

Onkolojide Yapay Zeka

YAPAY ZEKÂ KLİNİK KARAR VERMEYE HAZIR MI? PROSTAT KANSERİNDE BÜYÜK DİL MODELLERİNİN HEKİM REFERANSINA KARŞI PERFORMANSI

Gökhan Çolak¹, Ümit Bulut², İrem Aşık Bal³, Sercan Özbek¹, Dilara Özgül¹, Ayhan Açlan⁴, Muhammed Ali Kaypak⁴, Bilgin Demir¹, Merve Turan¹, Esin Oktay¹

¹Adnan Menderes Üniversitesi Tıp Fakültesi, Tıbbi Onkoloji Bilim Dalı, Aydın

²Ege Üniversitesi Fen Bilimleri Enstitüsü, Bilgisayar Mühendisliği Ana Bilim Dalı, İzmir

³Adnan Menderes Üniversitesi Tıp Fakültesi, İç Hastalıkları Ana Bilim Dalı, Aydın

⁴T.c. Sağlık Bakanlığı Aydın Atatürk Devlet Hastanesi, Tıbbi Onkoloji Kliniği, Aydın

Giriş-Amaç

Amaç: Yapay zekâ ve özellikle büyük dil modelleri (LLM), onkolojide tanı, prognostik değerlendirme ve tedavi planlamasına yönelik karar destek potansiyeli ile giderek daha fazla ilgi görmektedir. Bununla birlikte, gerçek yaşam klinik verileri karşısındaki doğrulukları ve uzman kararlarıyla uyumları henüz yeterince gösterilmemiştir. Bu çalışmada, metastatik hormon duyarlı prostat kanseri (mHSPC) yönetiminde dört LLM'nin (ChatGPT 5.2, Gemini 3, Claude, DeepSeek) klinik kararlarını 3 kişilik tıbbi onkoloji konseyi kararlarıyla karşılaştırarak uyum değerlendirmesi amaçlanmıştır.

Yöntem: Adnan Menderes Üniversitesi'nde takip edilen de novo mHSPC'li 193 hastanın verileri anonimleştirildi ve tedavi kararı için standart formlara aktarıldı. Modeller ve konseyden; volüm/risk sınıflaması, tedavi amacı (küratif/palyatif) ve tedavi önerisi değerlendirmesi istendi; konsey her hasta için tek ortak karar verdi. Uyum Cohen kappa ve yüzde uyumla analiz edildi; $p < 0,01$ anlamlı kabul edildi.

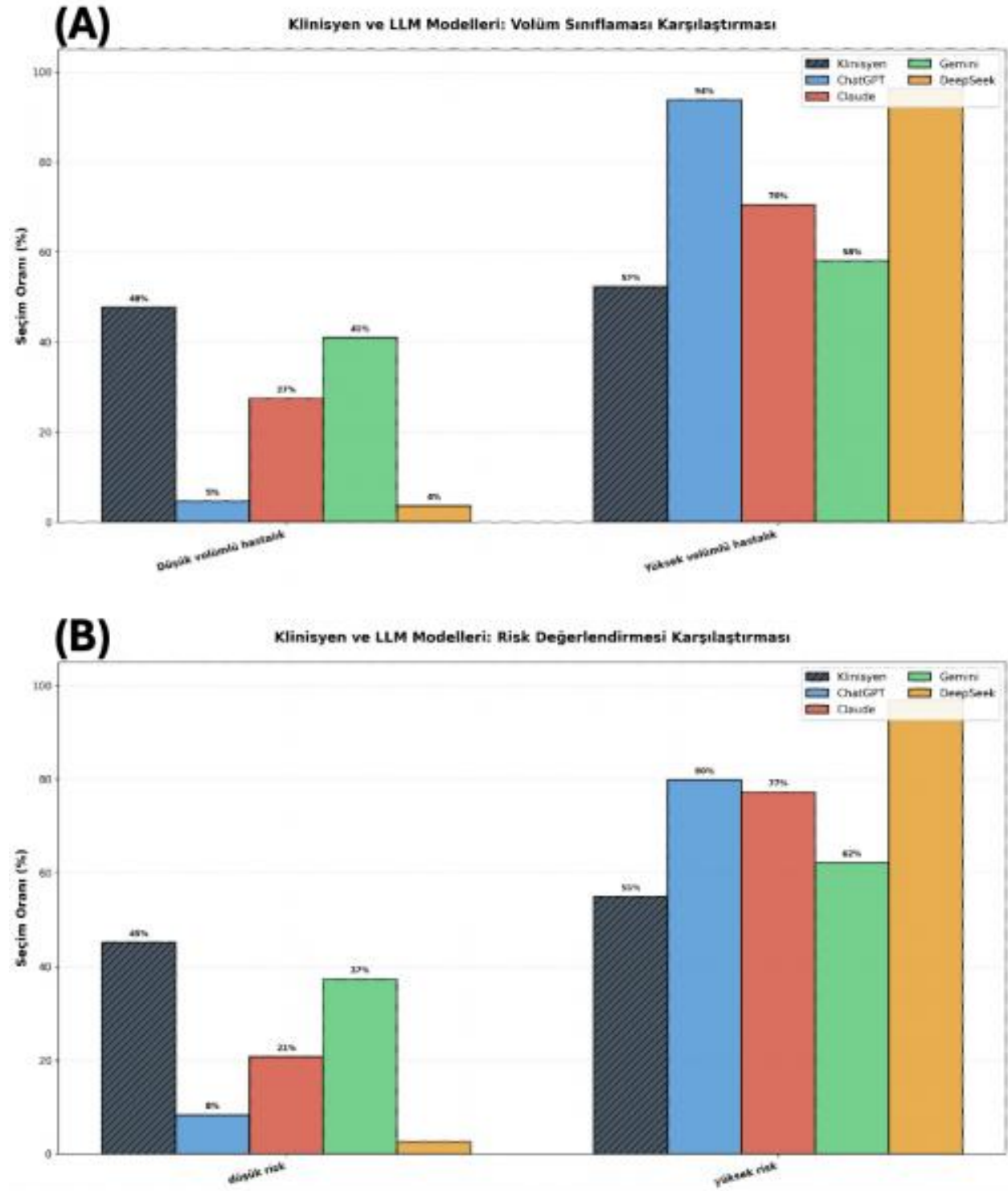
Bulgular: Klinisyen referansına göre volüm değerlendirmesinde (CHAARTED) en yüksek uyum Gemini'de saptandı ($\kappa = 0,814$; %90,7); Claude orta düzeyde kaldı ($\kappa = 0,544$; %77,2). ChatGPT ($\kappa = 0,096$; %55,9) ve DeepSeek ($\kappa = 0,058$; %54,9) düşük uyum gösterdi ve düşük volüm olgularını çoğunlukla yüksek volüm olarak sınıflandırdı. Risk sınıflamasında (LATITUDE) benzer şekilde Gemini yüksek uyum sağladı ($\kappa = 0,768$; %88,6), Claude orta düzeydeydi ($\kappa = 0,463$; %74,1); ChatGPT ($\kappa = 0,181$; %62,2) ve DeepSeek ($\kappa = 0,069$; %57,5) zayıf kaldı. Klinisyen 25 hastada küratif hedef belirlemişken, ChatGPT/Gemini/DeepSeek hiç küratif önermezken Claude yalnızca 3 olguda eşleşti. Sistemik tedavi seçiminde en yüksek uyum Claude'da ($\kappa = 0,532$; %73,6) olup Gemini ($\kappa = 0,412$; %65,8), ChatGPT ($\kappa = 0,302$; %59,1) ve DeepSeek ($\kappa = 0,150$; %58,5) izledi. Tüm modeller ADT tekli tedaviyi belirgin biçimde az tanıyıp ikili/üçlü rejimlere kayd; DeepSeek ikili, ChatGPT ve Gemini üçlü rejimleri orantısız sıklıkta önerdi. ARPI alt tip eşleşmesi Abirateron'da görece daha iyi, Apalutamid/Darolutamid'de ise yok denecek düzeydeydi.

Sonuç: mHSPC yönetiminde LLM–klinisyen uyumu heterojendi: Gemini volüm/riskte, Claude tedavi seçiminde daha iyi; ChatGPT/DeepSeek yanlış sınıflama ve rejim yoğunlaştırma eğilimindeydi. Bu çalışma; LLM'lerin klinik kararlarda bağımsız kullanım için henüz erken olduğunu, onkolojide ancak uzman denetiminde destekleyici araç olabileceklerini göstermektedir.

Anahtar Kelimeler : Büyük dil modelleri, Prostat kanseri, Klinik karar uyumu

Resimler :

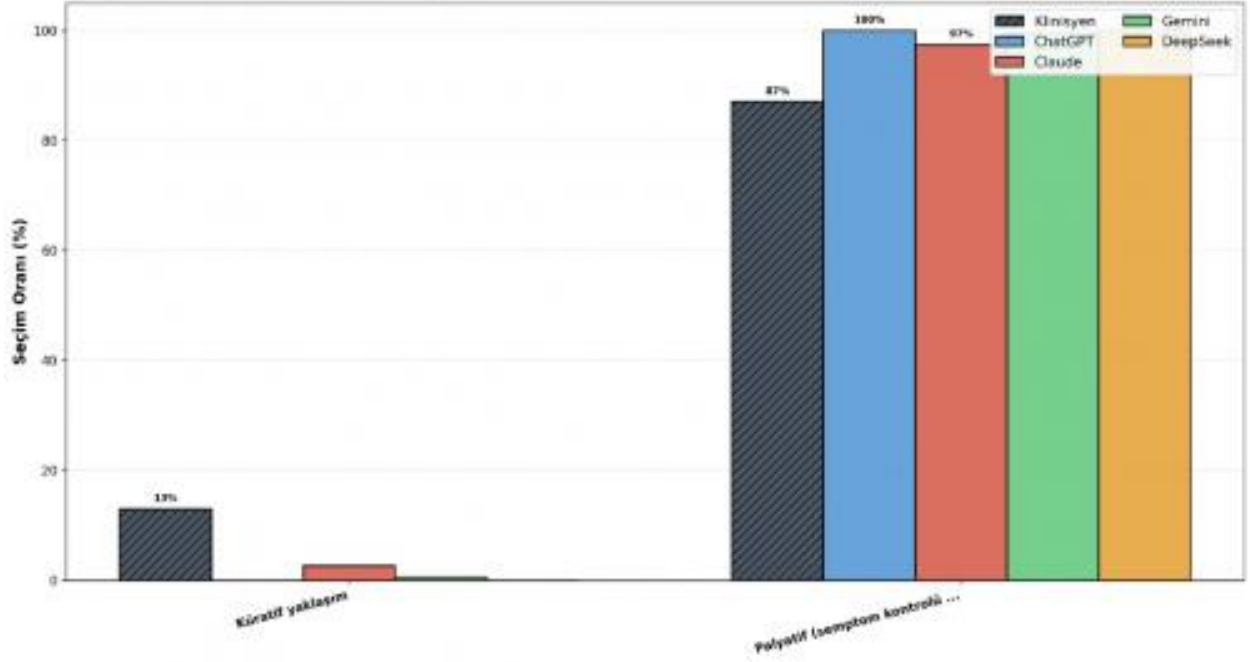
Resim Açıklaması: Şekil 1. Klinisyen ve LLM Modelleri (A) CHAARTED kriterlerine göre volüm değerlendirmesi seçim oranları, (B) LATITUDE kriterlerine göre risk değerlendirmesi seçim oranları.



Resim Açıklaması:

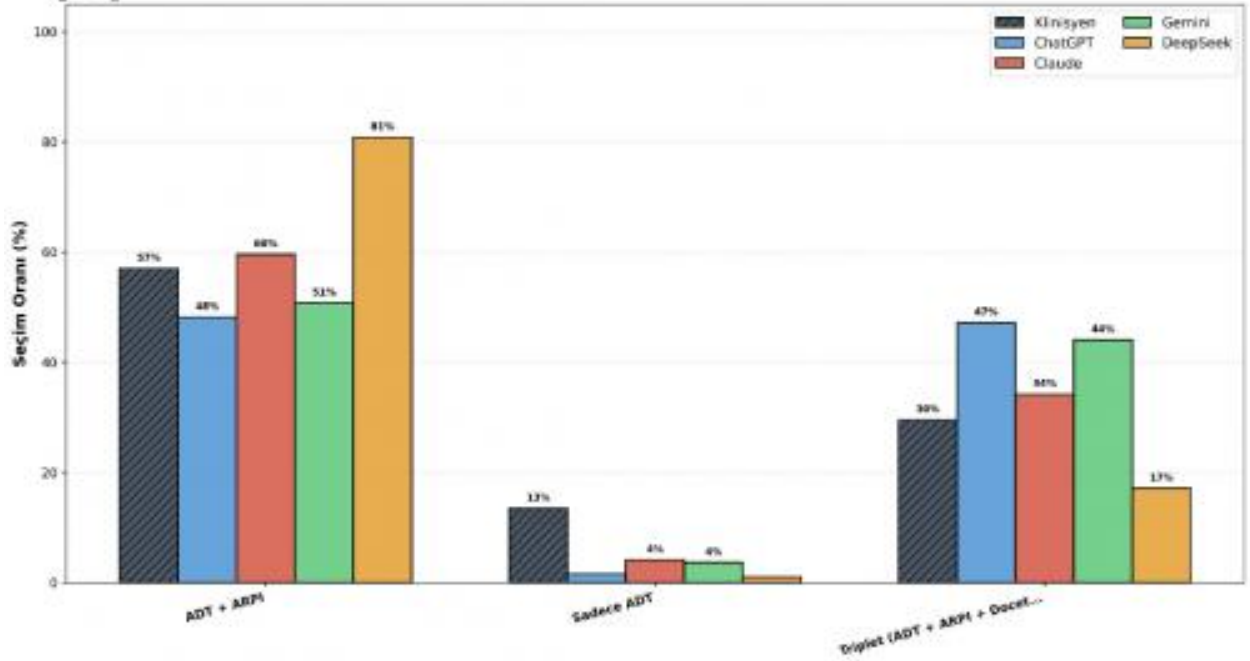
(A)

Klinisyen ve LLM Modelleri: Tedavi Hedefi Karşılaştırması

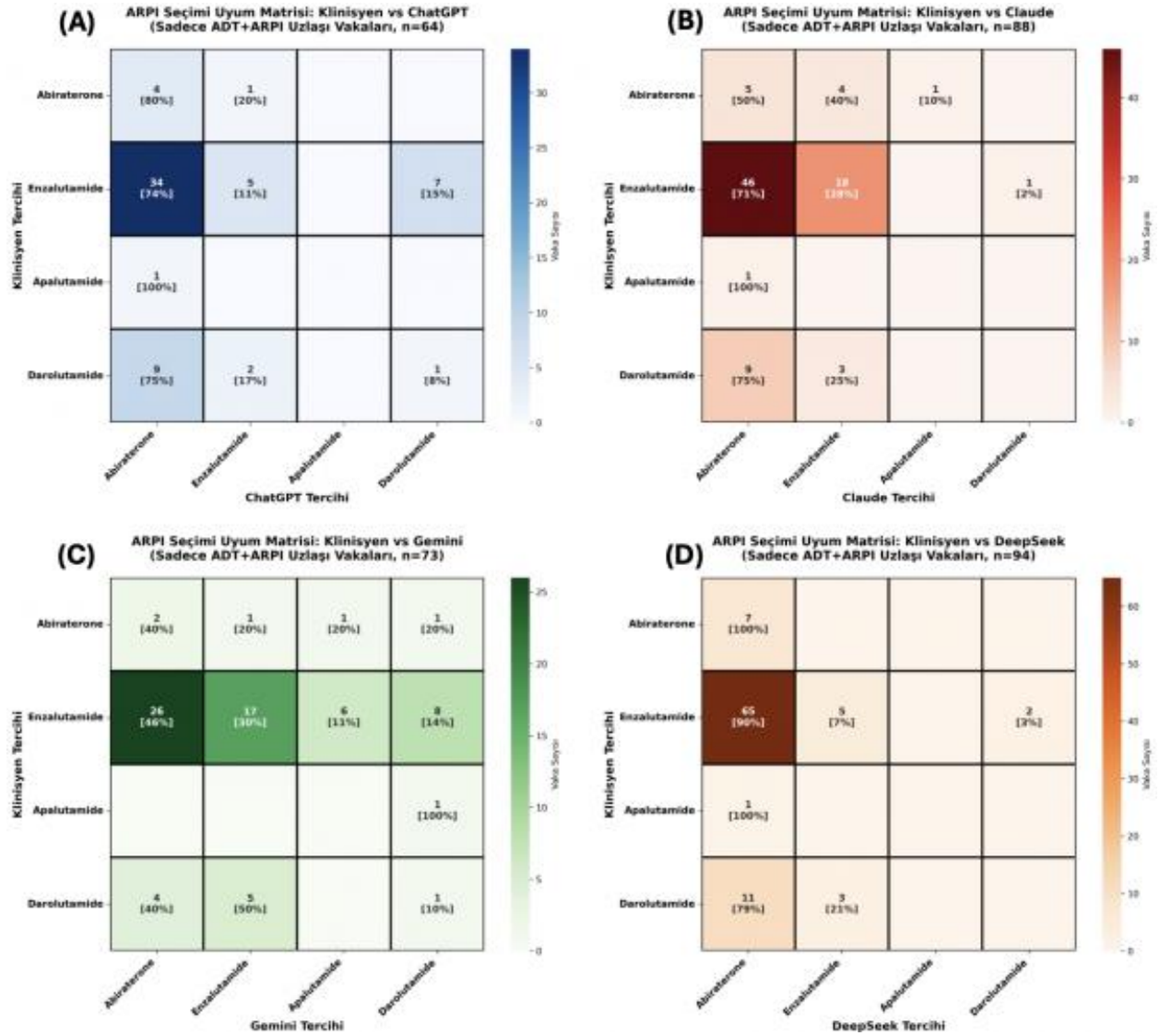


(B)

Klinisyen ve LLM Modelleri: Tedavi Seçimi Karşılaştırması



Resim Açıklaması:



Tables :

Tablo 1. LLM modellerinin hastalık volüm değeriendirilmesinde klinisyen referansına göre performansı

		Klinisyen	Klinisyen	Klinisyen			
		Düşük Volümlü Hastalık	Yüksek Volümlü Hastalık	Belirsiz/ Sınıflandırılmadı	Kappa	p değeri	Yüzde Uyum
		n(%)	n(%)	n(%)			
ChatGPT	Düşük Volümlü Hastalık	9 (9,8)	0(,0)	0(,0)	0,096	0,003	55,9

	Yüksek Volümlü Hastalık	82(89,1)	99(98,0)	0(,0)			
	Belirsiz/ Sınıflandırılmadı	0(,0)	1(1,0)	0(,0)			
Claude	Düşük Volümlü Hastalık	51(55,4)	2(2,0)	0(,0)	0,544	<0,001	77,2
	Yüksek Volümlü Hastalık	38(41,3)	98(97,0)	0(,0)			
	Belirsiz/ Sınıflandırılmadı	3(3,3)	1(1,0)	0(,0)			
Gemini	Düşük Volümlü Hastalık	77(83,7)	2(2,0)	0(,0)	0,814	<0,001	90,7
	Yüksek Volümlü Hastalık	14(15,2)	98(97,0)	0(,0)			
	Belirsiz/ Sınıflandırılmadı	1(1,1)	1(1,0)	0(,0)			
DeepSeek	Düşük Volümlü Hastalık	6(6,5)	1(1,0)	0(,0)	0,058	0,040	54,9
	Yüksek Volümlü Hastalık	86(93,5)	100(99,0)	0(,0)			
	Belirsiz/ Sınıflandırılmadı	0(,0)	0(,0)	0(,0)			

Tablo 2. LLM modellerinin hastalık risk değerlendirmesinde klinisyen referansına göre performansı

		Klinisyen	Klinisyen	Klinisyen			
		Düşük Risk	Yüksek Risk	Belirsiz/Sınıflandırılmadı	Kappa	P değeri	Yüzde Uyum
		n(%)	n(%)	n(%)			
ChatGPT	Düşük Risk	15(17,2)	1(,9)	0(,0)	0,181	<0,001	62,2
	Yüksek Risk	71(81,6)	105(99,1)	0(,0)			
	Belirsiz/Sınıflandırılmadı	1(1,1)	0(,0)	0(,0)			

[illegible]

[illegible]