

Израда синтетичког евалуативног скупа података за Српски SentiWordNet

УДК 811.163.41'322.2

САЖЕТАК: У раду се представља израда синтетичког скупа за евалуацију Српског SentiWordNet-а која користи велике језичке моделе, посебно модел *Мистрал*. Због недостатка ресурса за анализу сентимента на српском језику, циљ истраживања је премошћавање овог јаза генерисањем скупа за евалуацију и унапређење алата за анализу сентимента на српском. Вредности поларитета сентимента из енглеског SentiWordNet-а аутоматски су мапиране на Српски Ворднет. Како би се ове вредности прецизније прилагодиле српском језику, креиран је посебан скуп за евалуацију. Иницијално је одабрано 500 синсетова из Српског Ворднета, на основу њихове усклађености са senti-pol-sr лексиконом и мапираним вредностима из SentiWordNet-а. Ови синсетови су класификовани према поларитету сентимента коришћењем *Мистрал*-а. Избалансиран подскуп од 75 синсетова насумично је издвојен, додатно профињен финалном градацијом сентимента и ручно прегледан. Резултати показују високу прецизност, приближно 93%.

КЉУЧНЕ РЕЧИ: Анализа сентимента, велики језички модели, синтетички скуп података, Српски језик, SentiWordNet

РАД ПРИМЉЕН: 23. мај 2024.

РАД ПРИХВАЋЕН: 24. јун 2024.

Саша Петаљинкар

sasa5linkar@gmail.com

ORCID: 0009-0007-9664-3594

Универзитет у Београду

Мултидисциплинарни

докторске студије

Интелигентни системи

Београд, Србија

1. Увод

Анализа сентимента је процес рачунарског одређивања емоционалног тона иза текста како би се разумели ставови, мишљења и

емоције изражене у њему. Једна од метода за анализу сентимента заснива се на лексиконима сентимента. Лексикони сентимента су специјализовани речници који повезују речи и фразе са вредностима сентимента, омогућавајући аутоматизовану анализу поларитета емоција у тексту (Liu 2010).

Један од проблема који се идентификује код лексикона сентимента је тај што неки изрази могу имати више значења, при чему различита значења речи могу изазивати различите сентименте. Решавање овог проблема подразумева додељивање сентимента према контекстуалном значењу речи, а не према самој речи. Истакнути пример таквог лексикона је SentiWordNet (SWN), који проширује Princeton WordNet (PWN) речник додељивањем оцена сентимента сваком синсету (скуп когнитивних синонима) како би одражавао колективни емоционални тон појма. Овај процес укључује одабир малог броја синсетова са високим поларитетом, након чега следи проширење овог скупа кроз полуаутоматски метод, који идентификује односе унутар ворднета, а који или очувавају или преокрећу поларитет. Добијени проширени скуп се затим користи за обучавање класификатора заснованих на дефиницијама ових синсетова (Baccianella, Esuli, and Sebastiani 2010).

Ворднети за многе језике су међусобно повезани путем међународног језичког индекса (ILI) – синсет у ворднету одређеног језика добија вредност из ILI-а, која представља апстрактан концепт и додељује се синсетовима у другим језицима који лексикализују исти концепт. И EuroWordNet (Vossen 2004) и BalkaNet (Tufis, Cristea, and Stamou 2004) су на овај начин повезани са PWN-ом, а самим тим и са SWN-ом. EuroWordNet укључује ворднете за осам језика: енглески, холандски, шпански, италијански, немачки, француски, чешки и естонски, док BalkaNet проширује ову листу додавањем ворднета за пет додатних језика: бугарски, грчки, румунски, српски и турски (Krstev et al. 2004; Krstev, Koeva, and Vitas 2006; Stanković et al. 2018). Моруће је лако мапирати SWN на било који ворднет који има ILI, што је случај за све ворднете који припадају пројекту Отвореног вишејезичког ворднета (OMW) (Bond and Paik 2012) а који тренутно укључује 200 језика.

Истраживања су показала да се вредности сентимента из SWN-а могу применити на неенглеске језике коришћењем ILI-ја (Denecke 2008). Српски ворднет (СрпВН) садржи вредности сентимента добијене директним мапирањем синсетова коришћењем ILI-ја на SWN (Krstev et al. 2004; Mladenović, Mitrović, and Krstev 2014). СрпВН је коришћен

у креирању хибридног оквира за анализу сентимента на српском језику (Mladenović et al. 2015).

Наша хипотеза је да се овакав лексикон може побољшати заменом мапираних вредности сентимента онима које су репрезентативније за српски језик. Ово је већ спроведено за друге језике, као што су турски (Dehkharghani et al. 2016), арапски (Alhazmi and Black 2013), вијетнамски (Vu and Park 2014), одија (Mohanty, Kannan, and Mamidi 2017), индонежански (Wijayanti and Arisal 2021) и хинди (Bakliwal, Arora, and Varma 2012). Међутим, потребан је евалуациони скуп података – подскуп синсетоба из СрпВН-а већ анотираним са поларитетом сентимента – како би се проценила оваква побољшања за српски језик.

За SWN већ постоји такав евалуациони скуп података: Micro-WNO (Cerini et al. 2007), ручно означен подскуп синсетова из PWN-а, који је јавно доступан.¹ Креирање упоредивог евалуационог скупа података за српски језик ручним путем захтевало би значајан напор, укључујући ангажовање малог броја стручних анотатора или веће групе мање вештих анотатора. С обзиром на недостатак таквих ресурса, неопходно је размотрити алтернативни приступ.

Синтетички евалуациони скупови података представљају вештачки креиране колекције података, развијене за тестирање и валидацију рачунарских модела, посебно у областима где стварни подаци могу бити оскудни, пристрасни или превише осетљиви за употребу. Ови скупови података генеришу се путем алгоритама или симулација, са циљем да имитирају статистичке особине стварних података, омогућавајући истраживачима да спроводе робусне евалуације под контролисаним условима (Lu et al. 2024).

Појава ВЕЛИКИХ ЈЕЗИЧКИХ МОДЕЛА (ВЈМ) омогућила је креирање значајно бољих синтетичких скупова података. Показано је да ВЈМ-ови, за потребе задатака из области обраде природног језика (ОПЈ), укључујући анализу сентимента, могу успешно анотирати податке на основу само неколико познатих (виђених) примера (Amatriain 2024; Brown et al. 2020).

У сценарију УЧЕЊА БЕЗ ИЈЕДНОГ ПОКУШАЈА (енгл. Zero-shot learning), модел доноси предикције или анотације без експлицитно виђених примера задатка током тренинга. Ова способност је посебно корисна за анализу сентимента у језицима или контекстима где су анотирани

1. <https://github.com/aesuli/Sentiwordnet/blob/master/data/Micro-WNop-WN3.txt>

подаци ретки, јер омогућава ВЈМ-у да примени своје већ постојеће знање на нове, невиђене задатке (Amatriain 2024; Brown et al. 2020).

УЧЕЊЕ СА НЕКОЛИКО ПОКУШАЈА (енгл. Few-shot learning), с друге стране, омогућава моделу да из малог броја примера научи како да обавља одређени задатак. Овај приступ је посебно користан за побољшање разумевања модела и повећање његове тачности у специфичним применама, као што је разликовање нијансираних израза сентимента (Brown et al. 2020).

Коришћење УПИТА И ОДГОВОРА (енгл. prompts and response) са ВЈМ-овима омогућава да се ови приступи учења ефикасно примене. Креирањем упита који усмеравају модел ка жељеном излазу, истраживачи могу искористити способности ВЈМ-ова за анализу сентимента. Ово укључује дизајн упита који јасно дефинише задатак, као што је идентификација сентимента датог текста, и омогућавање моделу да генерише одговор на основу свог претходног тренинга и контекста који даје упит. Такав приступ помаже у коришћењу потенцијала ВЈМ-ова за детаљну анализу сентимента, нудећи флексибилан и ефикасан метод за анализу сентимента у различитим скуповима података (Amatriain 2024)

Ово поставља питање да ли би се ВЈМ-ови могли користити не само за креирање скупа за евалуацију података, већ и за анотацију целокупног СрпВН-а са вредностима поларитета сентимента. Одлука да се фокусира на креирање малог скупа за евалуацију проистиче из високих рачунарских трошкова повезаних са анотацијом целокупне мреже.

Примарни циљ овог приступа је побољшање постојећих вредности сентимента у СрпВН-у, које су добијене мапирањем из SWN-а. Фокус је на синсетовима који су у мапирању означени као објективни, иако садрже речи које у лексикону сентимента имају поларитет.

Да би се постигао балансиран узорак погодан за ефикасну евалуацију, посебно у каснијој примени модела машинског учења за класификацију сентимента, насумично је одабрано 500 синсетова. Ови синсетови су обрађени коришћењем модела *Mistral*. Први корак у обради био је категоризација синсетова у три групе: позитивни, негативни и објективни сентимент. Након тога, из сваке категорије је одабран једнак број синсетова за финесу у даљој градацији.

Резултати су прошли кроз процес ручне анотације, где је сваки синсет, заједно са вредностима које је вратио ВЈМ, процењен и

оцењен једноставном ознаком „успех“ или „неуспех“, на основу њихове усклађености са очекиваним вредностима сентимента.

Првобитно је методологија била осмишљена да укључи учење са неколико покушаја за финију градицију сентимента, користећи примере из синсетова који нису изабрани за примарни скуп података — конкретно, оних синсетова који су показивали очигледну усклађеност сентимента у српском и енглеском SWN-у. Међутим, прелиминарни резултати добијени приступом учења без иједног покушаја показали су се задовољавајућим, чиме је компонента учења са неколико покушаја постала непотребна. Због тога се истраживање наставило искључиво са парадигмом учења без иједног покушаја, при чему је *Mistral* модел примењен без претходних специфичних примера за задатак анализе сентимента.

Ручна анотација резултата потврдила је валидност овог поједностављеног приступа. Методологија учења без иједног покушаја показала је изванредну стопу успеха од преко деведесет процената, афирмативно показујући да чак и без укључивања учења са неколико покушаја и додатног контекста који оно пружа, ВЈМ може ефикасно да препозна и класификује сентимент у оквиру одабраних српских синсетова.

Главни ВЈМ коришћен у овом истраживању је *Mistral 7B – Instruct*.² То је фино подешена варијанта модела *Mistral 7B*, језичког модела са 7 милијарди параметара, дизајнираног за врхунске перформансе и ефикасност. *Mistral 7B* надмашује модел *Llama 2 13B* на разним бенчмарковима. Нарочито, превазилази *Llama 1 34B* у задацима расуђивања, математике и генерисања кода (Jiang et al. 2023). *Mistral 7B – Instruct* такође надмашује модел *Llama 2 13B – Chat* како на људским, тако и на аутоматизованим бенчмарковима (Jiang et al. 2023). Објављен под Apache 2.0 лиценцом, овај модел нуди флексибилно средство за истраживаче, са могућношћу да се користи на локалном рачунару без додатних трошкова (Jiang et al. 2023).

Овај рад је организован у неколико кључних поглавља како би се свеобухватно дискутовало о истраживању и резултатима. У поглављу 2 детаљно је описана методологија, укључујући одабир дефиниција синсетова за обраду, специфичне уште коришћене у истраживању, као и примену ових упита. Поред тога, укратко су поменути покушаји коришћења других језичких модела ради поређења. Поглавље 3 приказује резултате, описујући тачност модела са примерима синсетова

2. <https://huggingface.co/mistralai/Mistral-7B-Instruct-v0.2>

оцењених као исправни и неисправни, као и одговоре који нису припадали задатом скупу дефинисаном инструкцијама упита. На крају, поглавље 4 закључује рад са сажетком налаза и дискусијом о потенцијалним будућим радовима који би могли проширити ово истраживање.

2. Методологија

Senti-pol-sr је лексикон поларитета за српски језик, анотиран на нивоу речи, а не на нивоу значења речи (Stanković et al. 2022). У нашем истраживању, овај лексикон је коришћен за одабир узорка погодног за анотацију помоћу ВЈМ-а. Ово је постигнуто идентификацијом свих синсетова из СрпВН-а који садрже литерале присутне у лексикону *senti-pol-sr* и истовремено имају неутралну вредност сентимента (0,0), како је мапирано из *SWN*-а.

Детектована су четири главна разлога за неслагања између вредности сентимента у лексикону *senti-pol-sr* и оних изведених из *SWN*-а. Прво, док реч може преносити поларизујући сентимент, стварно значење у којем се користи можда не носи тај сентимент. На пример, синсет ENG30-06828389-n (деф. „карактер налик на звезду (*) који се користи у штампаним текстовима“) у Српском Ворднету садржи литерал „звезда“, која у другим синсетовима може имати позитивне конотације. Друго, вредности сентимента у *SWN*-у, генерисане методама машинског учења, могу бити нетачне. На пример, синсет ENG30-02585489-v (деф. „изазвати патњу“) који је очигледно веома негативан. Треће, ако синсет нема одговарајући синсет у *PWN*-у, као што је BILI-00000941 (деф. „Гласно изражавати жалост, запевати, тужити.“), уобичајено му се додељује поларитет (0,0).

Најинтересантнија могућност је да се концепт сматра објективним на енглеском, док на српском језику носи сентимент. Такви синсетови још нису пронађени.

Почетна анализа идентификовала је 2.956 синсетова од укупно 25.320 унутар СрпВН-а (Mladenović, Stanković, and Krstev 2017) који су садржали литерале анотирани са јасним поларитетом у лексикону *senti-pol-sr* и имали вредност поларитета (0,0), од чега 1.511 показује позитиван сентимент, а 1.445 негативан. С обзиром на велики обим, обрада свих ових синсетова помоћу ВЈМ-а није била изводљива, па је насумично одабран узорак од 500 синсетова за даљу анализу.

Ради креирања уравнотеженог скупа узорака, дефиниције из овог насумичног узорка су обрађене коришћењем *LangChain Python* библиотеке, моћног алата за креирање, експериментисање и анализирање језичких модела и агената (Chase 2022). *LangChain Python* библиотека функционише као интерфејс за модуле више великих језичких модела. Овај пројекат је користио омотаче пројектоване за Hugging Face Hub и Transformer библиотеке. Процедура коришћена у овом истраживању садржала је само један шаблон упита са једном улазном променљивом – дефиницијом синсета, и *Mistral 7B – Instruct* модел преузет са *Hugging Face Hub*-а. Модел је извршен на локалном рачунару. Библиотека *LangChain* омогућава да се упитни шаблони са променљивама, које су означене витичастим заградама, попуне вредностима тих променљивих и проследе као упит ВЈМ модулу, који затим враћа одговор.

Користећи одговарајући упит као што је приказано на слици 1, узорак је подељен на оне означене као позитивне, негативне, објективне и оне који нису исправно означени. Било је 290 објективних, 102 негативна, 33 позитивна и 75 погрешно означених синсетова.

Да би се искористиле пуне могућности ВЈМ-а за анализу сентимента, упит је пажљиво пројектован да обухвати три кључна елемента: доделу улога (енгл. *role-play*), јасне инструкције и очекиване исходе.

1. Додела улога: инструкција ВЈМ-у као експерту: Упит започиње сценаријом играња улога, где се ВЈМ-у даје инструкција да преузме улогу експерта за анализу сентимента (слика 1). Овај приступ је усвојен како би се модел усмерио ка аналитичком и фокусираном испитивању текста, ослањајући се на своје опсежно унапред обучено знање и разумевање нијанси анализе сентимента.
2. Јасне инструкције: суштина ефикасне комуникације са ВЈМ-ом лежи у јасноћи инструкција. Упит експлицитно описује задатак, водећи ВЈМ да идентификује и анализира сентимент изражен у датом делу текста (дефиницији синсета). Ова јасноћа осигурава да су аналитичке способности модела усмерене ка тачном процењивању сентимента, минимизујући двосмисленост у његовим одговорима.
3. Очекивани исходи: да би се додатно усавршио излаз модела, упит прецизира очекиване исходе анализе. Наводи жељени формат одговора, било да је то класификација сентимента (позитиван, негативан, неутралан) или нијансиранија оцена сентимента.

Упит коришћен за класификацију је усавршен путем експерименталног тестирања на мањем подскупу дефиниција синсета. Упоредне анализе инструкција упита на енглеском и српском језику откриле су да су упити на српском језику давали тачније резултате. Како би се обезбедило свеобухватно покривање потенцијалних одговора, број токена у излазу је постављен на пет.

Даље испитивање показало је постојање дупликата унутар скупа, због тога што неки синсетови садрже више литерала који су такође речи из senti-pol-sr лексикона, као што је синсет ENG30-01375831-a „славан, познат, чувен, значајан, реномиран, признат, прослављен: опште прихваћен и уважен“. Важно је напоменути да се сви осим два литерала („реномиран“, „признат“) из овог синсета појављују у лексикону, и означени су као позитивни. Након уклањања дупликата, остало је 279 објективних, 97 негативних и 27 позитивних синсетова. У овом кораку није било потребе за бројањем неправилно означених синсетова, јер уклањање дупликата из њих нема практичан значај.

За детаљнију анализу сентимента, насумично је одабран подскуп од по 25 синсетова из сваке категорије сентимента, формирајући такозвани „балансирани узорак“.

Kao ekspert za analizu sentimenta, analizirajte sledeći tekst na srpskom jeziku i odredite njegov sentiment.

Sentiment treba da bude striktno klasifikovan kao "pozitivan", "negativan", ili "objektivan".

Nijedan drugi odgovor neće biti prihvaćen.

Tekst: {text}

Sentiment:



Слика 1. Шаблон упита за одређивање поларитета.

Балансирани узорак је даље обрађен кроз два шаблона упита за нијансирану обраду сентимента, један за позитиван и један за негативан сентимент (слике 2 и 3). Ови шаблони упита су пројектовани експериментално, кроз итеративно тестирање разних опција и постепено прилагођавање текста, користећи мале скупове дефиниција синсетова из

СрпВН-а. На пример, понављање реченице „Ниједан други одговор неће бити прихваћен“ два пута значајно је смањило број одговора ван оквира.

Јединствени упит за одређивање обе вредности није одабран како би се очувала доследност са структуром SWN-а, где синсет може имати позитивне (POS) и негативне (NEG) вредности изнад нуле, све док њихов збир није већи од један. Да би се прилагодила очекивана дужина излазних стрингова, број излазних токена је повећан на девет.

Током ове фазе, тестирани су и други BJM модели на малом узорку – *GPT2-ORAO* (Škorić 2024), *Llama-2-7b* (Touvron et al. 2023), *alpaca-serbian-7b-base* и *WizardLM-1.0-Uncensored-Llama2-13b*. Идеја је била да се упореде различити резултати, симулирајући анотацију од стране више анотатора. Међутим, резултати су довели до њиховог искључења из даљег рада због превеликог броја неправилно означених синсетова и недостатка финије градације.

Kao ekspert za analizu sentimenta, analizirajte sledeći tekst na srpskom jeziku i odredite da li ima pozitivan sentiment.

Sentiment treba da bude striktno klasifikovan kao "nije pozitivan", "slabo pozitivan", "umereno pozitivan", "veoma pozitivan", ili "ekstremno pozitivan". Nijedan drugi odgovor neće biti prihvaćen.

Nijedan drugi odgovor neće biti prihvaćen.

 *Текст: {text}*

Pozitivan sentiment:

Слика 2. Шаблон упита за фину ознаку позитивног сентимента.

За сваку дефиницију у балансираном узорку, додељене су две вредности на тај начин. Интенција шаблона упита и дизајн ограничавају одговоре на следећи сет одговора:

— За позитиван сентимент:

- „није позитиван“
- „слабо позитиван“
- „умерено позитиван“

Kao ekspert za analizu sentimenta, analizirajte sledeći tekst na srpskom jeziku i odredite da li ima negativan sentiment.

Sentiment treba da bude striktno klasifikovan kao "nije negativan", "slabo negativan", "umereno negativan", "veoma negativan", ili "ekstremno negativan". Nijedan drugi odgovor neće biti prihvaćen.

Nijedan drugi odgovor neće biti prihvaćen.

Tekst: {text}

Negativan sentiment:

Слика 3. Шаблон упита за фину ознаку негативног сентимента.

- „веома позитиван“
- „екстремно позитиван“
- За негативан сентимент:
 - „није негативан“
 - „слабо негативан“
 - „умерено негативан“
 - „веома негативан“
 - „екстремно негативан“

Резултирајући скуп података је сачуван у форми датотеке са вредностима одвојеним зарезима (CSV).³ Колоне у тој датотеци су следеће:

ILI: Међујезички Индекс, који повезује синсет са његовим еквивалентима у WordNet-има других језика.

Definition: Глоса синсета, која пружа његово значење или објашњење.

Lemma_names: Литерали садржаних у синсету.

Sentiment_SWN: Вредност сентимента мапирана из SWN-а, која показује оригинални скор сентимента у енглеској верзији.

3. https://github.com/sasa5linkar/SWN-synth-eval-set/blob/main/balanced_sample2.csv

Sentiment_lexicon : Вредност сентимента изведена из лексикона сентимента senti-pol-sr креираног за српски језик, одражавајући локалне нијансе сентимента.

Sentiment_sa : Иницијална класификација од стране великог језичког модела, категоризујући сентимент као позитиван, негативан или неутралан.

Sentiment_sa_positive : Фино класификовање сентимента за позитиван сентимент, означавајући степен позитивности од „није позитиван“ до „екстремно позитиван“.

Sentiment_sa_negative : Фино класификовање сентимента за негативан сентимент, означавајући степен негативности од „није негативан“ до „екстремно негативан“.

С обзиром на мали узорак од 75, аутор је могао спровести ручну евалуацију целог скупа података. Овај процес је унапређен коришћењем *Data Wrangler*⁴ екстензије у *Visual Studio Code*-у, која је алатка за визуализацију и чишћење података и добро се интегрише са *VS Code*-ом и *VS Code Jupyter Notebooks*-ом. Алатка пружа интерактивни интерфејс за преглед и анализу података, нуди информативне статистике и визуелне приказе, и може убрзати чишћење података аутоматским генерисањем Pandas скрипти за модификацију података. Користила се за пажљиво испитивање дефиниција синсетова и њихово поређење са детаљним повратним информацијама о сентименту.

3. Резултати

Излаз	Број у sentiment_sa_negative
екстремно негативан	1
није негативан	60
умерено негативан	1
умјерено негативан	2
веома негативан	11
Укупно	75

Табела 1. Негативан сентимент.

4. [microsoft/vscode-data-wrangler \(github.com\)](https://github.com/microsoft/vscode-data-wrangler)

Израз	Број у sentiment _sa _positive
Није Ова из	1
Није позитиван	39
Ниједан О	6
Ниједан Текст	1
Учимо се држати	1
Умерено позитиван	3
Умерено позитиван The	2
Умјерено позитиван	6
веома позитиван	12
Веома позитиван О	1
Веома позитиван R	2
Веома позитиван Ov	1
Укупно	75

Табела 2. Позитиван сентимент

Резиме излаза генерисаних од стране ВЈМ-а коришћењем шаблона упита за фину градиацију сентимента приказан је у табелама 1 и 3.. Подаци у табелама показују да излаз ВЈМ-а није био ограничен како је назначено у шаблону упита, али је остао унутар граница које се лако могу исправити. Израз генерисан од негативног шаблона упита у великој мери се поклапао са предложеним излазом, уз само мање дијалекатске варијације. Насупрот томе, излаз из позитивног шаблона упита укључивао је један бесмислен одговор, „учимо се држати“, као и неке одговоре који се могу лабаво тумачити као не-позитивни. Током ручне провере, сви ови одговори су класификовани као не-позитивни.

Додатно, мање типографске грешке, као што су додатна слова, погрешни падежи или грешке у великим словима, занемарене су током ручне провере. Само једна особа је спровела евалуацију. Надаље, у излазима нису пронађени примери благо позитивних или благо негативних синсетова.

За садржај су два синсета очигледно неправилно означена, а три су била сумњива.

Очигледно неправилно су означени синсетови:

- BILI-00000941, ‘нарицати: Гласно изражавати жалост, запевати, тужити.’ означен је као потпуно објективан (нити позитиван нити

негативан), док очигледно носи негативан сентимент, штавише, изразито негативан.

- ENG30-04525038-n, 'баршун, сомот, плиш, кадифа: Свиленкаста, густа, чупава тканина, равна са полеђина.', који је означен као благо позитиван. Као врста текстила, требало би да буде објективан.

Синсетови који се сматрају сумњивим су:

- ENG30-01215137-v, 'ухапсити, затворити, скембати: Ставити у притвор, као осумњичене криминалце; о полицији.' означен као благо негативан, док га је евалуатор оценио као изразито негативан.
- ENG30-00309647-n, 'експедиција: Путовање организовано са специјалним циљем.' означен као веома позитиван, док га је евалуатор сматрао неутралним.
- ENG30-03135152-n, 'крст: крст као знамење хришћанства' означен као веома позитиван, док га је евалуатор сматрао неутралним.

Од већине синсетова који су тачно означени током нашег процеса анализе сентимента, овде представљамо избор илустративних примера. Ови примери би такође требало да демонстрирају критеријуме коришћене у евалуацији приликом разматрања исправне ознаке у класификацијама на позитивне, негативне и неутралне сентименте. Неки примери су:

- ENG30-01220336-n: Синсет повезан са акцијама попут „клевета“, „клеветање“ и „омаловажавање“, описан као „оштар напад на чију личност или добро име“. Сентимент овог синсета оцењен је као „Није позитиван“, што одражава одсуство позитивног сентимента, и прецизније, „Екстремно негативан“, тачно хватајући негативне конотације описаних акција.
- ENG30-06828389-n: Овај синсет се односи на „карактер налик на звезду (*) који се користи у штампаним текстовима“, са литеаралима „астериск“, „звездица“ и „звезда“. Сентимент за овај синсет је прецизно класификован као „Није позитиван“ и „Није негативан“, указујући на његову објективну природу без инхерентног позитивног или негативног сентимента.
- ENG30-10407310-n: Односи се на „онај који воли и брине о својој земљи“, са литеаралима „домољуб“, „патриота“ и „родољуб“. Сентимент овог синсета је фино оцењен као „Веома позитиван“ и „Није негативан“, ефикасно хватајући позитивне конотације повезане са патриотизмом.

4. Закључак

Анализа излаза генерисаног од стране ВЈМ-а, конкретно *Mistral* модела, показује да је 70 од 75 одговора испунило критеријуме прихватљивости дефинисане за ово истраживање. Овај резултат, који представља значајну већину скупа података, наглашава поузданост и ефикасност *Mistral* модела у извођењу класификације сентимента за српски језик. Са приближном стопом успеха од 93,3%, може се са сигурношћу закључити да је скуп података и употребљив и довољног квалитета за намеравану сврху евалуације предложених корекција поларитета сентимента унутар Српског WordNet-a.

Неуспех других тестираних модела може бити резултат упита оптимизованих за *Mistral* модел. Креирање упита за сваки модел појединачно могло би омогућити употребу више модела за постизање бољег квалитета.

Значајно подручје за будућа истраживања укључује проширење шаблона упита са специфичним примерима, прелазећи из тренутног приступа учења без иједног покушаја на парадигму учења са неколико покушаја. Ова модификација се очекује да побољша капацитет модела за нијансирану градацију сентимента, нудећи прецизнију и контекстуално свеснију анализу. Посебно се препоручује да се ово унапређење селективно примени на шаблоне за фину градацију сентимента, где је потенцијал за повећану тачност и осетљивост на језичке суптилности најизраженији (Brown et al. 2020).

Даље, истраживање напредних методологија за креирање упита, као што је приступ „ланац мисли“ (енгл. chain-of-thought), заслужује пажњу. Ова техника, олакшавајући структуриран и логичан ток мисли у оквиру обраде модела, може значајно побољшати способност модела да извади релевантне податке из одговора. Потенцијал таквих напредних стратегија за превазилажење инхерентних изазова у тачном идентификовању и класификацији израза сентимента у сложеним језичким контекстима представља обећавајући правац за истраживање (Amatriain 2024).

Ово истраживање представља иницијално испитивање могућности коришћења модела *Mistral* за генерисање додатних синтетичких скупова података за српски језик, посебно у областима где су ресурси оскудни или их је тешко прибавити из природних корпуса. Успешна примена модела *Mistral* за анализу сентимента наглашава његов потенцијал као вредног алата у превазилажењу ових изазова.

Литература

- Alhazmi, Samah, and William Black. 2013. “Arabic SentiWordNet in Relation to SentiWordNet 3.0.” *John McNaught International Journal of Computational Linguistics (IJCL)*, no. 4, 1.
- Amatriain, Xavier. 2024. *Prompt Design and Engineering: Introduction and Advanced Methods*. ArXiv:2401.14423 [cs], January. Accessed February 7, 2024. <https://doi.org/10.48550/arXiv.2401.14423>. <http://arxiv.org/abs/2401.14423>.
- Baccianella, Stefano, Andrea Esuli, and Fabrizio Sebastiani. 2010. “SENTIWORDNET 3.0: An Enhanced Lexical Resource.” In *7th LREC*, 2200–2204. <http://nmis.isti.cnr.it/sebastiani/Publications/LREC10.pdf>.
- Bakliwal, Akshat, Piyush Arora, and Vasudeva Varma. 2012. “Hindi Subjective Lexicon: A Lexical Resource for Hindi Adjective Polarity Classification.” In *International Conference on Language Resources and Evaluation*. <https://api.semanticscholar.org/CorpusID:1657748>.
- Bond, Francis, and Kyonghee Paik. 2012. “A Survey of Wordnets and Their Licenses.” In *Proceedings of the 6th Global WordNet Conference (GWC 2012)*, 64–71. Matsue, Japan.
- Brown, Tom B., Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, et al. 2020. *Language Models are Few-Shot Learners*. ArXiv:2005.14165 [cs], July. Accessed February 8, 2024. <https://doi.org/10.48550/arXiv.2005.14165>. <http://arxiv.org/abs/2005.14165>.
- Cerini, S., V. Compagnoni, A. Demontis, Maicol Formentelli, and C. Gandini. 2007. “Micro-WNOp: A gold standard for the evaluation of automatically compiled lexical resources for opinion mining.” *Language Resources and Linguistic Theory* (January): 200–210.
- Chase, Harrison. 2022. *LangChain*, October. <https://github.com/langchain-ai/langchain>.
- Dehkharghani, Rahim, Y. Saygin, B. Yanikoglu, and Kemal Oflazer. 2016. “SentiTurkNet: a Turkish polarity lexicon for sentiment analysis.” *Language Resources and Evaluation* 50:667–685. <https://doi.org/10.1007/s10579-015-9307-6>.

- Denecke, Kerstin. 2008. “Using SentiWordNet for multilingual sentiment analysis.” In *International Conference on Data Engineering*, 507–512. <https://doi.org/https://doi.org/10.1109/ICDEW.2008.4498370>.
- Jiang, Albert Q., Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, et al. 2023. *Mistral 7B*. ArXiv:2310.06825 [cs], October. Accessed August 1, 2024. <https://doi.org/10.48550/arXiv.2310.06825>. <http://arxiv.org/abs/2310.06825>.
- Krstev, Cvetana, Svetla Koeva, and Duško Vitas. 2006. “Towards the Global Wordnet.” In *Conference Abstracts of the First International Conference of Digital Humanities Organisations (ADHO) Digital Humanities 2006*, 114–117. Paris-Sorbonne, July. <http://poincare.matf.bg.ac.rs/~cvetana/biblio/DigHum06-KKV.pdf>.
- Krstev, Cvetana, Gordana Pavlović-Lažetić, Duško Vitas, and Ivan Obradović. 2004. “Using Textual and Lexical Resources in Developing Serbian Wordnet.” Publisher: Romanian Academy, Publishing House of the Romanian Academy, *Romanian Journal of Information Science and Technology* 7 (1-2): 147–161. <http://poincare.matf.bg.ac.rs/~cvetana/biblio/RJIST-MATF.pdf>.
- Liu, Bing. 2010. “Sentiment Analysis and Subjectivity.” In *Handbook of Natural Language Processing, Second Edition*, edited by Fred J. Damerau and Nitin Indurkha, 629–661. Chapman & Hall/CRC.
- Lu, Yingzhou, Minjie Shen, Huazheng Wang, Xiao Wang, Capucine van Rechem, Tianfan Fu, and Wenqi Wei. 2024. *Machine Learning for Synthetic Data Generation: A Review*. ArXiv:2302.04062 [cs] version: 6, June. Accessed August 1, 2024. <https://doi.org/10.48550/arXiv.2302.04062>. <http://arxiv.org/abs/2302.04062>.
- Mladenović, Miljana, Jelena Mitrović, and Cvetana Krstev. 2014. “Developing and Maintaining a WordNet: Procedures and Tools.” In *Proceedings of the Seventh Global WordNet Conference 2014*, edited by Heili Orav, Christiane Fellbaum, and Piek Vossen, 55–62. Tartu, Estonia: University of Tartu, January. ISBN: 978-9949-32-492-7.
- Mladenović, Miljana, Jelena Mitrović, Cvetana Krstev, and Duško Vitas. 2015. “Hybrid sentiment analysis framework for a morphologically rich language.” *Journal of Intelligent Information Systems* 46 (3): 599–620. Accessed January 1, 2023. <https://doi.org/10.1007/s10844-015-0372-5>.

- Mladenović, Miljana, Ranka Stanković, and Cvetana Krstev. 2017. "A WordNet Ontology in Improving Searches of Digital Dialect Dictionary." In *New Trends in Databases and Information Systems: ADBIS 2017 Short Papers and Workshops - SW4CH (Semantic Web for Cultural Heritage)*, 767:373–383. Lecture Notes in Computer Science. Nicosia, Cyprus: Springer International Publishing, September.
- Mohanty, Gaurav, Abishek Kannan, and Radhika Mamidi. 2017. "Building a SentiWordNet for Odia." In *Proceedings of the 8th Workshop on Computational Approaches to Subjectivity, Sentiment and Social Media Analysis*, edited by Alexandra Balahur, Saif M. Mohammad, and Erik van der Goot, 143–148. Copenhagen, Denmark: Association for Computational Linguistics, September. <https://doi.org/10.18653/v1/W17-5219>. <https://aclanthology.org/W17-5219/>.
- Škorić, Mihailo. 2024. "New Language Models for Serbian." Publisher: Zajednica biblioteka univerziteta u Srbiji, Beograd, *Infotheca - Journal for Digital Humanities* 24 (1). <https://doi.org/https://doi.org/10.48550/arXiv.2402.14379>. <https://arxiv.org/abs/2402.14379>.
- Stanković, Ranka, Miloš Košprdić, Milica Ikonić Nešić, and Tijana Radović. 2022. "Sentiment Analysis of Serbian Old Novels." In *2nd Workshop on Sentiment Analysis and Linguistic Linked Data (SALLD)*, 31–38. Marseille: ELRA, June. Accessed August 21, 2023. <https://aclanthology.org/2022.salld-1.6>.
- Stanković, Ranka, Miljana Mladenović, Ivan Obradović, Marko Vitas, and Cvetana Krstev. 2018. "Resource-based WordNet Augmentation and Enrichment." In *Proceedings of the Third International Conference on Computational Linguistics in Bulgaria (CLIB 2018)*, 104–114. Sofia, Bulgaria: Department of Computational Linguistics, Institute for Bulgarian Language, Bulgarian Academy of Sciences, May. <https://aclanthology.org/2018.clib-1.14/>.
- Touvron, Hugo, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, et al. 2023. *Llama 2: Open Foundation and Fine-Tuned Chat Models*. ArXiv:2307.09288 [cs], July. Accessed August 1, 2024. <https://doi.org/10.48550/arXiv.2307.09288>. <http://arxiv.org/abs/2307.09288>.

- Tufis, D., D. Cristea, and S. Stamou. 2004. “BalkaNet: Aims, Methods, Results and Perspectives. A General Overview.” Publisher: Institute for Artificial Intelligence, Romanian Academy, Bucharest, Romania, *Romanian Journal of Information Science and Technology* 7 (1-2): 9–43.
- Vossen, Piek. 2004. “EuroWordNet: A Multilingual Database of Autonomous and Language-Specific WordNets Connected via an Inter-Lingual Index.” Eprint: <https://academic.oup.com/ijl/article-pdf/17/2/161/2912694/ech041.pdf>, *International Journal of Lexicography* 17, no. 2 (June): 161–173. ISSN: 0950-3846. <https://doi.org/10.1093/ijl/17.2.161>. <https://doi.org/10.1093/ijl/17.2.161>.
- Vu, Xuan-Son, and Seong-Bae Park. 2014. “Construction of Vietnamese SentiWordNet by using Vietnamese Dictionary.” *ArXiv* abs/1412.8010.
- Wijayanti, Rini, and Andria Arisal. 2021. “Automatic Indonesian Sentiment Lexicon Curation with Sentiment Valence Tuning for Social Media Sentiment Analysis.” *ACM Transactions on Asian and Low-Resource Language Information Processing (TALLIP)* 20:1–16. <https://doi.org/10.1145/3425632>.