

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

Uncovering Language Disparity of ChatGPT on Retinal Vascular Disease Classification: Cross-Sectional Study

بررسی نابرابری زبانی در عملکرد ChatGPT برای طبقه‌بندی بیماری‌های
عروقی شبکیه: یک مطالعه مقطعی

۱۴۰۴/۲/۲۹

ژورنال کلاب

مبینا فیاض

مشخصات ژورنال

Indexing:
ISI, Scopus , PubMed, DOAJ

Impact Factor : 5.8

**Name : Journal of Medical
Internet Research**

Category : Q1



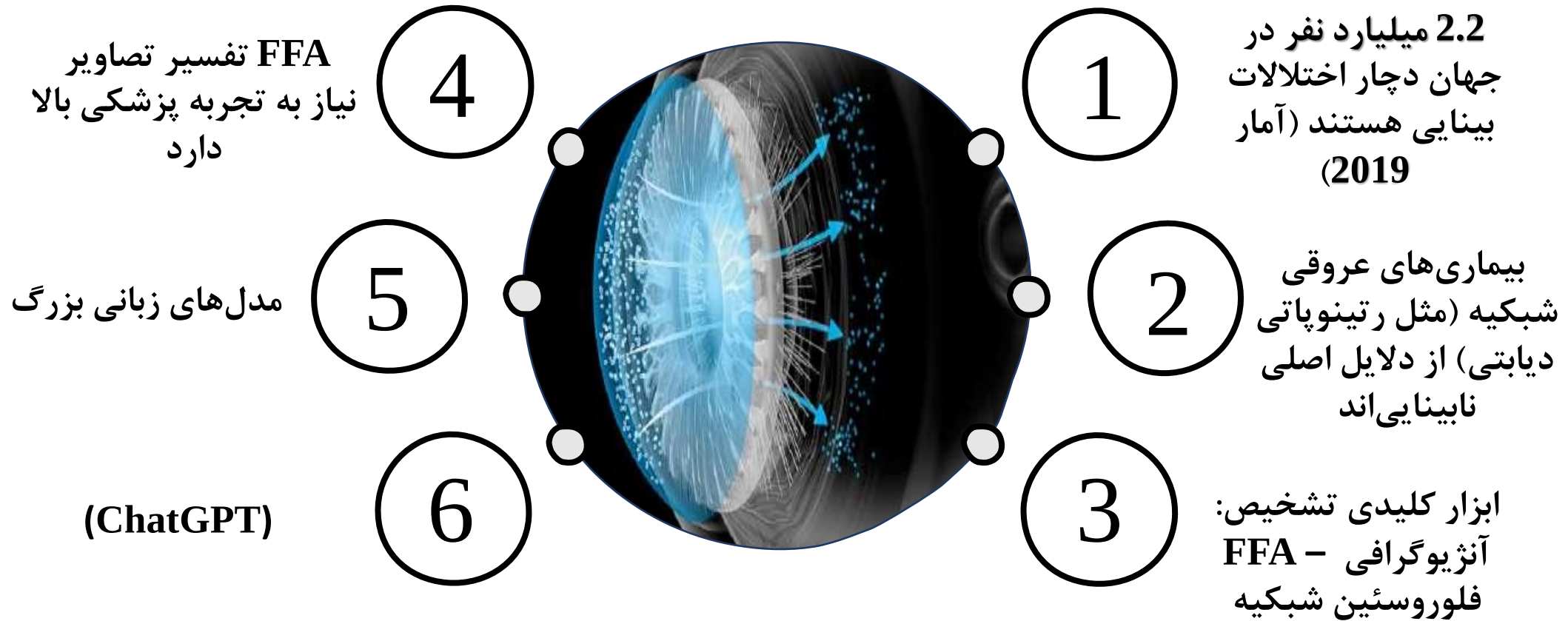
اختصارات

- FFA :Fundus Fluorescein Angiography..... آنژیوگرافی فلورسئین فوندوس
- LLMs :Large Language Models..... مدل های زبانی بزرگ
- API: Application Programming Interface..... رابط برنامه نویسی کاربردی
- AI: Artificial Intelligence هوش مصنوعی

مقدمه و بیان مسئله



مقدمه و بیان مسئله

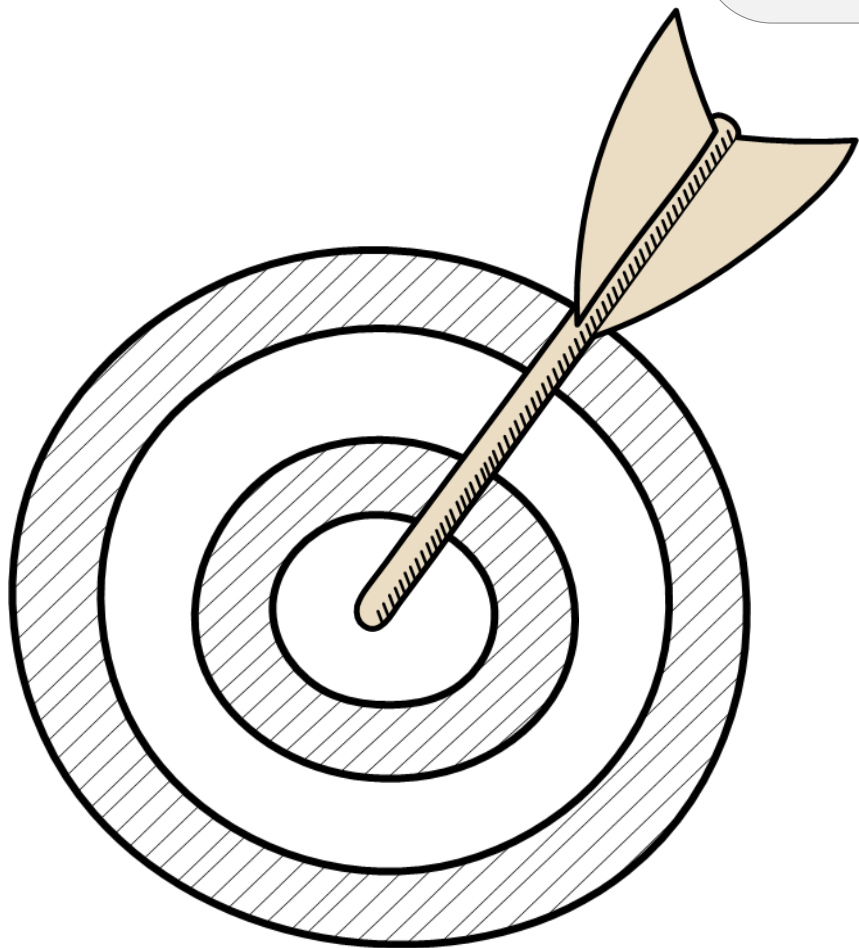


هدف مطالعه



۲

هدف مطالعه



هدف اصلی:

ارزیابی عملکرد و استدلال ChatGPT در تشخیص
بیماری‌های عروقی شبکه از روی گزارش‌های FFA به
زبان چینی

اهداف فرعی:

- مقایسه دقت ChatGPT با پزشکان و کارورزان
- تحلیل تفاوت عملکرد در زبان چینی و انگلیسی
- ارزیابی تأثیر نوع درخواست

روش اجرا



۳

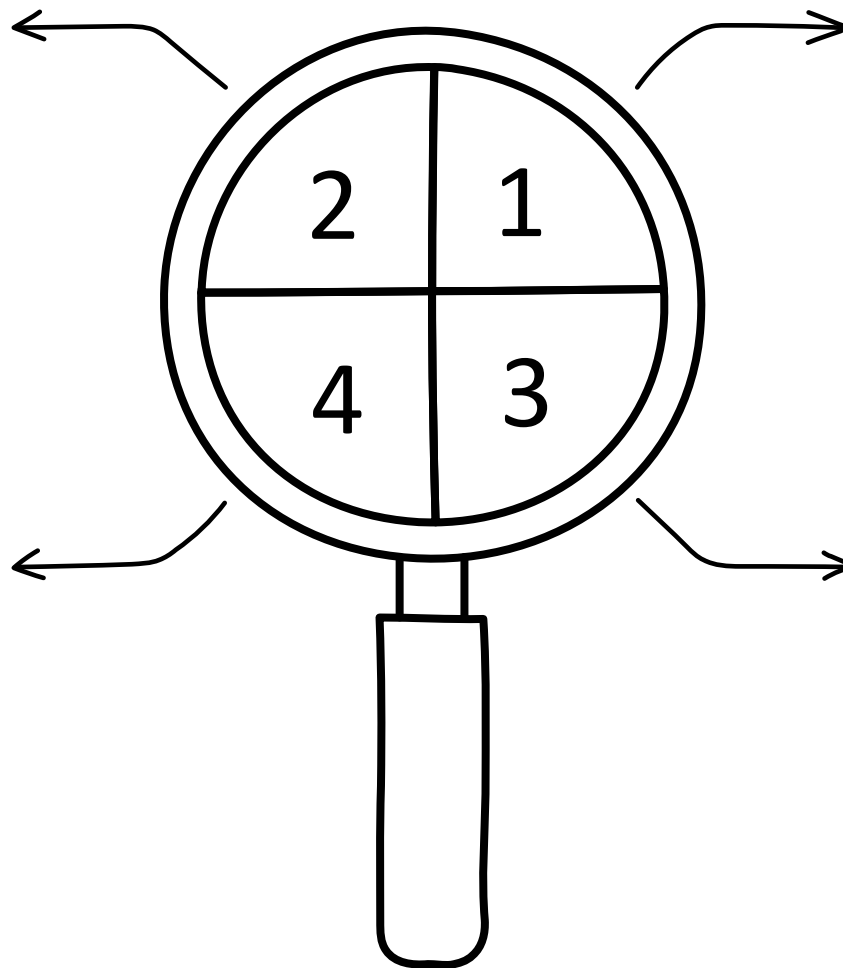
روش اجرا

ملاحظات اخلاقی

- تصویب اخلاقی
- حفظ حریم خصوصی

معیارهای ارزیابی

عملکرد تشخیصی
توانایی استنتاج
حذف اطلاعات
هالوکینیشن
اطلاعات نادرست
ناهماهنگی



داده‌گذاری و جمع‌آوری اطلاعات:
1226 گزارش FFA از 728 بیمار
جمع‌آوری شد.

تشخیص بیماری‌های شبکه‌ای با
استفاده از ChatGPT
پرامپت مستقیم و استدلالی به
دو زبان چینی و انگلیسی

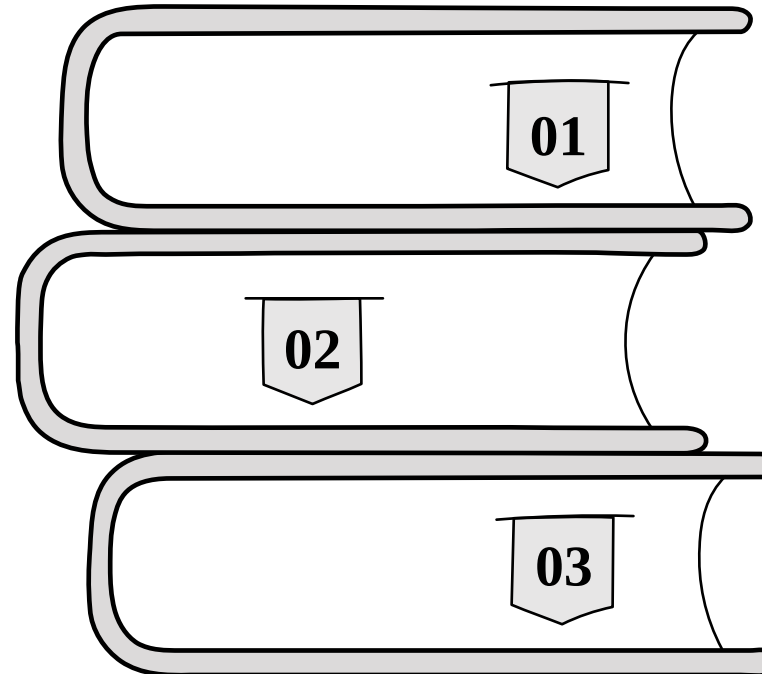


روش اجرا

معیارهای ارزیابی

توانایی استنتاج

- تعداد مراحل استدلال
- تعداد اشتباهات استدلال
- نقص‌های موجود در فرایند استدلال



عملکرد تشخیصی

ارزیابی دقت تشخیصی ChatGPT از سه معیار دقت (Precision)، یادآوری (Recall) و امتیاز F1 استفاده شده است.

حذف اطلاعات

نادیده گرفتن یا جا انداختن نکات مهم موجود در گزارش اصلی.



روش اجرا

معیارهای ارزیابی

هالوکینیشن

تولید اطلاعات یا نتایج پزشکی ساختگی که در گزارش اصلی وجود ندارند.

04

اطلاعات نادرست

ارائه دانش پزشکی اشتباه یا نقل قول ناصحیح از اطلاعات علمی.

05

06

ناهماهنگی

تضاد بین نتیجه گیری نهایی و مراحل استدلال ارائه شده توسط مدل.



یافته ها



یافته ها

Diagnostic Performance

Category	Direct-Chinese (%)			Direct-English (%)			Step-Chinese (%)			Step-English (%)		
	P ^a	R ^b	F ₁	P	R	F ₁	P	R	F ₁	P	R	F ₁
Normal	100	85.47	92.17	100	88.03	93.64	98.39	52.14	68.16	97.37	94.87	96.1
DR ^c	91.55	72.52	80.93	91.05	95.12	93.04	85.07	95.4	89.94	82.13	93.58	87.48
wetAMD ^d	44.72	87.98	59.3	59.92	80.87	68.84	63.58	60.11	61.8	60	34.42	43.75
CSC ^e	4.35	2.74	3.36	33.33	1.37	2.63	34.15	19.18	24.56	50	6.85	12.05
BRVO ^f	41.32	79.37	54.35	63.33	90.47	74.51	83.61	80.95	82.26	67.95	84.13	75.18
CRVO ^g	93.1	71.05	80.6	84.85	73.68	78.87	41.27	68.42	51.49	58.33	73.68	65.12
VKH ^h	0	0	0	0	0	0	0	0	0	0	0	0
Overall	75.03	70.15	70.47	79.61	83.12	80.05	76.24	77.16	75.61	74.56	75.94	73.46

Measurement	Step-Chinese	Step-English	<i>P</i> value ^a
Reasoning steps per report, mean (SD)	1.4 (0.8)	2.6 (1.5)	<.001
Reasoning errors per report, mean (SD)	0.4 (0.5)	0.4 (0.6)	0.88
Incompleteness (%)	63.53	44.31	<.001
Omission of information (%)	0.78	0.39	0.68
Hallucinations (%)	5.88	0.59	<.001
Misinformation (%)	7.84	1.96	<.001
Inconsistency (%)	0.59	0.39	>.99

یافته ها

Confusion Matrix

Direct-Chinese

Normal	100	16	0	0	0	0	0	1
DR	0	520	105	25	64	2	0	1
wetAMD	0	6	161	15	1	0	0	0
CSC	0	0	71	2	0	0	0	0
BRVO	0	6	7	0	50	0	0	0
CRVO	0	4	3	0	4	27	0	0
VKH	0	16	13	4	2	0	0	0
Undiag	0	0	0	0	0	0	0	0

Actual disease

Prediction

Direct-English

Normal	103	13	0	1	0	0	0	0
DR	0	682	13	1	18	3	0	0
wetAMD	0	31	148	0	4	0	0	0
CSC	0	1	70	1	1	0	0	0
BRVO	0	4	2	0	57	0	0	0
CRVO	0	3	0	0	7	28	0	0
VKH	0	15	14	0	3	2	0	1
Undiag	0	0	0	0	0	0	0	0

Actual disease

Prediction

Step-Chinese

Normal	61	17	3	0	0	35	0	1
DR	1	684	6	4	7	1	0	14
wetAMD	0	49	110	19	2	0	0	3
CSC	0	11	47	14	0	0	0	1
BRVO	0	9	1	1	51	1	0	0
CRVO	0	9	0	0	1	26	0	2
VKH	0	25	6	3	0	0	0	1
Undiag	0	0	0	0	0	0	0	0

Actual disease

Prediction

Step-English

Normal	111	2	0	0	0	0	0	4
DR	0	671	6	2	12	16	0	10
wetAMD	1	86	63	2	5	2	0	24
CSC	1	29	32	5	1	0	0	5
BRVO	0	6	0	0	53	0	0	4
CRVO	0	5	0	0	5	28	0	0
VKH	1	18	4	1	2	2	0	7
Undiag	0	0	0	0	0	0	0	0

Actual disease

Prediction

بحث و نتیجه گیری



5

نتایج

- ❖ ChatGPT ابزار کمکی مفید در تشخیص است، اما فقط در کنار نظارت پزشک و با دقت بالا قابل استفاده است.
- ❖ برای کاربرد گسترده تر، نیاز به آموزش با داده های واقعی، چندزبانه و بالینی دارد.



چالش ها :

- ❖ استدلال ناقص در زبان چینی می شود.
- ❖ تفاوت عملکرد بین زبان ها نشان دهنده ی نیاز به آموزش چندزبانه است.
- ❖ هنوز به سطح پزشکان متخصص نمی رسد.

منو



نتایج

محدودیت ها

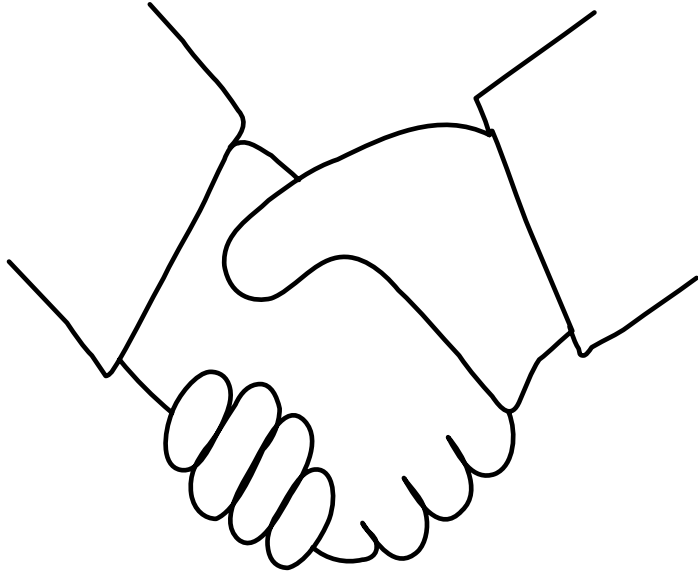


- دامنه زبانی و مکانی محدود
- عدم استفاده از تصاویر پزشکی
- (FFA)
- نبود ارزیابی در محیط واقعی بالینی

دیدگاه من



دیدگاه من



- استفاده محدود به نسخه GPT-3.5
- نبود تحلیل چندوجهی (Multimodal)
- ضعف عملکرد در زبان غیرانگلیسی



با تشکر از توجه شما