, ²Kemal Barut, ³Ramazan Giden, ³İbrahim Halil Yasak

¹Silvan Dr. Yusuf Azizoglu Devlet Hastanesi, Acil Tıp, Diyarbakır, Türkiye

²Gazi Yaşargil Eğitim ve Araştırma Hastanesi, Acil Tıp, Diyarbakır, Türkiye

³Harran Üniversitesi, Tıp Fakültesi Acil Tıp Anabilimdalı, Şanlıurfa, Türkiye

ORCID ID of the authors

EB. [0000-0003-0844-0468](https://orcid.org/0000-0003-0844-0468)

KB. [0000-0001-8758-726X](https://orcid.org/0000-0001-8758-726X)

RG. [0000-0003-2127-1056](https://orcid.org/0000-0003-2127-1056)

İHY. [0000-0002-6399-7755](https://orcid.org/0000-0002-6399-7755)

Correspondence / Sorumlu yazar:

Esat BARUT

Silvan Dr. Yusuf Azizoglu Devlet Hastanesi,
Acil Tıp, Diyarbakır, Türkiye

e-mail: esbarut94@gmail.com

Etik Kurul Onayı: Çalışmamıza hiçbir insan katılımcı veya hayvan dahil edilmediğinden etik kurul onayı gerekli değildir. ChatGPT sorgulamaları OpenAI'nin kullanım politikalarına uygun şekilde yürütülmüştür.

Onam: Çalışmamız için onam gerekli değildir.

Telif Hakkı Devir Formu: Tüm yazarlar tarafından Telif Hakkı Devir Formu imzalanmıştır.

Hakem Değerlendirmesi: Hakem değerlendirmesinden geçmiştir.

Yazar Katkı Oranları: Veri toplama: EB, KB. Konsept: EB, KB. Tasarım: EB, KB. Veri işleme: EB, KB, RG, İHY. Analiz veya Yorum: EB, KB, RG, İHY. Literatür taraması: EB, KB, RG, İHY. Yazma: EB, KB.

Çıkar Çatışması Bildirimi: Yazarlar çıkar çatışması olmadığını beyan etmişlerdir.

Destek ve Teşekkür Beyanı: Yazarlar bu çalışma için finansal destek almadıklarını beyan etmişlerdir.

Özet: Çalışmamızda, kuduz riskli hayvan teması vakalarının değerlendirilmesinde Türkiye Sağlık Bakanlığı “Kuduz Profilaksi Rehberi (2019)” ile ChatGPT (v.4)’ü karşılaştırmayı amaçladık. Hedefimiz, giderek kullanımı artan ve sorulara hızla yanıt verebilen yapay zekâ araçlarının, acil servis gibi yoğun ve zamanla yarışılan alanlarda rehberlere alternatif olup olamayacağını ve vakalara yaklaşımdaki performansını belirlemektir. Çalışmamızda, acil serviste sıklıkla karşılaşılabilecek 50 adet kuduz riskli hayvan teması içeren hasta senaryosu oluşturuldu. ChatGPT’den, vakaları temas kategorisi tayini (I-II-III-IV), kuduz aşı gerekliliği, kuduz immünoglobulini gerekliliği ve antibiyotik profilaksisi açısından değerlendirilerek yanıtlanması istendi. Ayrıca, rehberde sıkça sorulan sorular bölümünde yer alan 44 soru ChatGPT’ye yöneltildi. Yanıtlar; rehberdeki cevaplarla tam uyumlu, kısmi uyumlu, ve uyumsuz şeklinde sınıflandırıldı. ChatGPT yanıtları Türkiye Sağlık Bakanlığı “Kuduz Profilaksi Rehberi” ile karşılaştırılarak SPSS 26.0 programıyla analiz edilmiş olup; Ki-kare, Cohen’s Kappa ve oran testleri (Z-testi) uygulanmıştır. ChatGPT, 50 hasta senaryosunun 27’sinde kategori tayini, 11’inde aşı gerekliliği, 13’ünde immünoglobulin gerekliliği ve 7’sinde antibiyotik profilaksisi açısından rehberle uyumsuz öneride bulunmuştur. ChatGPT, rehberde sıkça sorulan 44 sorunun ise 21’ine (%47,73) tam uyumlu, 18’ine (%40,91) kısmi uyumlu ve 5’ine (%11,36) uyumsuz yanıt vermiştir. ChatGPT, kuduz profilaksisi rehberiyle kısmen uyumlu sonuçlar üretmiştir. Yapay zekâ araçları, kısa sürede sağlık alanında her ne kadar vazgeçilmez hale gelmiş olsa da, hasta güvenliği açısından mutlaka rehberlerle çapraz kontrollü ve yardımcı araçlar olarak kullanılmalıdır. Yapay zekâ araçlarının eksik veya uyumsuz yanıtlar üretme riski, hasta sağlığını tehlikeye atabileceği ve hekim açısından klinik karar süreçlerinde ciddi risk oluşturabileceği için, hasta yönetiminde güncel sağlık rehberlerine başvurulması önem arz etmektedir. Bu araçların hâlâ rehberlerin yerini alacak düzeye olmadığını düşünmekteyiz.

Anahtar Kelimeler: Kuduz, Profilaksi, ChatGPT, Yapay Zekâ, Acil Tıp

Abstract: In our study, we aimed to compare the Turkish Ministry of Health “Rabies Prophylaxis Guide (2019)” and ChatGPT (v.4) in the evaluation of rabies risk animal contact cases. Our aim was to determine whether artificial intelligence tools, which are increasingly used and can respond to questions quickly, can be an alternative to guides in busy and time-sensitive areas such as emergency services and their performance in approaching cases. In our study, 50 patient scenarios involving rabies risk animal contact that can be frequently encountered in emergency services were created. ChatGPT was asked to evaluate and answer the cases in terms of contact category determination (I-II-III-IV), rabies vaccine requirement, rabies immunoglobulin requirement and antibiotic prophylaxis. In addition, 44 questions in the frequently asked questions section of the guide were directed to ChatGPT. The answers were classified as fully compatible, partially compatible and incompatible with the answers in the guide. ChatGPT responses were compared with the Turkish Ministry of Health “Rabies Prophylaxis Guide” and analyzed with SPSS 26.0 program; Chi-square, Cohen’s Kappa and ratio tests (Z-test) were applied. ChatGPT made recommendations that were inconsistent with the guideline in 27 of the 50 patient scenarios in terms of category determination, in 11 in terms of vaccine requirements, in 13 in terms of immunoglobulin requirements, and in 7 in terms of antibiotic prophylaxis. ChatGPT gave fully compatible answers to 21 (47.73%), partially compatible answers to 18 (40.91%) and incompatible answers to 5 (11.36%) of the 44 frequently asked questions in the guide. ChatGPT produced partially compatible results with the rabies prophylaxis guide. Although artificial intelligence tools have become indispensable in the health field in a short time, they should definitely be cross-checked with the guides and used as auxiliary tools for patient safety. The risk of artificial intelligence tools producing incomplete or inconsistent responses could endanger patient health and pose a serious risk to physicians in their clinical decision-making processes. Therefore, it is important to refer to current health guidelines in patient management. We believe that these tools are still not at a level to replace the guidelines.

Keywords: Rabies, Prophylaxis, ChatGPT, Artificial Intelligence, Emergency Medicine

How to cite/ Atıf için: Barut E, Barut K, Giden R, Yasak İH, Acil Tıpta Yapay Zekâ Kullanımı: ChatGPT'nin Türkiye Sağlık Bakanlığı Kuduz Profilaksisi Rehberi ile Uyumunun Değerlendirilmesi, Osmangazi Journal of Medicine, 2025;47(6):865-870

1. Giriş

Kuduz, hem insan hem de hayvan sağlığını etkileyen, zoonotik karakterli ve bilinen en eski hastalıklardan biridir (1). Lyssavirus cinsine ait nörotropik bir virüs tarafından hastalık meydana gelmektedir. Genellikle ısırılma ile oluşan yaradan, nadiren de potansiyel virüs bulundurma ihtimali olan tırmalama veya virüs bulunan salgıların mukozayla ya da bütünlüğü bozulmuş deriyle teması ile bulaşmaktadır (2).

Tüm sıcakkanlı hayvanlar kuduz virüsü ile enfekte olabilirler; ancak virüse karşı aynı oranda hassas olmayıp temas sonrası bulaştırma riskleri farklılık göstermektedir. Virüs, evcil hayvanlardan en sık köpeklerde ve sığırlarda yabani hayvanlardan ise en sık tilkilerde tespit edilmektedir (1). Toplumda ise “hayvan ısırığı” denince akla ilk gelen kedi ve köpeklerdir (3). Pek çok hayvandan hastalık bulaşabilse de hâlâ kuduzla bağlı ölümlerin en önemli kaynağı köpeklerdir (4).

İnsan ve memeli hayvanların çoğunda ensefalit tablosu meydana getiren ölümcül bir hastalık olmasına rağmen uygun yara bakımı ve temas sonrası profilaksi ile de tamamen önlenabilir (5). Gelişmekte olan ülkeler başta olmak üzere tüm dünyada her yıl yaklaşık 40,000–70,000 arasında insanın kuduz nedeniyle hayatını kaybettiği tahmin edilmektedir (6). Türkiye, kuduz yönünden endemik bir bölge olup her yıl yaklaşık 180 bin kuduz riskli temas (KRT) bildirimi yapılmakta ve her yıl 1 ila 4 kuduzla bağlı insan ölümü gerçekleşmektedir (7).

Bu nedenle, kuduzu halk sağlığı sorunu olmaktan çıkarabilmek için KRT olgularının titizlikle değerlendirilmesi ve buna uygun stratejilerin planlanması büyük önem taşımaktadır (8). Ülkemizde temas öncesi ve temas sonrası profilaksi hizmetleri, bakanlığımız tarafından koruyucu sağlık hizmetleri kapsamında verilmekte olup; Dünya Sağlık Örgütü ve diğer ülkelerin konuyla ilgili rehberleri dikkate alınarak hazırlanan Kuduz Profilaksi Rehberi’nde belirtilen esaslar çerçevesinde uygulanmaktadır (1).

Son yıllarda, sağlık da dahil olmak üzere birçok alanda yapay zekâ kullanımı hızlı ve doğru bilgiye ulaşma ihtiyacı nedeniyle giderek yaygınlaşmaktadır (9). Herkesin kullanımına açık en yaygın uygulamalardan biri de ChatGPT’dir. ChatGPT, en büyük yapay zekâ araçlarından biri olup en gelişmiş sohbet robotlarından biridir. Bu gibi yapay zekâ modelleri, büyük miktarda veriyi barındıracak şekilde tasarlanmış olup hızlıca sorulan soruları

yanıtlayabilmektedir (10). Bu nedenle günümüzde yapay zekâ gerek sağlık çalışanlarının gerekse hastaların sık kullandıkları bir araç haline gelmiştir. Ancak, yapay zekâ sistemlerinin rehberlerle uyumu ve güvenilirliği tartışma konusudur (11).

Bu çalışmamızda ChatGPT’nin acil servislere sık karşılaşılan kuduz riskli hayvan teması olan hasta senaryolarının değerlendirilmesinde ve sıkça merak edilen soruların cevaplandırılmasında ne derece doğru sonuçlar verdiğini; temas öncesi ve sonrası profilaksi hizmetleri hakkındaki bilgilerinin TC Sağlık Bakanlığının “Kuduz Profilaksi Rehberi” (1) ile ne derece uyumlu olduğunu tespit etmeyi amaçladık.

2. Gereç ve Yöntem

Çalışmamız, yapay zekâ tabanlı bir karar destek aracının, kuduz profilaksi rehberiyle karşılaştırmalı performansını değerlendirmeyi amaçlayan, tanımlayıcı ve karşılaştırmalı kesitsel bir analizdir. Referans rehber olarak Türkiye Sağlık Bakanlığı “Kuduz Profilaksi Rehberi (2019)” kullanıldı; yapay zekâ kaynağı olarak ise ChatGPT-4o (v.4) (OpenAI) tercih edildi.

Acil uzmanları tarafından, acil serviste sıklıkla karşılaşılabilecek türden 50 farklı hasta senaryosu oluşturuldu. Vakalar hastanın yaşı, cinsiyeti, temasın gerçekleştiği vücut bölgesi, temas şekilleri (ısırık, tırmalama, salya teması vb.), temas edilen hayvan türü (köpek, kedi, yabani hayvanlar vs.), yaranın özellikleri, hastanın aşı/immünoglobulin öyküsü, hayvanın aşı durumu gibi değişkenleri içerecek şekilde gerçek klinik pratikten esinlenilerek hazırlandı. Karşılaştırma kriterleri olarak yapay zekâdan vakalara temas kategorisi tayini (I-II-III-IV) yapması ayrıca vakaları kuduz aşısı (Evet/Hayır), kuduz immünoglobulini gerekliliği (Evet/Hayır), antibiyotik profilaksisi (Evet/Hayır) açısından değerlendirmesi istendi. Örneğin, ‘15 yaşındaki bir erkek hasta, aşı durumu bilinmeyen bir sokak köpeği tarafından ısırılması nedeniyle acil servise başvurdu. Hastanın elinde kanamalı bir lezyon mevcut. Vakayı, temas kategorisi tayini (I-II-III-IV), kuduz aşısı gerekliliği (Evet/Hayır), kuduz immünoglobulini gerekliliği (Evet/Hayır) ve antibiyotik profilaksisi (Evet/Hayır) açısından değerlendir’; benzer şekilde, ‘8 yaşındaki bir kız hasta, aşıları bir evcil kedi tarafından tırmalanması nedeniyle acil servise başvurdu. Hastanın yüzünde kanamasız bir lezyon mevcut. Vakayı, temas kategorisi tayini (I-II-III-IV), kuduz aşısı

gerekliliği, kuduz immünoglobulini gerekliliği ve antibiyotik profilaksisi açısından değerlendir' şeklinde ChatGPT'ye sorular yöneltildi. Ayrıca, Türkiye Sağlık Bakanlığı Kuduz Profilaksi Rehberi'nin 34–39. sayfaları arasında yer alan 'Sıkça Sorulan Sorular' bölümündeki toplam 44 soru çalışmaya dahil edildi (1). ChatGPT'nin yanıtları; rehberdeki cevaplarla tam uyumlu (rehberle %100 aynı öneriler), kısmi uyumlu (rehberle büyük ölçüde uyumlu ancak eksik veya farklı detaylar içeren yanıtlar), uyum yok (rehberden tamamen farklı veya yanlış öneriler) şeklinde sınıflandırıldı.

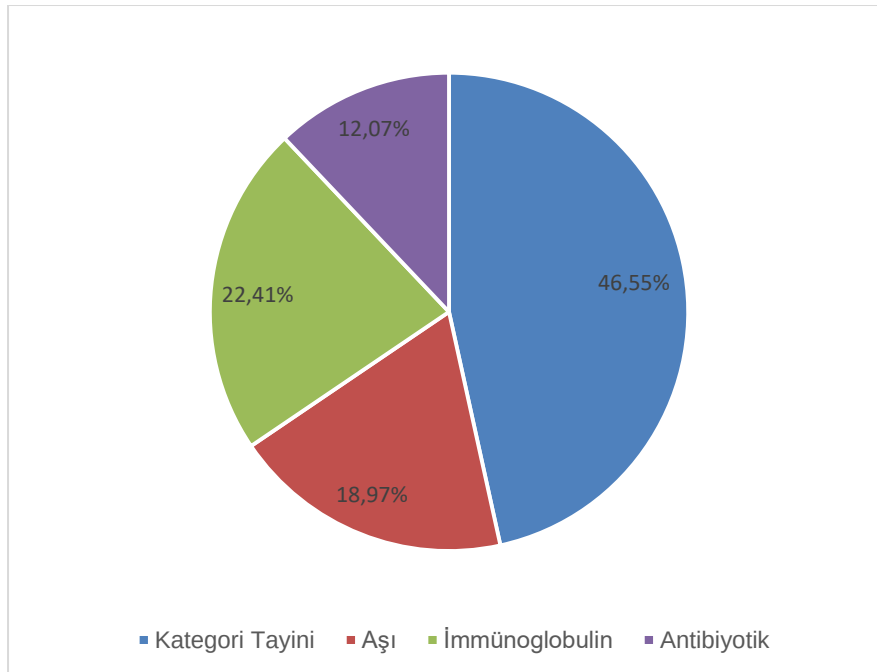
Yapay zekânın vaka senaryolarına ve sıkça sorulan sorulara verdiği yanıtlar, iki bağımsız acil uzmanı tarafından rehber baz alınarak değerlendirilmiş; uyuşmazlıklarda üçüncü bir acil uzmanının görüşüne başvurulmuştur. Yanıtların tekrarlanabilirliğini değerlendirmek amacıyla tüm sorular farklı zamanlarda iki kez sorulmuş ve tüm yanıtlar tutarlı olarak saptanmıştır. ChatGPT'nin yanıtları SPSS 26.0 kullanılarak analiz edilmiştir. Yanıtlarda dağılım farkları için ki-kare testi ($p < 0.05$ anlamlı), yanıtların uyum düzeyi için Cohen Kappa analizi (Kappa değerleri < 0 : uyumsuzluk, 0.01–0.20: hafif uyum, 0.21–0.40: orta uyum, 0.41–0.60: makul uyum, 0.61–0.80: yüksek uyum, 0.81–1.00:

mükemmel uyum), yanıtların tesadüfi olup olmadığını değerlendirmek için oran testi (Z-testi) ($p < 0.05$ anlamlı) uygulanmıştır.

Çalışmamıza hiçbir insan katılımcı veya hayvan dahil edilmediğinden literatürdeki benzer çalışmalarda olduğu gibi etik kurul onayı alınmamıştır (12,13). ChatGPT sorgulamaları da OpenAI'nin kullanım politikalarına uygun şekilde yürütülmüştür.

3. Bulgular

ChatGPT hazırlanan 50 hasta senaryosunun 27'sine rehberle uyumsuz kategori tayini yapmıştır. Ayrıca vakaların 11'ine aşı gerekliliği, 13'üne immünoglobulin gerekliliği ve 7'sine ise antibiyotik profilaksisi açısından rehberle uyumsuz öneride bulunmuştur. Toplam rehberle uyumsuz öneri sayısı 58 olup, bazı senaryolarda birden fazla uyumsuz öneride bulunduğu gözlenmiştir. En sık uyumsuz öneriyi kategori tayini; en az ise antibiyotik profilaksisi hakkında yapmıştır. Rehberle uyumsuz öneri oranları; kategori tayini için %46.55, aşı önerisi için %18.97, immünoglobulin önerisi için %22.41 ve antibiyotik profilaksisi için %12.07 olarak bulunmuştur (Şekil 1).



Şekil 1. ChatGPT'nin Hasta Senaryolarında Yaptığı Rehberle Uyumsuz Önerilerin Dağılımı

Rehberle uyumsuz önerilerin dağılımlarına Ki-kare testi uygulandığında anlamlı bir fark saptanmış olup ($\chi^2 = 13.45$, $df = 3$, $p = 0.004$), uyumsuz önerilerin en sık kategori tayininde olduğu ortaya konulmuştur.

ChatGPT'nin rehberle uyum düzeyini değerlendirmek amacıyla uygulanan Cohen Kappa analiziyle tüm kriterlerde düşük uyum saptanmış olup ($\kappa = 0.20$, 95% GA: 0.05–0.35), bu sonuç

ChatGPT'nin rehberle uyumunun sınırlı olduğunu göstermiştir. ChatGPT'nin senaryolarda yaptığı rehberle uyumsuz önerilerin dağılımına ve istatistiksel analizine Tablo 1'de yer verilmiştir. ChatGPT rehberle uyumsuz kategori tayinlerini genellikle

kuduz profilaksisi gerekmeyen hayvanlarla ilgili senaryolarda yapmış olup bunun sonucunda bu vakalara genellikle uyumsuz aşı ve immunoglobulin önerisinde bulunmuştur.

Tablo 1. ChatGPT'nin Hasta Senaryolarında Yaptığı Rehberle Uyumsuz Önerilerin Dağılımı

Öneri Türü	Sayı	Yüzde (%)	Ki-kare p-değeri	Cohen Kappa κ değeri
Kategori Tayini	27	46.55	0.004	0.20
Aşı	11	18.97		
İmmünoglobulin	13	22.41		
Antibiyotik	7	12.07		
Toplam	58	100		

Not: Yüzdelere, toplam 58 uyumsuz öneri sayısına göre hesaplanmıştır (birden fazla uyumsuz öneride bulunduğu senaryolar nedeniyle). Ki-kare testi ($\chi^2 = 13.45$, $df = 3$, $p = 0.004$), ChatGPT'nin rehberle uyumsuz önerilerinin kategori tayininde yoğunlaştığını gösterir ve sonuç anlamlıdır. Cohen Kappa ($\kappa = 0.20$, 95% GA: 0.05–0.35), ChatGPT'nin tüm kriterlerde rehberle düşük uyumunu belirtir.

ChatGPT 50 senaryonun sadece 15 (%30)'inde tüm karşılaştırma kriterlerine rehberle uyumlu cevap vermiştir. Kritik öneme sahip kuduz aşısı ve immünoglobulini gerekliliği açısından, iki öneriyi birlikte rehberle uyumlu yanıtladığı senaryo sayısı 36 (%72) olarak bulunmuştur. ChatGPT'nin yalnızca 15 senaryoda tüm kriterleri rehberle

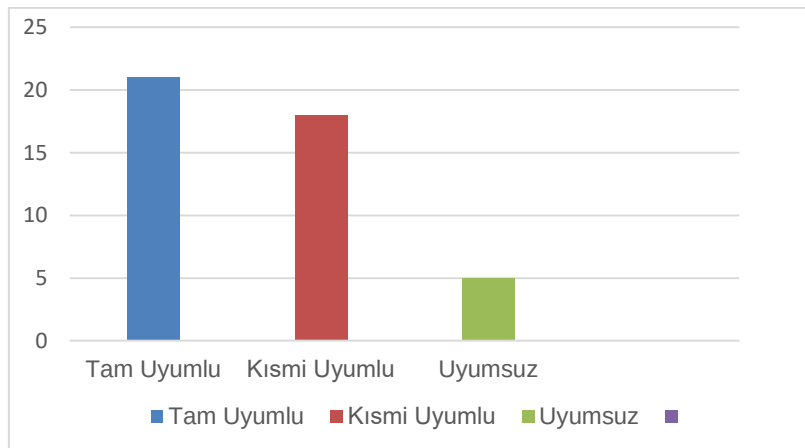
uyumlu yanıtlaması vakaları yönetirken sıklıkla rehberle uyumsuz öneriler yaptığını, her iki öneriyi 36 hastada rehberle uyumlu yanıtlaması da da yanıtlarının tesadüf olmaktan öte olduğunu göstermiştir. Bu sonuçlara oran testi uygulandığında anlamlı bulunmuştur ($z = 3.11$, $p = 0.001$ ve $p = 0.002$) (Tablo 2).

Tablo 2. Aşı Ve İmmünoglobulin Gerekliliği İle Tüm Önerilerde Chatgpt'nin Rehberle Uyumlu Yanıtlarının Dağılımı

Öneri	Sayı	Yüzde (%)	Oran Testi	p-değeri
Aşı ve İmmünoglobulin	36	72	$z = 3.11$	0.002
Tüm Öneriler (Tam Doğru)	15	30	$z = -2.83$	0.001

Not: Oran testi Aşı ve İmmünoglobulin için $z = 3.11$, $p = 0.002$ olup ChatGPT'nin iki öneriyi birlikte rehberle uyumlu yanıtlamada tesadüfi orandan daha iyi performans sergilediğini gösterir ve sonuç anlamlıdır. Oran testi Tüm öneriler için $z = -2.83$, $p = 0.001$ olup ChatGPT'nin tüm önerileri rehberle uyumlu yanıtlamada tesadüfi orandan daha düşük başarı sergilediğini gösterir ve sonuç anlamlıdır.

ChatGPT rehberde yer alan 44 sıkça sorulan sorunun ise 21'ine (%47.73) tam uyumlu, 18'ine (%40.91) kısmi uyumlu ve 5'ine (%11.36) uyumsuz yanıt vermiştir (Şekil 2).



Şekil 2. ChatGPT'nin Sıkça Sorulan Sorulara Verdiği Yanıtların Dağılımı

Uyum kategorilerinin dağılımına ki-kare testi uygulandığında anlamlı bulunmuştur ($\chi^2 = 9.55$, $df = 2$, $p = 0.008$), bu da tam ve kısmi uyumlu yanıtların beklenenden daha fazla olduğunu göstermektedir. Cohen Kappa analizi ChatGPT ile rehber arasında orta düzey uyum olduğunu ortaya koymuştur ($\kappa =$

0.38, 95% GA: 0.20–0.56). Z-testi ile ChatGPT'nin rehberle tam uyumlu yanıt vermede tesadüfi beklenen orandan daha iyi performans sergilediği gösterilmiş olup sonuç anlamlıdır ($z = 2.07$, $p = 0.038$) (Tablo 3).

Tablo 3.. ChatGPT'nin Sıkça Sorulan Sorulara Verdiği Yanıtların Dağılımı

Uyum Kategorileri	Sayı	Yüzde (%)	Ki-kare p-değeri	Cohen Kappa κ değeri	Oran Testi p-değeri
Tam Uyumlu	21	47.73	0.008	0.38	0.038
Kısmi Uyumlu	18	40.91			
Uyumsuz	5	11.36			
Toplam	44	100			

Not: Ki-kare testi ($\chi^2 = 9.55$, $df = 2$, $p = 0.008$), uyum kategorilerin eşit dağılmadığını, tam ve kısmi uyumlu yanıtların beklenenden daha fazla olduğunu göstermekte olup sonuç anlamlıdır. Cohen Kappa ($\kappa = 0.38$, 95% GA: 0.20–0.56), ChatGPT'nin rehberle orta düzey uyumunu belirtir. Oran testi ($z = 2.07$, $p = 0.038$), ChatGPT'nin rehberle tam uyumlu yanıt vermede tesadüfi beklenen orandan daha yüksek performans sergilediğini gösterir ve sonuç anlamlıdır.

Komplike vakalar ve nadir görülen hayvan temasları ile ilgili sıkça sorulan sorulara rehber net ve açıklayıcı cevaplar veriyorken, ChatGPT yanlış cevaplamış veya daha muğlak ve kısmen uyumlu önerilerde bulunmuştur. Kısmi uyumlu ve uyumsuz yanıtlar; profilaksi, özel hasta grupları (örneğin, immünosuprese, önceden profilaksi uygulanmış bireyler vs), kuduz riskli hayvanlar ve protokoller ile ilgili bazı soruları eksik ya da hatalı yorumlamasıyla ilişkilendirilmiştir.

4. Tartışma

Kuduz (rabies) hastalığı hâlâ dünyada ölümle sonuçlanan zoonotik enfeksiyonlar arasında ciddiyetini sürdürmektedir (14). Acil servislere sıklıkla hayvan ısırıkları nedeniyle hasta başvuruları olmaktadır. Acil servis gibi zorlayıcı bir alanda hastaların hızlıca yara bakımı, kuduz aşısı ve immünoglobulini, antibiyotik profilaksisi konularında titizlikle değerlendirilmesi ve hastaların doğru bilgilendirilmesi gerekmektedir.

Rehberlerin incelenmesi vakit alabildiğinden acil hekimleri hızlı ve doğru yanıtlar alma yolları aramaktadır. Bu yollardan biri olarak yapay zekâ araçları ön plana çıkmaktadır (15). Literatürde pek çok çalışmada ChatGPT'nin hastaların merak ettiği soruları ne ölçüde doğru yanıtladığı araştırılmıştır (16).

Biz de çalışmamızda kuduz riskli hayvan teması vakalarının değerlendirilmesinde ChatGPT'nin performansını değerlendirmek istedik. Çalışmamızda ChatGPT her ne kadar rehberle uyumlu öneriler sunsa da bazı vakalarda eksik ve uyumsuz yanıtlar verebildiği gösterilmiştir.

ChatGPT'nin özellikle temas kategorisi tayininde, kuduz profilaksi gerekmeyen hayvanları tanımda, kompleks vakalara ve özel hasta gruplarına yaklaşımda rehberle uyumsuz veya eksik sonuçlar verdiği görülmüştür. Wang ve ark. çalışmasında da ChatGPT'nin atipik prezentasyonlarda başarısız olduğu gösterilmiş olup, bizim bulgularımızı desteklemektedir (17).

Türkiye'de yapılan bir çalışmada, ChatGPT'nin refraktif cerrahi sorularında %52,5 doğruluk oranıyla performans gösterdiği rapor edilmiş olup çalışmamızı destekler niteliktedir (13). Çağlar ve ark. (18) ve Çolakoğlu ve ark.(12) yaptıkları çalışmalarda ChatGPT'nin sorulara bizim elde ettiğimiz sonuçların aksine tatmin edici derecede doğru yanıtlar verdiğini göstermiştir. Bu farklılık yapay zekâ araçlarının her alanda güncel bilgileri içermediğinden kaynaklanıyor olabilir.

ChatGPT'nin Türkçe yanıtlarının İngilizce yanıtlara kıyasla daha az tutarlı olabileceği, dil modeli eğitimindeki veri dengesizliklerinden kaynaklanabileceği çalışmalarda gösterilmiş olup (19), ChatGPT'nin sorularımıza verdiği rehberle uyumsuz yanıtlar kısmen bundan da kaynaklanıyor olabilir. Ozturk ve arkadaşlarının çalışmasında ChatGPT'nin İngilizce testlerde Türkçe testlere göre %28 daha yüksek doğruluk skoru elde ettiği gösterilmiştir (20). Yine de ChatGPT diğer sosyal medya platformlarından daha doğru ve kaliteli bilgi sağladığı çalışmalarla gösterilmiştir (21).

Kuduz her ne kadar ölümcül bir hastalık olsa da doğru yara bakımı ve profilaksi ile %100 önlenbilir. Bu nedenle bu hastalara yaklaşımda hata

payına yer olmayıp en ufak bir hatanın hastanın ölümüne ve klinik karar süreçlerinde ciddi risk oluşturabileceği unutulmamalıdır. Bu nedenle ulusal ve uluslararası düzeyde kuduz profilaksi rehberlerinin devamlı güncel tutulması gerekmektedir. ChatGPT gibi yapay zekâ araçlarının ise her konuda güncel bilgileri içermeyebileceği daima akılda tutulmalıdır.

ChatGPT gibi yapay zekâ araçlarının, sağlık alanında kullanımlarında mutlaka güncel rehberlerle çapraz kontrollü ve yardımcı araç olarak

kullanılması gerekmektedir. Bu çalışma yapay zekâ araçlarının klinik karar vermede tamamen rehberlerin yerini alamayacağını, ancak hızlı bilgiye erişim ve karar desteği açısından yardımcı bir araç olabileceğini göstermektedir. Gelecekte, yapay zekâ modelleri rehberler doğrultusunda eğitilirse ve düzenli validasyon çalışmaları yapılırsa, hasta güvenliğini arttıracığına ve sağlıkta etkin olarak kullanılarak daha büyük öneme sahip olacağına inanmaktayız. Ancak yine de o gün gelene kadar hasta yönetiminde güncel sağlık rehberlerine başvurulması önem arz etmektedir.

KAYNAKLAR

1. Aylan O, Baykam N, Güner R, ve ark. TC Sağlık Bakanlığı Kuduz Profilaksi Rehberi. Ankara: T.C. Sağlık Bakanlığı; 2019.
2. Zhao H, Zhang J, Cheng C, Zhou YH. Rabies acquired through mucosal exposure, China, 2013. *Emerg Infect Dis*. 2019 May;25(5):1028-9.
3. Aksoy M, Demirbaş B, Maden F, ve ark. Ankara ilinde 2005-2009 yılları arasında görülen şüpheli ısırıkların ve kuduz aşılamaının değerlendirilmesi. 3. EKMUD Kongresi; 2010 May 12-16; Ankara, Türkiye. Kongre Özet Kitabı. p.199.
4. World Health Organization. Rabies. Factsheet. Available from: <http://www.who.int/mediacentre/factsheets/fs099/en/> [Accessed 25 Apr 10]
5. O'Brien KL, Nolan T. The WHO position on rabies immunization - 2018 updates. *Vaccine*. 2019 Oct 3;37 Suppl 1(Suppl 1):A85-7.
6. Chin J. Control of communicable diseases manual. Washington DC: American Public Health Association; 2000.
7. Türkiye Halk Sağlığı Kurumu. Kuduz Saha Rehberi. Ankara: Türkiye Halk Sağlığı Kurumu; 2014.
8. Abela-Ridder B. Rabies: 100 per cent fatal, 100 per cent preventable. *Vet Rec*. 2015 Aug 8;177(6):148-9.
9. Deng J, Lin Y. The benefits and challenges of ChatGPT: an overview. *Front Comput Intell Syst*. 2022 Jan 5;2(2):81-3.
10. Kasneci E, Sessler K, Küchemann S, et al. ChatGPT for good? On opportunities and challenges of large language models for education. *Learn Individ Differ*. 2023 Apr;103:102274.
11. Biswas SS. Role of ChatGPT in public health. *Ann Biomed Eng*. 2023 May;51(5):868-9.
12. Çolakoğlu Y, Ayten A, Sertkaya C, Toksal K, Karadağ S. Evaluation of Chat Generative Pretrained Transformer (ChatGPT) performance in answering kidney transplant related questions. *New J Urol*. 2025;20(1):21-3.
13. Aydın FO, Aksoy BK, Ceylan A, ve ark. Refraktif cerrahide sık sorulan sorular için ChatGPT 3.5, ChatGPT 4.0, Gemini ve Microsoft Copilot tarafından oluşturulan yanıtların okunabilirliği ve uygunluğu. *Türk J Ophthalmol*. 2024 Dec;54(6):313-7.
14. Knobel DL, Cleaveland S, Coleman PG, et al. Re-evaluating the burden of rabies in Africa and Asia. *Bull World Health Organ*. 2005 May;83(5):360-8.
15. Zhang H. Artificial intelligence in healthcare: opportunities and challenges. *Theor Nat Sci*. 2023 Dec 20;21(1):130-4.
16. Secinaro S, Calandra D, Secinaro A, Muthurangu V, Biancone P. The role of artificial intelligence in healthcare: a structured literature review. *BMC Med Inform Decis Mak*. 2021 Apr 10;21(1):125.
17. Wang J, Shue K, Liu L, Hu G. Preliminary evaluation of ChatGPT model iterations in emergency department diagnostics. *Scientific Reports*. 2025;15:10426.
18. Caglar U, Yildiz O, Meric A, et al. Evaluating the performance of ChatGPT in answering questions related to benign prostate hyperplasia and prostate cancer. *Minerva Urol Nephrol*. 2023 Dec;75(6):729-33.
19. Yiğit S, Berşe S, Dirgar E. Yapay zekâ destekli dil işleme teknolojisi olan ChatGPT'nin sağlık hizmetlerinde kullanımı. *Eurasian J Health Technol Assess*. 2023 Jun 30;7(1):57-65.
20. Ozturk N, Yakak I, Ağ MB, Aksoy N. Is ChatGPT reliable and accurate in answering pharmacotherapy-related inquiries in both Turkish and English? *Currents in Pharmacy Teaching and Learning*. 2024;16(7):102101.
21. Antaki F, Touma S, Milad D, El-Khoury J, Duval R. Evaluating the performance of ChatGPT in ophthalmology: an analysis of its successes and shortcomings. *Ophthalmol Sci*. 2023 Dec 1;3(4).