



## تحلیل احساسات مبتنی بر جنبه در زبان فارسی با استفاده از مدل تنظیم دقیق شده LLaMA3

نیلوفر رنجبر آ\* حامد باغبانی آ

آ گروه مهندسی کامپیوتر، دانشکده مهندسی جم، دانشگاه خلیج فارس، بوشهر، ایران.

Persian  
Abstract

### چکیده

### اطلاعات مقاله

تاریخچه مقاله:

دریافت: 12 July 2025

اصلاح: 28 June 2026

پذیرش: 22 July 2026

انتشار آنلاین: 28 July 2026

کلمات کلیدی:

تحلیل احساسات مبتنی بر جنبه، زبان فارسی،

تحلیل احساسات، مدل LLaMA3، مدل های

زبانی بزرگ، پردازش زبان طبیعی.

تحلیل احساسات مبتنی بر جنبه در زبان فارسی، در مقایسه با زبان انگلیسی، کمتر توسعه یافته است؛ به ویژه هنگامی که هدف، استخراج ساختارهای کامل جنبه، دسته، عبارت نظر و احساس باشد، نه صرفاً طبقه بندی احساس برای جنبه های از پیش تعیین شده. این مقاله یک چارچوب مولد برای تحلیل احساسات مبتنی بر جنبه فارسی بر پایه مدل های دستورمحور تنظیم دقیق شده خانواده LLaMA3 ارائه می کند. مدل با دریافت یک نقد فارسی، آرایه ای JSON از چهارتایی های شامل عبارت جنبه، دسته، عبارت نظر و برچسب احساس تولید می کند. داده های آموزشی از یک پیکره تحلیل احساسات مبتنی بر جنبه فارسی با پیش برچسب گذاری GPT-3.5 گرفته شده اند که برچسب های آن توسط ارزیابان انسانی بررسی و پالایش شده اند. از این رو، در این مقاله از آن با عنوان پیکره ای کمک گرفته از مدل زبانی و اعتبارسنجی شده توسط انسان یاد می کنیم، نه یک مجموعه کاملاً انسان نویس. برای کاهش خطر ارزیابی دوری، یک مجموعه آزمون مستقل با برچسب گذاری صرفاً انسانی نیز از نقدهای خام فارسی ساخته شد. سه ارزیاب، ۲۴۰ نقد را از ابتدا برچسب گذاری کردند و توافق دست کم دو نفر از سه نفر، ۵۵۵ برچسب طلایی در سطح جنبه از ۲۳۹ نقد ایجاد کرد. ما LLaMA3 در حالت صفرنمونه، مدل های LLaMA3 تنظیم دقیق شده، و Dorna-LLaMA3-8B-Instruct را به عنوان یک خط مبنای مولد فارسی محور، با طرح خروجی، تجزیه گر و معیارهای یکسان ارزیابی می کنیم. در مجموعه آزمون صرفاً انسانی، LLaMA3 صفرنمونه به امتیاز F1 استخراج جنبه برابر با 0.1426 و امتیاز F1 جنبه و احساس برابر با 0.1185 می رسد. میانگین پنج اجرای LLaMA3-CG این مقادیر را به ترتیب به  $0.5477 \pm 0.0229$  و  $0.4394 \pm 0.0212$  افزایش می دهد، در حالی که Dorna-CG به  $0.5809 \pm 0.0123$  و  $0.4897 \pm 0.0194$  می رسد. افزودن ParsiNLU عملکرد درون دامنه ای ParsiNLU را به طور چشمگیری بهبود می دهد، اما عملکرد روی مجموعه صرفاً انسانی را کاهش می دهد و نوعی موازنه میان دامنه و سبک برچسب گذاری را آشکار می کند.

© 2026. تمامی حقوق محفوظ است.

\* نویسنده مسئول.

آدرس های رایانامه: nranjbar@pgu.ac.ir (ن. رنجبر)،

baghbani.hamed@gmail.com (ح. باغبانی)

تمامی حقوق محفوظ است. ISSN: 2322-4460 © 2026

