

SVEUČILIŠTE U SPLITU
MEDICINSKI FAKULTET

MARIANO KALITERNA

KORIŠTENJE UMJETNE INTELIGENCIJE PRI PISANJU ESEJA IZ
MEDICINSKE ETIKE TIJEKOM NASTAVE NA MEDICINSKOM FAKULTETU

Doktorski rad

Split, 2025.

SVEUČILIŠTE U SPLITU
MEDICINSKI FAKULTET

MARIANO KALITERNA

**KORIŠTENJE UMJETNE INTELIGENCIJE PRI PISANJU ESEJA IZ
MEDICINSKE ETIKE TIJEKOM NASTAVE NA MEDICINSKOM FAKULTETU**

Doktorski rad

Mentor: prof. dr. sc. Darko Duplančić, dr. med.

Split, 2025.

Doktorski rad izrađen je na Medicinskom fakultetu Sveučilišta u Splitu, Katedri za humanistiku pod voditeljstvom prof. dr. sc. Darka Duplančić. Istraživanje je dio projekta “Improving Research Ethics Expertise and Competencies to Ensure Reliability and Trust in Science“ iRECS.

Ova disertacija uključuje istraživanje objavljeno u časopisu Scientific Report:

Kaliterna M, Žuljević MF, Ursić L, Krka J, Duplančić D. Testing the capacity of Bard and ChatGPT for writing essays on ethical dilemmas: A cross-sectional study. Sci Rep. 2024 Oct 30;14(1):26046. doi: 10.1038/s41598-024-77576-3. (čimbenik odjeka 4,6)

ZAHVALA

Iskreno se zahvaljujem Prof dr.sc. Darku Duplančić na podršci, stručnim i korisnim savjetima prilikom izrade disertacije.

Zahvaljujem se mojim suradnicima Mariji, Luki, Jakovu na pomoći, strpljenju, savjetima te posebno na ugodnoj atmosferi prilikom rada na ovim istraživanjima.

Zahvaljujem se mojoj obitelji na razumjevanju i podršci u mom povratku u stručni i istraživački rad.

Posebna zahvala dragim prijateljima i kolegama koji su me podržavali u izradi ove disertacije, koji su imali vjeru u mene i u onim trenucima kada sam je ja sam počeo gubiti.

Sadržaj

POPIS OZNAKA I KRATICA.....	1
1. UVOD.....	2
1.1 Umjetna inteligencija	3
1.1.1 Uvod i osnove o umjetnoj inteligenciji.....	3
1.1.2 Ključne tehnologije i pojmovi	4
1.1.3 Povijest i razvoj umjetne inteligencije.....	5
1.1.4. Primjena umjetne inteligencije u društvu	6
1.1.5 Izazovi vezani uz umjetnu inteligenciju	8
1.2 Veliki jezični modeli	10
1.2.1 Uvod	10
1.2.2 Tehnologija LLM-ova	10
1.2.3 Primjene velikih jezičnih modela	12
1.2.4 Izazovi i ograničenja.....	12
1.2.5 Budućnost velikih jezičnih modela.....	13
1.3 Korištenje UI u pisanju eseja	14
2. CILJEVI RADA I HIPOTEZE	17
2.1 Lingvistička analiza i usporedba eseja studenata medicine i eseja generiranih ChatGPT umjetnom inteligencijom.....	18
2.2 Kako korištenje ChatGPT mijenja stavove o umjetnoj inteligenciji (UI).....	18
3. ISPITANICI I POSTUPCI.....	20
3.1 Lingvistička analiza i usporedba eseja studenata medicine i eseja generiranih umjetnom inteligencijom.....	21
3.1.1 Ustroj istraživanja.....	21
3.1.2 Ispitanici	21
3.1.3 Postupci	21
3.1.4 Statistička raščlamba	23

3.1.5 Etički aspekti	23
3.2 Povezanost korištenja ChatGPT-a i pristupa esejskom zadatku na nastavi iz medicinske humanistike	23
3.2.1 Ustroj istraživanja.....	23
3.2.2 Ispitanici	24
3.2.3 Postupci	24
3.2.4 Statistička raščlamba	28
3.2.5 Etički aspekti	28
4. REZULTATI.....	29
4.1 Lingvistička analiza i usporedba eseja studenata medicine i eseja generiranih umjetnom inteligencijom.....	30
4.1.1 Usporedba eseja koje su napisali studenti s esejima generiranim umjetnom inteligencijom	30
4.1.2 Eseji s djelomičnim u odnosu na potpune ključne riječi	34
4.1.3 Podanaliza – studentski eseji koji su potencijalno generirani uz pomoć umjetne inteligencije	35
4.2 Povezanost korištenja ChatGPT-a i pristupa esejskom zadatku na nastavi iz medicinske humanistike	45
5. RASPRAVA	49
5.1 Lingvistička analiza i usporedba eseja studenata medicine i eseja generiranih umjetnom inteligencijom.....	50
5.1.1 Sažetak glavnih nalaza istraživanja	50
5.1.2 Snaga i ograničenja istraživanja	52
5.1.3 Implikacije rezultata istraživanja i prijedlozi za daljnja istraživanja	53
5.2 Povezanost korištenja ChatGPT-a i pristupa esejskom zadatku na nastavi iz medicinske humanistike	54
5.2.1. Sažetak glavnih nalaza istraživanja.....	54
5.2.2 Snaga i ograničenja istraživanja	56

5.2.3 Implikacije rezultata istraživanja i prijedlozi za daljnja istraživanja	57
6. ZAKLJUČCI.....	58
7. SAŽETAK	60
8. LAIČKI SAŽETAK.....	62
9. SAŽETAK NA ENGLISKOM JEZIKU	64
10. LAIČKI SAŽETAK NA ENGLISKOM JEZIKU	67
11. LITERATURA	69
12. ŽIVOTOPIS.....	83
13. DODACI.....	88
Dodatak 1: Upute studentima o načinu pisanja eseja (prvo istraživanje)	89
Dodatak 2: Anketa povezanost korištenja ChatGPT-a i pristupa esejskom zadatku na nastavi iz medicinske humanistike (drugo istraživanje).....	89

POPIS OZNAKA I KRATICA

ATTARI-12 - Skala za mjerenje stavova prema umjetnoj inteligenciji i robotici (engl. *Attitudes Toward Artificial Intelligence and Robotics Indeks*)

BERT- Dvosmjerne enkoderske reprezentacije temeljem transformera (engl. *Bidirectional Encoder Representations from Transformers*)

ChatGPT - Razgovorni generativni unaprijed uvježbani transformator (engl. *Chat Generative Pre-trained Transformer*)

COMPAS - Profiliranje zatvorenika za upravljanje korektivnim sankcijama (engl. *Correctional Offender Management Profiling for Alternative Sanctions*)

GPT - generativni unaprijed uvježbani transformator (engl. *Generative Pre-trained Transformer*)

GPU - grafički procesor (engl. *Graphics Processing Unit*)

LIWC - Lingvistička analiza i brojanje riječi (engl. *Linguistic Inquiry and Word Count*)

LLM - veliki jezični model (engl. *Large Language Models*)

NLP - obrada prirodnog jezika (engl. *Natural Language Processing*)

OECD - Organizacija za ekonomsku suradnju i razvoj (engl. *Organisation for Economic Co-operation and Development*)

PaLM – Pathways jezični model (engl. *Pathways Language Model*)

UI - umjetna inteligencija

UNESCO - Organizacija Ujedinjenih naroda za obrazovanje, znanost i kulturu (engl. *United Nations Educational, Scientific and Cultural Organization*)

USSM - Medicinski fakultet Sveučilišta u Splitu (engl. *University of Split, School of Medicine*)

1. UVOD

1.1 Umjetna inteligencija

1.1.1 Uvod i osnove o umjetnoj inteligenciji

Jedno od zanimljivijih i najbrže rastućih područja znanosti i tehnologije današnjice je Umjetna inteligencija (UI) (1). Pojam UI obuhvaća razne tehnologije, metode i aplikacije koje omogućavaju strojevima da oponašaju ljudske kognitivne sposobnosti, a to pripadaju sve vrste razmišljanja, učenja, rješavanja problema i odlučivanja. Današnji UI se ne upotrebljava samo u naprednim tehnološkim sektorima poput informacijske tehnologije (IT), već sve više koristi i po specijalnim područjima poput zdravstvenih usluga, obrazovanja, prometa, financija, medija i mnogim drugim dijelovima društva. Cilj UI-ja nije samo oponašati ljudsku inteligenciju, već je i premašiti u mnogim zadacima kako bi se poboljšali procesi, povećala produktivnost i omogućila rješenja složenih problema koji su do sada bili izvan ljudskog dosega (2).

UI daje mogućnost strojevima da obavljaju zadatke koji su nekada bili isključivo ljudski, kao što su prepoznavanje govora, prepoznavanje lica, donošenje kompleksnih odluka i analiza velikih količina podataka. Jedan od najvažnijih aspekata korištenja UI-ja je strojno učenje, koje omogućava računalnim sustavima da uče autonomno iz podataka bez potrebe za detaljnim programiranjem svakog pojedinačnog zadatka (3).

Jedan od ključnih razloga zbog kojih se UI tako brzo širi u posljednjim desetljećima svakako je dostupnost velike količine podataka (engl. Big Data), koja omogućava sustavima strojnog učenja da učinkovito treniraju na različitim tipovima podataka. Uz to, povećanje računalne snage, posebice kroz razvoj specijaliziranih grafičkih procesora (engl. Graphics Processing Unit kratica GPU) i korištenje oblaka u računarstvu (engl. cloud computing), omogućilo je primjenu složenih algoritama koji su nekada bili tehnički neizvedivi ili preskupi za implementaciju (4, 5).

UI kao pojam nastao je sredinom 20. stoljeća, kada su znanstvenici pokušali stvoriti strojeve koji mogu oponašati ljudsko ponašanje, razmišljanje i učenje. Prvi put je pojam „umjetna inteligencija“ upotrijebljen 1956. godine tijekom konferencije u Dartmouthu, koju su organizirali John McCarthy, Marvin Minsky, Nathaniel Rochester i Claude Shannon (6). Tada je definirana kao odjel računarstva koji se bavi kreiranjem sustava koji će biti sposobni zamijeniti ljude u svim zadacima koji zahtijevaju inteligenciju. John McCarthy, često nazivan ocem umjetne inteligencije, je definirao UI kao disciplinu koja se bavi naukom i inženjerstvom

za izradu pametnih mašina te posebno inteligentnih softverskih programa (3). Prema njegovom opisu, UI obuhvaća zadatke koji se odnose na sposobnost strojeva da razmišljaju, uče, rješavaju probleme, percipiraju okolinu, razumiju prirodni jezik i izvršavaju zadatke koji bi inače zahtijevali ljudsku inteligenciju (3).

Tijekom godina, UI definicija se mijenjala u sklopu napretka tehnologija i novih spoznaja. Danas se UI obično opisuje kao sposobnost računalnih sustava da izvode zadatke koji tradicionalno zahtijevaju ljudsku inteligenciju, uključujući vizualnu percepciju, prepoznavanje govora, donošenje odluka i prevođenje između jezika (7).

Bitna razlika koja se često spominje u kontekstu definicije UI-ja jest podjela na usku (slabu) i opću (jaku) umjetnu inteligenciju (8). Uska umjetna inteligencija odnosi se na specijalizirane sustave dizajnirane za obavljanje određenih, jasno definiranih zadataka. To su, npr. algoritmi za prepoznavanje lica, sustava preporučivanja ili autonomnih vozila. Ovi sustavi su dizajnirani za jednu zadatak i nemaju opću sposobnost razmišljanja ili razumijevanja konteksta izvan postojeće zadaće. S druge strane, opća UI podrazumijeva sustave koji bi teoretski mogli razumjeti, naučiti ili primijeniti znanje na bilo koji intelektualni zadatak koji može obaviti ljudsko biće. Opća UI bila bi sposobna za samostalno razmišljanje, donošenje odluka o složenim problemima u različitim situacijama, a također bi imala sposobnost razumijevanja emocionalnih i društvenih aspekata ljudskog života. Opća UI trenutno je hipotetski koncept i predmet futurističkih projekcija. (9).

1.1.2 Ključne tehnologije i pojmovi

Osim definicije UI, važno je razumjeti i ključne koncepte koji se koriste u okviru UI-ja. Najvažniji od njih je svakako strojno učenje. Strojno učenje podrazumijeva pristupe u kojima sustavi UI-ja uče na temelju dostupnih podataka te na osnovi tog učenja donose odluke ili predviđaju ishode. Strojno učenje može biti nadzirano, nenadzirano i polunadzirano, ovisno o tome koliko su jasno označeni podaci na kojima sustav trenira (10,11).

Duboko učenje (engl. deep learning) jedna je od podgrana strojnog učenja, a temelji se na složenim strukturama nazvanim neuronske mreže. Ove mreže su inspirirane biološkim neuronima i sastoje se od velikog broja međusobno povezanih slojeva koji obrađuju informacije na vrlo apstraktan način, omogućavajući sustavima rješavanje izuzetno složenih zadataka poput prepoznavanja govora, obrade prirodnog jezika ili prepoznavanja vizualnih uzoraka (12,13).

Neuronske mreže predstavljaju temelj velikih jezičnih modela kao što su GPT (engl. Generative Pre-trained Transformer), koji su u posljednjih nekoliko godina napravili revoluciju u obradi prirodnog jezika. Zahvaljujući dubokom učenju, ovi modeli mogu generirati koherentne i kontekstualno primjerene tekstove koji se često teško razlikuju od onih koje pišu ljudi (14,15).

Iako definicija UI-ja može varirati ovisno o kontekstu i pristupu, zajednička osnova svih definicija jest sposobnost računalnih sustava da na neki način oponašaju ljudsku inteligenciju, bilo kroz specijalizirane zadatke ili opću inteligenciju koja se još uvijek razvija kao konceptualni cilj (16). Razumijevanje ovih definicija ključno je za pravilnu analizu i raspravu o etičkim aspektima korištenja UI-ja u različitim područjima, uključujući obrazovanje.

1.1.3 Povijest i razvoj umjetne inteligencije

Sama ideja inteligentnih strojeva postoji već stoljećima, nadahnućem ljudske mašte i filozofskog preispitivanja prirode inteligencije. Međutim, tek od ranog 20. stoljeća s razvojem računala i IT-a, inteligencija je postala praktična znanost (17). UI trenutno integrira discipline poput matematike, statistike, informatike, lingvističke neuroznanosti i robotike kako bi proizvela sustave sposobne razumjeti, interpretirati i reagirati na svijet na inteligentan način.

Povijest UI-a može se podijeliti na nekoliko ključnih faza. Od ranog 1950-ih, pioniri poput Alana Turinga postavili su temelje za razumijevanje što se podrazumijeva pojam „inteligentnosti“. Turingov test razvijen 1950. postao je klasični kriterij za procjenu mogućnosti stroja da oponaša ljudsku inteligenciju (18). Razvoj UI-ja doživljava svoj prvi značajniji zamah nakon, ranije spomenute, konferencije u Dartmouthu 1956. godine (6).

Tijekom toga stadija, istraživači su stvorili prve algoritme za rješavanje matematičkih problema i igranje šaha. U narednim desetljećima, razvijeni su algoritmi za rješavanje problema, igranje igara i otkrivanje uzoraka što čini novi podsticaj za UI (19).

Sredinom stoljeća fokus je bilo na razvoju simboličkih sustava i logičkih modela koji su mogli obavljati složene matematičke operacije i donositi odluke. Međutim, taj pristup ima nekoliko ograničenja u smislu prilagodljivosti i skalabilnosti (2). Iako su očekivanja bila velika, tehničke mogućnosti u to vrijeme bile su ograničene. To razdoblje je poznato kao „UI zima“ tijekom 70-ih i 80-ih godina prošlog stoljeća kada su ulaganja u UI i istraživanje dosegla najniže razine (2,20). Usprkos stagnaciji, neki važni projekti ostvarili su značajan napredak.

Primjerice, razvoj ekspertnog sustava MYCIN sredinom 70-ih godina prošlog stoljeća, koji je korišten za medicinsku dijagnostiku, pokazao je praktični potencijal UI-ja (21).

Krajem 20. stoljeća, s pojavom moćnijih računala i velikih količina podataka, dolazi do razvoja strojnog učenja, grane UI-ja koja koristi statističke metode za učenje iz podataka. Danas je duboko učenje, tehnika koja koristi slojevite neuronske mreže, postale srž naprednih UI sustava (2).

S početkom 21. stoljeća, UI doživljava pravu renesansu zahvaljujući dubokom učenju. Duboko učenje temelji se na složenim neuronskim mrežama dostatnim za učenje iz velikih količina podataka (13). Jedan od prvih većih uspjeha dubokog učenja bilo je prepoznavanje slika i glasova, što je omogućilo razvoj tehnologija virtualnih asistenata (Siri, Alexa, Google Assistant) (22). Godine 2011., IBM-ovo računalo „Watson“ prvi put je pobedio prvake u kvizu Jeopardy!, što je označilo veliki korak naprijed u razumijevanju prirodnog jezika i potencijal za složene zadatke (23). Ubrzo nakon toga, autonomna vozila poput Tesle i Googleovih samovozećih automobila dodatno potvrđuju potencijal UI-ja u svakodnevnoj primjeni, podižući interes javnosti i industrije (24). U posljednjih desetak godina, veliki jezični modeli (engl. Large Language Models – LLM) kao što su GPT-3 i GPT-4 postaju ključni dio UI-jevog razvoja. Njihova sposobnost generiranja teksta gotovo neprimjetno različito od ljudskog rada unaprijedila je mnoge industrije (25). Danas najveće svjetske kompanije poput Googlea, Amazona, Applea, Microsofta i mnoge druge aktivno ulažu enormna sredstva u razvoj UI-ja, smatrajući je ključnim čimbenikom budućeg poslovanja (26).

Razvoj UI-ja u zadnjih sedam desetljeća imao je izražene uspone i padove, no trenutno se nalazi u razdoblju stalnog i intenzivnog napretka (27).

1.1.4. Primjena umjetne inteligencije u društvu

Jedna od glavnih karakteristika modernog UI-ja je njegovog sveprisutnost. UI se koristi u gotovo svim aspektima našeg života, od industrije i ekonomije, preko obrazovanja i zdravstva, pa sve do pravosuđa i politike (28).

Područje u kojem UI pokazuje sve veću prisutnost je zdravstvo. Algoritmi dubokog učenja koriste se za dijagnosticiranje bolesti s točnošću koja u nekim slučajevima nadmašuje ljudske stručnjake, osobito u poljima poput radiologije, dermatologije i patologije (29).

Primjena sustava temeljenih na UI omogućuje bržu obradu medicinskih slika, preciznije prepoznavanje anomalija i personaliziranije terapijske pristupe, čime se ne samo povećava učinkovitost liječenja nego i smanjuje opterećenje zdravstvenih radnika (30). Uz to, umjetna inteligencija sve više pomaže u otkrivanju novih lijekova, optimiziranju kliničkih ispitivanja i predviđanju ishoda liječenja (31,32).

Primjenom UI-ja u obrazovanju postizemo mnoge koristi kao što su automatizacija administrativnih poslova te prilagođeno učenje putem razvoja pametnih sustava koji pomažu učenicima i nastavnicima u procesima poučavanja i ocjenivanja. Sustavi UI-ja omogućuju automatsko ocjenivanje eseja, kreiraju prilagođen edukacijske materijale prema specifičnim potrebama svakog pojedinog učenika te pružaju podršku učenicima s posebnim potrebama. (33,34). U obrazovanju, UI omogućuje personalizirano učenje prilagođeno individualnim potrebama učenika. Adaptivne platforme za učenje koriste analitiku podataka kako bi optimizirale kurikule i metode poučavanja u stvarnom vremenu (35). Ova tehnologija može doprinijeti i smanjenju obrazovnih nejednakosti, pružajući dodatnu podršku učenicima koji imaju poteškoće u učenju. Takvi inteligentni tutori, automatizirani sustavi za procjenu znanja i alati za automatsko generiranje povratnih informacija postali su standard u mnogim obrazovnim institucijama (36).

Industrijski sektor također intenzivno koristi umjetnu inteligenciju, osobito u područjima automatizacije, optimizacije lanaca opskrbe i prediktivnog održavanja. Roboti vođeni UI-jem obavljaju zadatke na proizvodnim linijama s razinom preciznosti i učinkovitosti koja je nekada bila nezamisliva (37). UI analitika pomaže poduzećima u donošenju informiranijih odluka temeljenih na obradi velikih skupova podataka, predviđanju potražnje i optimizaciji resursa (38).

U domeni pravosuđa, primjena algoritama za predviđanje recidivizma i automatizaciju procesa odlučivanja već je stvarnost u nekim zemljama. Sustavi poput COMPAS-a (engl. Correctional Offender Management Profiling for Alternative Sanctions) u Sjedinjenim Američkim Državama koriste se za procjenu rizika ponavljanja kaznenih djela, što utječe na odluke o jamčevini, uvjetnim otpustima i presudama (39). Nadalje, UI se koristi i za analizu sudskih presedana, prediktivno modeliranje ishoda sudskih procesa te pomoć u pisanju pravnih dokumenata (40).

Politička scena također sve više koristi UI, primjerice kroz analizu velikih količina podataka o biračima, prediktivno modelira izborne rezultate, generira političke poruka

usmjerene prema specifičnim demografskim skupinama (41). Studije su pokazale da mikrotargetiranje birača pomoću algoritama može značajno utjecati na ishod izbora (42). UI danas omogućuje i sofisticirane analize društvenih mreža kako bi se razumjele političke sklonosti i predviđjele izborne dinamike (43).

Posebno značajnu ulogu UI ima u području sigurnosti i nadzora. Sustavi za prepoznavanje lica, prediktivne sigurnosne analize i masovno prikupljanje podataka koriste UI za poboljšanje javne sigurnosti (44). Osim toga, UI se koristi za obradu i analizu videonadzornih snimki, identifikaciju sumnjivih obrazaca ponašanja i optimizaciju rasporeda sigurnosnih resursa (45).

Osim navedenih područja, UI nalazi svoju primjenu i u financijama, gdje se koristi za algoritamsko trgovanje, procjenu kreditne sposobnosti klijenata i otkrivanje financijskih prevara (46), u poljoprivredi za optimizaciju prinosa putem preciznog navodnjavanja i praćenja stanja usjeva (47), kao i u transportu, osobito u razvoju autonomnih vozila i pametnih prometnih sustava (48).

Široka primjena UI-ja pokazuje koliko je duboko ova tehnologija prodrla u sve pore modernog društva, redefiniirajući načine na koje radimo, učimo, liječimo se, komuniciramo i donosimo odluke.

1.1.5 Izazovi vezani uz umjetnu inteligenciju

Iako UI donosi značajne koristi, ona također otvara niz pitanja i izazova. Jedan od ključnih izazova je etika korištenja UI-ja. Kako osigurati da se UI koristi na odgovoran i pravedan način? Na primjer, algoritmi koji donose odluke u važnim područjima kao što su zapošljavanje, pravosuđe ili medicina mogu biti pristrani ako su podaci na kojima su trenirani nesavršeni. Nadalje, postavlja se pitanje privatnosti i sigurnosti podataka te odgovornosti u eri gdje UI sustavi obrađuju ogromne količine osjetljivih informacija (49).

Uz prednosti koje UI donosi, posebno u zdravstvu, pojavljuju se i zabrinutosti vezane uz transparentnost modela, etičke dileme oko donošenja medicinskih odluka te rizik od sustavne pristranosti u treniranju modela (50). Slično tome, u obrazovanju se otvaraju pitanja o privatnosti podataka učenika i mogućem smanjenju ljudskog aspekta, što su ključni čimbenici za buduća istraživanja i sam razvoj obrazovanja. U gospodarstvu inovacije vezane uz UI donose duboke ekonomske i društvene posljedice, osobito u obliku smanjenja radnih mjesta za

niskokvalificirane radnike, dok istovremeno otvara prilike za nova zanimanja, što potencijalno produbljuje društvene nejednakosti (51). U političkoj sferi otvaraju se ozbiljna pitanja o manipulaciji javnim mnijenjem, širenju dezinformacija i transparentnosti političkih procesa. Na primjer, masovna upotreba tehnologije prepoznavanja lica u javnim prostorima bez pristanka građana izazvala je burne rasprave o granicama legitimne upotrebe ovakvih tehnologija i potrebi da javnost ima uvid i kontrolu nad njima (52).

Na globalnoj razini, UI sve više utječe na oblikovanje međunarodnih odnosa. Koncept "UI utrke" između vodećih svjetskih sila, osobito SAD-a i Kine, postaje značajan faktor u geopolitici (53). Pitanje kontrole nad razvojem i primjenom naprednih UI sustava nije samo tehnološko, već i pitanje nacionalne sigurnosti, ekonomskog prosperiteta i društvene stabilnosti. Istovremeno, postoji zabrinutost da bi nejednak pristup UI-ju mogao dodatno produbiti globalne nejednakosti između razvijenih i nerazvijenih zemalja (54).

Ključno je pitanje kako će društvo upravljati ovom tranzicijom i osigurati da koristi UI-ja budu dostupne svima, a ne samo odabranoj manjini (55).

Tehnološki napredak u UI-ju također pokreće filozofska pitanja o budućnosti čovječanstva. Hoće li strojevi postati toliko napredni da nadmaše ljudsku inteligenciju? Ako da, što to znači za našu ulogu u svijetu? Ideja o superinteligenciji, koja bi mogla nastati od naprednih UI sustava, izaziva i divljenje i strah. Dok određeni futuristi zamišljaju budućnost u kojoj UI rješava svjetske izazove poput klimatskih promjena i epidemija, drugi upozoravaju na moguće rizike od nekontroliranog razvoja (48).

U konačnici, primjena UI-ja u društvu donosi ogromne potencijale za unapređenje kvalitete života, ali istovremeno nosi i značajne rizike koji zahtijevaju promišljeno upravljanje, interdisciplinarno istraživanje i međunarodnu suradnju. Ključna pitanja uključuju potrebu za transparentnošću algoritama, očuvanjem ljudskih prava, pravednim pristupom i osiguravanjem da koristi od razvoja UI-ja budu ravnomjerno raspoređene. Ako se razvoj UI-ja nastavi bez uključivanja marginaliziranih zajednica i razmišljanja o društvenim nejednakostima, postoji opasnost da će UI sustavi dodatno produbiti postojeće nejednakosti, koristeći tehnološki napredak u korist privilegiranih skupina, dok će ranjive populacije ostati još više zapostavljene i marginalizirane (56).

Unatoč izazovima, UI predstavlja neizmjernu priliku za transformaciju našeg društva. Kroz interdisciplinarni pristup i suradnju između znanstvenika, inženjera, filozofa i zakonodavaca, možemo osigurati da UI služi kao alat za poboljšanje kvalitete života, poticanje

inovacija i stvaranje bolje budućnosti. Kako nastavljamo istraživati granice ove tehnologije, važno je zadržati otvoren i kritičan pristup, prepoznajući njezin potencijal, ali i odgovornost koju nosimo kao kreatori i korisnici UI-ja (57).

Istodobno, razvoj UI tehnologija prate rastuće zabrinutosti oko etičkih pitanja kao što su privatnost, sigurnost podataka, algoritamska pristranost i utjecaj na zapošljavanje. Ova pitanja postaju sve izraženija u raspravama o društvenim implikacijama UI-ja (58). Stoga mnoge zemlje i međunarodne institucije rade na izradi regulativa i etičkih smjernica za primjenu UI-ja, kako bi se osigurala njezina sigurna i društveno odgovorna upotreba (59).

1.2 Veliki jezični modeli

1.2.1 Uvod

Veliki jezični modeli (engl. Large Language Models – LLM) predstavljaju najnoviji i najutjecajniji razvojni pravac u području UI-ja, osobito u domeni obrade prirodnog jezika (engl. Natural Language Processing – NLP) (55). Danas postoji veliki broj LLM-ova i njihov broj se stalno povećava jer različite kompanije i istraživačke institucije stvaraju svoje verzije. Njihova glavna svrha jest razumijevanje i generiranje prirodnog jezika na način koji oponaša ljudsku komunikaciju (57). Takvi modeli omogućuju računalima da pišu tekstove, odgovaraju na pitanja, sažimaju sadržaje, pa čak i vode dijaloge na raznim temama. Zbog svoje prilagodljivosti i široke primjene, postaju sve važniji alat u obrazovanju, znanosti, medijima, poslovanju i svakodnevnom životu. Koriste se za podršku u učenju, pomoć pri pisanju, automatizaciju zadataka, pa i za kreativne procese poput pisanja eseja ili priča (58). Njihova sposobnost da generiraju smislen i često vrlo uvjerljiv tekst dovodi do novih mogućnosti, ali i izazova, posebno u kontekstu autentičnosti i akademske čestitosti. Unatoč tome što ne „razumiju“ sadržaj kao ljudi, LLM-ovi su u stanju proizvesti tekst koji stilom, tonom i strukturom često nalikuje ljudskom. Time mijenjaju način na koji komuniciramo s tehnologijom, ali i kako promišljamo o znanju, izražavanju i kreativnosti (59).

1.2.2 Tehnologija LLM-ova

LLM-ovi predstavljaju jednu od najnaprednijih tehnologija u području UI-ja, s fokusom na razumijevanje i generiranje prirodnog jezika. Ovi modeli koriste tehnike strojnog učenja, posebice duboko učenje, kako bi analizirali ogromne količine tekstualnih podataka i oponašali ljudski način komunikacije (55).

Ovi modeli temeljeni su na sofisticiranim arhitekturama neuronskih mreža koje su sposobne analizirati, interpretirati i generirati ljudski jezik u visoko koherentnom i smislenom obliku (60–62). LLM-ovi su trenirani na masivnim količinama tekstualnih podataka, preuzetih s interneta, knjiga, znanstvenih članaka, foruma i mnogih drugih izvora (63). Oni se temelje na arhitekturama neuronskih mreža, od kojih je najpoznatija transformer. Transformeri su modeli dizajnirani za razumijevanje i generiranje jezika, a posebno su dobri u obradi redoslijeda riječi u rečenicama (64). Ključni element ove arhitekture je mehanizam pozornosti (engl. attention mechanism) (65), koji modelima omogućuje da istovremeno analiziraju i povezuju različite dijelove teksta. To je značajno poboljšalo razumijevanje složene semantike i konteksta jezika (66).

Korištenjem arhitektura kao što su transformeri, modeli poput GPT (Generative Pre-trained Transformer), BERT (Bidirectional Encoder Representations from Transformers), PaLM (Pathways Language Model), Claude i Bard sposobni su razumjeti kontekst rečenica, odgovarati na pitanja, prevoditi jezike, pa čak i stvarati kreativne narative ili rješavati kompleksne problemske zadatke (67).

LLM-ovi se oslanjaju na nekoliko ključnih tehnoloških principa koji omogućuju njihovu učinkovitost i široku primjenu (68). Prvi korak u razvoju ovih modela je predtreniranje (engl. pre-training), tijekom kojeg modeli uče statističke obrasce jezika iz golemih količina nepovezanog teksta, stječući tako temeljno razumijevanje jezika. Nakon toga slijedi fino podešavanje (engl. fine-tuning), gdje se modeli dodatno prilagođavaju za specifične zadatke ili domene, čime se poboljšava njihova učinkovitost u određenim kontekstima (69). Drugi važan aspekt je veličina modela, izražena brojem parametara. Što ih je više to su veće mogućnosti razumijevanja i generiranja složenog jezika. Na primjer, GPT-3 ima 175 milijardi parametara, dok se procjenjuje da GPT-4 ima oko 1,8 trilijuna parametara (70). Treći ključni princip je mehanizam pozornosti (engl. attention mechanism), koji omogućuje modelima da razumiju odnose između riječi u rečenicama, bez obzira na njihovu udaljenost. Ova sposobnost poboljšava razumijevanje složenih tekstualnih struktura i konteksta (71). Konačno, i skalabilnost je važan faktor, jer povećanje veličine modela i dostupnosti podataka može dovesti do poboljšanja kvalitete generiranog teksta. Međutim, veći modeli zahtijevaju više računalnih resursa za treniranje i implementaciju, što predstavlja izazov u pogledu učinkovitosti i održivosti (72,73).

Ovi tehnološki principi zajedno omogućuju LLM-ovima da generiraju koherentan i kontekstualno relevantan tekst, čime se otvaraju brojne mogućnosti primjene u različitim industrijama (74).

1.2.3 Primjene velikih jezičnih modela

Jedna od najvažnijih prednosti LLM-ova je njihova svestranost. Zahvaljujući sposobnosti generalizacije, isti model može se koristiti za mnogobrojne zadatke bez potrebe za dodatnim treniranjem. Na primjer, korisnik može postaviti pitanje na prirodnom jeziku, zatražiti sažetak teksta, prevođenje ili generiranje eseja i sve to u sklopu jedne platforme (75).

U području obrazovanja, LLM-ovi donose revoluciju. Oni omogućuju personaliziranu podršku učenicima, odgovaraju na njihova pitanja u realnom vremenu, pomažu pri učenju jezika, matematici, pa čak i programiranju. Nastavnici ih mogu koristiti za izradu kvizova, nastavnih planova i sadržaja prilagođenog razini razumijevanja učenika. Alati poput ChatGPT-a postali su popularni među studentima za izradu nacрта eseja, provjeru gramatike i formuliranje ideja (76).

LLM-ovi imaju širok spektar primjena u različitim područjima svakodnevnog života, obrazovanja i rada. Tehnologije poput Amazona Alexa, Google Assistanta i Appleove Siri koriste LLM-ove kako bi razumjele i odgovarale na glasovne naredbe korisnika, čime omogućuju prirodniju i učinkovitiju interakciju s uređajima (77). Sustavi za automatski prijevod, poput Google prevoditelja, oslanjaju se na velike jezične modele kako bi poboljšali točnost i razumijevanje konteksta prilikom prevođenja između jezika (78). LLM-ovi se sve češće koriste i za pisanje i kreiranje sadržaja. Platforme poput ChatGPT-a služe za stvaranje članaka, računalnog koda, marketinških materijala i mnogih drugih tekstualnih formi (79). Njihova sposobnost obrade velikih količina teksta omogućuje i naprednu analizu sadržaja, uključujući analizu sentimenta, prepoznavanje entiteta i sažimanje teksta, što je posebno korisno u poslovnim i istraživačkim okruženjima (80). U medicinskom i pravnom sektoru, LLM-ovi se koriste za obradu složenih dokumenata kao što su medicinski izvještaji ili pravni akti, čime se ubrzava donošenje odluka i povećava preciznost (81,82). U obrazovanju i znanstvenom istraživanju, studenti, istraživači i edukatori koriste velike jezične modele za učenje, sumiranje znanstvenih radova, rješavanje zadataka, provođenje simuliranih dijaloga te za mnoge druge svrhe koje podupiru razvoj znanja i vještina (83).

1.2.4 Izazovi i ograničenja

Uz svoju moć, LLM-ovi se suočavaju s nizom izazova i ograničenja koji prate njihov razvoj i primjenu (65, 66). Treniranje ovih modela zahtijeva izuzetno velike računalne resurse, što ih čini teško dostupnima manjim istraživačkim timovima i institucijama (86). Osim toga, modeli su izrazito osjetljivi na pristranosti prisutne u podacima na kojima su trenirani, što može dovesti do neželjenih posljedica poput održavanja ili pojačavanja postojećih stereotipa (87). Postoji i značajan rizik od zlouporabe, primjerice u svrhu generiranja dezinformacija, lažnih vijesti ili drugih zlonamjernih sadržaja (88). Iako su sposobni proizvoditi uvjerljiv i gramatički točan tekst, ovi modeli zapravo ne razumiju jezik na način kako to čini čovjek, što može rezultirati nelogičnim ili pogrešnim odgovorima u specifičnim situacijama (89). Uz sve to, etička pitanja postaju sve važnija, od odgovornosti za štetu uzrokovanu netočnim informacijama, do zaštite privatnosti korisnika i transparentnosti u korištenju UI-ja (90).

1.2.5 Budućnost velikih jezičnih modela

LLM-ovi imaju važnu ulogu u medicini, pravu, obrazovanju, novinarstvu, informatici i brojnim drugim industrijama. Ova široka primjena ističe njihovu korisnost, ali i nužnost odgovorne primjene (75).

U akademskom kontekstu postavlja se pitanje koliko je rad zaista rezultat znanja i truda studenta, a koliko je generirano od strane UI-ja. To potiče rasprave o akademskoj čestitosti, plagiranju i etičnosti upotrebe alata u procesu učenja (91,92).

Unatoč ogromnom potencijalu, LLM-ovi nose sa sobom i niz rizika. Jedan od najvećih izazova je već spomenuta pristranost modela. Budući da se treniraju na stvarnim internetskim podacima, LLM-ovi mogu preuzeti stereotipe, netočnosti i diskriminatorne obrasce prisutne u tim podacima. Time se otvara opasnost od širenja netočnih ili štetnih informacija (93).

Drugi važan problem je tzv. „halucinacija“ – pojava kada model generira informacije koje zvuče uvjerljivo, ali nisu točne. U kontekstu obrazovanja, medicine ili prava, takve pogreške mogu imati ozbiljne posljedice. Zato se sve više naglašava potreba za ljudskim nadzorom i provjerom rezultata koje generiraju LLM-ovi (94). Nadalje, tu su i pitanja vezana uz privatnost podataka. Iako LLM-ovi ne pohranjuju eksplicitne podatke o korisnicima, treniranje na nestrukturiranim internetskim podacima može uključivati osjetljive informacije, što otvara mogućnosti zlouporabe ili neželjenog izlaganja podataka (95). S pravnog aspekta, još uvijek nije jasno kako se regulira sadržaj koji generira UI. Primjerice, ako model generira tekst koji sadržava autorski zaštićen materijal, postavlja se pitanje tko snosi odgovornost.

Slično vrijedi i za generiranje štetnog ili dezinformativnog sadržaja. Tko je odgovoran: korisnik, tvrtka koja je izgradila model ili sam sustav (96)?

Zbog svih navedenih rizika, važno je primjenjivati načela odgovorne UI. To uključuje transparentnost modela, etički dizajn, inkluzivnost, mogućnost ljudske kontrole i pristup podacima koji poštuje privatnost. Organizacije poput UNESCO-a, OECD-a i Europske komisije već su objavile smjernice za etičku uporabu UI, uključujući i LLM-ove (97–99). Odgovorna i etička primjena LLM-ova u obrazovanju, znanosti i svakodnevnom životu zahtijeva suradnju znanstvenika, zakonodavaca, edukatora i tehnoloških kompanija. Pritom je važno razvijati kritičku pismenost korisnika za sposobnost razlikovanja korisnog i točnog sadržaja od potencijalno štetnog, što je naročito važno kod mladih generacija koje odrastaju uz UI (100).

Kako tehnologija napreduje, očekuju se učinkovitije arhitekture i manji resursni zahtjevi. Ključno je razvijati metode za smanjenje pristranosti i jačanje etičke primjene, uz suradnju znanstvenika, industrije i zakonodavaca radi odgovorne uporabe LLM-ova (101). Ovi modeli uz sve šire korištenje u društvu, trebaju, za maksimalnu korist i minimalne rizike, odgovoran razvoj i primjenu, da bi imali optimalno korištenje ove tehnologije. (82, 83).

1.3 Korištenje UI u pisanju eseja

Uključivanje UI-ja u proces pisanja eseja otvara nove mogućnosti, ali i zahtijeva kritičko sagledavanje potencijalnih rizika. Iako UI alati mogu ubrzati proces pisanja eseja, olakšati sintezu informacija i ponuditi stilističke smjernice, pretjerano oslanjanje na njih može umanjiti originalnost i kritičko razmišljanje autora. (104).

Jedno od najznačajnijih i ujedno najkontroverznijih područja primjene UI-ja u obrazovanju je generiranje pisanih radova, posebice eseja. LLM poput ChatGPT-a, omogućuju stvaranje sadržaja koji je često stilistički i sadržajno na razini prosječnog ili čak iznadprosječnog studentskog rada (105). Ova sposobnost donosi nove mogućnosti u obrazovnom procesu, ali također postavlja niz etičkih pitanja vezanih uz akademsko poštenje, originalnost i razvoj intelektualnih vještina (106).

U tradicionalnom obrazovanju, pisanje eseja služilo je ne samo kao način provjere znanja, već i kao alat za razvoj sposobnosti analitičkog mišljenja, strukturiranja argumenata, interpretacije izvora te izražavanja osobnog stava. U tom kontekstu, pisanje eseja predstavlja ključnu akademsku kompetenciju (107). Pojava UI-ja koja može pisati eseje brzo i uvjerljivo dovodi u pitanje relevantnost tog zadatka kao pedagoškog alata. Naime, studenti sada imaju

moгуćnost da putem nekoliko ključnih riječi ili kratkog opisa dobiju gotov tekst koji zadovoljava osnovne kriterije forme, sadržaja i koherencije. Iako se ovakvi alati mogu koristiti kao podrška u učenju, primjerice za strukturiranje ideja, lekturu ili pomoć pri izražavanju, njihova nekritička i neetična upotreba može dovesti do plagijata, smanjenja napora uloženog u učenje i degradacije akademskih standarda (104).

Empirijska istraživanja sve više ukazuju na to da studenti širom svijeta koriste UI alate u pisanju eseja (108,109). Po nekim istraživanjima, više od polovine studenata koristi UI za pisanje studentskih radova (110). Iako obrazovne institucije pokušavaju odgovoriti na ovaj trend uvođenjem detekcijskih alata za prepoznavanje UI-generiranog sadržaja, tehnologije se razvijaju paralelno i često nadilaze mogućnosti postojećih sustava za otkrivanje takvih radova (111).

S obzirom na to, sve je jasnije da detekcija više nije dovoljno učinkovita metoda. Potrebno je redefinirati način ocjenjivanja, zadavanja zadataka i procjene autentičnosti studentskog izraza. Neki stručnjaci predlažu povratak na usmene ispite ili pisanje eseja pod nadzorom, dok drugi zagovaraju etičku edukaciju studenata i integraciju umjetne inteligencije u obrazovni proces kroz jasno definirane i transparentne smjernice (104).

Upotreba UI-ja u pisanju eseja ne mora nužno biti negativna pojava. Kada se koristi odgovorno, UI može služiti kao mentor, vodič ili alat za unaprjeđenje pisanja. Studenti koji imaju poteškoćama u izražavanju ili oni kojima nije materinji jezik mogu imati koristi od jezične podrške i primjera strukture argumentacije. Važno je da se studenti educiraju o tome kako ispravno koristiti ovu tehnologiju, važnosti stvaranja originalnog rada i razvoju vlastitih vještina u akademskom pisanju (92, 93).

Akadska čestitost, autorska odgovornost i transparentnost u korištenju alata postaju novi izazovi za sve sudionike obrazovnog procesa; nastavnike, studente, administratore i zakonodavce. Pristup koji kombinira regulaciju, edukaciju i tehnološku podršku predviđa se kao najrealniji put prema odgovornoj integraciji UI-ja u akademski kontekst (114).

UI sustavi često djeluju na temelju unaprijed definiranih algoritama i obrazaca, što može rezultirati homogenizacijom sadržaja (115). Umjesto poticanja inovativnosti, oslanjanje na automatizirane alate može ograničiti dubinu analize i kritičko promišljanje, ključne elemente akademskog pisanja. Time se pojavljuje opasnost da se eseji pretvore u mehanički reproducirane tekstove, u kojima autor gubi osobni doprinos te stvarnu analitičku vrijednost (116). Dodatno, etički izazovi korištenja tehnologija UI postaju sve izraženiji. Transparentnost

u načinu primjene ovih alata, kao i jasno navođenje izvora i korištenih algoritama, postaju imperativ kako bi se izbjegle situacije plagiranja i devalvacije autentičnog autorstva (97, 98). Prekomjerna automatizacija može dovesti do situacija gdje se odgovornost za sadržaj prebacuje s autora na tehnologiju, što dodatno komplicira pitanja akademske integriteta i intelektualnog vlasništva (119).

Kritički osvrt na primjenu UI-ja u pisanju eseja također otkriva rizik od površne obrade tema. Iako sustavi za obradu prirodnog jezika mogu brzo analizirati velike količine informacija, oni ne posjeduju sposobnost dubokog razumijevanja konteksta i nijansi specifičnih akademskih rasprava. Kao rezultat toga, eseji izrađeni uz značajnu pomoć UI-ja mogu biti informativni, ali im nedostaje dubina argumentacije i sposobnost artikuliranja kompleksnih ideja na način koji bi poticao originalno istraživanje (120).

2. CILJEVI RADA I HIPOTEZE

2.1 Lingvistička analiza i usporedba eseja studenata medicine i eseja generiranih ChatGPT umjetnom inteligencijom

Glavni cilj ovog istraživanja:

Provesti lingvističku analizu i usporedbu reflektivnih eseja studenata 5. i 6. godine medicine i jednakog broja eseja na istu temu koji su generirani ChatGPT i Bard umjetnom inteligencijom.

Specifični ciljevi ovog istraživanja su:

1. Provesti lingvističku analizu tekstova koristeći softver Linguistic Inquiry and Word Count (LIWC) (121), unutar kojega ćemo analizirati sve njegove ponuđene varijable za različite kategorije jezika koje se pojavljuju u tekstu.
2. Ispitati svaki od tekstova softverom GPTZero (122) koji provjerava je li tekst napisao čovjek ili je tekst generiran umjetnom inteligencijom.

Hipoteza istraživanja:

1. Postoji razlika između lingvističkih parametara eseja studenata medicine i eseja koje je generirala ChatGPT umjetna inteligencija.
2. Očekujemo veću prevalenciju jezičnih varijabli koje ukazuju na izražavanje emocija u tekstovima studenata medicine nego u tekstovima koje je generirala umjetna inteligencija.
3. U tekstovima koje je napisala umjetna inteligencija očekujemo veću prevalenciju jezičnih varijabli koje ukazuju na kognitivne procese i razmišljanje.

2.2 Kako korištenje ChatGPT mijenja stavove o umjetnoj inteligenciji (UI)

Glavni cilj ovog istraživanja:

Provesti analizu i usporedbu stavova o UI-ju, koji će biti prikazan kao promjena prije i nakon intervencije za svaku grupu. Analize će se provesti usporedbom rezultata svakog sudionika prije i nakon intervencije. Specifični ciljevi ovog istraživanja su:

1. Ispitati da li će veća promjena u stavovima prema UI-ju bit predviđena većom intrinzičnom motivacijom,

2. Ispitati svaki od tekstova softverom GPTZero koji provjerava je li tekst napisao čovjek ili je tekst generiran UI-jem.

Hipoteza istraživanja:

Učenici koji ne koriste ChatGPT, tj. pišu samostalno esej o UI-ju u zdravstvu, pomnije će analizirati zadane materijale, što će dovesti do veće promjene u stavovima prema UI-ju od učenika koji pišu esej pomoću ChatGPT-a. Veća promjena u stavovima prema UI-ju bit će predviđena većom intrinzičnom motivacijom.

3. ISPITANICI I POSTUPCI

3.1 Lingvistička analiza i usporedba eseja studenata medicine i eseja generiranih umjetnom inteligencijom

3.1.1 Ustroj istraživanja

Ovo istraživanje je po načinu dobivanja podataka primarno, a po ustroju presječno s analizom između eseja koje su napisali studenti i onih koje su generirali LLM-ovi poput ChatGPT-a i Barda,

3.1.2 Ispitanici

Prikupljeni su eseji koje su napisali studenti 5. i 6. godine Medicinskog fakulteta Sveučilišta u Splitu, Medicina na engleskom jeziku (USSM, engl. University of Split, School of Medicine) 2023. godine. Kao dio obaveznog zadatka tijekom kolegija Medicinska humanistika V – Klinička etika IV (5. godina) i Medicinska humanistika VI – Medicinska etika V (6. godina), studenti su morali napisati esej u kojem su se osvrnuli na etičke, profesionalne ili moralne dileme s kojima su se susreli tijekom studija i kliničkih rotacija na USSM-u ili tijekom bilo koje kliničke prakse koju su možda imali prije upisa na USSM. Eseji su morali imati 400-500 riječi, bez naslova i podnaslova; morali su sadržavati promišljanje učenika o etičkoj dilemi, a nisu smjeli sadržavati nikakva osobna imena ili druge identifikatore za osobe uključene u opisanu situaciju.

3.1.3 Postupci

Studenti su bili upoznati sa istraživanjem tijekom nastave, gdje smo ih informirali o ciljevima istraživanja. Također su obaviješteni da će eseji biti potpuno anonimizirani i analizirani samo na grupnoj razini. Ove informacije su bile ponovljene u Obrascu za informirani pristanak koji im je bio ponuđen za ispunjavanje pri predaji esejskog zadatka putem sustava za e-učenje Merlin, verzija 2022/2023 (Srce, Zagreb, Hrvatska). Podsjetnici za ispunjavanje navedenog obrasca bili su poslani e-poštom na svaku studijsku godinu u dva navrata (dan prije roka za predaju eseja i ponovno tjedan dana nakon tog roka). Iz analize smo isključili eseje učenika koji nisu dali informirani pristanak ili su odbili sudjelovati.

Svaki od studentskih tekstova ispitan je softverom Originality.AI i GPTZero koji provjerava je li tekst napisao čovjek ili je tekst generiran UI-jem. Izdvojili smo 16 eseja moguće kreiranih putem UI („potencijalno UI eseji“) te su oni zasebno analizirani kako ne bi utjecali na rezultate. Koristili smo javno dostupne, ali plaćene verzije softvera Originality.AI, verzija

2.0 (Originality.AI, Ontario, Kanada) i GPTZero, verzija 2023 za otkrivanje jesu li eseji koje su napisali studenti djelomično ili u potpunosti napisani UI softverom. Eseje koje su napisali studenti proglasili smo kao UI-napisane samo ako su oba softvera za detekciju UI-a kategorizirala ih kao $\geq 50\%$ UI-generirane.

Dva istraživača neovisno su izdvojili do 14 ključnih riječi iz svakog eseja, nakon čega su raspravljali o odstupanjima i zajednički odabrali ključne riječi (raspon = 12-14) koje su se koristile za generiranje uputa za ChatGPT i Bard. Druga dva istraživača zatim su upotrijebili odabrane ključne riječi kako bi neovisno razvili dvije razine uputa: pune upute koje su uključivale sve ključne riječi izdvojene u prvoj fazi i djelomične upute koje su sadržavale polovicu tih ključnih riječi. Osim ovih ključnih riječi, sve upute sadržavale su točne smjernice koje su studenti dobili u okviru zadatka (Dodatak 1).

Nakon stvaranja potpunih i djelomičnih uputa za ključne riječi, dva su istraživača raspravljala o razlikama i odabrala konačne verzije uputa za generiranje eseja putem ChatGPT-a, verzija 3.5 i Bard-a, verzija 2023. Pristupili smo LLM-u preko njihovog otvoreno-dostupnog korisničkog sučelja (ne preko sučelja za programiranje aplikacija) kako bismo simulirali način na koji mu student pristupa u stvarnom životu. Nakon korištenja svake upute, izbrisali smo prethodnu povijest razgovora kako bismo izbjegli utjecaj prethodnih uputa ili generiranih eseja. Nastavili smo s ovim procesom sve dok nismo generirali broj eseja jednak broju eseja koje su napisali studenti.

Sljedeći korak je bila analiza eseja pomoću LIWC-a, izdvajajući glavne LIWC sažete mjere i varijable povezane s društvenim i psihološkim procesima. Zatim smo usporedili dobivene LIWC vrijednosti između eseja koje su napisali studenti i eseja koje je generirala UI.

LIWC je računalni alat za analizu teksta koji procjenjuje postotak riječi u tekstu koje pripadaju više od 100 lingvističkih, psiholoških i tematskih kategorija. Program se temelji na znanstvenim istraživanjima koja povezuju jezik s emocijama, mislima i društvenim ponašanjem. Svaku riječ iz teksta uspoređuje s tim rječnicima i izračunava učestalost po kategoriji. Softver uključuje četiri sažete mjere: analitičko razmišljanje, autoritet, autentičnost i ton. Riječ može pripadati više kategorija, a točnost analize raste s duljinom teksta. LIWC je učinkovit i, što je najvažnije, objektivan alat za procjenu psiholoških aspekata jezika i često se koristi u znanstvenim istraživanjima, posebice u psihologiji i komunikaciji (121).

Također smo proveli pod-analize uspoređujući: eseje generirane umjetnom inteligencijom i one koje su napisali studenti, a koji su označeni kao moguće napisani UI-jem. ("potencijalno UI

eseji”); "prave eseje" koje su napisali studenti (tj. nisu označeni kao moguća UI) s esejima koje je generirala umjetna inteligencija; i "prave eseje" koje su napisali studenti s "potencijalno UI esejima" studenata koji su moguće su koristili UI.

3.1.4 Statistička raščlamba

Svi prikupljeni podatci analizirani su korištenjem statističkog softvera Jamovi, verzija 2.6.17 (projekt jamovi, Sydney, Australija). Izračun minimalne veličine uzorka nije bio primjenjiv u našem slučaju, budući da smo pratili cijele dvije generacije učenika. Broj eseja generiranih putem ChatGPT-a ili Bard-a bio je ekvivalentan broju studentskih eseja.

Kategorijske podatke prikazali smo kao frekvencije i postotke, a varijable iz LIWC analize kao medijane s interkvartilnim rasponima. Za usporedbu ukupnih LIWC rezultata koristili smo Kruskal–Wallisov test, a za parne usporedbe između svake skupine eseja proveli smo Dwass-Steel-Critchlow-Flignerov test. Za usporedbu bilo koje dvije skupine koristili smo Mann–Whitneyjev U test, te smo primijenili Bonferronijevu korekciju za višestruke usporedbe.

3.1.5 Etički aspekti

Etičko povjerenstvo Medicinskog fakulteta Sveučilišta u Splitu odobrilo je ovo istraživanje (klasa: 003–08/23–03/0015, urbroj: 2181–198-03–04-23–0011). Osnovne informacije o istraživanju i što sudjelovanje podrazumijeva usmeno je predstavio jedan od autora studije na početku kolegija tijekom kojih su eseji zadani kao zadatak sudionicima. Naglasili smo da sudjelovanje u istraživanju ni na koji način neće utjecati na njihovu ukupnu ocjenu iz kolegija niti na našu procjenu njihovih eseja. Studentima su te informacije ponovno predstavljene, zajedno s anketom u kojoj su mogli dati ili uskratiti svoj informirani pristanak, na Merlin e-learning platformi Medicinskog fakulteta Sveučilišta u Splitu, gdje su predavali svoje eseje. Samo autori koji su ujedno bili nastavnici na kolegiju i ocjenjivali eseje imali su pristup neanonimiziranim esejima; ostali autori koji su provodili analizu pristupali su i koristili samo anonimizirane verzije. Sve metode istraživanja provedene su u skladu s Helsinškom deklaracijom.

3.2 Povezanost korištenja ChatGPT-a i pristupa esejskom zadatku na nastavi iz medicinske humanistike

3.2.1 Ustroj istraživanja

Ovo istraživanje je po načinu dobivanja podataka primarno, a po ustroju je kvaziekperimentalna studija s dizajnom predtesta i posttesta.

3.2.2 Ispitanici

Sudionici su bili studenti 5. i 6. godine medicine koji su pohađali kolegij Medicinske humanistike na Medicinskom fakultetu Sveučilišta u Splitu, bilo uživo u Splitu, Hrvatska (grupa 1) ili online iz Coburga u Njemačkoj (grupa 2). Studenti dolaze iz različitih zemalja, stoga se studijski program i oba kolegija izvode na engleskom jeziku.

3.2.3 Postupci

Studenti su bili podijeljeni u dvije skupine koje su dobile isti znanstveni članak za obradu. Obje grupe su analizirale pregledni članak koji opisuje različite načine na koje se UI može koristiti u medicini, s primjerima primjene UI u različitim medicinskim disciplinama poput radiologije, patologije i drugih (Topol, 2019). Njihov zadatak je bio pročitati članak i napisati esej u kojem su, na temelju jednog odabranog medicinskog područja, opisali potencijalnu primjenu UI u poboljšanju zdravstvene skrbi.

Prva skupina, koja je pohađala nastavu uživo, je dobila uputu da samostalno napiše esej, dok je drugoj, koja je pohađala nastavu online, bilo izričito naloženo da koriste ChatGPT za generiranje eseja. Studente koji su samostalno pisali eseje nadgledali su dva nastavnika kako bi se osiguralo da eseji nisu pisani uz pomoć LLM-ova. Eseji su dodatno bili provjereni pomoću GPTZero softvera, pri čemu su u analizu bili uključeni samo oni eseji za koje se procijenilo da su najmanje 70 % napisani od strane čovjeka (uz dopuštenu marginu pogreške).

Uoči pisanja eseja smo od studenata zatražili da ispune anketu prije intervencije (obrada članka), kojom su prikupljeni podaci o spolu, dobi, prethodnom iskustvu s korištenjem ChatGPT-a, njihovoj ukupnoj intrinzičnoj motivaciji i stavovima prema UI-ju. Nakon analiziranja članka i pisanja eseja, ponovno smo mjerili stavove obje grupe prema UI-ju te ih dodatno pitali sljedeće: jesu li koristili članak koji smo im dali kao literaturu za esej i koliko su ga pomno proučili, koliko im je zadatak bio težak te bi li radije koristili neki drugi način pisanja eseja (npr. bi li članovi skupine koja je pisala samostalno koristili ChatGPT i obrnuto).

Intrinzična motivacija mjerena je pomoću podskale Intrinsic Motivation iz upitnika Work Preference Inventory (123) (Tablica 1). Podskala se sastoji od 15 tvrdnji, pri čemu su

odgovori ocijenjeni na ljestvici od 1 (u potpunosti se ne slažem) do 5 (u potpunosti se slažem), s ukupnim rasponom bodova od 15 do 75. Viši rezultati ukazuju na višu razinu intrinzične motivacije.

Tablica 1. Popis radnih preferencija – Ljestvica intrinzičke motivacije (engl. Work Preference Inventory – Intrinsic Motivation Scale (Amabile, 1994))

1.	Uživam u rješavanju problema koji su mi potpuno novi.
2.	Više volim posao za koji znam da ga mogu dobro raditi nego posao koji zahtjeva nove sposobnosti.
3.	Uživam u pokušajima rješavanja složenih problema.
4.	Radije samostalno pronalazim rješenja.
5.	Uživam u radu koji me toliko zaokupi da zaboravljam na sve ostalo.
6.	Važno mi je da mogu raditi ono u čemu najviše uživam.
7.	Uživam u pokušajima rješavanja vrlo teškog problema
8.	Uživam u provođenju eksperimenata kako bih testirao svoje ideje.
9.	Radije radim na projektima koji me tjeraju da učim nove stvari.
10.	Uživam smišljati nove ideje.
11.	Uživam u prihvaćanju teškog izazova.
12.	Više volim aktivnosti koje mogu dobro raditi nego one koje izazivaju moje vještine.
13.	Znatiželja je pokretačka snaga iza većine onoga što radim.
14.	Što je problem teži, to više uživam u pokušaju da ga riješim.
15.	Najviše sam motiviran za rad na zadacima koji su mi zanimljivi.

Stavovi prema Ui-ju mjereni su pomoću upitnika Attitudes Towards AI: Measurement and Associations with Personality (ATTARI-12) (124) (Tablica 2). Upitnik se sastoji od 12 tvrdnji, pri čemu se odgovori ocjenjuju na ljestvici od 1 (u potpunosti se ne slažem) do 5 (u

potpunosti se slažem), s ukupnim rasponom bodova od 12 do 60. Viši rezultati ukazuju na pozitivnije stavove prema UI-ju.

Tablica 2. Upitnik o stavovima prema umjetnoj inteligenciji (engl. Attitudes towards artificial intelligence Scale (ATTARI-12))

1	UI će ovaj svijet učiniti boljim mjestom.
2	Imam snažne negativne emocije prema UI.
3	Želim koristiti tehnologije koje se oslanjaju na UI.
4	UI ima više nedostataka nego prednosti.
5	Radujem se budućem razvoju UI.
6	UI nudi rješenja za mnoge svjetske probleme.
7	Više volim tehnologije koje ne sadrže UI.
8	Bojim se UI.
9	Radije bih odabrao tehnologiju s UI nego onu bez nje.
10	UI stvara probleme, a ne rješava ih.
11	Kad razmišljam o UI, imam uglavnom pozitivne osjećaje.
12	Radije bih izbjegavao tehnologije koje se temelje na UI.

Oba upitnika inkorporirana su u anketu (Privitak 3) koju su ispunjavali studenti. Podatke iz ankete prikupili smo putem internetskog alata SurveyMonkey (SurveyMonkey Inc., San Mateo, CA, SAD). Studenti su anketu ispunjavali za vrijeme nastave, pri čemu su imali na raspolaganju najviše 135 minuta za ispunjavanje zadatka. Jedan nastavnik bio je cijelo vrijeme prisutan i nadgledao provedbu kako bi osigurao razumijevanje uputa i njihovo pravilno provođenje. Analizirani su samo u potpunosti ispunjeni upitnici.

3.2.4 Statistička raščlamba

Svi prikupljeni podatci analizirani su korištenjem statističkog softvera Jamovi, verzija 2.6.17 (projekt jamovi, Sydney, Australija). Pretpostavili smo da bi očekivana razlika između dvije skupine na ATTARI-12 bila sedam bodova na ljestvici od 12 do 60, sa standardnom devijacijom od 16. Na temelju tih parametara, s 80% snage i stope alfa pogreške od 5%, izračunali smo da će nam trebati 83 sudionika za promatrati pretpostavljenu razliku. Koristili smo online kalkulator veličine uzorka (<https://select-statistics.co.uk/calculators/sample-size-calculator-two-means/>).

Najprije smo provjerili postoje li razlike među studentskim skupinama na početku istraživanja koristeći hi-kvadrat test ili Mann–Whitneyjev U test. Zatim smo analizirali promjene u primarnom ishodu (rezultat na ljestvici ATTARI-12) prije i nakon intervencije unutar svake skupine. To smo mjerili usporedbom rezultata svakog sudionika na toj ljestvici prije i nakon intervencije, koristeći t-test za povezane uzorke. Također smo ispitali u kojoj mjeri dob, spol ili intrinzična motivacija predviđaju rezultat na ljestvici ATTARI-12.

3.2.5 Etički aspekti

Ovo istraživanje dobilo je odobrenje od Etičkog povjerenstva Medicinskog fakulteta Sveučilišta u Splitu (br. 2181-198-03-04-24-0056). Budući da je zadatak bio obvezan u sklopu kolegija, informirani pristanak prikupljen je tek nakon njegova završetka; tada su studenti obaviješteni da njihova odluka o sudjelovanju ni na koji način neće utjecati na kolegij te da sudjelovanje mogu odbiti iz bilo kojeg razloga. Studentski podatci bili anonimizirani prije analize, koja je provedena isključivo na skupnoj razini.

4. REZULTATI

4.1 Lingvistička analiza i usporedba eseja studenata medicine i eseja generiranih umjetnom inteligencijom

Od ukupno 69 studenata koji su pohađali kolegije na kojima su ovi eseji zadani kao završni zadatak, 47 (68 %) pristalo je sudjelovati u istraživanju; stoga smo uključili njihove eseje i generirali jednak broj eseja s djelomičnim i potpunim ključnim riječima koristeći ChatGPT i Bard. Na temelju provjere pomoću UI detektora, utvrđeno je da je 16 (34 %) eseja vjerojatno djelomično ili u potpunosti generirano uz pomoć umjetne inteligencije.

4.1.1 Usporedba eseja koje su napisali studenti s esejima generiranim umjetnom inteligencijom

Objе skupine eseja generiranih umjetnom inteligencijom imale su značajno više vrijednosti u kategoriji „Afekt“ u usporedbi s esejima koje su napisali studenti, pri čemu se ta razlika ponajprije odnosi na više vrijednosti podvarijabli „Pozitivan ton“ i „Emocije“. Suprotno tome, studentski eseji općenito su imali više rezultate u varijablama „Broj riječi po rečenici“ i „Kognicija“ (posebno u podvarijablama „Kognitivni procesi“ i „Sve-ili-ništa“ unutar kategorije „Kognicija“), iako je razlika u prvoj varijabli bila marginalna u odnosu na eseje generirane putem ChatGPT-a, a u drugoj varijabli u odnosu na eseje generirane putem Barda. Razlike u preostalim varijablama bile su promjenjive; primjerice, eseji koje je generirao ChatGPT imali su više rezultate u varijablama „Analitičko razmišljanje“, „Autentičnost“ i „Duge riječi“ u usporedbi s bilo kojom skupinom eseja. Također smo pronašli statistički značajne razlike između eseja koje je generirao ChatGPT i onih koje su napisali studenti u varijablama „Autentičnost“, „Ton“, „Duge riječi“, „Nagoni“ i „Riječi iz rječnika“ (s različitim smjerovima razlike), dok takve razlike nisu bile prisutne u usporedbi eseja koje je generirao Bard sa studentskim esejima. Ukupno gledano, nismo pronašli statistički značajnu razliku u varijabli „Autoritet“ (Tablica 3).

Tablica 3. Usporedba eseja koje su napisali studenti s esejima generiranim umjetnom inteligencijom na temelju svih ključnih riječi*

	LIWC rezultati, medijan (IQR)			Sveukupno			ChatGPT generirani vs. eseja koje su napisali studenti		Bard generirani vs. eseja koje su napisali studenti		Bard generirani vs. ChatGPT-generirani eseja	
	Eseji koje su napisali studenti (n = 47)	Eseji generirani ChatGPT-om (n = 47)	Eseji generirani Bard-om (n = 47)	χ^2	df	P	W	P	W	P	W	P
Analitičko razmišljanje	71.4 (59.2-84)	90.5 (82.6-93.5)	52.8 (46.4-66.1)	68.689	2	<0.001	-7.47	<0.001	5.22	<0.001	11.1	<0.001
Autoritet	41 (27.2-54.2)	36.5 (23.4-48.3)	40.8 (23.3-63.5)	2.619	2	0.27	2.069	0.309	-0.15	0.994	-1.888	0.376
Autentičnost	31.8 (24.3-53.3)	44.4 (34.6-62.5)	27.9 (12.1-42.3)	15.085	2	<0.001	-3.87	0.017	2.01	0.33	5.13	<0.001
Ton	22.2 (11.9-39.2)	34.3 (22.8-48.4)	33.9 (22-46.8)	8.133	2	0.017	-3.962	0.014	-2.887	0.103	0.594	0.907
Nagon	4.14 (3.31-5.32)	5.58 (4.64-6.8)	4.9 (3.85-5.89)	13.576	2	0.001	-5.15	<0.001	-2.13	0.289	3.08	0.075
Pripadnost	1.39 (0.885-2.25)	2.13 (1.64-3.08)	1.63 (0.87-2.67)	14.343	2	<0.001	-5.33	<0.001	-1.04	0.741	3.67	0.025
Postignuće	1.17 (0.83-1.72)	1.71 (1.22-2.25)	1.25 (0.705-1.94)	8.259	2	0.016	-3.7	0.024	0.278	0.979	3.326	0.049
Moć	1.29 (1.01-1.76)	1.67 (1.2-2.13)	1.47 (1.14-1.92)	4.315	2	0.116	-2.89	0.102	-1.71	0.447	1.27	0.643
Kognicija	15 (13.5-16.8)	11.1 (10-12.4)	14.4 (12.1-15.2)	47.203	2	<0.001	9.05	<0.001	3.55	0.032	-6.88	<0.001
Sve ili ništa	0.8 (0.52-1.18)	0.37 (0.2-0.43)	0.7 (0.51-1.03)	34.779	2	<0.001	6.867	<0.001	0.776	0.847	-7.559	<0.001
Kognitivni procesi	13.9 (12.6-15.8)	10.7 (9.65-11.6)	13.3 (11.2-14.3)	39.716	2	<0.001	8.54	<0.001	3.25	0.057	-5.95	<0.001
Uvid	3.44 (2.46-4.09)	4.04 (3.27-4.68)	4.51 (3.71-4.92)	12.939	2	0.002	-3.68	0.025	-4.69	0.003	-1.84	0.395

Uzročnost	1.96 (1.43-2.51)	1.24 (0.945-1.63)	1.47 (1.06-1.98)	17.411	2	<0.001	5.68	<0.001	3.69	0.025	-2.49	0.183
Razlika	1.76 (1.29-2.57)	1.11 (0.805-1.67)	2.05 (1.5-2.49)	29.566	2	<0.001	5.64	<0.001	-1.2	0.674	-7.43	<0.001
Pokušaj	2.17 (1.48-3.25)	1.37 (0.825-1.67)	1.67 (1.17-2.04)	27.568	2	<0.001	7.1	<0.001	4.34	0.006	-3.68	0.025
Izvjescnost	0.53 (0.2-0.895)	0.21 (0-0.41)	0.17 (0-0.46)	15.017	2	<0.001	4.783	0.002	4.646	0.003	0.645	0.892
Neslaganje	4.22 (3.08-4.87)	1.99 (1.63-2.36)	3.01 (2.59-3.69)	61.041	2	<0.001	9.76	<0.001	5.6	<0.001	-7.75	<0.001
Afekt	4.27 (3.42-5.22)	5.66 (4.71-6.79)	5.54 (4.57-6.71)	21.35	2	<0.001	-5.599	<0.001	-5.69	<0.001	0.46	0.943
Pozitivan ton	2.21 (1.75-2.65)	3.13 (2.52-3.59)	3.1 (2.25-3.66)	22.781	2	<0.001	-6.134	<0.001	-5.518	<0.001	0.267	0.981
Negativan ton	2.02 (1.19-2.81)	1.96 (1.44-2.62)	2.01 (1.56-2.76)	1.403	2	0.496	-0.492	0.936	-1.588	0.5	-1.198	0.674
Osjećaj	1.22 (0.9-2.2)	2.14 (1.56-3.15)	1.46 (0.945-2.12)	15.62	2	<0.001	-4.989	0.001	-0.62	0.9	4.631	0.003
Pozitivna emocija	0.4 (0.125-0.605)	0.55 (0.28-0.8)	0.36 (0.17-0.59)	5.116	2	0.077	-2.721	0.132	-0.15	0.994	2.808	0.116
Negativne emocije	0.83 (0.4-1.28)	1.2 (0.615-1.46)	0.79 (0.435-1.36)	3.811	2	0.149	-2.589	0.16	-0.626	0.898	2.075	0.307
Tjeskoba	0 (0-0.22)	0.36 (0.2-0.6)	0.2 (0-0.36)	16.198	2	<0.001	-5.46	<0.001	-3.13	0.069	3.05	0.079
Ljutnja	0 (0-0.195)	0 (0-0.195)	0 (0-0.195)	1.465	2	0.481	1.356	0.603	-0.301	0.975	-1.602	0.494
Tuga	0 (0-0.095)	0.2 (0-0.385)	0 (0-0.2)	9.397	2	0.009	-4.18	0.009	-2.07	0.308	2.51	0.178
Društveni procesi	12.4 (10.5-15)	12.1 (10.5-14)	14.5 (12.8-17)	16.27	2	<0.001	0.866	0.814	-4.374	0.006	-5.363	<0.001
Društveno ponašanje	4.83 (3.8-5.56)	7.04 (5.5-7.67)	5.94 (5.29-6.77)	24.402	2	<0.001	-6.26	<0.001	-4.79	0.002	3.49	0.036
Prosocijalno ponašanje	0.55 (0.29-1.29)	1.6 (1.06-2.33)	1.09 (0.845-1.73)	22.609	2	<0.001	-6.33	<0.001	-4.3	0.007	3.05	0.078
Uljudnost	0 (0-0.225)	0 (0-0.2)	0 (0-0.18)	1.291	2	0.524	1.5481	0.518	1.2101	0.669	- 0.0425	1

Međuljudski sukob	0.21 (0-0.51)	0.21 (0.09-0.43)	0.35 (0.145- 0.5)	0.517	2	0.772	-0.308	0.974	-0.96	0.776	-0.728	0.864
Moraliziranje	0.93 (0.4-1.33)	1.74 (1.4-2.04)	1 (0.67-1.31)	40.468	2	<0.001	-7.81	<0.001	-1.2	0.674	7.67	<0.001
Komunikacija	1.77 (1.19- 2.37)	1.28 (0.825- 1.65)	1.81 (1.27- 2.62)	13.079	2	0.001	3.89	0.016	-1.07	0.73	-4.78	0.002

*Kruskal-Wallis test with Dwass-Steel-Critchlow-Fligner pairwise comparisons.

4.1.2 Eseji s djelomičnim u odnosu na potpune ključne riječi

Nismo pronašli razlike pri usporedbi eseja generiranih pomoću ChatGPT-a na temelju djelomičnih i potpunih skupova ključnih riječi. Međutim, prilikom usporedbe eseja generiranih Bardom, uočili smo statistički značajne razlike u varijablama „Analitičko razmišljanje“, „Broj riječi po rečenici“, „Riječi iz rječnika“ i „Društveni procesi“. Iako su te razlike nestale nakon korekcije za višestruke usporedbe, proveli smo dodatnu podanalizu uspoređujući eseje s djelomičnim ključnim riječima s izvornim studentskim esejima. Rezultati su ostali isti kao u glavnoj analizi, što sugerira da su u početku utvrđene razlike između dviju vrsta eseja generiranih Bardom vjerojatno bile posljedica slučajnosti (Tablice 4 i 5).

Tablica 4. Usporedba eseja generiranih umjetnom inteligencijom na temelju djelomičnih i potpunih ključnih riječi

	LIWC rezultati, medijan (IQR)		Statistika	P*
	ChatGPT-generirani eseji s djelomičnim ključnim riječima (n = 47)	ChatGPT-generirani eseji s potpunim ključnim riječima (n = 47)		
Analitičko razmišljanje	87.5 (85.3-92.6)	90.5 (82.6-93.5)	1095	0.946
Autoritet	31.3 (19.1-45.7)	36.5 (23.4-48.3)	985	0.37
Autentičnost	47.3 (33.1-62.3)	44.4 (34.6-62.5)	1069	0.792
Ton	39.7 (29.3-57.7)	34.3 (22.8-48.4)	955	0.26
Nagoni	5.51 (4.63-6.82)	5.58 (4.64-6.8)	1097	0.958
Kognicija	11.8 (10.7-13.1)	11.1 (10-12.4)	861	0.066
Afekt	5.79 (4.68-7.19)	5.66 (4.71-6.79)	1049	0.675
Društveni procesi	12.4 (10.3-14.7)	12.1 (10.5-14)	1059	0.731
	Bard-generirani eseji s djelomičnim ključnim riječima (n = 47)	Bard-generirani eseji s potpunim ključnim riječima (n = 47)	Statistika	P*
Analitičko razmišljanje	43.6 (32.1-59.2)	52.8 (46.4-66.1)	766	0.010
Autoritet	46.5 (29.4-64.5)	40.8 (23.3-63.5)	979	0.345
Autentičnost	24.1 (9.71-44.2)	27.9 (12.1-42.3)	1039	0.623
Ton	29 (14.2-48.2)	33.9 (22-46.8)	1006	0.460
Nagoni	4.63 (3.67-5.36)	4.9 (3.85-5.89)	954	0.255
Kognicija	13.6 (12.8-15.7)	14.4 (12.1-15.2)	999	0.425
Afekt	5.89 (5.13-6.61)	5.54 (4.57-6.71)	977	0.337
Društveni procesi	17.4 (14-18.6)	14.5 (12.8-17)	839	0.045

*Mann-Whitney U. $p < 0,004$ smatra pragom značajnosti nakon Bonferronijeve korekcije.

Tablica 5. Usporedba izvornih eseja i eseja generiranih Bardom na temelju djelomičnih ključnih riječi za varijable u kojima je postojala razlika između eseja s potpunim i djelomičnim ključnim riječima (n = 47)

	LIWC vrijednost, medijan (IQR)		Statistika	P*
	Eseji koje su napisali studenti (n = 47)	Bard-generirani eseji (n = 47)		
Broj riječi	503 (465-593)	539 (465-591)	1050.0	0.683
Analitičko razmišljanje	71.4 (59.2-84)	43.6 (32.1-59.2)	438.0	<0.001
Društveni procesi	12.4 (10.5-15)	17.4 (14-18.6)	459.0	<0.001
Društveno ponašanje	4.83 (3.8-5.56)	6.38 (4.46-7.54)	711.0	0.003
Prosocijalno ponašanje	0.55 (0.29-1.29)	1.27 (0.735-1.94)	633.0	<0.001
Uljudnost	0 (0-0.225)	0 (0-0.19)	1027.5	0.520
Intepersonalni sukob	0.21 (0-0.51)	0.19 (0-0.41)	1003.0	0.437
Moraliziranje	0.93 (0.4-1.33)	0.91 (0.51-1.46)	1016.5	0.508
Komunikacija	1.77 (1.19-2.37)	1.94 (1.29-2.71)	934.5	0.200

*Mann-Whitney U. $p < 0,004$ smatra pragom značajnosti nakon Bonferronijeve korekcije.

4.1.3 Podanaliza – studentski eseji koji su potencijalno generirani uz pomoć umjetne inteligencije

Šesnaest studentskih eseja bilo je, prema detektorima UI, vjerojatno u cijelosti ili djelomično napisano pomoću umjetne inteligencije; ova podskupina imala je LIWC rezultate sličnije “UI-generiranim” esejima, ali se i dalje razlikovala od „pravih“ eseja koje su napisali studenti. Na primjer, u usporedbi s „pravim“ studentskim esejima, ova je skupina imala više rezultate u varijablama poput „Afekt“ i „Analitičko razmišljanje“, što je u skladu s glavnom analizom, ali i niže rezultate u varijabli „Autentičnost“, dok u varijabli „Ton“ nije bilo statistički značajnih razlika (Tablica 6).

Tablica 6. Usporedba eseja koje su napisali studenti u suradnji s umjetnom inteligencijom i "istinitih" eseja koje su napisali učenici*

	Eseji koje su napisali studenti u suradnji s UI	Pravi eseji koje su napisali studenti	Statistika	P
Analitičko razmišljanje	77.6 (71.1- 92.1)	67.4 (54-74.8)	129.5	0.008
Autoritet	51 (38.6-56.5)	34.5 (26.5-48)	162.5	0.056
Autentičnost	26 (14.1-29)	38.6 (29.2- 54.3)	102	0.001
Ton	34.8 (12.4-59)	20.2 (11.6- 32.8)	171	0.086
Nagon	4.28 (3.43- 6.05)	4.08 (3.29-4.9)	212	0.43
Pripadnost	1.32 (0.927- 2.24)	1.53 (0.845- 2.3)	247	0.991
Postignuće	1.46 (0.917- 2.35)	1.13 (0.83- 1.61)	200	0.286
Moć	1.49 (1.06- 2.12)	1.19 (1-1.67)	194	0.23
Kognicija	14.2 (12.5- 15.6)	15.6 (13.5- 17.4)	173	0.094
Sve ili ništa	0.485 (0-0.848)	0.88 (0.67- 1.21)	130	0.008
Kognitivni procesi	13.6 (12-15)	14.4 (12.9-16)	192	0.215
Uvid	3.71 (2.46-5)	3.44 (2.54- 3.81)	198	0.266
Uzročnost	2.14 (1.96- 2.65)	1.66 (1.29- 2.41)	173	0.094
Razlika	1.45 (1.1-2.46)	1.81 (1.36- 2.63)	195	0.238
Pokušaj	1.61 (1.26- 3.01)	2.28 (1.69- 3.35)	186.5	0.171
Izvjesnost	0.2 (0-0.762)	0.59 (0.315- 1.11)	149	0.026
Neslaganje	3.02 (2.4-3.98)	4.42 (3.67- 5.19)	108.5	0.002
Afekt	4.97 (4.32- 5.71)	3.94 (3.07- 4.92)	137.5	0.014
Pozitivan ton	2.72 (2.22- 3.43)	2 (1.5-2.3)	110.5	0.002
Negativan ton	2.08 (1.13- 2.65)	2 (1.21-2.81)	244.5	0.946
Osjećaj	1.17 (0.925- 1.8)	1.41 (0.86- 2.23)	228.5	0.67
Pozitivna emocija	0.465 (0-0.633)	0.4 (0.155- 0.565)	238	0.83

Negativna emocija	0.565 (0.363-0.992)	0.89 (0.435-1.39)	198.5	0.271
Tjeskoba	0.195 (0-0.292)	0 (0-0.205)	163	0.04
Ljutnja	0 (0-0.05)	0.1 (0-0.195)	190	0.152
Tuga	0 (0-0)	0 (0-0.195)	228.5	0.578
Društveni procesi	14.3 (11.6-15.8)	12.1 (9.96-13.9)	163	0.058
Društveno ponašanje	6.02 (4.64-8.59)	4.4 (3.77-5.39)	111	0.002
Prosocijalno ponašanje	1.51 (0.495-2.24)	0.44 (0.24-1.01)	135	0.011
Uljudnost	0.1 (0-0.255)	0 (0-0.21)	209	0.343
Međuljudski sukob	0.215 (0-0.547)	0.2 (0.075-0.495)	234.5	0.768
Moraliziranje	1.53 (1.09-1.79)	0.52 (0.28-0.985)	66	<0.001
Komunikacija	1.72 (0.942-2.11)	1.77 (1.21-2.4)	221.5	0.559

*Mann-Whitney U. $p < 0,004$ smatra pragom značajnosti nakon Bonferronijeve korekcije.

Ipak, eseji koje su studenti predali, a koji su potencijalno suautorski generirani uz pomoć UI, pokazali su veću sličnost s “UI-generiranim esejima”, budući da ranije utvrđene razlike u glavnoj analizi – u varijablama poput „Ton“, „Društveni procesi“ i „Afekt“ – više nisu bile statistički značajne. Također, nismo uočili statistički značajnu razliku u varijabli „Analitičko razmišljanje“ između studentskih i ChatGPT-generiranih eseja (Tablica 7).

Tablica 7. Usporedba studentskih eseja koji su potencijalno napisani uz pomoć UI s esejima generiranim UI na temelju potpunih ključnih riječi*

Variable	LIWC rezultati, medijan (IQR)			Sveukupno			ChatGPT generirani vs. eseja potencijalno napisanih uz pomoć UI		Bard generirani vs. eseja potencijalno napisanih uz pomoć UI		ChatGPT generirani vs. Bard generiranih eseja	
	UI/studentski eseji (n = 16)	Ekvivalentni Bard-generirani eseji (n = 16)	Ekvivalentni ChatGPT-generirani eseji (n = 16)	χ^2	df	P	W	P	W	P	W	P
Analitičko razmišljanje	77.6 (71.1-92.1)	61.9 (49.7-70.5)	89.5 (81.3-94.5)	22.1037	2	<0.001	-2.53	0.173	4.37	0.006	6.34	<0.001
Autoritet	51 (38.6-56.5)	44 (27.6-65.5)	32.6 (25.1-48.1)	4.9311	2	0.085	3.358	0.046	0.48	0.939	-1.866	0.384
Autentičnost	26 (14.1-29)	25.2 (19.1-40.5)	55.5 (34.9-65.2)	14.346	2	<0.001	-5.14	<0.001	-1.01	0.754	3.89	0.016
Ton	34.8 (12.4-59)	41.8 (30.1-48.1)	30.4 (24.7-54.6)	0.4534	2	0.797	-0.373	0.962	-0.746	0.858	-0.853	0.819
Nagon	4.28 (3.43-6.05)	5.07 (4.16-5.78)	5.18 (4.73-6.23)	1.9529	2	0.377	-1.919	0.364	-0.933	0.787	1.119	0.709
Pripadnost	1.32 (0.927-2.24)	2 (0.867-2.7)	1.87 (1.46-2.74)	3.6549	2	0.161	-2.958	0.092	-1.306	0.626	0.8	0.839
Postignuće	1.46 (0.917-2.35)	1.72 (0.973-2.22)	2 (1.25-2.9)	0.7676	2	0.681	-1.039	0.743	-0.32	0.972	1.066	0.732
Moć	1.49 (1.06-2.12)	1.25 (0.94-1.46)	1.6 (1.27-2.08)	3.887	2	0.143	-0.533	0.925	1.866	0.384	2.799	0.117
Kognicija	14.2 (12.5-15.6)	14 (12.4-16)	11.3 (10.5-12.5)	12.6393	2	0.002	4.371	0.006	0.133	0.995	-4.291	0.007
Sve ili ništa	0.485 (0-0.848)	0.67 (0.37-0.86)	0.2 (0-0.42)	8.7037	2	0.013	1.52	0.528	-1.77	0.425	-4.57	0.004

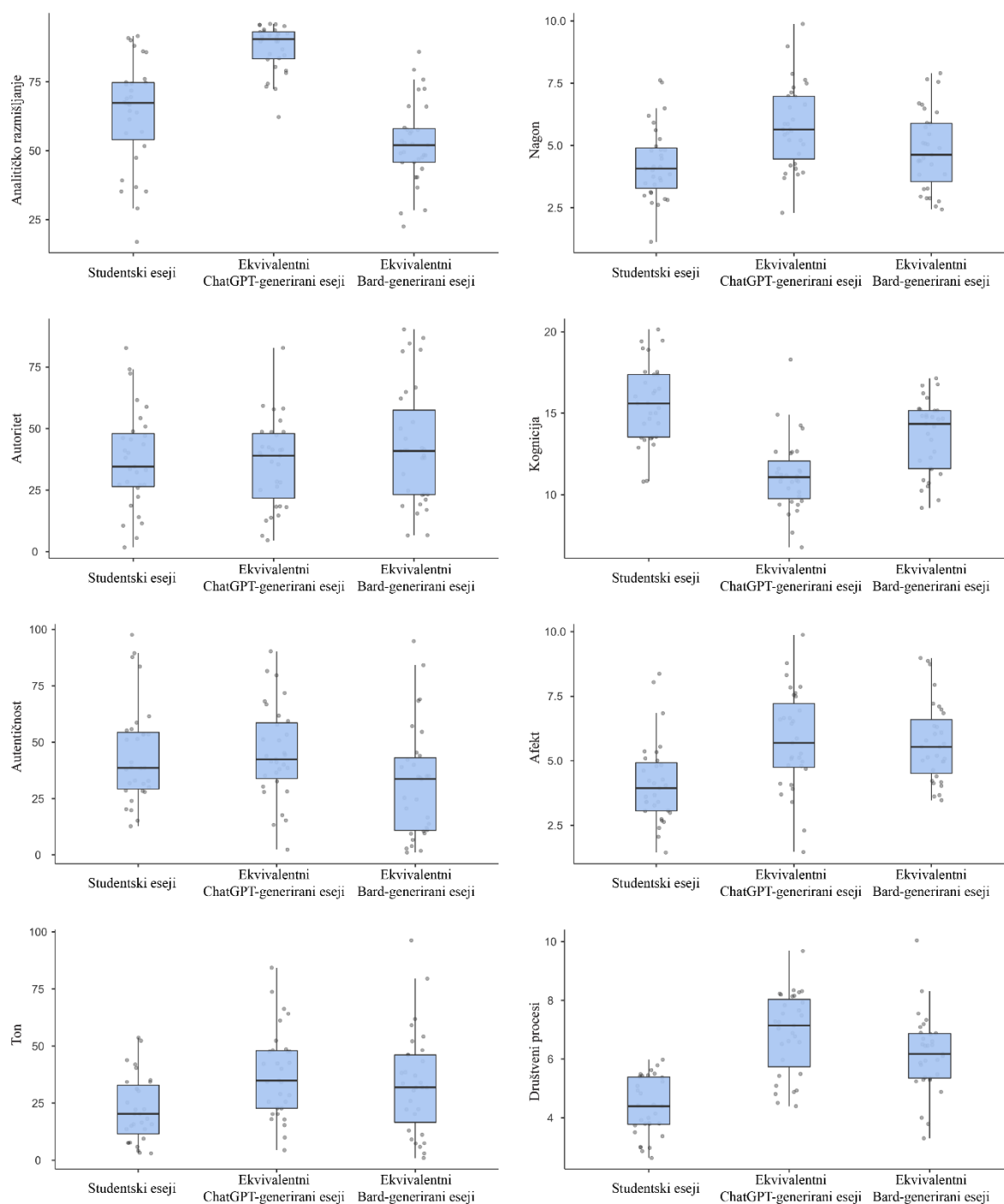
Kognitivni procesi	13.6 (12-15)	13.3 (11.8-15.2)	11 (10.3-11.9)	9.1198	2	0.01	3.731	0.023	0.16	0.993	-3.625	0.028
Uvid	3.71 (2.46-5)	4.44 (2.82-5.87)	4.68 (3.9-5.49)	1.9066	2	0.385	-2.106	0.296	-1.093	0.72	0.4	0.957
Uzročnost	2.14 (1.96-2.65)	1.53 (1.27-2.21)	1.23 (1.09-1.71)	9.3623	2	0.009	4.16	0.009	2.48	0.186	-2.11	0.296
Razlika	1.45 (1.1-2.46)	2.13 (1.66-2.52)	0.905 (0.695-1.37)	16.497	2	<0.001	3.04	0.08	-2.21	0.261	-5.92	<0.001
Pokušaj	1.61 (1.26-3.01)	1.86 (1.47-2.34)	1.19 (0.77-1.69)	7.8108	2	0.02	3.305	0.051	0.48	0.939	-3.492	0.036
Izvjesnost	0.2 (0-0.762)	0.4 (0-0.573)	0.22 (0.135-0.412)	0.209	2	0.901	0.325	0.971	0.219	0.987	-0.785	0.844
Neslaganje	3.02 (2.4-3.98)	2.99 (2.48-3.43)	2.15 (1.82-2.34)	12.1214	2	0.002	4.238	0.008	0.666	0.885	-4.211	0.008
Afekt	4.97 (4.32-5.71)	5.54 (4.79-6.68)	5.63 (4.7-6.3)	3.7359	2	0.154	-2.399	0.207	-2.292	0.237	0.32	0.972
Positive tone	2.72 (2.22-3.43)	3.18 (2.97-3.56)	2.92 (2.57-3.35)	2.2081	2	0.332	-1.09	0.72	-1.95	0.354	-1.28	0.637
Negativan ton	2.08 (1.13-2.65)	1.98 (1.48-2.52)	1.87 (1.64-2.66)	0.2454	2	0.885	-0.6664	0.885	-0.533	0.925	0.0533	0.999
Osjećaj	1.17 (0.925-1.8)	1.23 (0.963-1.94)	1.99 (1.64-2.83)	6.7666	2	0.034	-3.225	0.059	-0.4	0.957	3.091	0.074
Pozitivna emocija	0.465 (0-0.633)	0.415 (0.195-0.55)	0.57 (0.35-0.68)	1.2154	2	0.545	-1.233	0.658	0.107	0.997	1.442	0.565
Negativne emocije	0.565 (0.363-0.992)	0.61 (0.4-0.978)	1.02 (0.617-1.22)	2.8448	2	0.241	-2.08	0.305	-0.4	0.957	1.999	0.334
Tjeskoba	0.195 (0-0.292)	0.275 (0.135-0.5)	0.415 (0.235-0.612)	4.6418	2	0.098	-3.08	0.075	-1.43	0.57	1.5	0.54
Ljutnja	0 (0-0.05)	0 (0-0.13)	0 (0-0.045)	0.2926	2	0.864	0.456	0.944	-0.47	0.941	-0.673	0.883
Tuga	0 (0-0)	0 (0-0.225)	0 (0-0.408)	2.6381	2	0.267	-2.11	0.296	-1.16	0.692	1.44	0.566

Društveni procesi	14.3 (11.6-15.8)	14.5 (12.5-16.4)	12 (11.1-13.9)	5.3796	2	0.068	2.079	0.306	-0.88	0.808	-3.305	0.051
Društveno ponašanje	6.02 (4.64-8.59)	5.61 (4.88-6.42)	6.54 (5.41-7.49)	1.0645	2	0.587	0.213	0.988	0.853	0.819	1.652	0.472
Prosocijalno ponašanje	1.51 (0.495-2.24)	1.35 (0.895-1.65)	1.36 (1.06-1.66)	0.0378	2	0.981	-0.08	0.998	0.1866	0.99	0.2665	0.981
Uljudnost	0.1 (0-0.255)	0 (0-0.135)	0 (0-0.203)	3.1738	2	0.205	1.858	0.388	2.308	0.232	0.519	0.929
Međuljudski sukob	0.215 (0-0.547)	0.185 (0-0.387)	0.205 (0-0.265)	0.6647	2	0.717	1.057	0.735	0.925	0.79	0.163	0.993
Moraliziranje	1.53 (1.09-1.79)	1.1 (0.848-1.3)	1.5 (1.26-1.88)	5.6028	2	0.061	-0.426	0.951	2.532	0.173	3.172	0.064
Komunikacija	1.72 (0.942-2.11)	1.46 (0.99-2.02)	1.34 (0.98-1.64)	1.1153	2	0.573	1.5464	0.519	0.0267	1	1.0127	0.754

*Kruskal-Wallisov test s Dwass-Steel-Critchlow-Fligner usporedbama u paru.

Te su razlike ponovno bile prisutne kada smo usporedili „prave“ studentske eseje s UI-generiranima (Slika 1 i Tablica 8).

Slika 1. Usporedba "istinitih" eseja koje su napisali studenti s esejima generiranim umjetnom inteligencijom s punim ključnim riječima*



Tablica 8. Usporedba "istinitih" eseja koje su napisali studenti s esejima generiranim umjetnom inteligencijom s punim ključnim riječima*

	LIWC rezultati, medijan (IQR)			Sveukupno			ChatGPT generirani vs. istinitih eseja koje su napisali studenti		Bard generirani vs. istinitih eseja koje su napisali studenti		Bard generirani vs. ChatGPT-generirani eseja	
	Eseji koje su napisali studenti (n = 31)	Ekvivalentni Bard-generirani eseji (n = 31)	Ekvivalentni ChatGPT-generirani eseji (n = 31)	χ^2	df	P	W	P	W	P	W	P
Analitičko razmišljanje	67.4 (54-74.8)	52 (45.9-58)	90.5 (83.5-93.2)	49.54 1	2	<0.001	-7.36	<0.001	3.61	0.029	9.05	<0.001
Autoritet	34.5 (26.5-48)	40.8 (23.2-57.4)	38.9 (21.7-48)	0.601	2	0.741	0.119	0.996	-0.806	0.836	-1.065	0.732
Autentičnost	38.6 (29.2-54.3)	33.7 (10.8-43.1)	42.4 (33.9-58.5)	7.561	2	0.023	-0.727	0.865	2.997	0.086	3.614	0.029
Ton	20.2 (11.6-32.8)	31.9 (16.6-46.1)	35 (22.8-47.9)	10.30 5	2	0.006	-4.61	0.003	-2.79	0.119	1.34	0.609
Nagon	4.08 (3.29-4.9)	4.63 (3.56-5.89)	5.65 (4.46-6.97)	12.33 6	2	0.002	-4.96	0.001	-1.93	0.359	2.92	0.098
Pripadnost	1.53 (0.845-2.3)	1.54 (0.87-2.48)	2.23 (1.83-3.22)	12.20 5	2	0.002	-4.569	0.004	-0.418	0.953	3.922	0.015
Postignuće	1.13 (0.83-1.61)	1.15 (0.635-1.65)	1.66 (1.21-2.06)	9.038	2	0.011	-3.754	0.022	0.737	0.861	3.554	0.032
Moć	1.19 (1-1.67)	1.68 (1.23-2.4)	1.71 (1.11-2.17)	6.867	2	0.032	-2.987	0.087	-3.405	0.042	-0.11	0.997
Kognicija	15.6 (13.5-17.4)	14.4 (11.6-15.2)	11.1 (9.76-12.1)	35.95	2	<0.001	7.89	<0.001	4.15	0.009	-5.35	<0.001
Sve ili ništa	0.88 (0.67-1.21)	0.7 (0.535-1.05)	0.39 (0.21-0.43)	33.93 6	2	<0.001	7.63	<0.001	1.92	0.363	-6.29	<0.001
Kognitivni procesi	14.4 (12.9-16)	13.3 (11-14.2)	10.6 (9.38-11.5)	30.91 7	2	<0.001	7.5	<0.001	3.83	0.018	-4.66	0.003

Uvid	3.44 (2.54-3.81)	4.66 (3.83-4.89)	3.9 (3.15-4.38)	17.09 6	2	<0.001	-3.04	0.081	-5.61	<0.001	-3.25	0.057
Uzročnost	1.66 (1.29-2.41)	1.42 (1.01-1.94)	1.25 (0.785-1.63)	8.449	2	0.015	3.89	0.016	2.65	0.147	-1.77	0.422
Razlika	1.81 (1.36-2.63)	2.02 (1.21-2.45)	1.23 (1.02-1.67)	14.92 7	2	<0.001	4.779	0.002	0.0398	1	-4.6599	0.003
Pokušaj	2.28 (1.69-3.35)	1.62 (1.15-1.89)	1.37 (0.89-1.65)	23.57 9	2	<0.001	6.27	<0.001	5.22	<0.001	-2.03	0.322
Izvjescnost	0.59 (0.315-1.11)	0.17 (0-0.375)	0.21 (0-0.405)	21.79 3	2	<0.001	5.67	<0.001	5.59	<0.001	1.36	0.601
Neslaganje	4.42 (3.67-5.19)	3.01 (2.65-3.75)	1.86 (1.61-2.35)	52.85 4	2	<0.001	8.88	<0.001	6.56	<0.001	-6.44	<0.001
Afekt	3.94 (3.07-4.92)	5.54 (4.52-6.59)	5.7 (4.75-7.21)	17.96 9	2	<0.001	-4.888	0.002	-5.466	<0.001	0.468	0.942
Positive tone	2 (1.5-2.3)	2.93 (2.1-3.66)	3.21 (2.47-3.59)	22.97	2	<0.001	-6.223	<0.001	-5.416	<0.001	0.707	0.872
Negativan ton	2 (1.21-2.81)	2.02 (1.64-2.89)	1.96 (1.42-2.53)	1.541	2	0.463	-0.289	0.977	-1.493	0.542	-1.523	0.529
Osjećaj	1.41 (0.86-2.23)	1.53 (0.975-2.12)	2.14 (1.44-3.23)	9.254	2	0.01	-3.873	0.017	-0.448	0.946	3.524	0.034
Pozitivna emocija	0.4 (0.155-0.565)	0.35 (0.17-0.695)	0.43 (0.28-0.82)	3.889	2	0.143	-2.2169	0.26	0.00999	1	2.59492	0.158
Negativne emocije	0.89 (0.435-1.39)	0.88 (0.47-1.42)	1.22 (0.53-1.69)	1.629	2	0.443	-1.733	0.438	-0.428	0.951	1.294	0.631
Tjeskoba	0 (0-0.205)	0.17 (0-0.34)	0.21 (0.095-0.495)	12.12 7	2	0.002	-4.62	0.003	-3.17	0.064	2.41	0.205
Ljutnja	0.1 (0-0.195)	0.13 (0-0.21)	0 (0-0.205)	1.299	2	0.522	1.319	0.62	-0.285	0.978	-1.473	0.551
Tuga	0 (0-0.195)	0 (0-0.195)	0.2 (0-0.295)	6.851	2	0.033	-3.6	0.029	-1.68	0.462	2.16	0.28

Društveni procesi	12.1 (9.96-13.9)	14.5 (12.9-17.7)	12.2 (10.3-14.4)	13.71 9	2	0.001	-0.627	0.897	-4.689	0.003	-4.311	0.007
Društveno ponašanje	4.4 (3.77-5.39)	6.17 (5.36-6.87)	7.14 (5.73-8.04)	39.53	2	<0.001	-7.83	<0.001	-6.94	<0.001	3.23	0.059
Prosocijalno ponašanje	0.44 (0.24-1.01)	1.05 (0.845-1.83)	1.64 (1.06-2.42)	30.27	2	<0.001	-7.34	<0.001	-5.27	<0.001	3.01	0.085
Uljudnost	0 (0-0.21)	0 (0-0.18)	0 (0-0.2)	0.118	2	0.943	0.483	0.938	-0.229	0.986	-0.341	0.969
Međuljudski sukob	0.2 (0.075-0.495)	0.42 (0.195-0.515)	0.21 (0.185-0.595)	2.061	2	0.357	-1.214	0.667	-2.001	0.333	-0.82	0.831
Moraliziranje	0.52 (0.28-0.985)	0.88 (0.67-1.31)	1.88 (1.58-2.21)	45.77 7	2	<0.001	-8.7	<0.001	-3.43	0.041	7.09	<0.001
Komunikacija	1.77 (1.21-2.4)	1.93 (1.34-2.67)	1.28 (0.61-1.65)	13.43 4	2	0.001	3.81	0.019	-1.52	0.529	-4.83	0.002

*Kruskal-Wallisov test s Dwass-Steel-Critchlow-Fligner usporedbama u paru.

4.2 Povezanost korištenja ChatGPT-a i pristupa esejskom zadatku na nastavi iz medicinske humanistike

Ukupno je 85 studenata bilo upisano u dva kolegija Medicinske humanistike. Troje studenata 5. godine i jedan student 6. godine nisu prisustvovali nastavi kada je provedeno istraživanje zbog bolesti, tako da je ukupno 81 student prisustvovao sesiji u kojoj je provedena intervencija. Svi su pristali sudjelovati u našem istraživanju.

Prikupili smo 86 odgovora, od kojih smo isključili pet duplikata koji su nastali ili zato što su studenti započeli ispunjavanje upitnika, prekinuli ga i zatim ponovno ispunili u novom unosu, ili zato što su slučajno odabrali ime pogrešnog sakupljača podataka iz druge skupine eseja, pa izašli i zatim ispunili ispravan upitnik. Također smo iz analize isključili jedan nepotpuno ispunjen upitnik, što znači da smo na kraju imali 80 valjanih odgovora, 42 iz skupine koji su generirali eseje ChatGPT-om i 38 iz skupine koji su eseje pisali samostalno, što daje stopu dovršenosti od 94,1 % (Tablica 9).

Tablica 9. Karakteristike ispitanika

	Grupa eseja pisanih samostalno (n = 38)	Grupa eseja generirana ChatGPT- om (n = 42)	Svi (n = 80)	<i>P</i>
Spol, n (%)				
Muški	17 (44.7)	12 (28.6)	29 (36.3)	0.133*
Ženski	21 (55.3)	30 (71.4)	51 (63.7)	
Godina studija, n (%)				
5 th	18 (47.4)	23 (54.8)	41 (51.2)	0.509*
6 th	20 (52.6)	19 (45.2)	39 (48.8)	
Dob u godinama, MD (IQR)	26.0 (24.0, 27.8)	25.0 (24.0, 28.0)	25.0 (24.0, 28.0)	0.922†
Intrinzična motivacija na početku (raspon: 15.0–75.0), MD (IOR)	58.5 (52.0, 60.8)	54.0 (51.0, 57.8)	55.0 (51.0, 60.0)	0.043†

IQR – interkvartilni raspon, MD – medijan *Hi-kvadrat test asocijacije. †Mann-Whitney U-test.

Studenti su u najvećem broju izvijestili da ChatGPT koriste povremeno (37,5 %) ili često (30,0 %). Nijedan student iz nijedne skupine nije naveo da mu je zadatak bio vrlo težak. Nešto više od polovice studenata izjavilo je da su pročitali samo dijelove članka koji su bili relevantni za njihov esej. Njihov stav prema umjetnoj inteligenciji uglavnom se nije promijenio prije i nakon intervencije ($P = 0,310$), s medijanom pozitivne promjene u ATTARI-12 rezultatu od 1 boda (raspon ljestvice: 12–60).

Iako su studenti općenito izrazili preferenciju za korištenje ChatGPT-a pri pisanju eseja, veći udio studenata u skupini onih koji su eseje pisali samostalno je želio pisati samostalno (31,6 % prema 9,5 %, $P < 0,001$). Većina studenata koja je generirala eseje ChatGPT-om smatrala je zadatak lakim (42,9 %), dok su oni koji su eseje pisali samostalno uglavnom ocijenili da zadatak nije bio ni lak ni težak (44,7 %). Studenti koji su generirali eseje ChatGPT-om također su češće izjavili da uopće nisu pročitali članak ili su ga samo površno preletjeli, u usporedbi sa skupinom koja je pisala samostalno, gdje je gotovo tri četvrtine studenata izjavilo da su pročitali dijelove članka relevantne za svoj esej ($P < 0,001$) (Tablica 10).

Tablica 10. Rezultati vezani uz zadatak pisanja eseja

	Svi (n = 80)	Grupa eseja pisanih samostalno (n = 38)	Grupa eseja generirana ChatGPT-om (n = 42)	P
Prijašnja razina korištenja ChatGPT-a, n (%)				0.357*
Nikada	5 (6.3)	3 (7.9)	2 (4.8)	
Rijetko	21 (26.3)	13 (34.2)	8 (19.0)	
Ponekad	30 (37.5)	13 (34.2)	17 (40.5)	
Često	24(30.0)	9 (23.7)	15 (35.7)	
Percepcija težine zadatka, n (%)				0.061*
Vrlo lagano	19 (23.8)	5 (13.2)	14 (33.3)	
Lagano	32 (40.0)	14 (36.8)	18 (42.9)	
Ni lagano ni teško	26 (32.5)	17 (44.7)	9 (21.4)	
Teško	3 (3.8)	2 (5.3)	1 (2.4)	
Vrlo teško	0 (0.0)	0 (0.0)	0 (0.0)	
Preferirani način pisanja eseja, n (%)				<0.001*
ChatGPT	51 (63.7)	26 (68.4)	38 (90.5)	
Pisanje samostalno	29 (36.3)	12 (31.6)	4 (9.5)	
U kojoj je mjeri učenik pročitao dodijeljeni rad, n (%)				<0.001*
Uopće nije pročitao rad	14 (17.5)	0 (0.0)	14 (33.3)	
Površno je prelistao rad	19 (23.8)	5 (13.2)	14 (33.3)	
Pročitao dijelove rada relevantne za esej	42 (52.5)	28 (73.7)	14 (33.3)	
Read whole paper	5 (6.3)	5 (13.2)	0 (0.0)	
ATTARI-12 rezultati (raspon rezultata 12,0–60,0), medijan (interkvartilni raspon)				

Polazna vrijednost	39.0 (36.0, 43.0)	38.0 (35.0, 41.0)	40.0 (37.0, 43.8)	0.059†
Nakon pisanja eseja	40.0 (36.0, 44.0)	39.0 (35.0, 44.0)	41.0 (38.0, 44.0)	0.347†
Promjena prije i poslije pisanja eseja	1.0 (-1.0, 2.25)	1.0 (-1.0, 3.0)	0.0 (-1.0, 2.0)	0.310†

*Hi-kvadrat test povezanosti.

†Mann-Whitney U-test.

Naša analiza linearnom regresijom nije pronašla povezanost između promjena u stavovima prema UI-ju prije i nakon eseja i vrste kreiranja eseja (samostalno ili ChatGPT), dobi, spola, godine studija, prethodne razine korištenja ChatGPT-a, percipirane težine zadatka, opsega u kojem je student pročitao članak te preferiranog načina pisanja eseja (Tablica 11).

Tablica 11. Rezultati analize linearne regresije prediktora promjene stavova studenata prema umjetnoj inteligenciji prije i nakon pisanja eseja. N = 80.

Varijable ishoda	Kovarijanta	B	P	Prilagođen R²*
Promjena stavova o umjetnoj inteligenciji prije i poslije pisanja eseja	Grupa	-0.408	0.201	0.083
	Dob	0.098	0.428	
	Spol	0.093	0.449	
	Godina studija	-0.028	0.816	
	Intrizična motivacija	-0.042	0.741	
	Prethodna razina korištenja ChatGPT-a	-0.156	0.216	
	Percipirana težina zadatka	0.117	0.359	
	U kojoj je mjeri učenik pročitao rad	-0.032	0.837	
	Preferirani modalitet pisanja eseja	-0.311	0.311	

Kratice: β =Standardizirani regresijski koeficijent, R²=Prilagođeni koeficijent determinacije

Ipak, pronašli smo umjerenu povezanost između intrinzične motivacije i opsega u kojem su studenti pročitali članak ($r = 0,339$, $P = 0,002$), kao i s preferiranim načinom pisanja eseja ($r = 0,262$, $P = 0,019$). Kada smo ovu povezanost dodatno istražili primjenom gore spomenute linearne regresije, utvrdili smo da je jedino opseg u kojem su studenti pročitali

članak bio marginalno značajan prediktor rezultata intrinzične motivacije ($P = 0,049$). Nijedna druga kovarijabla nije bila statistički značajna u ovoj analizi (Tablica 12).

Tablica 12. Rezultati analize linearne regresije prediktora intrinzične motivacije studenata. $N = 80$.

Varijable ishoda	Kovarijanta	B	P	Prilagođen R^2 *
Intrizička motivacija	Grupa	0.265	0.384	0.175
	Dob	0.174	0.137	
	Spol	0.052	0.655	
	Godina studija	0.012	0.920	
	Promjena stavova o umjetnoj inteligenciji prije i poslije pisanja eseja	-0.036	0.741	
	Prethodna razina korištenja ChatGPT-a	0.010	0.931	
	Percipirana težina zadatka	-0.064	0.599	
	U kojoj je mjeri učenik pročitao rad	0.282	0.049	
	Preferirani modalitet pisanja eseja	0.379	0.191	

Kratice: β = standardizirani regresijski koeficijent, R^2 = prilagođeni koeficijent determinacije.

5. RASPRAVA

5.1 Lingvistička analiza i usporedba eseja studenata medicine i eseja generiranih umjetnom inteligencijom

5.1.1 Sažetak glavnih nalaza istraživanja

U ovom istraživanju usporedili smo psihometrijska jezična obilježja između eseja koje su studenti sami napisali o osobnim etičkim dilemama i eseja koje je generirala UI koristeći upute izrađene na temelju ključnih riječi iz originalnih eseja. Utvrdili smo da su eseji generirani UI obično sadržavali više riječi povezanih s emocijama, osobito onima koje izražavaju pozitivne emocije, te su pokazivali višu razinu analitičkog razmišljanja i korištenje tzv. "velikih riječi". S druge strane, eseji koje su napisali studenti sadržavali su više jezičnih izraza povezanih s kognicijom te više riječi po rečenici, tj. duljim rečenicama, što može ukazivati na veći stupanj misaonog procesa tijekom pisanja te spontano i slobodnije izražavanje.

Koliko nam je poznato, naša studija je prva koja je kvantitativno uspoređuje psihometrijska obilježja ljudskih i eseja generiranih UI na temelju osobnih iskustava medicinskih studenata s realnim etičkim dilemama u obrazovnom ili profesionalnom kontekstu. Druge studije koje su uspoređivale ljudski i UI-generirani tekst uglavnom su se fokusirale na argumentativne ili reflektivne tekstove o raznim temama (100, 103) ili na sposobnost LLM-ova da pišu u određenim stilovima osobnosti, bez obzira na konkretnu temu (126).

Dva LLM-a koja su korištena za generiranje eseja, Bard i ChatGPT, također su se međusobno razlikovala: ChatGPT eseji su pokazivali viši stupanj autentičnosti, analitičkog izraza i upotrebe kompleksnih riječi, dok su Bard eseji imali viši rezultat u "Društvenim procesima", što bi moglo upućivati na drugačiji stil generiranja sadržaja kod ova dva LLM-a. To sugerira da nije samo važna činjenica koristi li se umjetna inteligencija, već i koji model se koristi, budući da svaki ima svoj vlastiti jezični profil i stil.

Međutim, otkrili smo da je trećina eseja koje su studenti predali već bila u cijelosti ili djelomično napisana uz pomoć umjetne inteligencije. To je potvrđeno dodatnim analizama, u kojima su razlike između tih eseja i eseja potpuno generiranih UI bile manje izražene nego razlike između "pravih" studentskih eseja i onih koje je generirala umjetna inteligencija. Ta skupina eseja pokazala je sličnosti s esejima generiranim u potpunosti UI u nizu psiholingvističkih kategorija. Kao i u istraživanju Weidinger et al. ti su eseji imali više emocionalnih izraza i analitičkog jezika nego studentski, ali su bili manje autentični i nisu

pokazivali statistički značajne razlike u tonu (103). To upućuje na to da su neki studenti koristili UI kao alat, bilo u pisanju nacrtu, strukturiranju misli ili završnom uređivanju. Iako su ovi eseji i dalje u određenoj mjeri odstupali od "pravih" studentskih eseja, njihova bliskost generiranima UI pokazuje koliko je teško jasno razgraničiti ljudsko i strojno autorstvo, osobito kada se UI koristi kao podrška umjesto kompletne zamjene. Također, budući da su suautorski eseji imali nižu autentičnost, ali više rezultate u analitičkom razmišljanju i češću upotrebu "velikih riječi", vjerojatno su doista bili generirani umjetnom inteligencijom, ali su ih studenti naknadno uređivali.

Subanaliza eseja generiranih temeljem djelomičnih i potpunih ključnih riječi pokazala je minimalne razlike kod ChatGPT eseja, ali određene statistički značajne razlike kod eseja generiranih Bardom. Te su razlike, međutim, nestale nakon korekcije za višestruke usporedbe, što sugerira da su mogle biti rezultat slučajnosti. To ukazuje na to da je Bard osjetljiviji na promjene u uputama, odnosno da struktura i kvaliteta ulaznih podataka imaju veći utjecaj na krajnji tekst.

Potvrdili smo da eseji koje je generirala UI sadrže više izraza povezanih s emocijama, autentičnošću i analitičkim razmišljanjem u usporedbi s esejima koje su studenti sami napisali, nakon što su iz analize uklonjeni suautorski UI-eseji. Rezultati da eseji napisani UI koriste više analitičkog jezika nije iznenađujući, budući da su Herbold i suradnici prethodno otkrili da eseji napisani pomoću ChatGPT-a nadmašuju one koje pišu ljudi u smislu korištenja akademski poželjnije strukture i naracije (120). Nadalje, Jiang i suradnici pokazali su da ChatGPT može uspješno oponašati određene osobine ličnosti unutar modela „Velikih pet“, što može djelomično objasniti zašto su LLM-ovi bili vrlo kompetentni u pisanju o nijansiranim, emocionalnim temama, kao što je bio slučaj u našoj studiji (126). Iako nismo provodili kvalitativnu ni kvantitativnu analizu tog aspekta, primijetili smo da eseji generirani UI često imaju šablonsku strukturu, vjerojatno pod utjecajem uputa koji upućuje model da napiše "esej". To uključuje, među ostalim, upotrebu ustaljenih fraza poput "Tijekom studija medicine" za početak eseja i "Zaključno" za kraj. Takav šablonski jezik, ako ga čitatelj eseja susreće više puta, mogao bi upućivati na upotrebu alata UI.

S druge strane, rezultati da eseji napisani pomoću UI teže emocionalno pozitivnijem i autentičnijem jeziku u odnosu na one koje su napisali studenti pomalo je neočekivan. To se može protumačiti rezultatima istraživanja Ovsyannikove i suradnika u kojem je internetska usluga za emocionalnu podršku, temeljena na GPT-3, bila percipirana kao emocionalno

podržavajuća više od odgovora koje su davali ljudi. Međutim, kada su sudionici saznali da odgovori nisu generirani od strane ljudi, taj učinak je nestao (127). Pretpostavljamo (iako to ne možemo sa sigurnošću potvrditi) da su se studenti možda suzdržavali u izražavanju svojih emocija i stavova zbog etičke dvosmislenosti stvarnih situacija koje su opisivali, što je rezultiralo nižim rezultatima za autentičnost i emocionalni ton. Nadalje, eseji koje je pisala UI prikazivali su idealizirani "najbolji scenarij" i moralno poželjnije odgovore, što je razlog zašto su imali pozitivniji ton i nudili općenito pozitivnija rješenja etičkih dilema. Iako to može djelovati kontraintuitivno, može se objasniti ograničenjima ljudskog pisanja u stvarnim, etički dvosmislenim kontekstima, gdje studenti možda oklijevaju zauzeti čvrst stav ili izraziti osobne emocije. Modeli UI, s druge strane, nude idealizirane, jasno strukturirane i moralno poželjne odgovore, što se očituje u višim rezultatima za ton i afektivni izraz.

I dok Sandler i suradnici nisu pronašli razliku u izraženosti emocija između ljudskog i UI-generiranog dijaloga, našli su da je potonji imao pozitivniji emocionalni ton i više riječi povezanih sa socijalnim procesima, što podupire naše rezultate (128). Ovaj pozitivni ton u UI-esejima sugerira da su prikazivali idealizirane "najbolje scenarije" s pozitivnijim ishodima etičkih dilema.

Naši rezultati pokazuju da su alati temeljeni na umjetnoj inteligenciji vrlo učinkoviti u stvaranju eseja nalik onima koje pišu ljudi, osobito kada je riječ o etici, osobnim iskustvima i stavovima. Uočen psihometrijski razmak između UI-eseja i onih koje su pisali studenti postao je znatno suptilniji u podanalizama suautorskih eseja, što sugerira da su studenti te eseje vjerojatno mijenjali nakon što su ih generirali LLM-ovi.

5.1.2 Snaga i ograničenja istraživanja

Glavna snaga našeg istraživanja leži u korištenju LIWC softvera, koji je prethodno višestruko validiran u znanstvenim studijama, što nam je omogućilo pouzdanu kvantitativnu usporedbu različitih vrsta eseja. Uz to, prikupili smo velik broj eseja s visokom stopom odziva studenata, svi su eseji napisani u stvarnom obrazovnom kontekstu, a ne preuzeti iz gotovih baza podataka ili internetskih izvora. Time naši nalazi ne predstavljaju teorijske pretpostavke već stvarne obrasce pisanja i potencijalnu upotrebu UI-ja alata u praksi. Posebnu vrijednost donosi i usporedba različitih izvora eseja (studentskih, potpuno UI-generiranih te tzv. hibridnih eseja), što nam je omogućilo slojevitije razumijevanje načina na koji studenti koriste velike jezične modele u akademskom pisanju.

Važno je naglasiti i određena ograničenja našeg istraživanja. Studenti koji su sudjelovali u istraživanju nisu izvorni govornici engleskog jezika, no pohađali su program u cijelosti na engleskom jeziku, što implicira zadovoljavajuću jezičnu kompetenciju. Također, koristili smo verzije modela ChatGPT 3.5 i Bard, iako su u međuvremenu postale dostupne i novije, naprednije inačice poput ChatGPT 4.0. Odluka da koristimo upravo te verzije temeljila se na njihovoj besplatnoj dostupnosti studentima, budući da smo smatrali da bi upotreba plaćenih verzija mogla umanjiti vanjsku valjanost rezultata, s obzirom na to da većina studenata vjerojatno ne bi koristila modele koji zahtijevaju pretplatu.

Osim toga, iako smo pomoću UI detekcijskih alata uspjeli identificirati dio eseja koji su djelomično ili u cijelosti generirani umjetnom inteligencijom, ne možemo sa sigurnošću isključiti mogućnost postojanja lažno pozitivnih ili lažno negativnih rezultata. Također, ostaje pitanje jesu li neki studenti u dovoljnoj mjeri prilagodili UI-generirane eseje kako bi izbjegli detekciju. Kako bismo ublažili ove potencijalne nedostatke, koristili smo dva detekcijska alata te dodatno ljudsku validaciju.

5.1.3 Implikacije rezultata istraživanja i prijedlozi za daljnja istraživanja

Nalazi ove studije otvaraju niz smjerova za buduća istraživanja o ulozi UI-ja u akademskom pisanju i obrazovanju. Prije svega, preporučuje se provođenje kvalitativnih analiza sadržaja eseja, poput tematske ili narativne analize, koje bi omogućile dublje razumijevanje načina na koji studenti izražavaju osobna iskustva, moralne dileme i stavove, za razliku od standardiziranih i često idealiziranih odgovora koje proizvode alati UI.

Istraživanje učinaka edukacije o etici korištenja UI pokazuje se kao izuzetno značajno. Kroz edukativne intervencije moglo bi se pratiti na koji način informiranost o akademskoj čestitosti i pravilima korištenja UI utječe na stavove, ponašanje i razinu angažiranosti studenata. Nadalje, korisno bi bilo proširiti istraživanje na različite akademske discipline i razine obrazovanja, kako bi se ispitalo postoje li razlike u načinu korištenja UI, stupnju prihvaćanja tehnologije i percepciji autentičnosti rada među studentima prve i završne godine studija, kao i među strukama koje se razlikuju po prirodi pismenog izražavanja.

Jednako važan smjer budućih istraživanja uključuje longitudinalno praćenje promjena u stavovima, motivaciji i navikama korištenja UI alata tijekom duljeg obrazovnog perioda. Takve bi studije mogle pokazati kako se s vremenom razvijaju studentske vještine pisanja i njihovo kritičko promišljanje, osobito u kontekstu stalnog napretka tehnologije. Dodatno,

korisno bi bilo istražiti mogućnost korištenja UI kao alata za razvoj pisanih vještina – primjerice kao podrške studentima kojima engleski nije materinji jezik ili onima s nižom motivacijom, pod uvjetom da se upotreba UI odvija na jasan i pedagoški opravdan način.

5.2 Povezanost korištenja ChatGPT-a i pristupa esejskom zadatku na nastavi iz medicinske humanistike

5.2.1. Sažetak glavnih nalaza istraživanja

Rezultati ovog istraživanja nude uvid u načine na koje studenti percipiraju i koriste umjetnu inteligenciju, konkretno ChatGPT, u obrazovnom kontekstu te kako te odluke utječu na njihov angažman, motivaciju i odnos prema zadatku. Uz kvantitativne podatke, otkriveni su i obrasci ponašanja koji mogu imati implikacije za oblikovanje obrazovne politike i didaktičke prakse u visokom obrazovanju.

Jedan od važnijih nalaza odnosi se na razliku u razini intrinzične motivacije između studenata koji su pisali eseje bez korištenja umjetne inteligencije i onih koji su koristili ChatGPT. Studenti u koji su samostalno pisali eseje pokazali su višu razinu intrinzične motivacije u usporedbi s kolegama koji su eseje kreirali ChatGPT-om. Ovaj rezultat može se tumačiti kroz teoriju o motivaciji, prema kojima autonomni angažman u zadatku potiče osjećaj kompetencije i unutarnje motivacije (107, 108). Mogućnost oslanjanja na vanjski alat kao što je ChatGPT, iako olakšava tehničku izvedbu zadatka, možda smanjuje osjećaj osobne uključenosti i procesom pisanja kao osobnom djelu.

Zanimljiv je i nalaz da su studenti koji su koristili ChatGPT za kreiranje eseja zadatak češće percipirali kao lak, dok su studenti koji su samostalno pisali eseje izražavali stav da zadatak nije bio ni težak ni lagan. Iako se razlike nisu pokazale statistički značajnima, one ipak ukazuju na moguću razliku u doživljaju kognitivnog opterećenja. Iako se olakšan pristup pisanju može činiti korisnim, dugoročno može rezultirati smanjenjem kognitivne angažiranosti, što potvrđuju nalazi istraživanja Borji et al. koja upozoravaju na "kognitivno oslanjanje" na UI alate (131). Olakšavanje zadatka može tako smanjiti prilike za razvoj važnih akademskih vještina kao što su kritičko mišljenje, analiza i sinteza informacija.

Značajnu razliku nalazimo u načinu na koji su studenti pristupili znanstvenom članku koji je bio zadana literatura za esej. Dok je većina studenata koji su samostalno pisali eseje pročitala relevantne dijelove članka, čak trećina onih koji su koristili ChatGPT nije pročitala članak uopće, a dodatna trećina ga je samo površno pregledala. Ovi rezultati upućuju na potencijalni negativni efekt upotrebe UI alata na razinu angažmana s izvornim znanstvenim

materijalima. Ovaj nalaz potvrđuje zabrinutosti izražene u literaturi, gdje se ističe da korištenje generativnih UI sustava može smanjiti duboko učenje i motivaciju za samostalnu analizu izvora (132). Ako studenti percipiraju da im je UI sposoban pružiti „gotov proizvod“, moguće je da će zanemariti ključne korake akademskog procesa poput čitanja, kritičke analize i sinteze informacija.

Važno je naglasiti da, unatoč izraženoj općoj preferenciji studenata za korištenje ChatGPT-a pri pisanju eseja, značajan broj studenata koji su samostalno pisali eseje (31,6 %) i dalje preferira pisanje eseja samostalno. U skladu s rezultatima Chan i Ho ovo može ukazivati na prisutnu svijest o važnosti autentičnosti u akademskom izražavanju, ali i na potencijalni sukob između koristi koje nudi tehnologija i osobnih uvjerenja ili vrijednosti koje studenti vezuju uz vlastiti rad (133).

S obzirom na rezultate ATTARI-12 skale, nije zabilježena značajna promjena u stavovima studenata prema umjetnoj inteligenciji nakon intervencije, što može upućivati na to da jednokratni zadatak nije dovoljan da bi se oblikovali ili promijenili dublji stavovi prema novim tehnologijama. Potrebni su sveobuhvatniji obrazovni pristupi, uključujući rasprave, primjere dobre prakse i etičke dileme, kako bi se studenti mogli kritički suočiti s kompleksnošću upotrebe UI u obrazovanju.

Dodatne regresijske analize nisu pokazale značajne prediktore promjene stavova prema UI-ju, no identificirana je umjerena povezanost između intrinzične motivacije i razine angažmana s pročitanim materijalom. Studenti koji su temeljitije pročitali znanstveni članak pokazali su višu motivaciju, što potvrđuje tezu da dubinski angažman u obrazovnom sadržaju pozitivno utječe na unutarnju motivaciju (130,134). Iako se preferirani način pisanja eseja također pokazao povezan s motivacijom, ta povezanost nije bila statistički značajna nakon uključivanja u višestruku regresijsku analizu.

Općenito, rezultati ukazuju na to da upotreba alata kao što je ChatGPT mijenja način na koji studenti pristupaju akademskom pisanju i učenju, olakšavajući tehničku izvedbu zadatka, ali potencijalno smanjujući dubinsku angažiranost i intrinzičnu motivaciju. To potvrđuje teze o nužnosti pažljivo osmišljenih obrazovnih strategija koje će studentima omogućiti da koriste prednosti umjetne inteligencije, ali bez gubitka ključnih akademskih vještina i osobne uključenosti.

Nadalje, ovi nalazi podupiru tvrdnje autora kao što su Mollick i Mollick (2023), koji upozoravaju da će obrazovni sustavi trebati redefinirati zadatke i metode ocjenjivanja kako bi

sačuvali autentičnost učenja u doba sveprisutne umjetne inteligencije (135). U konačnici, rezultati ovog istraživanja naglašavaju potrebu za kritičkim, strukturiranim i pedagoški promišljenim integriranjem UI alata u obrazovne kontekste.

5.2.2 Snaga i ograničenja istraživanja

Jedna od glavnih snaga ove studije jest visok odaziv ispitanika, što značajno doprinosi pouzdanosti i reprezentativnosti prikupljenih podataka. Također, studija je provedena u stvarnim nastavnim uvjetima unutar dva kolegija iz medicinske humanistike, čime se povećava valjanost i omogućuje uvid u stvarno ponašanje studenata, a ne samo u deklarativne stavove. Podjela ispitanika u dvije jasno definirane skupine – one koja koristi ChatGPT i one koja piše eseje samostalno – omogućila je izravnu usporedbu obrazaca ponašanja i preferencija. Dodatno, uključivanje psihometrijskih instrumenata poput ATTARI-12 i skale intrinzične motivacije pružilo je dublji uvid u stavove i motivacijske mehanizme studenata, što je od osobite važnosti u kontekstu sve češće primjene alata UI u obrazovanju.

Unatoč tim prednostima, studija ima i nekoliko važnih ograničenja. Prvo, studenti koji su sudjelovali u istraživanju nisu bili izvorni govornici engleskog jezika, iako su pohađali program na engleskom jeziku. To bi moglo utjecati na način interpretacije teksta, izražavanja u eseju i razumijevanja zadatka, što se može reflektirati i na rezultate povezanih mjerenja.

Nadalje, korištena verzija umjetne inteligencije – ChatGPT 3.5 nije najnovija dostupna verzija LLM-a, što ograničava generalizaciju rezultata na novije modele koji imaju naprednije jezične mogućnosti. Izbor tih verzija bio je opravdan njihovom dostupnošću studentima bez financijskog opterećenja, ali može predstavljati tehničko ograničenje u pogledu relevantnosti nalaza za najnovije tehnologije.

Dodatno, iako je korišten nadzor prilikom samostalno pisanja eseja i alata za detekciju UI-generiranog teksta te njihova evaluacija potvrđena od strane ljudskih ocjenjivača, mogućnost postojanja lažno pozitivnih ili negativnih detekcija ne može se isključiti. Nije moguće sa sigurnošću tvrditi da studenti nisu dodatno uređivali UI-generirane eseje na način koji bi ih učinio manje prepoznatljivima alatima za detekciju, što potencijalno utječe na valjanost analize korelata stavova i motivacije.

Budući da je studija provedena u okviru jednog sveučilišta, rezultati su primjenjivi prvenstveno na slične akademske kontekste i kulturne okvire, te ih treba s oprezom generalizirati na druge obrazovne sustave ili razine obrazovanja.

Nadalje, promjena stavova prema umjetnoj inteligenciji nije bila statistički značajna, što može ukazivati na to da jedan zadatak, čak i uz refleksiju, nije dovoljan da izazove dublje promjene u uvjerenjima – sugerirajući potrebu za dugoročnijim i višedimenzionalnim pristupom u obrazovnom oblikovanju stavova prema tehnologiji.

5.2.3 Implikacije rezultata istraživanja i prijedlozi za daljnja istraživanja

Dobiveni rezultati otvaraju niz važnih pitanja i praktičnih implikacija za oblikovanje obrazovne politike u kontekstu sve prisutnije integracije umjetne inteligencije u akademski proces. Uočena razlika u razini intrinzične motivacije između studenata koji su eseje pisali samostalno i onih koji su koristili ChatGPT može ukazivati na potencijalni negativni utjecaj prekomjernog oslanjanja na alate UI na osobni angažman, osjećaj postignuća i razvoj ključnih akademskih vještina. Također, činjenica da su studenti u UI-skupini u manjoj mjeri čitali zadanu literaturu otvara pitanje kako dostupnost automatiziranih alata mijenja pristup učenju i informiranju, što može dugoročno utjecati na sposobnost kritičkog razmišljanja i interpretacije izvora.

S obzirom na sveprisutnost alata poput ChatGPT-a, potrebno je jasno definirati granice njihove prihvatljive uporabe u obrazovnim zadacima, uz istovremenu izgradnju studentske svijesti o važnosti autentičnosti, etičnosti i intelektualnog integriteta. U tom smislu, rezultati istraživanja sugeriraju potrebu za razvojem novih strategija vrednovanja znanja koje bi mogle uključivati više usmenih ispita, zadataka pisanih pod nadzorom ili kombinaciju automatiziranih i tradicionalnih metoda kako bi se očuvala relevantnost i vjerodostojnost obrazovnog procesa.

Daljnja istraživanja trebala bi se usmjeriti na longitudinalno praćenje promjena stavova i ponašanja studenata kroz više semestara, kako bi se ispitao trajni učinak korištenja alata UI na akademski razvoj. Također, bilo bi korisno ispitati kako različite vrste zadataka (npr. argumentativni vs. refleksivni eseji) utječu na način upotrebe UI i koje pedagoške metode mogu potaknuti odgovorno korištenje tehnologije. Poseban fokus treba staviti na istraživanje strategija koje mogu povećati razinu intrinzične motivacije kod studenata u digitalnom okruženju, osobito kod zadataka koji uključuju asistenciju UI.

U konačnici, važno je razviti instrumente i alate koji ne samo da detektiraju UI-generirani sadržaj, već koji omogućuju procjenu stupnja ljudske intervencije i doprinosa. Time bi se omogućila pravednija evaluacija studentskog rada u eri u kojoj granice između ljudskog i strojno generiranog sadržaja postaju sve nejasnije.

6. ZAKLJUČCI

1. Eseji generirani s UI-jem razlikuju se od studentskih eseja na psiholingvističkoj razini. UI eseji koriste više emocionalnih izraza, analitičkog jezika i velikih riječi, dok eseji studenata imaju dulje rečenice i izraženiju kognitivnu obradu.
2. Stil i sadržaj eseja ovise o vrsti korištenog modela UI-ja. Eseji generirani ChatGPT-om bili su autentičniji i analitičniji, dok su oni napisani pomoću Barda naglašavali veću orijentaciju na društvene procese.
3. Korištenje UI-ja među studentima bilo je češće od očekivanog. Oko jedne trećine eseja bilo je u potpunosti ili djelomično napisano uz pomoć UI-ja, što otežava razlikovanje između ljudski pisanog i strojno generiranog sadržaja.
4. Primjena UI alata utječe na obrazovne procese. Studenti koji su koristili ChatGPT pokazali su manju intrinzičnu motivaciju i slabiji angažman s izvornim znanstvenim materijalima, što sugerira nižu razinu dubinske obrade zadatka.
5. Sklonost korištenju UI-ja ne isključuje svijest o važnosti autentičnosti. Značajan broj studenata i dalje preferira pisanje bez UI-ja, što ukazuje na prisutnost unutarnjih vrijednosnih dilema.
6. Iskustva kratkotrajne upotrebe UI-ja nisu dovoljna za promjenu stavova prema toj tehnologiji. Promjene stavova prema UI-ju nakon pisanja zadatka nisu bile značajne, što ukazuje da su potrebni dugoročniji i temeljitiji pristupi edukaciji studenata o UI-ju.
7. Tehnološki napredak i etičke dileme moraju se razmatrati istovremeno. Nužna je jasna edukacija o odgovornom korištenju UI-ja u akademskom kontekstu, uz razvijanje novih metoda ocjenjivanja studentskog rada koji će smanjiti ovisnost o UI alatima.
8. Buduća istraživanja na ovu temu trebala bi biti longitudinalna i multidisciplinarna. Fokus trebaju biti promjene u ponašanju ili stavovima tijekom vremena, u raznim disciplinama, uz dodatni razvoj alata za prepoznavanje i mjerenje ljudskog doprinosa u esejima.

7. SAŽETAK

Ciljevi: Cilj prvog istraživanja bio je provesti lingvističku analizu i usporedbu eseja studenata medicine s esejima generiranim pomoću UI-ja (ChatGPT i Bard) na temu etičke, profesionalne ili moralne dileme. Jezične varijable su analizirane uz pomoć LIWC softvera, a izvor eseja je ispitan pomoću softvera za prepoznavanje UI-ja. Hipoteze su pretpostavljale razlike u izražavanju emocija i kognitivnim procesima između studentskih i UI-generiranih eseja. U drugom istraživanju cilj je bio ispitati utjecaj korištenja ChatGPT-a na stavove studenata prema UI-ju, uz pretpostavku da će studenti koji sami pišu eseje kao i oni s višom intrinzičkom motivacijom pokazati snažniju promjenu postojećih stavova.

Ispitanici i postupci: U prvom istraživanju sudjelovali su studenti 5. i 6. godine Medicinskog fakulteta u Splitu koji su napisali reflektivne eseje. Uz pomoć LIWC programa eseji su potom uspoređeni s istim brojem eseja generiranih pomoću ChatGPT-a i Barda. Drugo istraživanje uključivalo je također studente medicinskog fakulteta podijeljenih u dvije skupine: jedna je pisala eseje koristeći ChatGPT, a druga samostalno. Prije i poslije pisanja, studenti su ispunjavali upitnike o stavovima prema UI-ju i o svojoj intrinzičnoj motivaciji.

Rezultati: Lingvistička analiza ukazala je da su eseji koje je generirao UI imali višu razinu analitičkog mišljenja, autentičnosti i veću upotrebu dugih riječi u usporedbi s esejima studenata. Studentski eseji pokazivali su više izraza kognitivnih procesa i duže rečenice, kao i veću raznolikost izražavanja neslaganja i kognitivnih procesa. Eseji potencijalno napisani uz pomoć UI-ja pokazivali su mješovite karakteristike, bliže UI-generiranim esejima nego originalnim studentskim esejima. U drugom istraživanju, promjena stavova prema UI-ju nije bila statistički značajna između grupa. Međutim, studenti koji su pisali bez pomoći ChatGPT-a temeljitije su proučili dodijeljeni članak i posvetili više napora u izvršavanju zadatka. Intrinzična motivacija bila je povezana s pozitivnijim stavovima prema UI-ju.

Zaključci: Rezultati istraživanja pokazuju da, iako UI vrlo uspješno stvara stilski kvalitetne i sadržajno bogate eseje, i dalje postoje uočljive lingvističke razlike između eseja napisanih od strane studenata i onih kreiranih pomoću UI. Te razlike posebno se očituju u područjima analitičkog razmišljanja, izražavanja emocija i kognitivnih procesa. Iako UI pokazuje izuzetnu sposobnost pisanja, ljudsko pisanje zadržava složenije izražavanje misli i emocija. Korištenje UI alata bilo je povezanom sa načinom na koji studenti pristupaju zadacima, ali nije značajno promijenilo njihove stavove prema UI-ju. Budući izazovi uključuju razvoj novih načina ocjenjivanje studentskog rada, uz etičku upotrebu UI-ja u obrazovanju.

8. LAIČKI SAŽETAK

Ova istraživanja su proučavala korištenje UI-ja pri pisanju eseja iz medicinske humanistike tijekom nastave na medicinskom fakultetu. Prvo istraživanje je uz pomoć računalne analize jezika usporedilo jezične karakteristike eseja koji su napisali studenti samostalno s esejima koje je generirala UI. U drugom istraživanju smo analizirali kako korištenje UI-ja utječe na stavove studenata prema toj tehnologiji. Rezultati su pokazali da eseji generirani UI-jem imaju višu razinu analitičkog izražavanja, veći broj dugih riječi i višu ocjenu "autentičnosti" stila. S druge strane, eseji koje su napisali studenti pokazali su veće korištenje izraza koji odražavaju razmišljanje, sumnju i neslaganje, kao i duže i složenije rečenice. Neki eseji koje su studenti predali nosili su obilježja oba stila, što upućuje na moguće korištenje UI-ja tijekom pisanja. U drugom istraživanju, studenti su podijeljeni u dvije skupine te su jedni pisali uz pomoć ChatGPT-a, a drugi bez njega. Iako su svi studenti pokazali slične stavove prema UI-ju prije i nakon pisanja, oni koji su pisali bez UI-ja uložili su više truda i dublje proučili temu. Također, studenti s višom unutarnjom motivacijom pokazali su pozitivnije stavove prema UI-ju. Zaključno, UI može stvoriti vrlo kvalitetne tekstove, ali ljudsko pisanje i dalje pokazuje veću dubinu u izražavanju misli i emocija. Iako UI može pomoći u pisanju, njegovo korištenje otvara važne rasprave o etici i vrednovanju studentskog rada u obrazovanju.

9. SAŽETAK NA ENGLESKOM JEZIKU

Title: Using artificial intelligence in writing medical ethics essays during medical school education

Objectives: The first study aimed to conduct a linguistic analysis and comparison of essays written by medical students with essays generated by AI (ChatGPT and Bard) on the topic of an ethical, professional, or moral dilemma. Linguistic variables were analyzed using the LIWC software, and the origin of the essays was examined using AI detection tools. The hypotheses assumed differences in the expression of emotions and cognitive processes between student-written and AI-generated essays. The second study aimed to examine how the use of ChatGPT affects students' attitudes toward AI, with the hypothesis that students who wrote essays independently and those with higher intrinsic motivation would show a stronger change in their existing attitudes.

Participants and procedures: In the first study, participants were 5th- and 6th-year medical students from the University of Split School of Medicine who wrote reflective essays. Using the LIWC software, these essays were compared with an equal number of essays generated by ChatGPT and Bard. The second study involved 80 students divided into two groups: one group wrote essays using ChatGPT, while the other wrote independently. Before and after the writing task, students completed questionnaires on their attitudes toward artificial intelligence and their intrinsic motivation.

Results: The linguistic analysis showed that AI-generated essays exhibited higher levels of analytical thinking, authenticity, and greater use of long words compared to student essays. Student essays demonstrated greater expression of cognitive processes and longer sentence lengths. Furthermore, student essays showed greater variation in the expression of disagreement and cognitive processes. Essays potentially written with AI assistance exhibited mixed characteristics, being more similar to AI-generated essays than to original student essays. In the second study, changes in attitudes toward AI were not statistically significant between groups. However, students who wrote essays without the help of ChatGPT engaged more thoroughly with the assigned article and demonstrated greater effort in completing the task. Intrinsic motivation was associated with more positive attitudes toward AI.

Conclusions: The results indicate that although AI is highly successful in producing stylistically polished and content-rich essays, recognizable linguistic differences still exist between student-written and AI-generated texts. These differences are particularly evident in aspects of analytical thinking, emotional expression, and cognitive processing. While AI

demonstrates exceptional writing capabilities, human writing retains a more complex expression of thoughts and emotions. The use of AI tools influenced the way students approached academic tasks but did not significantly alter their attitudes toward AI. Future challenges include developing new methods for assessing student work, alongside promoting the ethical use of AI in education.

10. LAIČKI SAŽETAK NA ENGLESKOM JEZIKU

These studies explored the use of artificial intelligence (AI) in writing essays on medical humanities during coursework at a medical school. The first study used computer-based language analysis to compare linguistic features of essays written independently by students with those generated by AI. The second study examined how the use of AI affects students' attitudes toward the technology. The results showed that AI-generated essays exhibited higher levels of analytical expression, a greater number of long words, and higher ratings of stylistic "authenticity." In contrast, student-written essays demonstrated more frequent use of language reflecting thought, doubt, and disagreement, along with longer and more complex sentences. Some student-submitted essays displayed characteristics of both styles, suggesting the possible use of AI during the writing process. In the second study, students were divided into two groups—one wrote essays using ChatGPT, while the other wrote without it. Although all students showed similar attitudes toward AI before and after writing, those who worked without AI engaged more deeply with the topic and put in more effort. Additionally, students with higher intrinsic motivation exhibited more positive attitudes toward AI. In conclusion, while AI is capable of producing high-quality and well-structured texts, human writing still offers greater depth in expressing thoughts and emotions. Although AI can assist in the writing process, its use raises important questions about ethics and the evaluation of student work in education.

11. LITERATURA

1. Jaakkola H, Henno J, Makela J, Thalheim B. Artificial Intelligence Yesterday, Today and Tomorrow. In: 2019 42nd International Convention on Information and Communication Technology, Electronics and Microelectronics (MIPRO) [Internet]. Opatija, Croatia: IEEE; 2019 [citirano 14. travnja 2025.]. p. 860–7. Dostupno na: <https://ieeexplore.ieee.org/document/8756913/>
2. Russell S, Norvig P. Artificial Intelligence: A Modern Approach, 4th US ed.; 2021.
3. McCarthy J. What is artificial intelligence? [Internet]. Stanford University; 2007. Dostupno na: <https://www-formal.stanford.edu/jmc/whatisai.pdf>
4. Toews R. Transformers Revolutionized AI. What Will Replace Them? [Internet]. Forbes; 2023. Dostupno na: <https://www.forbes.com/sites/robtoews/2023/09/03/transformers-revolutionized-ai-what-will-replace-them/>
5. Copeland BJ. Methods and goals in AI [Internet]. Encyclopaedia Britannica; 2025. [citirano 12. svibnja 2025.]. Dostupno na: <https://www.britannica.com/technology/artificial-intelligence/Methods-and-goals-in-AI>
6. McCarthy J, Minsky ML, Rochester N, Shannon CE. A proposal for the Dartmouth summer research project on artificial intelligence. AI Magazine Archive 1955. Vol. 27, No. 4.
7. Schuett J. A Legal Definition of AI. SSRN Electronic Journal. 2019. doi: 10.2139/ssrn.3453632
8. Bory P, Natale S, Katzenbach C. Strong and weak AI narratives: an analytical framework. AI & Soc. 2024. doi.org/10.1007/s00146-024-02087-8
9. Thorn PD, Bostrom N. Superintelligence: Paths, Dangers, Strategies. Minds & Machines. 2015. 25:285–289. doi: 10.1007/s11023-015-9377-7
10. Mahesh B. Machine Learning Algorithms - A Review. IJSR. 2020. 5;9(1):381–6. doi: 10.21275/ART20203995
11. Jain N, Kumar R. A Review on Machine Learning & It's Algorithms. IJSCE. 2022. 12(5):1-5. doi: 10.35940/ijscce.E3583.1112522
12. LeCun Y, Bengio Y, Hinton G. Deep learning. Nature. 2015. 521:436–44. doi: 10.1038/nature14539
13. Mathew A, Amudha P, Sivakumari S. Deep Learning Techniques: An Overview. U: Hassanien AE, Bhatnagar R, Darwish A, urednici. Advanced Machine

- Learning Technologies and Applications. Singapore: Springer Singapore; 2021. str. 599–608.
14. Dongare AD, Kharde RR, Kachare AD. Introduction to artificial neural network. IJEIT. 2012;2(1):189–94.
 15. Wu YC, Feng JW. Development and Application of Artificial Neural Network. Wireless Pers Commun. 2018;102(2):1645–56.
 16. Oliveira E. Artificial Intelligence: An Overview. U: Hsueh TT urednik. Cutting Edge Technologies and Microcomputer Applications for Developing Countries. New York: Routledge; 2019. str. 61–5.
 17. Haenlein M, Kaplan A. A Brief History of Artificial Intelligence: On the Past, Present, and Future of Artificial Intelligence. California Management Review. 2019;61(4):5–14.
 18. Turing AM. Computing machinery and intelligence. Mind. 1950;59(236):433–60.
 19. Newell A, Simon HA, Shaw JC. Report on a General Problem-Solving Program. Proceedings of the International Conference on Information Processing. 1959;256–64.
 20. Umbrello S. AI Winter. U: Michael Klein M, Frana P, urednici. Encyclopedia of Artificial Intelligence: The Past, Present, and Future of AI. Bloomsbury Publishing; 2021. str. 7-8.
 21. Van Melle W. MYCIN: a knowledge-based consultation program for infectious disease diagnosis. International Journal of Man-Machine Studies. 1978;10(3):313–22.
 22. Farrell N, Hall P. The Best Smart Speakers With Alexa, Google Assistant, and Siri [Internet]. Wired. 2025 [citirano 22. travnja 2025.]. Dostupno na: <https://www.wired.com/story/best-smart-speakers/>
 23. Watson, ‘Jeopardy!’ champion [Internet]. IBM. [citirano 22. travnja 2025.]. Dostupno na: <https://www.ibm.com/history/watson-jeopardy>
 24. Taeihagh A, Lim HSM. Governing autonomous vehicles: emerging responses for safety, liability, privacy, cybersecurity, and industry risks. Transport Reviews. 2019;39(1):103–28.
 25. OpenAI. GPT-4 Technical Report. [Internet]. arXiv; 2024 [citirano 18. travnja 2025.]. Dostupno na: <https://arxiv.org/abs/2303.08774v6>

26. Schroer A. 76 Artificial Intelligence (AI) Companies to Know [Internet]. BuiltIn; 2025 [citirano 24. travnja 2025.]. Dostupno na: <https://builtin.com/artificial-intelligence/ai-companies-roundup>
27. Zhang L. Artificial Intelligence: 70 Years Down the Road [Internet]. arXiv; 2023 [citirano 24. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2303.02819>
28. McCarthy, J. What is Artificial Intelligence? [Internet]. Stanford University. 2007 [citirano 24. travnja 2025.]. Dostupno na: <https://www-formal.stanford.edu/jmc/whatisai.pdf>
29. Topol EJ. High-performance medicine: the convergence of human and artificial intelligence. *Nat Med*. 2019;25(1):44–56.
30. Litjens G, Kooi T, Bejnordi BE, Setio AAA, Ciompi F, Ghafoorian M, i sur. A survey on deep learning in medical image analysis. *Medical Image Analysis*. 2017;42:60–88.
31. Esteva A, Robicquet A, Ramsundar B, Kuleshov V, DePristo M, Chou K, i sur. A guide to deep learning in healthcare. *Nat Med*. 2019;25(1):24–9.
32. Rajpurkar P, Chen E, Banerjee O, Topol EJ. AI in health and medicine. *Nat Med*. 2022;28(1):31–8.
33. Anderson J. Educating in a World of Artificial Intelligence [Internet]. Harvard Graduated school of education. 2023 [citirano 24. travnja 2025.]. Dostupno na: <https://www.gse.harvard.edu/ideas/edcast/23/02/educating-world-artificial-intelligence>
34. Writers of UoPeople. AI In Education: Where Is It Now And What Is The Future [Internet]. UoPeople 2024 [citirano 24. travnja 2025.]. Dostupno na: <https://www.uopeople.edu/blog/ai-in-education-where-is-it-now-and-what-is-the-future/>
35. Luckin R, Holmes W, Griffiths M, Forcier LB. *Intelligence Unleashed: An argument for AI in Education*. London: Pearson; 2016.
36. Holmes W, Bialik M, Fadel C. *Artificial intelligence in education: promises and implications for teaching and learning*. Boston, MA: The Center for Curriculum Redesign; 2019.
37. Chui M, Manyika J, Miremadi M. Where machines could replace humans—and where they can't (yet) [Internet]. McKinsey & Company. 2016 [citirano 26. travnja 2025.]. Dostupno na: <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/where-machines-could-replace-humans-and-where-they-cant-yet#/>

38. Bughin J, Seong J, Manyika J, Chui M, Joshi R. Notes from the AI frontier: Modeling the impact of AI on the world economy [Internet]. McKinsey & Company. 2018 [citirano 26. travnja 2025.]. Dostupno na: <https://www.mckinsey.com/featured-insights/artificial-intelligence/notes-from-the-ai-frontier-modeling-the-impact-of-ai-on-the-world-economy>
39. Martin K. Machine Bias. U Martin K urednik. Ethics of data and analytics: concepts and cases. Prvo izdanje. New York: Auerbach Publications; 2022.
40. Ashley KD. Artificial Intelligence and Legal Analytics: New Tools for Law Practice in the Digital Age [Internet]. 1st ed. Cambridge University Press; 2017 [citirano 26. travnja 2025.]. Dostupno na: <https://www.cambridge.org/core/product/identifier/9781316761380/type/book>
41. Papakyriakopoulos O, Tessono C, Narayanan A, Kshirsagar M. How Algorithms Shape the Distribution of Political Advertising: Case Studies of Facebook, Google, and TikTok. Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society 2022; p. 532–46. doi: 10.1145/3514094.3534166
42. Isaak J, Hanna MJ. User Data Privacy: Facebook, Cambridge Analytica, and Privacy Protection. Computer. 2018;51(8):56–9.
43. Howard PN, Woolley S, Calo R. Algorithms, bots, and political communication in the US 2016 election: The challenge of automated political communication for election law and administration. Journal of Information Technology & Politics. 2018;15(2):81–93.
44. Garvie C, Bedoya AM, Frankle J. The Perpetual Line-Up: Unregulated Police Face Recognition in America [Internet]. Washington DC: Georgetown Law Center on Privacy & Technology. 2016. [citirano 12. travnja 2025.]. Dostupno na: <https://www.perpetuallineup.org/>
45. Brey P. The strategic role of technology in a good society. Technology in Society. 2018;52:39–45.
46. Krauss C, Do XA, Huck N. Deep neural networks, gradient-boosted trees, random forests: Statistical arbitrage on the S&P 500. European Journal of Operational Research. 2017;259(2):689–702.
47. Liakos K, Busato P, Moshou D, Pearson S, Bochtis D. Machine Learning in Agriculture: A Review. Sensors. 2018;18(8):2674.
48. Grigorescu S, Trasnea B, Cocias T, Macesanu G. A Survey of Deep Learning Techniques for Autonomous Driving. Journal of Field Robotics. 2020;37(3):362–86.

49. Brundage M. Taking superintelligence seriously. *Futures*. 2015;72:32–5.
50. Char DS, Shah NH, Magnus D. Implementing Machine Learning in Health Care — Addressing Ethical Challenges. *N Engl J Med*. 2018 Mar;378(11):981–3.
51. Brynjolfsson E, McAfee A. The second machine age: work, progress, and prosperity in a time of brilliant technologies. 1. izdanje. New York, London: W. W. Norton & Company; 2016. 306 p.
52. Zalnieriute M. Facial recognition surveillance and public space: protecting protest movements. *International Review of Law, Computers & Technology*. 2025;39(1):116–35.
53. Graham A. Destined for War: Can America and China Escape Thucydides’s Trap? 1. izdanje. Dublin: Houghton Mifflin Harcourt; 2014.
54. Casilli AA, Tubaro P, Cornet M, Ludec CL, Torres-Cierpe J, Braz MV. Global Inequalities in the Production of Artificial Intelligence: A Four-Country Study on Data Work [Internet]. arXiv; 2024 [citirano 28. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2410.14230>
55. Zubiaga A. Natural language processing in the era of large language models. *Front Artif Intell*. 2024;6:1350306.
56. Barrett M. The dark side of AI: algorithmic bias and global inequality. Cambridge: Cambridge Judge Business School; 2023 [citirano 28. travnja 2025.]. Dostupno na: <https://www.jbs.cam.ac.uk/2023/the-dark-side-of-ai-algorithmic-bias-and-global-inequality/>
57. IBM. What are large language models (LLMs) [Internet]? IBM; 2023 [citirano 14. travnja 2025.]. Dostupno na: <https://www.ibm.com/topics/large-language-models>
58. Pixelplex. Top 10 Real-Life Applications of Large Language Models [Internet]. Pixelplex; 2024 [citirano 14. travnja 2025.]. Dostupno na: <https://pixelplex.io/blog/llm-applications/>
59. Lockhart ENS. Creativity in the age of AI: the human condition and the limits of machine generation. *J Cult Cogn Sci*. 2025;9(1):83–8.
60. Parthasarathy VB, Zafar A, Khan A, Shahid A. The Ultimate Guide to Fine-Tuning LLMs from Basics to Breakthroughs: An Exhaustive Review of Technologies, Research, Best Practices, Applied Research Challenges and Opportunities [Internet]. arXiv; 2024 [citirano 14 travnja 2025.]. Dostupno na: <https://arxiv.org/abs/2408.13296>

61. Euronews. ChatGPT a year on: 3 ways the AI chatbot has completely changed the world in 12 months [Internet]. Euronews; 2023 [citirano 14 travnja 2025.]. Dostupno na: <https://www.euronews.com/next/2023/11/30/chatgpt-a-year-on-3-ways-the-ai-chatbot-has-completely-changed-the-world-in-12-months>
62. Kerner SM. 25 of the best large language models in 2025 [Internet]. Techtarger; 2025 [citirano 14 travnja 2025.]. Dostupno na: <https://www.techtarger.com/whatis/feature/12-of-the-best-large-language-models>
63. Gao L, Biderman S, Black S, Golding L, Hoppe T, Foster C, i sur. The Pile: An 800GB Dataset of Diverse Text for Language Modeling [Internet]. arXiv; 2020 [citirano 24. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2101.00027>
64. Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, i sur. Attention is All you Need [Internet]. arXiv; 2023 [citirano 24. travnja 2025.]. Dostupno na: <https://arxiv.org/abs/1706.03762>
65. Niu Z, Zhong G, Yu H. A review on the attention mechanism of deep learning. *Neurocomputing*. 2021;452:48–62.
66. Naveed H, Khan AU, Qiu S, Saqib M, Anwar S, Usman M i sur. A Comprehensive Overview of Large Language Models [Internet]. arXiv; 2024 [citirano 18. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2307.06435>
67. Shen Z, Tao T, Ma L, Neiswanger W, Liu Z, Wang H i sur. SlimPajama-DC: Understanding Data Combinations for LLM Training [Internet]. arXiv; 2023 [citirano 14. travnja 2025.]. Dostupno na: <https://arxiv.org/abs/2309.10818>
68. Brown TB, Mann B, Ryder N, Subbiah M, Kaplan J, Dhariwal P i sur. Language Models are Few-Shot Learners [Internet]. arXiv; 2020 [citirano 15. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2005.14165>
69. Sun K, Dredze M. Amuro and Char: Analyzing the Relationship between Pre-Training and Fine-Tuning of Large Language Models [Internet]. arXiv; 2025 [citirano 24. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2408.06663>
70. Getgenie. GPT 3 vs. GPT 4: Differences Between ChatGPT Models [Internet]. Getgenie; 2025 [citirano 24. travnja 2025.]. Dostupno na: https://getgenie.ai/gpt-3-vs-gpt-4/?utm_source=chatgpt.com
71. Hong SK, Jang JS, Kwon HY. Enhancing performance of transformer-based models in natural language understanding through word importance embedding. *Knowledge-Based Systems*. 2024;304:112404.

72. Singh A, Patel NP, Ehtesham A, Kumar S, Khoei TT. A Survey of Sustainability in Large Language Models: Applications, Economics, and Challenges [Internet]. arXiv; 2025 [citirano 24. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2412.04782>
73. Bai G, Chai Z, Ling C, Wang S, Lu J, Zhang N i sur. Beyond Efficiency: A Systematic Survey of Resource-Efficient Large Language Models [Internet]. arXiv; 2024 [citirano 24. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2401.00625>
74. Kaur P, Kashyap GS, Kumar A, Nafis MT, Kumar S, Shokeen V. From Text to Transformation: A Comprehensive Review of Large Language Models' Versatility [Internet]. arXiv; 2024 [citirano 22. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2402.16142>
75. Bommasani R, Hudson DA, Adeli E, Altman R, Arora S, Arx S i sur. On the Opportunities and Risks of Foundation Models [Internet]. arXiv; 2022 [citirano 15. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2108.07258>
76. Wang S, Xu T, Li H, Zhang C, Liang J, Tang J i sur. Large Language Models for Education: A Survey and Outlook [Internet]. arXiv; 2024 [citirano 12 travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2403.18105>
77. Mindy-support. How llms are powering the next generation of smart assistants [Internet]. Mindy-support; 2025 [citirano 15. travnja 2025.]. Dostupno na: <https://mindy-support.com/news-post/how-llms-are-powering-the-next-generation-of-smart-assistants/>
78. Zhu W, Liu H, Dong Q, Xu J, Huang S, Kong L i sur. Multilingual Machine Translation with Large Language Models: Empirical Results and Analysis [Internet]. arXiv; 2023 [citirano 12 travnja 2025.]. Dostupno na: <https://arxiv.org/abs/2304.04675>
79. Devenport TH, Mittal N. How Generative AI Is Changing Creative Work [Internet]. Harvard business review; 2022 [citirano 12 travnja 2025.]. Dostupno na: <https://hbr.org/2022/11/how-generative-ai-is-changing-creative-work>
80. Stallbaumer C. Introducing Copilot for Microsoft 365—A whole new way to work [Internet]. Microsoft; 2023 [citirano 12 travnja 2025.]. Dostupno na: <https://www.microsoft.com/en-us/microsoft-365/blog/2023/03/16/introducing-microsoft-365-copilot-a-whole-new-way-to-work/>
81. Med-PaLM: A Medical Large Language Model [Internet]. Dostupno na: <https://sites.research.google/med-palm/>

82. Harvey AI [Internet]. Dostupno na: <https://www.harvey.ai/>
83. Tools such as ChatGPT threaten transparent science; here are our ground rules for their use. *Nature*. 2023;613(7945):612–612.
84. Hager P, Jungmann F, Holland R, Bhagat K, Hubrecht I, Knauer M i sur. Evaluation and mitigation of the limitations of large language models in clinical decision-making. *Nat Med*. 2024;30(9):2613–22.
85. Ullah E, Parwani A, Baig MM, Singh R. Challenges and barriers of using large language models (LLM) such as ChatGPT for diagnostic medicine with a focus on digital pathology – a recent scoping review. *Diagn Pathol*. 2024;19(1):43.
86. Rajaram S. The Ultimate Guide to Hardware Requirements for Training and Fine-Tuning Large Language Models (LLMs) [Internet]. Medium: Towards AI; 2025 [citirano 12 travnja 2025.]. Dostupno na: <https://pub.towardsai.net/the-ultimate-guide-to-hardware-requirements-for-training-and-fine-tuning-large-language-models-7b5fe3884f64>
87. Jacob C, Kerrigan P, Bastos M. The chat-chamber effect: Trusting the AI hallucination. *Big Data & Society*. 2025;12(1):20539517241306345.
88. Sallami D, Chang YC, Aïmeur E. From Deception to Detection: The Dual Roles of Large Language Models in Fake News [Internet]. arXiv; 2024 [citirano 18. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2409.17416>
89. Bender EM, Gebru T, McMillan-Major A, Shmitchell S. On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? U: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency [Internet]. Virtual Event Canada: ACM; 2021 [citirano 18. travnja 2025.]. Dostupno na: <https://dl.acm.org/doi/10.1145/3442188.3445922>
90. Zhui L, Fenghe L, Xuehu W, Qining F, Wei R. Ethical Considerations and Fundamental Principles of Large Language Models in Medical Education: Viewpoint. *J Med Internet Res*. 2024;26:e60083.
91. Pudasaini S, Miralles-Pechuán L, Lillis D, Salvador ML. Survey on Plagiarism Detection in Large Language Models: The Impact of ChatGPT and Gemini on Academic Integrity [Internet]. arXiv; 2024 [citirano 15. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2407.13105>
92. Chemaya N, Martin D. Perceptions and detection of AI use in manuscript preparation for academic journals. *PLoS One*. 2024. doi: 10.1371/journal.pone.0304807

93. Bender EM, Gebru T, McMillan-Major A, Shmitchell S. On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? U: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency [Internet]. Virtual Event Canada: ACM; 2021 [citirano 15 travnja 2025.]. p. 610–23. Dostupno na: <https://dl.acm.org/doi/10.1145/3442188.3445922>
94. Huang L, Yu W, Ma W, Zhong W, Feng Z, Wang H, i sur. A Survey on Hallucination in Large Language Models: Principles, Taxonomy, Challenges, and Open Questions. *ACM Trans Inf Syst.* 2025;43(2):1–55.
95. Miranda M, Ruzzetti ES, Santilli A, Zanzotto FM, Bratières S, Rodolà E. Preserving Privacy in Large Language Models: A Survey on Current Threats and Solutions [Internet]. arXiv; 2025 [citirano 24. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2408.05212>
96. Novelli C, Casolari F, Hacker P, Spedicato G, Floridi L. Generative AI in EU Law: Liability, Privacy, Intellectual Property, and Cybersecurity [Internet]. arXiv; 2024 [citirano 15 travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2401.07348>
97. Fengchun M, Wayne H. Guidance for Generative AI in Education and Research. [Internet]. UNESCO; 2025 [citirano 24. travnja 2025.]. Dostupno na: <https://www.unesco.org/en/articles/guidance-generative-ai-education-and-research>
98. OECD. Recommendation of the Council on Artificial Intelligence. [Internet]. OECD Legal instruments; 2021 [citirano 24. travnja 2025.]. Dostupno na: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>
99. European Commission. Proposal for a Regulation on a European approach for Artificial Intelligence. *EUR-Lex*; 2021 [citirano 24. travnja 2025.]. Dostupno na: <https://eur-lex.europa.eu/>.
100. The Nerd Academy. AI Literacy for Students: Cultivating Critical Thinking in the Age of Intelligent Machines. The Nerd Academy; 2024 [citirano 24. travnja 2025.]. Dostupno na: <https://thenerdacademy.com/ai/ai-literacy-for-students-cultivating-critical-thinking-in-the-age-of-intelligent-machines/>
101. Ferdaus MM, Abdelguerfi M, Ioup E, Niles KN, Pathak K, Sloan S. Towards Trustworthy AI: A Review of Ethical and Robust Large Language Models [Internet]. arXiv; 2024 [citirano 24. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2407.13934>
102. Bommasani R, Hudson DA, Adeli E, Altman R, Arora S, Arx S von, i sur. On the Opportunities and Risks of Foundation Models [Internet]. arXiv; 2022 [citirano 15. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2108.07258>

103. Weidinger L, Mellor J, Rauh M, Griffin C, Uesato J, Huang PS i sur. Ethical and social risks of harm from Language Models [Internet]. arXiv; 2021 [citirano 15. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2112.04359>
104. Cotton D, Cotton P, Shipway JR. Chatting and Cheating. Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*. 2023;61(2), 228–239.
105. Herbold S, Hautli-Janisz A, Heuer U, Kikteva Z, Trautsch A. A large-scale comparison of human-written versus ChatGPT-generated essays. *Sci Rep*. 2023;13(1):18617.
106. Kasneci E, Sessler K, Küchemann S, Bannert M, Dementieva D, Fischer F, i sur. ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*. 2023;103:102274.
107. Quitadamo IJ, Kurtz MJ. Learning to Improve: Using Writing to Increase Critical Thinking Performance in General Education Biology. *LSE*. 2007;6(2):140–54.
108. Jelson A, Manesh D, Jang A, Dunlap D, Lee SW. An Empirical Study to Understand How Students Use ChatGPT for Writing Essays [Internet]. arXiv; 2025 [citirano 18. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2501.10551>
109. Von Garrel J, Mayer J. Artificial Intelligence in studies—use of ChatGPT and AI-based tools among students in Germany. *Humanit Soc Sci Commun*. 2023;10(1):799.
110. Lucariello K. Turnitin: More than Half of Students Continue to Use AI to Write Papers [Internet]. *Campus technology*; 2024 [citirano 15. travnja 2025.]. Dostupno na: https://campustechnology.com/Articles/2024/05/08/Turnitin-More-than-Half-of-Students-Continue-to-Use-AI-to-Write-Papers.aspx?utm_source=chatgpt.com
111. Weber-Wulff D, Anohina-Naumeca A, Bjelobaba S, Foltýnek T, Guerrero-Dib J, Popoola O i sur. Testing of detection tools for AI-generated text. *Int J Educ Integr*. 2023;19(1):26.
112. Yang K, Raković M, Liang Z, Yan L, Zeng Z, Fan Y i sur. Modifying AI, Enhancing Essays: How Active Engagement with Generative AI Boosts Writing Quality. U: *Proceedings of the 15th International Learning Analytics and Knowledge Conference* [Internet]. Dublin Ireland: ACM; 2025 [citirano 16. travnja 2025.]. Dostupno na: <https://dl.acm.org/doi/10.1145/3706468.3706544>

113. Zheldibayeva R, Nascimento AK de O, Castro V, Kalantzis M, Cope B. The Impact of AI-Driven Tools on Student Writing Development: A Case Study From The CGScholar AI Helper Project [Internet]. arXiv; 2025 [citirano 16. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2501.08473>
114. Foltynnek T, Bjelobaba S, Glendinning I, Khan ZR, Santos R, Pavletic P i sur. ENAI Recommendations on the ethical use of Artificial Intelligence in Education. *Int J Educ Integr.* 2023;19(1):12, s40979-023-00133–4.
115. Liu Q, Zhou Y, Huang J, Li G. When ChatGPT is gone: Creativity reverts and homogeneity persists [Internet]. arXiv; 2024 [citirano 16. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2401.06816>
116. Anderson BR, Shah JH, Kreminski M. Homogenization Effects of Large Language Models on Human Creative Ideation. *Creativity and Cognition.* 2024. doi: 10.1145/3635636.3656204
117. Wahle JP, Ruas T, Mohammad SM, Meuschke N, Gipp B. AI Usage Cards: Responsibly Reporting AI-generated Content [Internet]. arXiv; 2023 [citirano 16. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2303.03886>
118. Mishra T, Sutanto E, Rossanti R, Pant N, Ashraf A, Raut A, i sur. Use of large language models as artificial intelligence tools in academic research and publishing among global clinical researchers. *Sci Rep.* 2024;14(1):31672.
119. Kasneci E, Sessler K, Küchemann S, Bannert M, Dementieva D, Fischer F, i sur. ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences.* 2023;103:102274.
120. Herbold S, Hautli-Janisz A, Heuer U, Kikteva Z, Trautsch A. AI, write an essay for me: A large-scale comparison of human-written versus ChatGPT-generated essays [Internet]. arXiv; 2023 [citirano 15. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2304.14276>
121. Boyd RL, Ashokkumar A, Seraj S, Pennebaker JW. The Development and Psychometric Properties of LIWC-22. University of Texas. Austin. 2022.
122. Tian E, Cui A. GPTZero: Towards detection of AI-generated text using zero-shot and supervised methods. *GPTZero*; 2023. Dostupno na: <https://gptzero.me>
123. Amabile TM, Hill KG, Hennessey BA, Tighe EM. The Work Preference Inventory: Assessing intrinsic and extrinsic motivational orientations. *Journal of Personality and Social Psychology.* 1994;66(5):950–67.

124. Stein JP, Messingschlager T, Gnambs T, Hutmacher F, Appel M. Attitudes towards AI: measurement and associations with personality. *Sci Rep*. 2024;14(1):2909.
125. Li Y, Sha L, Yan L, Lin J, Raković M, Galbraith K, i sur. Can large language models write reflectively. *Computers and Education: Artificial Intelligence*. 2023;4:100140.
126. Jiang H, Zhang X, Cao X, Breazeal C, Roy D, Kabbara J. PersonaLLM: Investigating the Ability of Large Language Models to Express Personality Traits [Internet]. arXiv; 2024 [citirano 26. travnja 2025.]. Dostupno na: <http://arxiv.org/abs/2305.02547>
127. Ovsyannikova D, De Mello VO, Inzlicht M. Third-party evaluators perceive AI as more compassionate than expert humans. *Commun Psychol*. 2025;3(1):4.
128. Sandler M, Choung H, Ross A, David P. A Linguistic Comparison between Human and ChatGPT-Generated Conversations. U: Wallraven C, Liu CL, Ross A, urednici. *Pattern Recognition and Artificial Intelligence. ICPRAI 2024. Lecture Notes in Computer Science*, vol 14892. Singapore: Springer. doi: 10.1007/978-981-97-8702-9_25
129. Deci EL, Ryan RM. *Intrinsic Motivation and Self-Determination in Human Behavior*. Boston, MA: Springer US; 1985.
130. Ryan RM, Deci EL. Self-determination theory and the facilitation of intrinsic motivation, social development, and well-being. *American Psychologist*. 2000;55(1):68–78.
131. Borji A, Mohammadian M. Battle of the wordsmiths: Comparing ChatGPT, GPT-4, Claude, and bard. *SSRN Electronic Journal*. 2023;12,1–10.
132. Yusuf A, Pervin N, Román-González M. Generative AI and the future of higher education: a threat to academic integrity or reformation? Evidence from multicultural perspectives. *Int J Educ Technol High Educ*. 2024;21(1):21.
133. Chan CKY, Hu W. Students' voices on generative AI: perceptions, benefits, and challenges in higher education. *Int J Educ Technol High Educ*. 2023;20(1):43.
134. Wigfield A, Guthrie JT. Engagement and motivation in reading. U Kamil ML, Mosenthal PB, Pearson PD, Barr R, urednici. 2000. *Handbook of reading research*. Lawrence Erlbaum Associates Publishers; 2000. str. 403–22.

135. Mollick ER, Mollick L. Using AI to Implement Effective Teaching Strategies in Classrooms: Five Strategies, Including Prompts. The Wharton School Research Paper. 2023. doi: 10.2139/ssrn.4391243

12. ŽIVOTOPIS

Osobni podaci

Ime i prezime: Mariano Kaliterna

Elektronička pošta: mariano.kaliterna@gmail.com

Državljanstvo: hrvatsko

Datum i mjesto rođenja: 29. listopada 1968., Split, Hrvatska

Obrazovanje

1988. - 1995. Sveučilište u Zagrebu, Studij u Splitu, Medicinski fakultet, Doktor medicine

2009. - 2010. Sveučilište u Splitu, Medicinski fakultet, Poslijediplomski specijalistički studij
Klinička epidemiologija, Magistar kliničke epidemiologije

2011. - 2015. Specijalizacija iz psihijatrije, Specijalist psihijatrije

2016. - 2017. Sveučilište u Zagrebu, Poslijediplomski specijalistički studij Psihoterapija,
Subspecijalist psihoterapije

2017. - 2018. Sveučilište u Splitu, Medicinski fakultet, Doktorski studij Medicina utemeljena
na dokazima

Radno iskustvo:

1996.-1997., 1998.-1999. Pripravnički staž, Nastavni zavod za javno zdravstvo, SDŽ

1997.-1998. Istraživački novak, University of Connecticut health center, Department of
developmental biology, Farmington, SAD

1999.-2011. Marketinški predstavnik u Merck Sharp & Dohme

2011.-2015. Specijalizant psihijatrije, KBC Split

2015.- Specijalist psihijatar, KBC Split

2017.- -Subspecijalist psihoterapije

2017.- Asistent na Medicinskom fakultetu, Sveučilište u Splitu, na Katedri Psihijatrija i
suradnik na Katedri Medicinska humanistika

Poznavanje stranih jezika:

Engleski jezik

Talijanski jezik

Članstva:

Hrvatska liječnička komora (HLK).

Od 2024. godine sam zastupnik u Skupštini HLK.

Hrvatski liječnički zbor

Hrvatsko psihijatrijsko društvo (HPD).

Od 2023. godine obnašam funkciju dopredsjednika HPD.

Nagrade i priznanja

Zahvala Hrvatskog liječničkog zbora za zdravstvenu i humanitarnu djelatnost.

Tečajevi i projekti

Sudionik sam dva međunarodna projekta vezano uz etiku istraživanja: “Bridging the Gap: Research Ethics Training and Education“ (BriGRETE) i “Improving Research Ethics Expertise and Competencies to Ensure Reliability and Trust in Science“ (iRECS).

Publikacije

1. **Kaliterna M**, Žuljević MF, Ursić L, Krka J, Duplančić D. Testing the capacity of Bard and ChatGPT for writing essays on ethical dilemmas: A cross-sectional study. Sci Rep. 2024 Oct 30;14(1):26046. doi: 10.1038/s41598-024-77576-3.
2. Žuljević MF, Hren D, Storman D, **Kaliterna M**, Duplančić D. Attitudes of European psychiatrists on psychedelics: a cross-sectional survey study. Sci Rep. 2024 Aug 12;14(1):18716. doi: 10.1038/s41598-024-69688-7.

3. Mastelić T, Borovina Marasović T, Žuljević MF, Sučević Ercegovac M, **Kaliterna M**, Pleić N, Vukorepa D, Topić J, Žuljan Cvitanović M, Lasić D, Uglešić B, Kozina S, Glavina T. Attitudes on Psychedelics in a Sample of Croatian Mental Health Professionals: A Cross-Sectional National Survey Study. *J Psychoactive Drugs*. 2024 Jun 27;1-10. doi: 0.1080/02791072.2024.2370343.
4. Žuljević MF, Breški N, **Kaliterna M**, Hren D. Attitudes of European psychiatrists on psychedelics: a qualitative study. *Front Psychiatry*. 2024 May 24;15:1411234. doi: 10.3389/fpsyt.2024.1411234. eCollection 2024.
5. Žuljević MF, Buljan I, Leskur M, **Kaliterna M**, Hren D, Duplančić D. Validation of a new instrument for assessing attitudes on psychedelics in the general population. *Sci Rep*. 2022 Oct 29;12(1):18225. doi: 10.1038/s41598-022-23056-5.
6. Uglešić L, Glavina T, Lasić D, **Kaliterna M**. Postinjection Delirium/Sedation Syndrome (PDSS) Following Olanzapine Long-Acting Injection: A Case Report. *Psychiatr Danub*. 2017 Mar;29(1):90-91.
7. Britvić D, Antičević V, **Kaliterna M**, Lušić L, Beg A, Brajević-Gizdić I, Kudrić M, Stupalo Ž, Krolo V, Pivac N. Comorbidities with Posttraumatic Stress Disorder (PTSD) among combat veterans: 15 years postwar analysis. *Int J Clin Health Psychol*. 2015 May-Aug;15(2):81-92. doi: 10.1016/j.ijchp.2014.11.002. Epub 2014 Dec 25.
8. Kaliterna V, **Kaliterna M**, Hrenovic J, Barisic Z, Tonkic M, Goic-Barisic I. *Acinetobacter baumannii* in Southern Croatia: Clonal lineages, biofilm formation, and resistance patterns. *Infectious Diseases*. 2015;47(12):902-907.
9. Kaliterna V, **Kaliterna M**, Pejković L, Barišić Z, Karin Ž. Is there a need for the introduction of »Screening« for chlamydia trachomatis in women younger than 25 years? *Gynaecologia et Perinatologia*. 2012;21(3):96–9
10. Culic V, Balarin L, Zierler H, Sertic J, Culic S, Relic B, et al. Two rare mutations in cystic fibrosis. *Paediatrica Croatica*. 2000;44(4):161–5.
11. Kalajzic I, Kalajzic Z, **Kaliterna M**, Gronowicz G, Clark SH, Lichtler AC, et al. Use of Type I Collagen Green Fluorescent Protein Transgenes to Identify Subpopulations of Cells at Different Stages of the Osteoblast Lineage. *Journal of Bone and Mineral Research*. 2002 Jan 1;17(1):15–25.

Priopćenja na stručnim ili znanstvenim skupovima

1. Prvi hrvatski kongres o psihotraumi s međunarodnim sudjelovanjem, Split, Hrvatska, 9.-11. studeni 2017. "Može li izostanak podrške okoline u traumatskim situacijama potaknuti razvoj PTSP-a? / Can the Absence of Support in the Traumatic Situations Induce PTSP Development?" **M. Kaliterna**, D. Lasić, M. Žuljan-Cvitanović, knjiga sažetaka 2017, 82-83 (poster)
2. 7. psihijatrijski kongres s međunarodnim sudjelovanjem, Opatija, 24.-27. listopada, 2018. "Brintellix - omogućuje bolesnicima da se osjećaju, razmišljaju i funkcioniraju bolje - primjer iz kliničke prakse". **Kaliterna M** (predavanje)
3. 16. mostarska psihijatrijska subota, Ličnost u psihijatriji: Značenje u prevenciji, nastanku i liječenju duševnih smetnji, Mostar, BiH, 1. lipanj 2019. "Psihopatizacija: poremećaj ili odgovor na društvena zbivanja" **Kaliterna M**, Krnić S, Glavina T (predavanje)
4. 16. psihijatrijski dani s međunarodnim sudjelovanjem, virtualno izdanje, 25.-27. studeni, 2020. "Pregled rada dnevne bolnice klinike za psihijatriju u periodu od 2004. do 2019. godine". **Mariano Kaliterna**, Marija Žuljan-Cvitanović, Joško Topić, Duška Krnić, Lea Kustura, Romilda Roje, knjiga sažetaka 2020, 26-26 (poster)
5. 16. hrvatski Cochrane simpozij s međunarodnim sudjelovanjem „Moderni pogledi na mentalno zdravlje: uključivanje medicine temeljene na dokazima za učinkovitu zdravstvenu skrb“, Split, 18. listopada, 2024. "Testing the capacity of Bard and ChatGPT for writing essays on ethical dilemmas: A cross-sectional study". **Mariano Kaliterna**, Marija Franka Žuljević, Luka Ursić, Jakov Krka, Darko Duplančić, Ana Marušić (poster)
6. Poslijediplomski tečaj II. Kategorije: Aktualni izazovi u dijagnostici i liječenju poremećaja hemostaze, Split, 13.-14. prosinca 2024. "Psihijatrijski lijekovi i poremećaj hemostaze". **Kaliterna M** (predavanje)
7. Međužupanijski stručni skup: Izazovi u radu stručnih suradnika psihologa i nastavnika psihologije, Imotski, 2.05.2024. "Suradnja psihologa i psihijatra". **Kaliterna M** (predavanje)
8. 6. dani koagulacije, multidisciplinarni kongres s međunarodnim sudjelovanjem, Zagreb, 8.-10. svibnja, 2025. "Psychiatric treatment and hemostatic risk". **Kaliterna M** (predavanje)

13. DODACI

Dodatak 1: Upute studentima o načinu pisanja eseja (prvo istraživanje)

Vaš zadatak je napisati esej o nekom Vašem stvarnom iskustvu s pacijentom/pacijentima ili kolegama/suradnicima koje ste doživjeli za vrijeme studija i kliničkih rotacija ili tijekom prakse prije upisa na fakultet, a u kojem ste se susreli sa etičkim, moralnim ili profesionalnim problemom. Esaj treba imati 400 - 500 riječi, bez naslova i podnaslova. Niste obavezani stavljati reference, ali možete ako želite. Bitno je da navedete zbog čega je situacija koju ste opisali etički, moralno ili profesionalno upitna. Možete predstaviti situaciju u kojoj ste bili sudionik ili samo promatrač, a po želji možete dati vlastiti osvrt (profesionalan, osobni, itd.), uz izneseno vlastito mišljenje koje se neće ocjenjivati. Niste obavezani navoditi nikakve stvarne osobe ili mjesta ako ne želite.

Dodatak 2: Anketa povezanost korištenja ChatGPT-a i pristupa esejskom zadatku na nastavi iz medicinske humanistike (drugo istraživanje)

Ime i prezime: _____

1. Dob (unesite kao broj): _____

2. Spol:

- a) Muški
- b) Ženski

3. Koliko često koristite ChatGPT?

- a) Nikada
- b) Rijetko
- c) Ponekad
- d) Često

4. Molimo vas da procijenite svoj stupanj slaganja ili neslaganja sa sljedećim tvrdnjama:

- a) Potpuno se ne slažem
- b) Ne slažem se
- c) Ni slažem se ni ne slažem
- d) Slažem se
- e) Potpuno se slažem

- 1. Uživam u rješavanju problema koji su mi potpuno novi.
- 2. Uživam u pokušajima rješavanja složenih problema.
- 3. Što je problem teži, to više uživam pokušavajući ga riješiti.
- 4. Želim da mi posao pruža prilike za povećanje znanja i vještina.
- 5. Radoznalost je glavni pokretač mnogih mojih aktivnosti.
- 6. Želim otkriti koliko dobar stvarno mogu biti u svom poslu.
- 7. Više volim samostalno otkrivati rješenja.

8. Najvažnije mi je uživati u onome što radim.
 9. Važno mi je imati priliku za samoizražavanje.
 10. Više volim posao za koji znam da ga mogu dobro obaviti nego posao koji traži pomicanje mojih sposobnosti.
 11. Bez obzira na ishod projekta, zadovoljan/na sam ako osjećam da sam stekao/la novo iskustvo.
 12. Ugodnije mi je kada mogu samostalno postavljati ciljeve.
 13. Uživam u radu koji me toliko zaokupi da zaboravim na sve ostalo.
 14. Važno mi je moći raditi ono u čemu najviše uživam.
 15. Uživam u relativno jednostavnim i izravnim zadacima.
5. U nastavku nas zanima vaš stav prema umjetnoj inteligenciji (UI).
- a) Potpuno se ne slažem
 - b) Ne slažem se
 - c) Ni slažem se ni ne slažem
 - d) Slažem se
 - e) Potpuno se slažem
1. UI će učiniti ovaj svijet boljim mjestom.
 2. Imam snažne negativne emocije prema UI.
 3. Želim koristiti tehnologije koje se oslanjaju na UI.
 4. UI ima više nedostataka nego prednosti.
 5. Radujem se budućem razvoju UI.
 6. UI nudi rješenja za mnoge svjetske probleme.
 7. Više volim tehnologije koje ne uključuju UI.
 8. Bojim se UI.
 9. Radije bih odabrao/la tehnologiju s UI nego onu bez nje.
 10. UI stvara probleme umjesto da ih rješava.
 11. Kada razmišljam o UI, pretežno imam pozitivne osjećaje.
 12. Radije bih izbjegavao/la tehnologije temeljene na UI.
6. Koliko vam je bio težak zadatak pisanja eseja?
- a) Vrlo težak
 - b) Djelomično težak
 - c) Ni težak ni lagan
 - d) Djelomično lagan
 - e) Vrlo lagan
7. U kojoj mjeri ste pročitali znanstveni članak koji smo vam dali za ovaj zadatak?
- a) Pročitao/la sam cijeli članak
 - b) Pročitao/la sam dijelove članka relevantne za moj esej
 - c) Površno sam pregledao/la članak
 - d) Nisam pročitao/la članak
8. (Za grupu studenata koja je koristila ChatGPT) Biste li radije koristili drugi način pisanja eseja?
- Da – radije bih sam/a napisao/la esej
- Ne – svidio mi se ovaj način pisanja eseja

(Za studente koji su sami pisali esej) Biste li radije koristili drugi način pisanja eseja?

Da – radije bih koristio/la ChatGPT za pisanje eseja
Ne – svidio mi se ovaj način pisanja eseja

9. (Za grupu studenata koja je koristila ChatGPT) Jeste li koristili članak koji smo vam predložili za izradu prompta za ChatGPT?

- a) Da
- b) Ne

10. (Ako da) Možete li ukratko opisati kako ste koristili članak u kontekstu izrade prompta za ChatGPT? Dovoljan je kratak tekstualni odgovor.