

PAPER DETAILS

TITLE: Okülofasiyal Plastik ve Orbital Cerrahide İngilizce ve Türkçe Dil Çesitliliğinin Yapay Zeka Chatbot Performansına Etkisi: ChatGPT-3.5, Copilot ve Gemini Üzerine Bir Çalıřma

AUTHORS: Eyüpcan Sensoy,Mehmet Çitirik

PAGES: 781-786

ORIGINAL PDF URL: <https://dergipark.org.tr/tr/download/article-file/4089242>

Research Article / Araştırma Makalesi

Okülofasiyal Plastik ve Orbital Cerrahide İngilizce ve Türkçe Dil Çeşitliliğinin Yapay Zeka Chatbot Performansına Etkisi: ChatGPT-3.5, Copilot ve Gemini Üzerine Bir Çalışma
Impact of Language Variation English and Turkish on Artificial Intelligence Chatbot Performance in Oculofacial Plastic and Orbital Surgery: A Study of ChatGPT-3.5, Copilot, and Gemini

Eyüpcan Şensoy, Mehmet Çitirik

Ankara Etlik Şehir Hastanesi, Göz Hastalıkları Kliniği, Ankara, Türkiye

Abstract: The aim is to investigate the effects of applying the same questions in different languages related to oculofacial plastic and orbital surgery to ChatGPT-3.5, Copilot, and Gemini artificial intelligence chatbots, which are freely accessible, on the performance of these programs. English and Turkish versions of 30 questions related to oculofacial plastic and orbital surgery were applied to ChatGPT-3.5, Copilot, and Gemini chatbots. The answers given by the chatbots were compared with the answer key at the back of the book and grouped as correct and incorrect. Their superiority over each other was compared statistically. While ChatGPT-3.5 answered 43.3% of the English questions correctly, it answered 23.3% of the Turkish questions correctly ($p=0.07$). While Copilot answered 73.3% of the English questions correctly, it answered 63.3% of the Turkish questions correctly ($p=0.375$). While Gemini answered 46.7% of the English questions correctly, it answered 33.3% of the Turkish questions correctly ($p=0.344$). Copilot showed higher performance than other programs in answering Turkish questions ($p<0.05$). In addition to improving the knowledge level of chatbots, their performance in different languages also needs to be examined and improved. Correcting these disadvantages in chatbots will pave the way for more widespread and reliable use of these programs.

Keywords: ChatGPT-3.5, Copilot, Gemini, İngilizce ve Türkçe, Okülofasiyal plastik ve orbita cerrahisi.

Özet: Ücretsiz olarak erişim sağlanabilen ChatGPT-3.5, Copilot ve Gemini yapay zeka sohbet botlarına okülofasiyal plastik ve orbita cerrahisi ile ilişkili farklı dillerdeki aynı soru uygulamalarının bu programların performanslarına olan etkilerini araştırmaktır. Okülofasiyal plastik ve orbita cerrahisi ile ilişkili 30 sorunun İngilizce ve Türkçe versiyonları ChatGPT-3.5, Copilot ve Gemini sohbet botlarına uygulandı. Sohbet botlarının verdikleri cevaplar kitap arkasında yer alan cevap anahtarı ile karşılaştırıldı, doğru ve yanlış olarak gruplandırıldı. Birbirlerine üstünlükleri istatistiksel olarak karşılaştırıldı. ChatGPT-3.5 İngilizce soruların %43,3'üne doğru cevap verirken, Türkçe soruların %23,3'üne doğru cevap verdi ($p=0,07$). Copilot İngilizce soruların %73,3'üne doğru cevap verirken, Türkçe soruların %63,3'üne doğru cevap verdi ($p=0,375$). Gemini İngilizce soruların %46,7'sine doğru cevap verirken, Türkçe soruların %33,3'üne doğru cevap verdi ($p=0,344$). Copilot, Türkçe soruları cevaplama da diğer programlardan daha yüksek performans gösterdi ($p<0,05$). Sohbet botlarının bilgi düzeylerinin geliştirilmesinin yanında farklı dillerdeki performanslarının da incelenmeye ve geliştirilmeye ihtiyacı vardır. Sohbet botlarındaki bu dezavantajların düzeltilmesi, bu programların daha yaygın ve güvenilir bir şekilde kullanılmasına zemin hazırlayacaktır.

Anahtar Kelimeler: ChatGPT-3.5, Copilot, Gemini, English and Turkish, Oculofacial plastic and orbital surgery.

ORCID ID of the authors: ES. [0000-0002-4401-8435](https://orcid.org/0000-0002-4401-8435), MÇ. [0000-0002-0558-5576](https://orcid.org/0000-0002-0558-5576)

Received 22.07.2024

Accepted 03.09.2024

Online published 04.09.2024

Correspondence: Eyüpcan ŞENSOY– Ankara Etlik Şehir Hastanesi, Göz Hastalıkları Kliniği, Ankara, Türkiye
e-mail: dreyupcansensoy@yahoo.com

1. Giriş

Teknolojideki son gelişmeler ile birlikte yapay zeka uygulamaları da hayatımızın her alanında yaygınlaşmaya başlamış, etkisinin izlendiği bir alan da tıp olmuştur(1). Yapay zeka uygulamalarının tıpta yaygınlaşması sonrasında oftalmoloji de sıklıkla etkilenmiş özellikle 2015 yılı sonrasında çok çeşitli hastalıkların tanı ve tedavi takiplerinde kendisine önemli bir yer bulmuştur(2-4). Oftalmolojide kullanılan bu ilk programlar sıklıkla kalıpları öğrenerek yanıtlar verebilen derin öğrenme temelli yapay zeka uygulamaları idi. Son zamanlarda gitgide popülerlikleri artan Büyük Dil Modeli (BDM) temelli sohbet botları ise derin öğrenme temelli uygulamalardan farklı olarak verilen girdileri anlayabilme, yorumlayabilme ve çeşitli olasılıklarla tahminler yürütebilme yetisine sahip, insan düşünce sistemini taklit etmeyi amaçlayan güncel programlardır(5). Üç farklı üretici tarafından ücretsiz olarak kullanıma açılmış ChatGPT-3,5 (Open AI), Copilot (Microsoft) ve Gemini (Google AI) yapay zeka sohbet botları bu alanın önemli temsilcilerini oluşturmaktadır(5-7).

Büyük Dil Modeli temelli bu sohbet botlarının performansları, avantajları ve dezavantajları son zamanlarda oftalmoloji alanında sıkça araştırılan önemli bir konu haline gelmiştir. Çalışmamızın amacı ChatGPT-3,5, Copilot ve Gemini yapay zeka sohbet botlarının okülofasiyal plastik ve orbita cerrahisi hakkındaki çoktan seçmeli sorulardaki başarısına, soruların farklı dillerde sorulmasının etkilerini araştırmaktır.

2. Gereç ve Yöntemler

Amerikan Akademi ve Oftalmoloji 2023-2024 Basic and Clinical Science Course (BCSC) Okülofasiyal Plastik ve Orbita Cerrahisi adlı kitabın çalışma soruları kısmında yer alan 30 sorunun tamamı araştırmaya dahil edildi(8). Bu sorular metin tabanlı sorular olup tablo, görsel veya figür içermemekte idi. Soruların Türkçe çevirileri sertifikasyonlu çevirmen (native speaker) tarafından gerçekleştirildikten sonra soruların İngilizce ve Türkçe versiyonları soru metninde herhangi bir değişiklik yapılmadan ChatGPT-3,5 (OpenAI; San Francisco, CA), Copilot (Microsoft, Redmond, WA) ve Gemini

(Google, Mountain View, California, United States) yapay zeka sohbet botlarına 20 Temmuz 2024 tarihinde soruldu. Her soru uygulanmadan önce ‘Sana çoktan seçmeli sorular soracağım. Lütfen bana doğru cevap şikkını belirt.’ komutu verildi ve her soru sonrasında oturum sonlandırılarak oturum tekrar başlatıldı. Sohbet botlarının sorulara verdikleri cevaplar kitap arkasında yer alan cevap anahtarı ile karşılaştırıldı. Doğru ve yanlış olarak gruplandırıldı. Çalışmamızda insan ve hayvan kaynaklı veriler mevcut olmadığından etik kurul onayına ihtiyaç yoktur.

İstatistiksel Analiz

Verilerin analizinde Statistical Package for the Social Sciences sürüm 23 (SPSS Inc., Chicago, IL, USA) programı kullanıldı. Yüzdelik değerler hesaplandı. Bağımsız gruplardaki nominal verilerin istatistiksel analizinde Pearson ki-kare testi ve Yates ki-kare testi kullanıldı. Bağımlı gruplardaki nominal verilerin istatistiksel analizinde McNemar testi kullanıldı. Anlamlılık düzeyi $p < 0,05$ olarak kabul edildi.

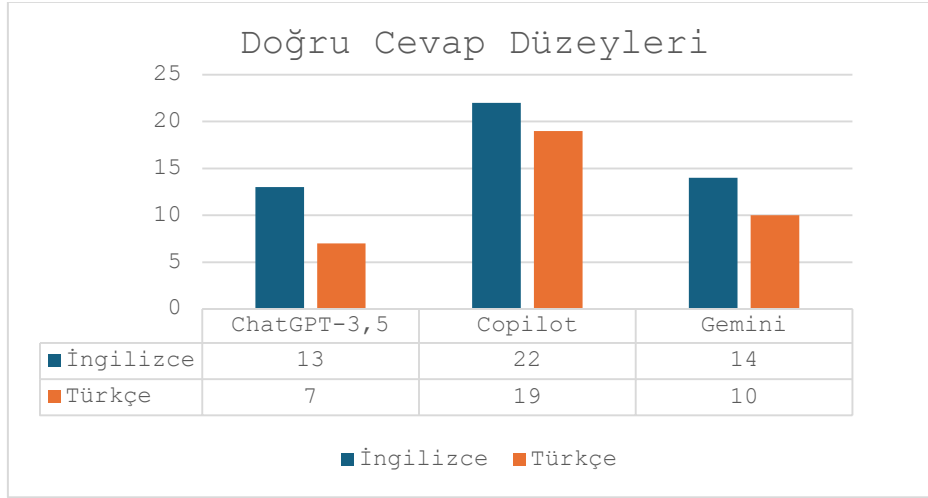
3. Bulgular

Okülofasiyal plastik ve orbita cerrahisi ile ilişkili 30 İngilizce soru sohbet botlarına uygulandı. ChatGPT-3,5 soruların 13’üne (%43,3) doğru cevap, 17’sine (%56,7) yanlış cevap verdi. Copilot sorulan soruların 22’sine (%73,3) doğru cevap, 8’ine (%26,7) yanlış cevap verdi. Gemini sorulan soruların 14’üne (%46,7) doğru cevap, 16’sına (%53,3) yanlış cevap verdi (Figür 1). Soruları cevaplamada Copilot, ChatGPT-3,5’a göre daha yüksek başarı gösterdi ($p=0,036$, Yates ki-kare test). ChatGPT-3,5 ve Gemini arasında ayrıca Copilot ve Gemini arasında istatistiksel olarak anlamlı düzeyde başarı farkı tespit edilmedi (sırası ile $p=1,0$, $p=0,065$ Yates ki-kare test).

Soruların Türkçe versiyonları yapay zeka sohbet botlarına uygulandı. ChatGPT-3,5 soruların 7’sine (%23,3) doğru cevap, 23’üne (%76,7) yanlış cevap verdi. Copilot sorulan soruların 19’una (%63,3) doğru cevap, 11’ine (%36,7) yanlış cevap verdi. Gemini sorulan soruların 10’una (%33,3) doğru cevap, 20’sine

(%66,7) yanlış cevap verdi (Figür 1). Copilot, ChatGPT-3,5 ve Gemini ye göre Türkçe sorularda daha yüksek başarı gösterdi (sırası ile $p=0,004$, $p=0,039$, Yates ki-kare).

ChatGPT-3,5 ve Gemini'nin soruların Türkçe versiyonlarını cevaplamada istatistiksel olarak anlamlı düzeyde bir fark tespit edilmedi ($p=1,0$ Yates ki kare test).



Figür 1. ChatGPT-3,5, Copilot ve Gemini'nin okülofasyal plastik ve orbita cerrahisi ile ilişkili İngilizce ve Türkçe soruları doğru cevaplama düzeyleri

ChatGPT-3,5; soruların İngilizce ve Türkçe versiyonlarının 22'sine (%73,3) aynı cevap, 8'ine (%26,7) farklı cevap verdi. Farklı cevap verdikleri soruların 1'i (%12,5) Türkçe sorulduğunda doğru olarak cevaplanmış iken 7'si (%87,5) Türkçe sorulduğunda yanlış cevaplandı. ChatGPT-3,5, İngilizce ve Türkçe soruları cevaplamada benzer performanslar gösterdi ($p=0,07$, McNemar test) (Tablo 1).

Copilot; soruların İngilizce ve Türkçe versiyonlarının 25'ine (%83,3) aynı cevap, 5'ine (%16,7) farklı cevap verdi. Farklı cevap verdikleri soruların 1'i (%20) Türkçe sorulduğunda doğru olarak cevaplanmış iken

4'ü (%80) Türkçe sorulduğunda yanlış cevaplandı. Copilot, İngilizce ve Türkçe soruları cevaplamada benzer performanslar gösterdi ($p=0,375$, McNemar test) (Tablo 1).

Gemini; soruların İngilizce ve Türkçe versiyonlarının 20'sine (%66,7) aynı cevap, 10'una (%33,3) farklı cevap verdi. Farklı cevap verdikleri soruların 3'ü (%30) Türkçe sorulduğunda doğru olarak cevaplanmış iken 7'si (%70) Türkçe sorulduğunda yanlış cevaplandı. Gemini, İngilizce ve Türkçe soruları cevaplamada benzer performanslar gösterdi ($p=0,344$, McNemar test) (Tablo 1).

Tablo 1. Yapay zeka sohbet botlarının aynı sorulara verdikleri cevaplar ve değişimleri

Cevaplar	ChatGPT-3,5 (İngilizce)	ChatGPT-3,5 (Türkçe)	Copilot (İngilizce)	Copilot (Türkçe)	Gemini (İngilizce)	Gemini (Türkçe)
Doğru	13 (%43,3)	7 (%23,3)	22 (%73,3)	19 (%63,3)	14 (%46,7)	10 (%33,3)
Yanlış	17 (%56,7)	23 (%76,7)	8 (%26,7)	11 (%36,7)	16 (%53,3)	20 (%66,7)
P değeri	0,07*		0,375*		0,344*	
Aynı cevabı üretme	22 (%73,3)		25 (%83,3)		20 (%66,7)	
Farklı cevabı üretme	8 (%26,7)		5 (%16,7)		10 (%33,3)	
Doğru-yanlış değişimi	7 (%87,5)		4 (%80)		7 (%70)	
Yanlış-doğru değişimi	1 (%12,5)		1 (%20)		3 (%30)	

*: McNemar testi

4. Tartışma

Büyük Dil Modelleri son zamanlardaki teknolojik gelişmelerin bir ürünü olarak karşımıza çıkmış, tıp eğitimi içerisinde doğru bilgiye hızlı erişim, özel öğrenme planı belirleyebilme, dil çevirileri gibi çok çeşitli konularda kullanılmaya başlanmış ve bu konularda fayda sağlamış yapay zekanın bir koludur(9). Ayrıca BMD temelli bu yapay zeka programları çok çeşitli hastalıkların ayırıcı tanısını yapabilme, tanının konulabilmesine yardımcı olabilme ve tedavisi ile ilgili bilgi verebilme gibi çok çeşitli faydalı özellikleri de içerisinde barındırmaktadır. Sohbet botlarının bu avantajları tıp eğitimi alanlar öğrencilerden, sağlık profesyonellerine kadar çok geniş bir kitlenin dikkatini çekmiş ve bu faydalı etkileri sıklıkla araştırılan bir konu haline gelmiştir(10).

Oftalmoloji BDM temelli yapay zeka sohbet botlarının avantajlarının ve dezavantajlarının sıkça araştırıldığı ve kullanım alanlarının yaygın bir şekilde test edildiği bir alandır.

Oftalmolojideki etkileri araştırılırken sıklıkla incelenen bir bilimi de okülofasiyal plastik ve orbita cerrahisi alanı olmuştur. Güncel bir çalışmada hastaların oküloplastik cerrahilerine yönelttikleri sorular toplanmış ve bu sorular ChatGPT-3,5 ve Bard sohbet botlarına uygulanmış. Daha sonra bu programların yanıtları oküloplastik cerrahlar tarafından değerlendirilmiştir. Hastalardan alınan 112 soruya verilen cevaplar kapsamlı, doğru ancak yetersiz, karışık ve tamamen yanlış olarak 4 grupta incelenmiş, ChatGPT-3,5 sorulara bu kategorilerde sırası ile %71,4, %12,9, %10,5 ve %5,1 oranında başarı göstermişlerdir. Bard ise sırası ile %53,1, %18,3, %18,1 ve %10,5 oranında başarı göstermiştir. Bu veriler sonucunda yazarlar bu sohbet botlarının geliştirilmesi ile birlikte doktorlara yardımcı olarak görev alabileceğini ve hastalara doğru bilgiler verebilmede etkin bir araç olarak kullanılabileceğini belirtmişlerdir(11). Oftalmolojideki çoktan seçmeli sorulardaki başarıları da bir başka incelenen konu olmuştur. Haddad ve ark.'nın yaptıkları bir

çalışmada 380 soru ChatGPT-3,5 ve ChatGPT-4,0 sohbet botlarına uygulanmış, sırası ile bu programlar %55 ve %70 oranında başarı göstermişlerdir. Oküloplastik ve orbita cerrahisi ile ilişkili sorulardaki başarıları ise ChatGPT-3,5'da %29, ChatGPT-4,0'da %50 düzeyinde olmuştur(12). ChatGPT-3,5 ve Bing'in oftalmoloji sorularındaki başarılarının test edildiği başka bir çalışmada ChatGPT-3,5 soruların %59,69'unu, Bing ise soruların %73,6'sını doğru cevapladığı belirtilmiştir. Orbita, göz kapağı ve lakrimal sistem ile ilgili sorularda ise ChatGPT-3,5'un %46,77, Bing'in ise %70,9 oranında başarı gösterdikleri ifade edilmiş ve Bing'in çevrimiçi internet erişiminin varlığı, alıntılama özelliği sayesinde oftalmoloji eğitimi alanlara ek faydalar sağlayabileceği ifade edilmiştir(13). Türkçe oftalmoloji sorularında sohbet botlarının başarısını inceleyen bir başka çalışmada ChatGPT-3,5, Bing ve Bard'ın soruları sırası ile %51, %63 ve %45,5 oranında doğru cevapladıkları; adneksler, üvea ve oküloplasti ile ilişkili sorularda ise sırası ile %62,1, %69 ve %58,6 oranında başarı gösterdikleri belirtilmiş ve sohbet botlarının geliştirilmeye ihtiyacının vurgusu yapılmıştır(14). Bir farklı konuda sohbet botlarının sorulduğu ülkeye göre performanslarının etkilenip etkilenmediğinin araştırılmasıdır. Oftalmoloji ile ilişkili 150 soru 4 farklı ülke internet erişimi kullanılarak Gemini sohbet botuna uygulanmış ve soruların erişildiği bölgeye göre değişen doğruluklarda cevaplama oranlarının ortaya çıktığı belirtilmiştir. Ayrıca orbita ve oküloplastik cerrahi ile ilişkili sorularda da bulunduğu bölgeye göre başarı düzeylerinin değiştiği, %70 ile %90 arasında değişen düzeylerde başarı oranlarının olduğu belirtilmiştir. Araştırmacılar bu verileri sohbet botlarındaki performansların bulunduğu ülkeye göre farklılık gösterdiği ve bu konunun daha geliştirilip araştırılması gereken bir konu olması gerektiğini belirtmişlerdir(15).

Copilot hem İngilizce hem Türkçe sorularda en başarılı program idi. Bu başarısının güncel internet erişimi sağlayabilmesi, soruları anlama yorumlama ve dil çeviri kabiliyetinin farklılığı ve üstünlüğü sonucunda gelişmiş olabileceğini düşünmekteyiz. Ama esas vurguladığımız konu ise aynı sorularda sohbet

botlarının gösterdikleri başarılarıdır. Her ne kadar sohbet botları kendi içlerinde İngilizce ve Türkçe soruları cevaplarken benzer performanslar göstermiş olsalar da verilen cevaplar soru bazlı incelendiğinde aynı sorulara farklı cevaplar ürettiği gözlemlenmektedir. Ayrıca dikkat çekici bir başka konu da her ne kadar Türkçe sorulduğunda doğru olarak cevaplanmış, İngilizce sorulduğunda ise yanlış cevaplanmış sorular olsa da soruların çoğunlukla Türkçe sorulduğunda doğru cevabın yanlış cevaba dönmüş olmasıdır. Bu durumun çok çeşitli sebepleri olabilir. Literatürün sıklıkla İngilizce olması soruların İngilizce versiyonlarının daha yüksek doğru cevaplanmasına yardımcı olmuş olabilir. Ayrıca bu durum sohbet botlarının farklı dilleri anlama, yorumlama ve çözüm üretebilme kabiliyetlerinin dillere göre farklılık gösterebilmesi ile ilişkili olabilir. Bu konudaki farklı düşüncelerin doğruluğu ileri çalışmalar ile desteklenmesi gerekmektedir.

Önemli bilgileri barındıran ve temel kitaplar arasında yer alan Amerikan Akademi ve Oftalmoloji 2023-2024 Basic and Clinical Science Course (BCSC) Okülofasiyal Plastik ve Orbita Cerrahisi kitabının çalışma soruları 30 soruyu içeriyordu ve sohbet robotlarına bu soruları sorduk. Soru sayısının az olmasının istatistiksel sonucun anlamlı çıkıp çıkmamasına etki edebileceğini ön görmekteyiz. Fakat, temel bilgileri ölçen bu kitabın sorularına ilave soru eklemeyi uygun bulmadık. Daha fazla soru içeren testlerde farklı değerlerin çıkıp çıkmayacağının araştırılması gerektiğini ön görebiliriz.

Soru sayımızın azlığı, tek merkez çalışması oluşu ve alt grupların ayrı değerlendirilememiş olması çalışmanın kısıtlılık yönlerini oluşturmaktadır.

Sonuç olarak her ne kadar sohbet botlarının daha aktif ve doğru bir şekilde kullanılabilmesi için bilgi düzeylerinin artırılması gereklilik gösterse de farklı dillerdeki performansları araştırılmayı ve geliştirilmeyi hak eden diğer önemli bir konuyu oluşturmaktadır. Bilgi ve dil performanslarının geliştirilmesi bu programlara olan güveni artıracak ve programların yaygın etkinliklerini önemli düzeyde geliştirecektir.

KAYNAKLAR

1. Rahimy E. Deep learning applications in ophthalmology. *Curr Opin Ophthalmol*. 2018;29(3):254-60.
2. Ting DSW, Pasquale LR, Peng L, et al. Artificial intelligence and deep learning in ophthalmology. *Br J Ophthalmol*. 2019;103(2):167-75.
3. Antaki F, Coussa RG, Kahwati G, Hammamji K, Sebag M, Duval R. Accuracy of automated machine learning in classifying retinal pathologies from ultra-widefield pseudocolour fundus images. *Br J Ophthalmol*. 2023;107(1):90-5.
4. Schmidt-Erfurth U, Sadeghipour A, Gerendas BS, Waldstein SM, Bogunović H. Artificial intelligence in retina. *Prog Retin Eye Res*. 2018;67:1-29.
5. Mikolov T, Deoras A, Povey D, Burget L, Černocký J. Strategies for training large scale neural network language models. 2011 IEEE Workshop on Automatic Speech Recognition and Understanding, ASRU 2011, Proceedings. Published online 2011:196-201.
6. Google AI updates: Bard and new AI features in Search. Accessed July 4, 2024. <https://blog.google/technology/ai/bard-google-ai-search-updates/>
7. Bing Chat | Microsoft Edge. Accessed July 4, 2024. <https://www.microsoft.com/en-us/edge/features/bing-chat?form=MT00D8>
8. Korn BS, Burkat CN, Couch SM, et al., eds. *Oculofacial Plastic and Orbital Surgery*. American Academy of Ophthalmology; 2023.
9. Khan RA, Jawaid M, Khan AR, Sajjad M. ChatGPT - Reshaping medical education and clinical management. *Pak J Med Sci*. 2023;39(2):605.
10. Jeblick K, Schachtner B, Dextl J, et al. ChatGPT Makes Medicine Easy to Swallow: An Exploratory Case Study on Simplified Radiology Reports. Published online December 30, 2022. Accessed June 10, 2023. <https://arxiv.org/abs/2212.14882v1>
11. Al-Sharif EM, Penteado RC, Dib El Jalbout N, et al. Evaluating the Accuracy of ChatGPT and Google BARD in Fielding Oculoplastic Patient Queries: A Comparative Study on Artificial versus Human Intelligence. *Ophthalmic Plast Reconstr Surg*. 2024;40(3):303-11.
12. Haddad F, Saade JS. Performance of ChatGPT on Ophthalmology-Related Questions Across Various Examination Levels: Observational Study. *JMIR Med Educ*. 2024;10:e50842.
13. Tao BKL, Hua N, Milkovich J, Micieli JA. ChatGPT-3.5 and Bing Chat in ophthalmology: an updated evaluation of performance, readability, and informative sources. *Eye* 2024. Published online March 20, 2024:1-6.
14. Canleblebici M, Dal A, Erdağ M. Evaluation of the Performance of Large Language Models (ChatGPT-3.5, ChatGPT-4, Bing and Bard) in Turkish Ophthalmology Chief-Assistant Exams: A Comparative Study. *Türkiye Klinikleri J Ophthalmol*. Published online June 11, 2024.
15. Mihalache A, Grad J, Patil NS, et al. Google Gemini and Bard artificial intelligence chatbot performance in ophthalmology knowledge assessment. *Eye* 2024. Published online April 13, 2024:1-6.

Etik Bilgiler

Etik Kurul Onayı: Çalışmamızda insan ve hayvan kaynaklı veriler mevcut olmadığından etik kurul onayına ihtiyaç yoktur.

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

Yazar Katkı Oranları: Cerrahi ve Tıbbi Uygulamalar:EŞ,MÇ. Konsept: EŞ. Tasarım: EŞ, MÇ. Veri Toplama veya İşleme: EŞ. Analiz veya Yorum: EŞ, MÇ. Literatür Taraması: EŞ. Yazma: EŞ.

Hakem Değerlendirmesi: Hakem değerlendirmesinden geçmiştir.

Çıkar Çatışması Bildirimi: Yazarlar çıkar çatışması olmadığını beyan etmişlerdir.

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