

MEGATREND UNIVERZITET
PRAVNI FAKULTET

Doktorska disertacija

IMPLEMENTACIJA VEŠTAČKE INTELIGENCIJE I RAZMATRANJE
KORPORATIVNE BEZBEDNOSTI

Mentor:

Doc. dr Vladimir Lović

Kandidat

Marko Stanojević

Beograd, 2024.

Apstrakt

Veštačka inteligencija (AI) je alat koji se koristi za podršku ljudskoj radnoj snazi u optimizaciji radnih tokova i efikasnijim poslovnim operacijama. Primena AI obuhvata različite modifikacije poput korišćenja AI za automatizaciju zadataka koji su rutinski i koji se često ponavljaju, generisanje informacija zasnovanih na algoritmima mašinskog učenja, zatim za brzu obradu ogromne količine skupova podataka i izvlačenje zaključaka koji omogućavaju predviđanje budućih pravaca delovanja i rezultata baziranih na analizi prikupljenih podataka. Sistemi AI usmeravaju razvijanje automatizacije poslovanja, uključujući automatizaciju samih poslovnih subjekata ali i automatizaciju pojedinačnih procesa, pri čemu štite poslovne subjekte od ljudskih grešaka i obezbeđuju dodatni prostor ljudima da pređu na obavljanje poslovnih zadataka višeg nivoa.

Međutim, pristup korišćenju AI i dalje se karakteriše podeljenim mišljenjima kada su u pitanju stručnjaci i kreatori politike naročito kada se radi o tome kako formulisati politiku veštačke inteligencije i kako je primeniti u odnosu na to kakav balans ona treba da postigne između različitih vrednosti, kao što je proaktivan rad da spreči štetu kroz vladinu intervenciju i zauzimanje popustljivijeg pristupa za podsticanje eksperimentisanja i inovacija.

Stalni napredak AI dovodi do veće upotrebe i većeg uticaja na to kako radimo, živimo i doživljavamo svet. Pored toga, Kako programeri čine značajne korake unapređujući AI, važno je razmotriti implikacije politika koje se danas donose na budući napredak i primenu tehnologije.

S obzirom na to da pojava AI i delovanje u crnoj ili sivoj zoni što se tiče regulative, predstavlja bojazan za izlazak AI izvan zadovoljavajućih okvira bezbednosti za ljudsko društvo, Evropska komisija je predložila širi regulatorni okvir za harmonizaciju politike veštačke inteligencije u državama članicama, poznat kao Zakon o veštačkoj inteligenciji (AI Act), koji prevazilazi samo standarde za stvaranje zakonskih i tehničkih obaveza za potencijalno štetnu veštačku inteligenciju.

Očigledno je da pravna regulativa još uvek nije u dovoljnoj meri regulisala primenu veštačke inteligencije i da je neophodna pravna intervencija u ovoj oblasti, u cilju zaštite korporativne bezbednosti.

Ključne reči: veštačka inteligencija (AI), poslovni zadaci, pravna regulativa, primena tehnologije, korporativna bezbednost

Abstract

Artificial intelligence (AI) is a tool used to support the human workforce in optimizing workflows and making business operations more efficient. The application of AI includes various modifications such as using AI to automate tasks that are routine and frequently repeated, generating information based on machine learning algorithms, then quickly processing huge amounts of data sets and drawing conclusions that allow predicting future courses of action and results based on analysis . collected data. AI systems guide the development of business automation, including the automation of business entities themselves, but also the automation of individual processes, while protecting business entities from human errors and providing additional space for people to move on to higher-level business tasks.

However, the approach to the use of AI is still characterized by divided opinion when it comes to experts and policy makers, especially when it comes to how to formulate an AI policy and how to implement it in relation to what balance it should achieve between different values, such as proactively working to prevent harm through government intervention and taking a permissive approach to encourage experimentation and innovation.

The constant advancement of AI is leading to greater use and greater impact on how we work, live and experience the world. In addition, as developers make significant strides in advancing AI, it is important to consider the policy implications that are being brought today to the future advancement and application of the technology.

Considering that the emergence of AI and its operation in the black or gray zone as far as regulation is concerned, represents a fear of AI going beyond the satisfactory security framework for human society, the European Commission has proposed a broader regulatory framework for the harmonization of artificial intelligence policy in member states, known as The Artificial Intelligence Act (AI Act), which goes beyond just standards to create legal and technical obligations for potentially harmful artificial intelligence.

It is obvious that legal regulations have not yet sufficiently regulated the application of artificial intelligence and that legal intervention in this area is necessary, in order to protect corporate security.

Keywords: artificial intelligence (AI), business tasks, legal regulation, technology application, corporate security

SADRŽAJ

Apstrakt	2
Abstract	3
SADRŽAJ	4
UVOD	6
PRVI DEO: NAUČNO ISTRAŽIVANJE.....	13
1.1. Predmet istraživanja	13
1.2. Cilj istraživanja	16
1.3. Definisane hipoteze.....	16
1.4. Metode istraživanja	17
1.5. Naučna i društvena opravdanost istraživanja	17
DRUGI DEO: VEŠTAČKA INTELIGENCIJA	19
2.1. Određivanje pojma	19
2.2. Ekspanzija veštačke inteligencije	21
2.3. Pristup definisanju veštačke inteligencije.....	22
2.4. Veštačka inteligencija i upravljanje poslovanjem	24
2.5. Aspekti mašinskog učenja	26
2.5.1. Duboko učenje	29
2.5.2. Obrada prirodnog jezika	30
2.5.3. Kompjuterski vid	32
TREĆI DEO: UTICAJ VEŠTAČKE INTELIGENCIJE NA DRUŠTVO I POSLOVNI SVET	34
3.1. Potencijal veštačke inteligencije.....	34
3.2. Poslovni menadžment.....	36
3.2.1. Podrška menadžmenta i otpor promenama	40
3.2.2. Strateško usklađivanje organizacije sa digitalizacijom.....	43
3.2.3. Izvršno digitalno liderstvo i organizacioni kapacitet za digitalizaciju.....	44
3.3. Uticaj novih tehnologija na korporativnu bezbednost	47

GLAVA ČETVRTA: IMPLEMENTACIJA AI I PRISTUP PRAVNOJ REGULATIVI	50
4.1. Najava regulative.....	50
4.2. Razmatranje regulatornog pristupa.....	51
4.2.1. Čvrsto pravo vs meko pravo	53
4.3. Procena Opšte regulative	56
4.4. Regulativa veštačke inteligencije i poređenja između SAD i EU	58
GLAVA PETA: PRAVNA REGULATIVA U DOMENU SAJBER BEZBEDNOSTI.....	62
5.1. Pravna rešenja i strategije sajber bezbednosti	62
5.2. Haktivizam i kriminalizacija	63
5.3. Postupak, dokazi i harmonizacija zakona.....	64
GLAVA ŠESTA: AI U PRAVNIČKOJ STRUCI - DEFINISANI I TRENUTNO KORIŠTENI PROGRAMI	70
6.1. Programi veštačke inteligencije kreirani za pravnu industriju	70
6.1.1. Napredne platforme za pravno istraživanje	71
6.1.2. Aplikacije za pravnu samopomoć	72
6.1.3. Pregled ugovora	72
6.2. Korišćenje AI za unapređenje kvaliteta pravnih usluga	73
6.3. Korišćenje AI za poboljšanje pristupa pravdi.....	77
6.4. Pravna etika	77
6.5. Pravci buduće regulative	78
6.6. Sajber bezbednost na nacionalnom nivou	81
6.8. Bezbednosni aspekti unutar Evropske unije.....	83
6.8.1. Odgovornost za štetnu upotrebu AI	84
6.8.2. Pristrasnost u korišćenju AI i Anti-diskriminacija.....	86
SEDMI DEO: REZULTATI EMPIRIJSKOG ISTRAŽIVANJA	88
ZAKLJUČNA RAZMATRANJA.....	179
LITERATURA	187

UVOD

U krugovima donosilaca odluka na globalnom nivou prisutne su dileme o tome koja zakonska regulativa je prihvatljiva za obuhvat AI: Odnosno, kada bi trebalo primeniti čvrsti, a kada meki zakon za upravljanje AI. Evidentno je da postoje različiti pristupi kada se razmatra ovo pitanje. Postoji opravdana bojazan da bi tvrdi zakon mogao usporiti tehnološki napredak. S druge strane pojedinci se zalažu za kontrolu sistema veštačke inteligencije jer smatraju da primena AI može da predstavlja pretnju po zdravlje ili bezbednost korisnika. Sledeća, ne tako beznačajna dilema odnosi se na rešavanje toga da li i kada politika treba da primeni čvrsti ili meki zakon u slučajevima kada je nivo rizika nepoznat ili neizvestan.

Razmatranje ovih otvorenih pitanja jesu glavni fokus mnogih zainteresovanih strana iz različitog društvenog miljea: vlade, civilnog društva, akademske zajednice i industrije. Razmatranje različitih pristupa dovodi do suprotstavljenih argumenata koji se sučeljavaju prilikom izrade i primene AI alata ili politike veštačke inteligencije i direktno utiču na identifikovanje načina za suzbijanje štetnih pristrasnosti. Rešavanje štetne pristrasnosti veštačke inteligencije zahteva višestruku strategiju.

Zakon o veštačkoj inteligenciji u okviru EU prati niz drugih politika o tehnologiji i podacima koje je Evropska komisija predložila i usvojila u protekloj deceniji. EU je osmislila ove zakone s ciljem da omogućí da se regulišu i usmeravaju neke operacije kompanija, i to tako da svi deluju po ovim pravilima na nivou cele Evropske unije. Mnogi od donetih akata imaju indirektan uticaj na razvoj AI.

AI zavisi od velikog obima podataka da bi efikasno funkcionisala, tako da ograničenja u prikupljanju i korišćenju podataka mogu da ugroze funkcionisanje AI sistema. S obzirom na to da različiti zakoni i propisi utiču jedni na druge, donosioci zakona ne bi trebalo da zanemare dupliranje zakonskih rešenja, zatim različite nedoslednosti i da nastoje da u zakonskim regulativama otklone praznine.

U aprilu 2021. Evropska komisija je predložila prvi regulatorni okvir EU za AI. Zakon kao polaznu osnovu za definisanje koristi klasifikaciju prema riziku koji AI sistemi predstavljaju za korisnike i u odnosu na takvo određenje opredeljuju se za više ili manje regulacije sistema veštačke inteligencije koji se mogu koristiti u različitim aplikacijama. Šta Prioritet EU

parlamenta usmeren je na to da se donese takva regulativa koja bi trebalo da osigura da sistemi veštačke inteligencije, koji se koriste u EU, budu bezbedni, transparentni, sledljivi, nediskriminatorni i ekološki prihvatljivi. Istovremeno, posebno je značajan stav da sisteme veštačke inteligencije treba da nadgledaju ljudi, a ne da se to odvija automatski, jer ljudska kontrola može da spreči štetne ishode. Parlament je takođe svestan da je teoretski pristup definisanju AI neujednačen i želi da uspostavi tehnološki neutralnu, jednoobraznu definiciju za veštačku inteligenciju koja bi se mogla primeniti na buduće sisteme veštačke inteligencije.

Realnost načina na koji su sistemi veštačke inteligencije dizajnirani (u suštini to je „crna kutija“) i implementirani znači da se postojeći standardi odgovornosti moraju promeniti.

AI će nesumnjivo imati transformativne efekte na ekonomiju, jer su čak su i računari, kada su prvi put predstavljeni, značajno promenili mnoge uslove zapošljavanja. Postoje ekonomski i društveni aspekti kojima se mora upravljati odgovarajućom regulativom AI.

Ovi rizici i štete, između ostalog, razlog su zašto se vlade širom sveta bore da uspostave odgovarajuće kontrole i smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju.

Regulisanje veštačke inteligencije je teško zbog potrebe da se uspostavi ravnoteža između omogućavanja inovacije i evolucije, dok se AI koristi na bezbedan i odgovoran način. Jedna stvar je jasna a to je da omogućavanje nespontanog razvoja veštačke inteligencije može biti neodrživo.

Postoji niz pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI. Vrlo malo je rečeno o tome koja organizacija u lancu razvoja AI treba da popravi takva pitanja kada se pojave, otvarajući jaz u odgovornosti. Životni ciklus veštačke inteligencije je složen i uključuje mnogo različitih operatera, od kojih svi imaju potencijal da utiču na kvalitet tehnologije.

Postoje, takođe, problemi sa samim konceptom vlasništva nad podacima. Zakonski ne postoji jedinstvena definicija ili okvir za vlasništvo nad podacima, a različite jurisdikcije mogu usvojiti različite pristupe u zavisnosti od prirode, izvora i upotrebe podataka. Dopisivanje prava vlasništva podataka podacima koje koristi ili generiše AI predstavlja brojne pravne izazove, uključujući: 1. Podaci koje obrađuje, generiše ili koristi AI često uključuju preklapanje interesa različitih strana. To takođe nije jasna linearna podela prava, već zavisi od prirode podataka, doprinosa strana koje su obezbedile, kreirale ili kontrolisale podatke i ugovornih ili pravnih

aranžmana između tih strana. Na primer, ko poseduje podatke ili sadržaj koji generiše sistem veštačke inteligencije koji je sakupio podatke iz javnog izvora? Ko poseduje podatke ili sadržaj koji se generiše korišćenjem nejavnih podataka, npr. platforma zasnovana na pretplati? Ovo će se nesumnjivo onda okrenuti razmatranjima licenciranja, kao i pravima na privatnost i poverljivost. Kao i kod većine pitanja veštačke inteligencije, još uvek nema utvrđenog stava po ovom pitanju.

2. Kada se utvrdi da podaci koje koristi ili kreira AI zapravo mogu biti predmet vlasničkih prava, onda se postavlja pitanje kako se ta prava mogu sprovesti. Uopšteno govoreći, ugovorni podaci i poverljivost i pravične obaveze poverljivosti i prava intelektualne svojine pružaju osnovu i okvir za uspostavljanje prava vlasnika podataka da ograniči upotrebu i otkrivanje podataka (osim ličnih podataka). Ali čak i tada, ostaje otvoreno pitanje kako ovi vlasnici podataka mogu da kontrolišu pristup, otkrivanje, prenos, brisanje ili modifikaciju podataka koji se nalaze u „crnoj kutiji“?

AI je definisan na nekoliko načina tokom svog postojanja, ali generalno AI je „sposobnost mašine da imitira inteligentno ljudsko ponašanje.“¹ Ovo uključuje „prepoznavanje govora i objekata, donošenje odluka na osnovu podataka i prevođenje jezici.“² AI može da oponaša ljudsku inteligenciju na dva načina: prvo, AI programi mogu biti obučeni unosom podataka, što je istorijski bio jedini način na koji su AI programi mogli da oponašaju ljudsku inteligenciju; drugo, napredniji AI programi mogu sami da uče putem pokušaja i grešaka. Iako koncept veštačke inteligencije postoji već duže vreme, upotreba naprednijih programa u pravnoj oblasti beleži tek neznatan razvoj. Početak regulisanja AI u pravnoj oblasti još uvek je skroman.³ Bilo je samo nekoliko kompanija koje su razvile AI programe prilagođene pravnim poslovima, a programske funkcije su uglavnom bile ograničene na eDiscoveri, pregled ugovora i dužnu pažnju.

¹ *Artificial intelligence*, MERRIAM-WEBSTER, <https://www.merriam-webster.com/dictionary/artificial%20intelligence> [https://perma.cc/3ZBP-PLAJ]

² Lauri Donahue, *A Primer on Using Artificial Intelligence in the Legal Profession*, HARV. J. L. & TECH. (Jan. 3, 2018), <http://jolt.law.harvard.edu/digest/a-primer-on-using-artificial-intelligence-in-the-legal-profession> [https://perma.cc/H65H-6A5A]

³ Keith Mullen, *Artificial Intelligence: Shiny Object? Speeding Train?*, AM. BAR ASS'N RPTE EREPORT 2018 FALL ISSUE, https://www.americanbar.org/groups/real_property_trust_estate/publications/ereport/rpte-2018/artificial-intelligence/ [https://perma.cc/HCE7-WNVV]

Tehnološki razvoj povezan sa sajber bezbednošću i sajber kriminalom znači da, dok se tradicionalni zakoni mogu donekle primeniti, zakonodavstvo takođe mora da se nosi sa novim konceptima i objektima, kao što su nematerijalni kompjuterski podaci, koji se tradicionalno ne obrađuju zakonom.

Ova dela u velikoj meri karakterišu novi nematerijalni objekti, kao što su podaci ili informacije. Iako se krivično pravo često smatra najrelevantnijim kada je u pitanju sajber kriminal, pravni odgovori na širu zabrinutost u vezi sa sajber bezbednošću takođe uključuju upotrebu drugih grana prava, kao što su građansko pravo i upravno pravo.

Pitanje kriminalizacije nepoželjnog ponašanja na internetu ima dvostruki efekat:

(a) stvaranje pravne osnove za retributivno suzbijanje ponašanja, i (b) stvaranje klime društvene neprihvatljivosti sajber kriminala, deromantizovanje i stigmatizovanje takvog ponašanja. Oni koji koriste internet za činjenje zločina odrasli su i bili su socijalizovani u klimi u kojoj je preovlađujući vid nezakonitih aktivnosti bio kriminal u stvarnom svetu, u svojim tradicionalnim obličjima,⁴ dok hakovanje, na primer, ne izaziva osećaj društvene neprihvatljivosti. Trenutno, hakersko ponašanje karakteriše „laissez-faire“ stav prema odgovornosti i zakonitosti u mnogim jurisdikcijama na globalnom nivou. U tom smislu, konceptualizacija sajber kriminala kao krivičnog dela unapređuje i odmazdu i odvratanje.

Prekogranična priroda sajber kriminala i činjenje sajber kriminala u elektronskom okruženju su glavne poteškoće sa kojima se suočavaju organi za sprovođenje zakona.

Izazovi u istrazi sajber kriminala proizilaze iz kriminalnih inovacija počinitelaca, poteškoća u pristupu elektronskim dokazima i iz internih resursa, kapaciteta i logističkih ograničenja. Osumnjičeni često koriste tehnologije anonimizacije i prikrivanja, a nove tehnike brzo prolaze do široke kriminalne publike preko onlajn tržišta kriminala.⁵

Dakle, na proceduralnom nivou, glavni problem za nacionalnu primenu zakona je pomirenje istorijske činjenice da je policija operativno i organizaciono lokalizovana unutar granica jurisdikcije (inherentno povezana sa državnim suverenitetom), dok su sajber bezbednost i sajber kriminal globalizovani i poremećene jurisdikcije.

⁴ Brenner Susan (2004). Toward a Criminal Law for Cyberspace: Distributed Security, BOSTON UNIVERSITY JOURNAL OF SCIENCE & TECHNOLOGY LAW. Vol. 10, Issue 1, pp. 1-109.

⁵ Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

Istrage kibernetičkog kriminala, od strane organa za sprovođenje zakona, zahtevaju spajanje tradicionalnih i novih policijskih tehnika. Dok se neke istražne radnje mogu postići tradicionalnim ovlašćenjima, mnoge proceduralne odredbe se ne prevode dobro sa prostornog, objektno orijentisanog pristupa na pristup koji uključuje elektronsko skladištenje podataka i tokove podataka u realnom vremenu.⁶

Dokaz je sredstvo kojim se utvrđuju činjenice relevantne za krivicu ili nevinost pojedinca na suđenju. Elektronski dokazi su sav takav materijal koji postoji u elektronskom, ili digitalnom, obliku, koji može biti uskladišten ili prolazan i može postojati u obliku računarskih datoteka, prenosa, evidencije, metapodataka ili mrežnih podataka.

Dodatni izazov je upotreba tehnologije od strane organa za sprovođenje zakona – čini se da u ovom trenutku sajber prestupnici bolje koriste tehnološke mogućnosti koje imaju. Prisustvo relevantnog tela sa posebnim znanjem i stručnošću u okviru policijskih snaga je ključni element u efikasnom regulisanju sajber prostora.⁷ Tužioci moraju da vide i razumeju elektronske dokaze kako bi izgradili slučaj na suđenju. Mnoge zemlje u razvoju, na globalnom nivou, nemaju dovoljno resursa kako bi tužioci to i učinili. Kompjuterske veštine tužilaštva su obično niže nego kod istražitelja. Isto važi i za sudije koje se bave visoko specijalizovanim predmetima sajber kriminala. Obuka pravosuđa o zakonu o sajber kriminalu, prikupljanju dokaza i osnovnom i naprednom poznavanju računara predstavlja poseban prioritet.⁸ Mnoge jurisdikcije ne prave pravnu razliku između elektronskih dokaza i fizičkih dokaza.⁹ Iako se pristupi razlikuju, mnoge zemlje smatraju ovo dobrom praksom, jer obezbeđuje pravičnu prihvatljivost pored svih drugih vrsta dokaza. Brojne zemlje van Evrope uopšte ne prihvataju elektronske dokaze, zbog čega je krivično gonjenje sajber kriminala, kao i bilo kog drugog zločina dokazanog elektronskim informacijama, neizvodljivo. Iako zemlje generalno nemaju posebna pravila o elektronskim dokazima, brojne zemlje su se pozvale na principe kao što su: pravilo o najboljim dokazima,

⁶ Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

⁷ Više o tome: WALL DAVID (2007). CYBERCRIME: THE TRANSFORMATION OF CRIME IN THE INFORMATION AGE. Polity Press.

⁸ Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

⁹ Brenner Susan (2004). Toward a Criminal Law for Cyberspace: Distributed Security, BOSTON UNIVERSITY JOURNAL OF SCIENCE & TECHNOLOGY LAW. Vol. 10, Issue 1, pp. 1-109.

relevantnost dokaza, pravilo iz druge ruke, autentičnost i integritet, od kojih svi mogu imati posebnu primenu na elektronske dokaze.¹⁰

Mnoge zemlje imaju elemente pravnog okruženja koji se bavi sajber bezbednošću i sajber kriminalom, ali se ovi nacionalni pravni okviri uveliko razlikuju u pogledu načina na koji se ova pitanja rešavaju.¹¹ U današnjem globalizovanom svetu, zakon se sastoji od mnoštva nacionalnih, regionalnih i međunarodnih pravnih sistema. Interakcije između ovih sistema se javljaju na više nivoa. Kao rezultat toga, odredbe su ponekad u suprotnosti jedna sa drugom, što dovodi do sukoba zakona, ili se ne preklapaju u dovoljnoj meri, ostavljajući praznine u nadležnostima. Ove razlike između nacionalnih zakona dovode do pitanja da li, i ako jesu, koliko daleko se nacionalne pravne razlike u zakonima o sajber kriminalu mogu i kako ih smanjiti? Drugim rečima, koliko je važno uskladiti zakone o sajber kriminalu? Ovo se može preduzeti na više načina, uključujući i obavezujuće i neobavezujuće međunarodne ili regionalne inicijative. Osnova harmonizacije može biti jedinstveni nacionalni pristup.¹²

Harmonizacija takođe može omogućiti globalno prikupljanje dokaza. Usklađivanje procesnog prava je drugi neophodan uslov za efikasnu međunarodnu saradnju. Eksterna validacija podataka i jasnost o nameravanoj upotrebi veštačke inteligencije pomaže da se obezbedi bezbednost i olakša regulacija. Posvećenost kvalitetu podataka, kao što je rigorozno procenjivanje sistema pre objavljivanja, je od vitalnog značaja da se obezbedi da sistemi ne pojačavaju pristrasnosti i greške.

Sajber prostor je često suočen sa bezbednosnim opasnostima. Poseban problem predstavlja pravna regulativa koja pokriva analognu oblast i susreće se sa nerešivom problematikom kada se zahteva njena primena u vreme digitalizacije. Rešavanje ovih problema je pred izazovima posebno kada se radi o sukobu nacionalne i međunarodne regulative, kada države reaguju na dva načina: ili se suprotstavljaju primeni međunarodnih odredbi jer brane

¹⁰ Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

¹¹ Više o tome: Satola David & Judy Henry (2011). Towards a Dynamic Approach to Enhancing International Cooperation and Collaboration in Cybersecurity Legal Frameworks: Reflections on the Proceedings of the Workshop on Cybersecurity Legal Issues at the 2010 United Nations Internet Governance Forum 37 WILLIAM MITCHELL LAW REVIEW.

¹² Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

svoju suverenost ili prihvataju intervenciju koju smatraju pogodnom za rešavanje nastalog problema. Činjenica je da virtuelni prostor ne može da bude usko vezan za teritorijalnost.¹³

Još uvek nije postignut zadovoljavajući nivo izgrađenosti međunarodnih pravnih normi koje bi garantovale bezbednost virtuelne teritorije u okviru koje dolazi do povrede bezbednosti. Svi dosadašnji pokušaji doveli su samo do parcijalnih rešenja što se pokazalo kao nedovoljno.¹⁴

Zbog nedostatka potpune i delotvorne međunarodne regulacije države često pribegavaju samostalnoj primeni odredbi domaćeg zakonodavstva i delu međunarodnog prava, i to u oblastima koje se posredno odnose na ovu problematiku.¹⁵ Ista situacije je evidentna i u oblasti međunarodnog ugovornog prava.¹⁶ Činjenica je da se i u ovoj ugovornoj oblasti radi samo o parcijalnim rešenjima.¹⁷

Ono što je evidentno u sadašnjem ugovornom međunarodm pristupu jeste nedostatak strateškog pristupa sveobuhvatnom rešavanju problema.

Pošto je pretnja sajber napadima na KI strateškog obima, nacionalni odgovor mora biti jednak zadatku, a to podrazumeva podizanje svesti javnosti, ulaganje u obrazovanje, naučna istraživanja, razvoj sajber prava i međunarodne saradnje. Nauka i tehnologija igraju važnu ulogu u međunarodnoj bezbednosti, ali iako moderna informaciona i komunikaciona tehnologija (IKT) nudi civilizaciji „najšire pozitivne mogućnosti“, IKT je ipak podložna zloupotrebi od strane kriminalaca i terorista.

¹³ Ibid

¹⁴ Ibid

¹⁵ Više o tome: Schmitt Michael (2013). Tallinn Manual on the International Law Applicable to Cyber Warfare. Cambridge University Press.

¹⁶ Stahl, William M. "The Uncharted Waters of Cyberspace: Applying the Principles of International Maritime Law to the Problem of Cybersecurity." Georgia Journal of International and Comparative Law. Vol. 40. (2011): 247-274.

¹⁷ Više o tome: Satola David & Judy Henry (2011). Towards a Dynamic Approach to Enhancing International Cooperation and Collaboration in Cybersecurity Legal Frameworks: Reflections on the Proceedings of the Workshop on Cybersecurity Legal Issues at the 2010 United Nations Internet Governance Forum 37 WILLIAM MITCHELL LAW REVIEW.

PRVI DEO: NAUČNO ISTRAŽIVANJE

1.1. Predmet istraživanja

Ne postoji univerzalno prihvaćena definicija AI. Termin koji je skovao kompjuterski naučnik Džon Makarti, nastao je 1956. Termin se često odnosi na inteligentne tehnologije sa mogućnostima koje oponašaju kognitivne sposobnosti slične ljudima. Neka vladina tela su pokušala da definišu veštačku inteligenciju,¹⁸ ali industrija, akademici i grupe civilnog društva često koriste veoma različite definicije kako bi zadovoljile svoje potrebe. AI često uključuje softver koji se kreće od razumevanja ljudskog govora (kao što su Aleka i Siri) do predviđanja koje slike ili sadržaja želite da vidite (kao što je Facebook News Feed) do obrade medicinskih rendgenskih zraka kako bi se pomoglo u otkrivanju bolesti. Definisanje šta čini AI sistem je ključna tačka AI ACT-a. Mnoge zainteresovane strane koje su davale povratne informacije o početnom nacrtu Zakona o veštačkoj inteligenciji izrazile su zabrinutost da definicija AI u predlogu (koji sadrži kompletnu listu tehnika i pristupa navedenih u Aneksu I,¹⁹ uključujući tehnologije „mašinskog učenja” zasnovane na softveru, „logički- i sistemi zasnovani na znanju, i „statistički pristupi“)²⁰ obuhvata previše sistema i može dovesti do određene nesigurnosti.²¹ Zainteresovane strane nastavljaju debatu o definiciji AI.

Generalno, zagovornici uže definicije smatraju da su jednostavni algoritmi ili tehnologija sortiranja izvan dometa AI i stoga bi se suočio sa neopravdanim regulatornim opterećenjima prema ovom zakonu.²² Zagovornici šire definicije tvrde da bi trebalo da obuhvati svaki

¹⁸ Institute of Electrical and Electronics Engineers, IEEE Position Statement of Artificial Intelligence, June 24, 2019. Available at: <https://globalpolicy.ieee.org/wp-content/uploads/2019/06/IEEE18029.pdf>.

¹⁹ European Commission, *Annexes to the Proposal for Regulation of the European Parliament and Council*, April 21, 2021. Available at: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52021PC0206>.

²⁰ Nick Gondek, “Learning about Machine Learning,” Bipartisan Policy Center, August 4, 2021. Available at: <https://bipartisanpolicy.org/explainer/learning-about-machine-learning>.

²¹ Big Data Value Association, *BDVA Position Paper: Response to the European Commission’s proposal for AI Regulation*, August 4, 2021. Available at: https://www.bdva.eu/sites/default/files/BDVA_DAIRO%20response-feedback%20AI%20Regulation_Final.pdf.

²² Centre for Data Ethics and Innovation, United Kingdom Government, “The European Commission’s Artificial Intelligence Act Highlights the Need for an Effective AI Assurance Ecosystem,” May 11, 2021. Available at: <https://cdei.blog.gov.uk/2021/05/11/the-european-commissions-artificial-intelligence-act-highlights-the-need-for-an-effective-ai-assurance-ecosystem>.

„automatski sistem zasnovan na softveru“ koji bi „mogao da utiče na osnovna prava“.²³ Kako Evropska komisija i drugi kreatori politike vremenom traže kompromis o definiciji sistema veštačke inteligencije, treba da imaju na umu složenost AI i dugovečnost zakonodavstva. AI tehnologija će se razvijati i možda se neće uklopiti u definicije koje su napravljene za AI sisteme danas.

Stručnjaci se slažu da je pristup zasnovan na riziku neophodan za regulisanje AI, ali se ne slažu oko detalja i implementacije. Pristup zasnovan na riziku uključuje kategorizaciju AI sistema ili slučajeva upotrebe na osnovu nivoa rizika koji predstavljaju i nametanje različitih propisa i standarda u skladu sa tim. Prema ovom pristupu, kreatori politike i regulatori mogu imati stroža pravila za sprečavanje štete u oblastima visokog rizika i dozvoljena pravila za oblasti sa niskim rizikom kako bi smanjili opterećenje usklađenosti i promovisali eksperimentisanje. Pristup zasnovan na riziku se ogleda u predlozima propisa i smernica u EU i Sjedinjenim Državama.

Zakon o veštačkoj inteligenciji EU razlikuje regulatorne obaveze na osnovu rizika koji sistem veštačke inteligencije predstavlja, bilo da stvara neprihvatljiv rizik, visok rizik, ograničen rizik ili nizak ili minimalan rizik. Kao što je navedeno u Zakonu o veštačkoj inteligenciji, prakse koje potpadaju pod neprihvatljiv rizik su potpuno zabranjene. Oni uključuju sisteme veštačke inteligencije koji krše osnovna prava, podsvesno manipulišu ljudima, iskorišćavaju ranjive grupe kao što su deca, verovatno izazivaju psihičke ili fizičke povrede, vrše društveno bodovanje za korišćenje od strane javnih vlasti ili prikupljaju u realnom vremenu udaljene biometrijske identifikacione podatke u javnim prostorima radi zakona svrhe izvršenja. AI sistemi klasifikovani pod visoki rizik se suočavaju sa strogim obavezama i primenom prema Zakonu o veštačkoj inteligenciji, dok se sistemi ograničenog, niskog i minimalnog rizika suočavaju sa manje strogim obavezama ili bez njih. U Sjedinjenim Državama, prethodne verzije predloženih zakona, koje su predstavili brojni demokratski članovi, klasifikovali su visokorizične sisteme veštačke inteligencije i predložili propise specifične za tu kategoriju.²⁴ Pored toga, NIST, vladina

²³ AlgorithmWatch and Access Now, “EU Policy Makers: Protect People’s Rights, Don’t Narrow Down the Scope of the AI Act!” November 23, 2021. Available at: <https://algorithmwatch.org/en/statement-scope-of-eu-ai-act/>

²⁴ U.S. Congress,, *Algorithmic Accountability Act of 2019*, H.R.223, 116th Congress, introduced in House April 11, 2019. Available at: <https://www.congress.gov/bill/116th-congress/house-bill/2231/text>.

agencija SAD odgovorna za razvoj okvira za upravljanje rizikom od veštačke inteligencije, radi na uspostavljanju okvira zasnovanog na riziku za upravljanje sistemima veštačke inteligencije.²⁵

Tehnološki razvoj povezan sa sajber bezbednošću i sajber kriminalom znači da, dok se tradicionalni zakoni mogu donekle primeniti, zakonodavstvo takođe mora da se nosi sa novim konceptima i objektima, kao što su nematerijalni kompjuterski podaci, koji se tradicionalno ne obrađuju zakonom.

Ova dela u velikoj meri karakterišu novi nematerijalni objekti, kao što su podaci ili informacije. Iako se krivično pravo često smatra najrelevantnijim kada je u pitanju sajber kriminal, pravni odgovori na širu zabrinutost u vezi sa sajber bezbednošću takođe uključuju upotrebu drugih grana prava, kao što su građansko pravo i upravno pravo.

Pitanje kriminalizacije nepoželjnog ponašanja na internetu ima dvostruki efekat:

(a) stvaranje pravne osnove za retributivno suzbijanje ponašanja, i (b) stvaranje klime društvene neprihvatljivosti sajber kriminala, deromantizovanje i stigmatizovanje takvog ponašanja. Oni koji koriste internet za činjenje zločina odrasli su i bili su socijalizovani u klimi u kojoj je preovlađujući vid nezakonitih aktivnosti bio kriminal u stvarnom svetu, u svojim tradicionalnim obličjima,²⁶ dok hakovanje, na primer, ne izaziva osećaj društvene neprihvatljivosti. Trenutno, hakersko ponašanje karakteriše „laissez-faire“ stav prema odgovornosti i zakonitosti u mnogim jurisdikcijama na globalnom nivou. U tom smislu, konceptualizacija sajber kriminala kao krivičnog dela unapređuje i odmazdu i odvratanje.

Prekogranična priroda sajber kriminala i činjenje sajber kriminala u elektronskom okruženju su glavne poteškoće sa kojima se suočavaju organi za sprovođenje zakona.

Izazovi u istrazi sajber kriminala proizilaze iz kriminalnih inovacija počinitelaca, poteškoća u pristupu elektronskim dokazima i iz internih resursa, kapaciteta i logističkih ograničenja. Osumnjičeni često koriste tehnologije anonimizacije i prikrivanja, a nove tehnike brzo prolaze do široke kriminalne publike preko onlajn tržišta kriminala.²⁷

²⁵ National Institute of Science and Technology, *AI Risk Management Framework*, 2022. Available at: <https://www.nist.gov/itl/ai-risk-management-framework>

²⁶ Brenner Susan (2004). *Toward a Criminal Law for Cyberspace: Distributed Security*, BOSTON UNIVERSITY JOURNAL OF SCIENCE & TECHNOLOGY LAW. Vol. 10, Issue 1, pp. 1-109.

²⁷ Više o tome: Sandage, John et al. eds. (2013). *Comprehensive Study on Cybercrime* (United Nations Office on Drugs and Crime).

Dakle, na proceduralnom nivou, glavni problem za nacionalnu primenu zakona je pomirenje istorijske činjenice da je policija operativno i organizaciono lokalizovana unutar granica jurisdikcije (inherentno povezana sa državnim suverenitetom), dok su sajber bezbednost i sajber kriminal globalizovani i poremećene jurisdikcije.

1.2. Cilj istraživanja

Iz predmeta istraživanja proizilaze naučni i društveni cilj istraživanja.

Primarni cilj istraživanja odnosi se na istraživanje i analizu uzroka i posledica veštačke inteligencije u savremenom svetu.

Naučni cilj istraživanja ogleda se u naučnoj deskripciji, naučnoj klasifikaciji i utvrđivanju empirijskih varijabli uticaja uzroka i posledica primena veštačke inteligencije u kontekstu pozitivnih i negativnih aspekata s kojima se ljudsko društvo suočava i ukazivanje na pravce mogućih zakonskih regulativa u ovoj oblasti.

Društveni cilj istraživanja ogleda se u analizi stanja o pojavama i procesima, kao što su: aktuelna primena veštačke inteligencije, razmatranje parametara i ukazivanje na smer za buduće regulatorne aktivnosti, kao i primene AI.

1.3. Definisane hipoteze

U toku pripreme za sprovođenje empirijskog istraživanja doktorske disertacije postavljena je jedna glavna i tri pomoćne hipoteze.

H0: Ukoliko je zakonodavstvo uspostavilo odgovarajuću kontrolu i dalo jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju utoliko je manja mogućnost da upotreba AI dovede do štetnih posledica.

H1: Ukoliko je regulisanje veštačke inteligencije uspostavljeno tako da je postignuta ravnoteža između regulisanja kontrole AI i omogućavanja inovacije i evolucije, utoliko je veća mogućnost da se AI koristi na bezbedan i odgovoran način.

H2: Ukoliko podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana utoliko je veća mogućnost da se otvore pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI.

H3: Ukoliko je zakonodavni pristup definisan oslanjanjem na procenu stepena rizika od upotrebe AI utoliko je veća mogućnost da se AI sistemi visokog rizika suoče sa strogo definisanim načinima delovanja.

1.4. Metode istraživanja

Tokom izrade disertacije korišćiće se sledeće istraživačke metode:

Metod analize sadržaja se primenjuje u cilju analize svih pisanih tragova koji su u vezi sa međunarodnim aspektom uvođenja pravne regulative u oblast primene AI. Glavna prednost ove metode je mogućnost primene dostupnih podataka o problemima, nedostatak je mali izvor podataka u vezi sa međunarodnim aspektom zaštite prava izbeglica iz ugla bezbednosti države.

Metod sinteze primenjivaće se kako bi se sistematizovala znanja po zakonitostima formalne logike, kao proces izgradnje teorijskog znanja u pravcu od posebnog ka opštem, u vezi sa međunarodnim aspektom zaštite prava u oblasti primene AI.

Induktivna metoda primenjivaće se kako bi se na osnovu analize pojedinačnih činjenica došlo do zaključka o opštem sudu, kao i zapažanja konkretnih pojedinačnih slučajeva primenom, koja je dovela do opštih zaključaka. Glavna prednost metode je što je moguće analizirati različite činjenice koje se odnosi na uticaj primene AI na multikulturizam i bezbednost. Možda i glavni nedostatak ove metode je činjenicu da postoji veliki broj pojedinačnih slučajeva, i da je važno odabrati pravi, kako bi se došlo do adekvatnih zaključaka.

1.5. Naučna i društvena opravdanost istraživanja

Rezultati istraživanja imaju, kako naučni, tako i društveni doprinos. U radu su prikazani efekti primene AI na ljudsko društvo i to s aspekta korisnosti, ali i ukazivanjem na moguće štetne posledice, kao i analiza mogućih regulatornih rešenja ovog savremenog dostignuća, ali i problema sa kojim se suočava savremeni svet. To će biti glavni doprinos ovog istraživanja.

Naučni doprinos se ogleda u potvrđivanju i proširivanju dosadašnjih naučnih saznanja o korišćenju AI, ali i o posledicama koja su moguće bez adekvatne zakonske regulative, koja bi trebalo da pomiri zahteve koji se odnose na bezbednost korisnika, ali i društva u celini, a da se istovremeno ne onemogući dalji razvoj ove tehnologije.

Društveni doprinos se ogleda u realnom sagledavanju uzroka, posledica i mogućih rešenja do kojih se dolazi primenom AI.

DRUGI DEO: VEŠTAČKA INTELIGENCIJA

2.1. Određivanje pojma

Kao što smo u uvodnom delu već naveli, ne postoji univerzalno prihvaćena definicija AI. Ovaj termin prvi put pominje kompjuterski naučnik Džon Makarti, još 1956. Termin se oslanja na inteligentne tehnologije sa mogućnostima koje oponašaju kognitivne sposobnosti slične ljudima. Neka vladina tela su pokušala da definišu veštačku inteligenciju,²⁸ međutim, industrija, akademici i grupe civilnog društva često koriste veoma različite definicije kako bi zadovoljile svoje potrebe. AI se često posmatra kao softver koji se kreće od razumevanja ljudskog govora do predviđanja koje slike ili sadržaje želite da vidite, ili do obrade medicinskih rendgenskih zraka kako bi se pomoglo u otkrivanju bolesti. Očigledno je da je precizno definisanje AI sistema neophodno. U toku rasprave o nacrtu Zakona o veštačkoj inteligenciji povratne informacije ukazale su na zabrinutost da definicija AI, kako se predlaže u Aneksu I,²⁹ proširuje polje delovanja AI na tehnologije „mašinskog učenja” zasnovane na softveru, logičke sisteme i sisteme zasnovane na znanju, kao i na statistički pristup³⁰ i da samim tim obuhvata previše sistema i može dovesti do određene nesigurnosti, jer se udaljava od osnovnog koncepta.³¹

Zbog iznetih neslaganja debata zainteresovanih strana o definiciji AI se nastavlja. Uglavnom su određena dva pristupna gledišta i to: zagovornici uže definicije koji smatraju da su jednostavni algoritmi ili tehnologija sortiranja izvan dometa AI i stoga su nepotrebni u definisanju AI i predstavljali bi nepotrebno regulatorno opterećenje;³² Zagovornici šire definicije

²⁸ Institute of Electrical and Electronics Engineers, IEEE Position Statement of Artificial Intelligence, June 24, 2019. Available at: <https://globalpolicy.ieee.org/wp-content/uploads/2019/06/IEEE18029.pdf>.

²⁹ European Commission, *Annexes to the Proposal for Regulation of the European Parliament and Council*, April 21, 2021. Available at: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52021PC0206>.

³⁰ Nick Gondek, “Learning about Machine Learning,” Bipartisan Policy Center, August 4, 2021. Available at: <https://bipartisanpolicy.org/explainer/learning-about-machine-learning>.

³¹ Big Data Value Association, *BDVA Position Paper: Response to the European Commission’s proposal for AI Regulation*, August 4, 2021. Available at: https://www.bdva.eu/sites/default/files/BDVA_DAIRO%20response-feedback%20AI%20Regulation_Final.pdf.

³² Centre for Data Ethics and Innovation, United Kingdom Government, “The European Commission’s Artificial Intelligence Act Highlights the Need for an Effective AI Assurance Ecosystem,” May 11, 2021. Available at: <https://cdei.blog.gov.uk/2021/05/11/the-european-commissions-artificial-intelligence-act-highlights-the-need-for-an-effective-ai-assurance-ecosystem>.

tvrdi da bi definisanje AI trebalo da obuhvati svaki „automatski sistem zasnovan na softveru“ koji bi „mogao da utiče na osnovna prava“.³³

Ovi suprotstavljeni pristupi naveli su Evropsku komisiju i sve ostale donosiocce odluka da traže kompromis o definiciji sistema veštačke inteligencije, insistirajući na tome da ne sme da se zanemari složenost AI. Činjenica je da će bilo kakvo definisanje AI vrlo brzo biti prevaziđeno razvojem ove tehnologije. To bi značilo da definisanje AI mora da bude podvrgnuto periodičnom redefinisaniu. Periodično redefinisanie podrazumevalo bi prilagodljivost tehnološkom napretku AI.

Ono u čemu se sve zainteresovane strane slažu odnosi se na da je pristup zasnovan na riziku neophodan za regulisanje AI, ali se ne slažu oko detalja i implementacije. Pristup zasnovan na riziku kategorizuje AI sisteme i upotrebu na osnovu nivoa rizika koji predstavljaju i smatraju da propisi i standardi moraju da se definišu u skladu sa tim. To podrazumeva definisanje strožih pravila za sprečavanje štete u oblastima visokog rizika i liberalnija pravila za oblasti sa niskim rizikom. Potpora ovakvom stavu je argumentacija da bi se tako smanjilo opterećenje usklađenosti i omogućilo dalje eksperimentisanje u ovoj oblasti. Pristup zasnovan na riziku primenjen je u predlozima propisa i smernica u EU i Sjedinjenim Državama. Zakon o veštačkoj inteligenciji EU implementira regulatorne obaveze na osnovu rizika koji sistem veštačke inteligencije predstavlja, bilo da stvara neprihvatljiv rizik, visok rizik, ograničen rizik ili nizak ili minimalan rizik, pri čemu su prakse koje potpadaju pod neprihvatljiv rizik zabranjene. Pod ovakvim praksama smatra se delovanje veštačke inteligencije koje dovodi do kršenja osnovnih prava, podsvesna manipulacija ljudima, zloupotreba ranjivih kategorija, poput dece i sl. Prakse sa neprihvatljivim rizikom mogu da dovedu do psihičkih ili fizičkih povreda, ali i do neovlašćenog prikupljanja biometrijskih identifikacionih podataka od strane javnih vlasti.

AI sistemi visokog rizika se suočavaju sa strogo definisanim načinima delovanja prema Zakonu o veštačkoj inteligenciji, dok se sistemi ograničenog, niskog i minimalnog rizika suočavaju sa manje strogim obavezama ili bez njih. U Sjedinjenim Državama, prethodne verzije predloženih zakona, koje su predstavili brojni demokratski članovi, klasifikovali su visokorizične

³³ AlgorithmWatch and Access Now, “EU Policy Makers: Protect People’s Rights, Don’t Narrow Down the Scope of the AI Act!” November 23, 2021. Available at: <https://algorithmwatch.org/en/statement-scope-of-eu-ai-act/>

sisteme veštačke inteligencije i predložili propise adekvatne za tu kategoriju.³⁴ Pored toga, NIST, vladina agencija SAD, odgovorna za razvoj okvira za upravljanje rizikom od veštačke inteligencije, i dalje radi na uspostavljanju okvira zasnovanom na riziku za upravljanje sistemima veštačke inteligencije.³⁵

2.2. Ekspanzija veštačke inteligencije

Pojava alternativne inteligencije koja ne potiče od ljudi već je produkt tehnologije ne prolazi vez uznemirenosti u smislu kako će se odraziti na same ljude. Ne mogu da se anuliraju pozitivne strane ove surogat inteligencije koja je posebno pozitivna kao podrška i pomoć upravljačkim procesima posebno u oblastima koji se bave proizvodnjom, kao i u uslužnim oblastima.³⁶

Veštačka inteligencija u ovim oblastima doprinosi kvalitetnije donošenju odluka na načinin koji se karakteriše saradnjom ljudi i digitalnih alata. Takva saradnja daje pozitivne efekte kada je u pitanju napredak i preciznost u delovanju. Međutim, to definitivno vodi preusmeravanju pojedinih poslovnih funkcija sa ljudi na mašine.³⁷

Veštačka inteligencija uspešno akumulira bezbroj podataka tražeći zakonitost u ponavljanju operacija i tako doprinosi analizi i reviziji poslovnih poteza. Postoji tendencija da veštačka inteligencija preuzme komunikacijski kanal na liniji preduzeće-klijent, ali je ova aktivnost još uvek u začetku.³⁸

³⁴ U.S. Congress,, *Algorithmic Accountability Act of 2019*, H.R.223, 116th Congress, introduced in House April 11, 2019. Available at: <https://www.congress.gov/bill/116th-congress/house-bill/2231/text>.

³⁵ National Institute of Science and Technology, *AI Risk Management Framework*, 2022. Available at: <https://www.nist.gov/itl/ai-risk-management-framework>

³⁶ Stevenson, W. J. (2020). *Operations Management*. McGraw Hill.

³⁷ Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard business review*, 96(1), 108-116.

³⁸ Ibid

2.3. Pristup definisanju veštačke inteligencije

Veštačka inteligencija su mašine koje su stvorile sposobnost da rade same i uče na prirodan način kao što to čine ljudi.³⁹ AI tehnologija može se koristiti za rad u različitim oblastima, koja ne deli sličnosti, kao što su zdravstvo, poslovanje, autonomna vozila, prepoznavanje slika, povećanje piksela i još mnogo toga.⁴⁰ Procenjuje se da će tržište veštačke inteligencije porasti za 50% do 2025. Mašinsko učenje obuhvata niz različitih metoda stvaranja veštačke inteligencije, različite složenosti i mogućnosti upotrebe. Jedna od osnovnih metoda mašinskog učenja je statističko učenje. Statističko učenje se gradi direktno na osnovu podataka i modela predviđanja i direktno je povezano sa podacima koje analizira.⁴¹

Model predviđanja se može promeniti kako bi odgovarao različitim potrebama, kao što su filtriranje e-pošte i predviđanje srčanih bolesti, sa velikom razlikom što se nefiltrirana e-pošta može prihvatiti, u poređenju sa srčanom bolešću koja se mora otkriti. Prilikom analize kvantitativnih podataka koristi se regresivna analiza, a pri analizi kvalitativnih koristi se klasifikaciona analiza.⁴² To znači da se statističko učenje može obaviti na brojne načine, sa lakšim i složenijim modelima predviđanja. Jedan od najčešće korišćenih modela statističkog predviđanja su linearni modeli i model najmanjeg kvadrata. Linearni model se sastoji od jednostavne formule, koja postoji dva puta, prvo kao samostalna pristrasnost od ulaznih podataka, koja je konstanta, a zatim i druga pristrasnost koja je povezana sa analizom ulaznih podataka.⁴³ Da bi se pojednostavio linearni model, može se koristiti model najmanjeg kvadrata, koji svodi na minimum razlike u ulaznim podacima, gde je pristrasnost odabrana da bi odgovarala određenim kriterijumima. Ove metode učenja osiguravaju složenije statističke metode, kao što su neuronske mreže.

³⁹ Sajja, P. S. (2021). Introduction to Artificial Intelligence. Illustrated Computational Intelligence, Springer: 1-25.

⁴⁰ Shaw, J., et al. (2019). "Artificial intelligence and the implementation challenge." Journal of medical Internet research 21(7): e13659.

⁴¹ Više o tome: Hastie, T., et al. (2009). The elements of statistical learning: data mining, inference, and prediction, Springer Science & Business Media.

⁴² Ibid

⁴³ Ibid

Mašina zasnovana na neuronskoj mreži je algoritam koji prati način na koji mozak uči iz različitih inputa, zašto se naziva neuronskim mrežama.⁴⁴ Neuralne mreže su postale naglašeni metod učenja, koji su mnogi skloni da definišu kao složen i misteriozan metod učenja. Najjednostavnije verzije neuronskih mreža se sastoje od jednog skrivenog sloja sa propagacijom unazad, ili perceptrona jednog sloja, što znači da podaci koji idu napred kroz algoritam mogu ići i unazad, što je način na koji neuronska mreža može poboljšati svoj način rada..⁴⁵ Neuronske mreže su nelinearni statistički modeli,⁴⁶ odnosno dvostepena klasifikacija ili regresiona analiza, što znači da može primeniti i regresiju i klasifikaciju, zbog čega su pogodne za kvalitativnu i za kvantitativnu analizu i široku upotrebu. Delovanje se odvija kao linija od dna ka vrhu, pri čemu postoji nelinearna funkcija, sa specifičnom pristrasnošću za tu funkciju.⁴⁷

U procesu implementacije, bez ikakvog povratnog širenja, kreirana neuronska mreža će dati odgovore nasumično, koji neće biti zadovoljavajući. Da bi se poboljšala zadovoljavajuća stopa odgovora, potrebno je širenje unazad. Povratno širenje hrani neuronsku mrežu tačnim odgovorima. Tačni odgovori se vraćaju unazad u neuronsku mrežu i neuronska mreža menja svoje funkcije, uključujući i pristrasnost, kako bi se uklopila u nove izlaze ispravnog materijala, čineći je sposobnom da daje bolje odgovore kada se primenjuju novi podaci. Ovde zaključujemo da će rezultati biti bolji, što se više podataka unese u neuronsku mrežu.⁴⁸ Metod mašinskog učenja zahteva kvalitetne podatke, jer u suprotnom ispravna implementacija je otežana.⁴⁹ Nasuprot tome, zadovoljavajuća količina podataka omogućava da implementacija mašinskog učenja može postati uspešna. U slučaju da je prisutan nedostatak podataka metod mašinskog učenja uči od nule, sa mogućim pozitivnim ishodom, ali sa evidentiranim nizom nedostataka. 1. Potrebno je prikupljanje mnogo podataka da bi metod ujednačeno funkcionisao. 2. Potrebno je mnogo vremena za postavljanje baze podataka i dolazi do grešaka, zbog čega je neophodna kontrola od strane iskusnih inženjera. 3. Mala razlika, odnosno odstupanje u obrascu ulaznih

⁴⁴ Više o tome: 3Blue1Brown (2018). Neural Networks & Deep Learning. 3Blue1Brown.

⁴⁵ Ibid

⁴⁶ Više o tome: Hastie, T., et al. (2009). The elements of statistical learning: data mining, inference, and prediction, Springer Science & Business Media.

⁴⁷ Više o tome: 3Blue1Brown (2018). Neural Networks & Deep Learning. 3Blue1Brown.

⁴⁸ Ibid

⁴⁹ Više o tome: Holzinger, A. (2018). From machine learning to explainable AI. 2018 world symposium on digital intelligence for systems and machines (DISA), IEEE.

podataka može da vrati ceo proces učenja na početak. 4. Za ove operacije neophodno je mnogo memorije i računarske snage. 5. Ovakva vrsta procesa karakteriše se odsustvom transparentnosti, tako da iako mašina može da radi kao uvučena, postaje teško osigurati da mašina radi kako je predviđeno.⁵⁰ Zbog toga je podučavanje ljudi kako AI funkcioniše važno za uspešnu primenu AI. Prvi korak za uspešnu implementaciju AI rešenja kroz mašinsko učenje je strategija podataka. U ovoj strategiji potrebna je jedna oblast etičke upotrebe AI.⁵¹

2.4. Veštačka inteligencija i upravljanje poslovanjem

Dobro je poznato da će uspon AI imati snažan uticaj na upravljanje poslovima i organizacijama u bliskoj budućnosti, a to je dovelo do interesovanja analitičara u ovoj oblasti.⁵² Veštački menadžment podrazumeva primenu mašina ili unapred programiranih logičkih sklopova u administrativnom, menadžerskom ili organizacionom rešavanju problema, i omogućava donošenje odluka na nižim, srednjim, ali i na najvišim nivoima korporativne hijerarhije.⁵³ To ukazuje da su kompjuteri, ekspertski sistemi i roboti zamenjeni kriterijumima podržanim, unapred programiranim i logičnim odlučivanjem.⁵⁴

Nisu retka predviđanja da će implementacija AI pomoći u efikasnom i sistematičnom upravljanju organizacijama, sa ciljem da se poboljša njihov učinak, rezultati i profit. U tom kontekstu, očekuje se da napredak u veštačkoj inteligenciji značajno reformisati radno okruženje, i uticati na redefinisane istraživačkog iskustva menadžera.⁵⁵

Vrhunski menadžeri provode većinu svog vremena u organizovanju i planiranju, dok menadžeri srednjeg i nižeg nivoa provode više vremena na kontroli i vođenju. Istraživanje o

⁵⁰ Ibid

⁵¹ Ibid

⁵² Chopra, K. N. (2018). Overview and application of artificial intelligence concepts and some important coevolving modern issues on management of organization and commerce. *Journal of Internet Banking and Commerce*, 23(3), 1-22. <https://search.proquest.com/docview/2216881012?accountid=27513>

⁵³ Shim, J. K., & Rice, J. S. (1988). Expert systems applications to managerial accounting. *Journal of Systems Management*, 39(6), 6. <https://search.proquest.com/docview/199850554?accountid=27513>

⁵⁴ Ibid

⁵⁵ Chopra, K. N. (2018). Overview and application of artificial intelligence concepts and some important coevolving modern issues on management of organization and commerce. *Journal of Internet Banking and Commerce*, 23(3), 1-22. <https://search.proquest.com/docview/2216881012?accountid=27513>

uticaju veštačke inteligencije u menadžmentu koju je sproveo Institut Accenture za visoke performanse (AIHP) i Accenture Strategy, gde je učestvovalo 1770 menadžera sa niskog, srednjeg i najvišeg nivoa, pokazalo je da će računarski sistemi AI najverovatnije preuzeti menadžerske funkcije i to posebno u oblasti koju čine zadaci, kao što su koordinacija, kontrola, saradnja i rešavanje problema. Benefit se u ovom slučaju ostvaruje jer menadžmentu tako ostaje više vremena koje mogu da usmere na inovacije, razvoj strategije i na ljudski faktor unutar organizacije.⁵⁶

Nesporno je da će veštačka inteligencija ući u poslovne subjekte u ekstremnim razmerama i da će AI promeniti sve nivoe upravljanja. AI će automatizovati zakazivanje, izveštavanje i dodelu resursa, dok će analitika, testiranje hipoteza i simulacija uz pomoć veštačke inteligencije biti izuzetno efikasni za strateško donošenje odluka i inovacije u celom poslovnom sistemu. Iako AI predstavlja mogućnosti za stvaranje vrednosti, ona takođe ima izazove za rukovodioce jer će biti primorani da preispitaju svoje uloge i preispitaju osnovne principe rada koji trenutno vode njihovu organizaciju. AI upućuje na neophodnost povećanja saradnje ljudi sa mašinama što neminovno dovodi do promena u podeli rada i do potrebe da kompanije moraju da pojačaju svoje strategije učinka, obuke i sticanja talenata kako bi se preusmerile na rad koji zavisi od ljudske procene i veština, zajedno sa eksperimentisanjem i saradnjom sa mašinama.⁵⁷

Kako je objasnio profesor MIT Sloan Erik Brinolfson, poznati stručnjak za menadžment, teškoća sa kojom se danas suočavamo nije objašnjena kao svet bez dela, već kao svet u kome se rad stalno menja. Shodno tome, menadžeri će morati da ostanu na nogama i da nastave da preduzimaju napredne korake kako bi kontinuirano evoluirali kompanije.⁵⁸ Ovde dolazimo do zaključka da će AI promeniti sve nivoe upravljanja, jer direktno utiče na koordinaciju, kontrolu, ali i na automatizaciju saradnje i rešavanje problema.

Predviđa se da će prodor AI biti mnogo veći do 2030. Međutim, zamena ljudi se smatra lošom stranom. Beker i saradnici su uvereni da će rešenja veštačke inteligencije postati široko

⁵⁶ Više o tome: Kolbjørnsrud, V., Amico, R., Thomas, R. J. (2015). The promise of artificial intelligence: Redefining management in the workforce of the future. Accenture.

⁵⁷ Ibid

⁵⁸ Chopra, K. N. (2018). Overview and application of artificial intelligence concepts and some important coevolving modern issues on management of organization and commerce. Journal of Internet Banking and Commerce, 23(3), 1-22. <https://search.proquest.com/docview/2216881012?accountid=27513>

primenjena i korišćena u upravljanju rizicima i usklađenosti sa propisima.⁵⁹ Čopra je zaključio da se u bliskoj budućnosti očekuju nova važna dostignuća u oblasti veštačke inteligencije u menadžmentu. gde se postiže novi napredak u razvoju neuronskih mreža, kvantnog računanja i tehnika simulacije. Nalazi koji su rezultirali istraživanjima u ovoj kategoriji uključuju činjenicu da je AI u menadžmentu važno i neophodno polje za organizacije koje treba uzeti u obzir, posebno ako žele da budu u stanju da se takmiče sa kompanijama kao što su Google, Apple, Microsoft, IBM itd.⁶⁰ Zanimljivo je da studije u ovoj oblasti potiču iz Kanade, Indije i Rumunije i da je svim studijama zajedničko da se bave efektima veštačke inteligencije na menadžment najvišeg nivoa.

2.5. Aspekti mašinskog učenja

Protekla decenija obeležena je stvaranjem i korišćenjem umreženih i mobilnih računarskih uređaja zahvaljujući masovnoj upotrebi pametnih telefona i drugih moćnijih računarskih uređaja. Pored toga, internet stvari (IOT), je stvorio značajan uticaj, proširivši se na skoro svaki uređaj koji probija put do sitnih komponenti različitih faza u proizvodnom procesu. Moderni primeri uključuju Soni i njegove digitalizovane igračke konzole i televizore, bežično kontrolisane Nest termostate i GE-ovu upotrebu senzora u medicinskim proizvodima, između ostalog.⁶¹

Sa ovim poslovičnim morem informacija dolazi i pitanje analize i korišćenja podataka. Tu dolaze veliki podaci. Definicija velikih podataka je sledeća: „veliki podaci se odnose na naprednu analizu podataka i brzo znanje iz ogromnih količina različitih skupova podataka, uključujući i strukturirane i nestrukturirane podatke. Ova tehnologija predstavlja analitičku digitalnu platformu koju karakterišu tri varijacije: zapremina, raznovrsnost i brzina.” Ove ogromne količine podataka se zatim šalju u oblak. Računarstvo u oblaku je model koji

⁵⁹ Becker, M., & Buchkremer, R. (2018). Ranking of current information technologies by risk and regulatory compliance officers at financial institutions – a German perspective. *The Review of Finance and Banking*, 10(1) <https://search.proquest.com/docview/2208135738?accountid=27513>

⁶⁰ Chopra, K. N. (2018). Overview and application of artificial intelligence concepts and some important coevolving modern issues on management of organization and commerce. *Journal of Internet Banking and Commerce*, 23(3), 1-22. <https://search.proquest.com/docview/2216881012?accountid=27513>

⁶¹ Više o tome: Spremić, M., Ph.D. (2017). *Enterprise information systems in digital economy*. Zagreb.

omogućava sveprisutan, zgodan i pristup podacima na zahtev preko mreže. Za rad su potrebne mreže, serveri, gomila skladišta, aplikacije i usluge. Krajnji rezultat je fleksibilan, rekonfigurabilan, efikasan i pristupačan rad.⁶²

Mašinsko učenje podrazumeva kodiranje računara da se ponašaju kao ljudski mozak umesto da ih učimo onome što znamo. On omogućava računarima pristup velikim podacima i omogućava im da izdvoje važne karakteristike kako bi rešili komplikovane probleme. Iz ovoga proizilazi važno pitanje: kako možemo kreirati računarske sisteme koji se automatski poboljšavaju kroz iskustvo? Takve metode su poznate kao sistemi nadgledanog učenja. U suštini, takva obuka se sastoji od prikazivanja primera ponašanja input-output kroz upotrebu ogromnih količina visoko specifičnih i relevantnih označenih podataka. Računar otkriva razlike između primera i koristi ih kao merila za procenu daljih primera. Ovaj pristup se koristi u dubokom učenju. Koristi milijarde parametara koji se zatim obučavaju na velikim zbirkama slika, uzoraka govora ili video zapisa. Profesori Džordan i Mičel u svom radu na mašinskom učenju dali su primer korišćenja ogromnog broja istorijskih kupovina kreditnim karticama. Neki su bili lažni, a drugi su bili legitimni.. Na sistemu je bilo da uči iz primera i da na kraju donese sopstvene odluke. Ovaj jednostavan problem binarne klasifikacije gde su moguće kombinacije ograničene na „prevaru“ i „legitimno“ nije jedini tip označavanja, jer postoje i klasifikacije sa više oznaka sa uključivanjem problema rangiranja i opštih problema strukturisanog predviđanja. Ovakve metode se koriste u prepoznavanju govora uz istovremeno označavanje reči u rečenicama.⁶³

Drugi metod se zove učenje bez nadzora i podrazumeva analizu neobebeženih podataka uz pretpostavku o strukturnim svojstvima podataka. Mašina je predstavljena sa ogromnom količinom nesortiranih informacija koje postepeno sortira prema njihovim osnovnim sličnostima, karakteristikama ili obrascima. Može se izvršiti grupisanje velikog uzorka podataka u manje grupe sličnih podataka, ili se može analizirati putem asocijacije, gde mašina nastoji da otkrije pravila koja opisuju velike delove podataka. Primer povezivanja AI u maloprodaji je lista „često kupuju zajedno“. Takav mehanizam za preporuke koristi koncepte u realnom vremenu kao što su kolaborativno filtriranje, filtriranje zasnovano na sadržaju ili hibridni sistemi preporuka za

⁶² Ibid

⁶³ Jordan, M. I., & Mitchell, T. M. (2015). Machine learning: Trends, perspectives, and prospects. *Science*, 349(6245), 255-260.

analizu određenih kupovina i povezivanje proizvoda. Ova vrsta veštačke inteligencije je sveprisutna za skoro svaku e-trgovinu, ili čak običnu radnju, kao što je Ikea, gde su često kupljeni proizvodi postavljeni blizu jedan drugom.⁶⁴

Treći metod se zove učenje uz pomoć. Odnosi se na algoritme orijentisane ka složenom cilju ili maksimiziranjem određene dimenzije, na primer maksimiziranjem poena osvojenih u igri. Polazna tačka je prazna tabla, slična detetu, a njegovo delovanje se metodom „šargarepe i štapa“ podstiče na postizanje određenog cilja kroz model ponašanja. U psihologiji, takve akcije se nazivaju učenjem potkrepljenja. Ova metoda je korišćena u kombinaciji sa dubokim učenjem za obuku Google-ovog AlphaGo AI sistema koji je na kraju pobedio najboljeg Go igrača na svetu. Učenje sa pojačanjem je prilično slično učenju pod nadzorom, ali sa retkim povratnim informacijama.⁶⁵

Mašinsko učenje je oblast sa mnoštvom neistraženih ili nedovoljno istraženih mogućnosti delovanja. Njegova metodologija učenja zasnovana je na interdisciplinarnim otkrićima koji kombinuju znanja iz oblasti sociologije, psihologije, inženjerstva i drugih kako bi se postavili temelji veštačke inteligencije. Međutim, mašinsko učenje je fokusirano na veoma specifične poduhvate učenja suprotno širokim i sveobuhvatnim naporima koje postiže ljudski mozak. Postoje sporovi oko prikupljanja, korišćenja i vlasništva ličnih podataka koji se prikupljaju u svrhu analize. Poslednjih godina neki od titana informacionih tehnologija prikupljaju, obrađuju i čuvaju mnoštvo novih korisničkih podataka radi ekonomskog profita. Primer za ovu je delovanje kompanija poput Gugla i Fejsbuka koje prodaju korisničke podatke kako bi finansirale svoje proizvode. Dobar primer korisnika koji imaju koristi od ugradnje različitih tehnologija u isporuku poruke prikazan je u sledećoj konstataciji: „Kombinovanjem podataka o lokaciji iz onlajn izvora (na primer, podataka o lokaciji sa mobilnih telefona, iz transakcija kreditnim karticama u maloprodajnim objektima i sa sigurnosnih kamera na javnim mestima i privatnim zgradama) sa onlajn medicinskim podacima (npr. prijem u hitnu pomoć), omogućava se implementacija sistema koji bi odmah telefonirao pojedincima ako je osoba sa kojom su juče bili u bliskom kontaktu upravo primljena u hitnu pomoć sa zaraznom bolešću, upozoravajući ih na simptome na koje treba da paze i mere predostrožnosti koje treba da preduzmu.” Takva

⁶⁴ Ibid

⁶⁵ Ibid

razmatranja nas navode da verujemo u pozitivan, transformativni efekat takve tehnologije u 21. veku i dalje.⁶⁶

2.5.1. Duboko učenje

Duboko učenje je jedna od vrsta mašinskog učenja, odnosno to je tehnika koja omogućava računarskim sistemima da se unaprede kroz iskustvo i prikupljene podatke. Autori Goodfellow, Bengio i Courville tvrde da je mašinsko učenje jedina održiva opcija za stvaranje sposobnih AI mašina koje mogu da rade u složenom okruženju u stvarnom svetu, jer se duboko učenje, kao posebna vrsta mašinskog učenja, najbolje uklapa između velike snage i dovoljno fleksibilnosti. Duboko učenje je sposobno da predstavi svet u obliku ugneždenih hijerarhija koncepata. To znači da je svaki koncept definisan u odnosu na jednostavnije koji čine celinu. Ove operacije se izvode kroz mnoštvo vidljivih i nevidljivih računskih slojeva. Zovu se tako što su varijable lako uočljive, dok se ostale izdvajaju iz subjekata na sve apstraktniji način.⁶⁷

Ključ njegovog modusa operandi leži u Stohastičkom Gradijentnom spuštanju (SGD). Sastoji se od predstavljanja ulaznog vektora na nekoliko primera, izračunavanja izlaza i grešaka, izračunavanja prosečnog gradijenta i prilagođavanja pondera. Ovaj proces se ponavlja za veliku količinu malih skupova primera i daje procenu prosečnog gradijenta u svim primerima. Da bi izračunali gradijent funkcije, istraživači su koristili tehnike povratnog širenja što se pokazalo uspešnim. Kombinacija rezultira višeslojnom neuronskom mrežom koja ima nekoliko vidljivih varijabli koje interaguju sa brojnim skrivenim varijablama kroz međusobno povezanu mrežu. Duboko učenje je danas najčešće korišćeni način primene AI. Njegov potencijal se najbolje vidi na primerima kompjuterskog vida, prepoznavanja govora i mašinskog prevođenja u primerima aplikacije TensorFlow. TensorFlow je najrasprostranjenija platforma za primenu AI u praksi. Osim ove platforme, značajne su i druge platforme za duboko učenje koje uključuju Caffe istraživača veštačke inteligencije Berklija i H2O.ai. Caffe je platforma za duboko učenje

⁶⁶ Ibid

⁶⁷ Više o tome: Goodfellow, I., Bengio, Y., & Courville, A. (2017). Deep learning. Cambridge, MA: MIT Press.

napravljena brzinom, ekspresijom i modularnošću kao svetla vodilja na putu razvoja veštačke inteligencije. H2O je AI platforma otvorenog koda sa misijom demokratizacije inteligencije.⁶⁸

2.5.2. Obrada prirodnog jezika

Obrada prirodnog jezika je metoda mašinskog učenja koja daje uputstva računaru da identifikuje znakove i jezik. Duboko učenje je moćniji način da se veštačka inteligencija dovede u komunikaciju sa ljudima putem prirodnog jezika, govornog ili kucanog. Kao značajna sfera veštačke inteligencije, obrada prirodnog jezika (NLP) prvenstveno stimuliše računare da čuvaju, otkrivaju i prevode podatke koji se odnose na jezik u tekst i reči, što vodi u prepoznavanju prirodnih reči nalik čoveku. Tradicionalni ručni pristup ispitivanja podataka slobodnog teksta ima sklonost da ostavi mnogo korisnih informacija neiskorišćenim. U tu svrhu, NLP se postepeno koristi da zameni dugotrajnu ljudsku analizu, jer može da obradi ogromne količine tekstualnih podataka po niskoj ceni i u kratkom vremenskom periodu.

NLP ima mogućnost da prođe kroz veliki broj tekstualnih datoteka prikupljenih u CEM domenu kako bi pomogao u upravljanju građevinskim projektima i inženjerskim podacima. Korisne i pravne informacije mogu se efikasno izvući iz ogromne količine nestrukturiranih tekstualnih dokumenata korišćenjem obrade prirodnog jezika (NLP). Kao posledica toga, nadzornici mogu imati koristi od prethodnika incidenta kako bi poboljšali praćenje i evaluaciju bezbednosti izgradnje.

Sa jednog aspekta, opasne radnje i uzroci mogu se lako otkriti i kategorisati u preliminarnoj fazi, omogućavajući blagovremeno delovanje ljudi na gradilištima i praksama kako bi se smanjila opasnost. Takođe može dati tačne pretpostavke o mogućim pretnjama koje će se verovatno pojaviti u budućnosti i na taj način omogućava administratorima da preduzmu preventivne mere kako bi izbegli da se takvi incidenti ponovo dogode. Na primer, Zhong je razvio novu platformu koja meša duboko učenje i rudarenje teksta za otkrivanja opasnosti, koja bi takođe mogla da konstruiše mreže zajedničkog pojavljivanja reči kako bi sumirala moguće

⁶⁸ Ibid

trendove opasnosti i pratila dinamičku evoluciju opasnosti tokom vremena radi izbegavanja nezgoda.⁶⁹

Budući da je glavni cilj svakog poslovanja ostvarivanje prihoda, NLP direktno optimizovanjem poslovnih procesa, ili indirektno, povećanjem produktivnosti ili poboljšanjem zadovoljstva kupaca, omogućava povećanje prihoda. NLP je koristan jer nudi unosnije podatke i rešenja za poslovanje nego što zapravo košta, međutim da bi se to postiglo mora da se dizajnira sekvenca mašinskog učenja koja je specifična za konkretne poslovne procese. Neke kompanije su skeptične prema NLP-u i mašinskom učenju jer neznaju kako da pravilno iskoriste tehnologiju u svoju korist. Za poslovanje je ključno fokusiranje na specifične oblasti u kojima NLP može pomoći u uvođenju stvarnih, pozitivnih promena u smislu smanjenja troškova, povećanja produktivnosti i profita. NLP se uspešno koristi za preduzeća koja postavljaju složene lance snabdevanja za predviđanje potražnje, upravljanje zalihama, izbor dobavljača, ispunjenje porudžbina, sve faze korisničke podrške i još mnogo toga. Obučeni AI program koji pokreće lanac snabdevanja je jedan od sjajnih slučajeva poslovne upotrebe za NLP. Rezultati mogu biti apsolutno pozitivni u poslovanju i mogu da doprinesu kreiranju sofisticiranijih proizvoda koji će zadovoljiti zahteve potrošača koji se stalno menjaju.

Većina preduzeća koristi zastarele tehnologije pretraživanja zasnovane na podudaranju ključnih reči kako bi pomogla da pronađu nešto za svoje zaposlene, klijente i partnere. Međutim, ovakav pristup ne može da obezbedi dobijanje optimalnih rezultata. Primenom NLP-a na pretrage, rezultati se mogu očekivati ako postoji dobro razumevanje osnovnog značenja upita, što znači da samo unošenje ključnih reči nije dovoljno. Jedno od najjednostavnijih i najefikasnijih rešenja za bolju pretragu je Elasticsearch, koji može da se kombinuje sa modelom mašinskog učenja po izboru u Python-u. Tamo gde je potrebno nijansiranije u potrazi za informacijama, koristili se baza podataka uz ugradnju vektora. Četbotovi obučeni za NLP mogu pomoći u smanjenju troškova koji su obično povezani sa zadacima koji se ponavljaju. Mašinsko učenje nastavlja da poboljšava kapacitet čet-botova, i ljudi se sve više opredeljuju da koriste ove sisteme. Četbotovi obučeni za NLP imaju sposobnost da razumeju, analiziraju i daju prioritet pitanjima klijenata na osnovu konteksta i namere. Ovo im omogućava da brzo i tačno odgovaraju

⁶⁹ Zhong, B. X. Pan, P.E. Love, J. Sun, C. Tao, (2020). Hazard analysis: a deep learning and text mining framework for accident prevention, Adv. Eng. Inform. 46, 101152, <https://doi.org/10.1016/j.aei.2020.101152>.

na upite, i to znatno brže od redovnog predstavnika korisničke službe. Vremenom, ovo može izgraditi puno poverenja, vrednosti i kredibiliteta za poslovanje.⁷⁰

2.5.3. Kompjuterski vid

Kompjuterski vid se može definisati kao naučna oblast koja izdvaja informacije iz digitalnih slika. Tip informacija dobijenih od slike može da varira od identifikacije, merenja prostora za navigaciju ili aplikacija proširene stvarnosti. Drugi način da se definiše kompjuterski vid je kroz njegove aplikacije. Kompjuterski vid gradi algoritme koji mogu da razumeju sadržaj slika i koriste ga za druge aplikacije. Kompjuterski vid objedinjuje veliki skup disciplina.

Kompjuterski vid se može posmatrati kao deo kompjuterskih nauka, a teorija algoritama ili mašinsko učenje su od suštinskog značaja za razvoj algoritama kompjuterskog vida. Kompjuterski vid je još uvek je veoma težak problem, kome tek predstoji rešavanje. To je nešto što mi ljudi radimo nesvesno, ali je predstavlja težak zadatak za računare. Danas se još uvek borimo da kreiramo algoritme koji daju dobro oblikovane rečenice. Ono što ljudi nazivaju inteligencijom često nije dobar kriterijum za procenu težine računarskog zadatka. Danas je lakše napraviti 3D model objekta do milimetarske preciznosti nego izgraditi algoritam koji prepoznaje objekte. Prepoznavanje objekata je i dalje veoma težak problem, iako se približavamo ljudskoj tačnosti. Kompjuterski vid je težak jer postoji ogroman jaz između piksela i značenja. Šta računar vidi u a 200×200 RGB slika je skup od 120 000 vrednosti. Put od ovih brojeva do značajnih informacija je veoma težak.⁷¹

Kompjuterski vid je jedna od oblasti veštačke inteligencije koja omogućava računarima da razumeju vizuelni svet. Računari mogu da koriste digitalne slike i modele dubokog učenja da precizno identifikuju i klasifikuju objekte i reaguju na njih. Kompjuterski vid u AI je posvećen razvoju automatizovanih sistema koji mogu da interpretiraju vizuelne podatke (kao što su fotografije ili filmske slike) na isti način kao i ljudi. Ideja koja stoji iza kompjuterskog vida je da se kompjuteri upute da tumače i razumeju slike na bazi piksel po piksel. Ovo je osnova polja kompjuterskog vida. Što se tiče tehničke strane stvari, računari će nastojati da izvuku vizuelne

⁷⁰ Wood, T. (2024). Business uses of natural language processing - the 2024 capabilities of NLP. <https://fastdatascience.com/business-uses-natural-language-processing/>

⁷¹ Više o tome: Krishna, R. (2017). Computer Vision: Foundations and Applications. Stanford University.

podatke, upravljaju njima i analiziraju rezultate koristeći sofisticirane softverske programe. Za kompjuterski vid su potrebne ogromne količine informacija. Ponovljene analize podataka se vrše sve dok sistem ne može da razlikuje objekte i identifikuje vizuelne elemente.

Duboko učenje, specifična vrsta mašinskog učenja i konvolucione neuronske mreže su dve ključne tehnike koje se koriste za postizanje ovog cilja. Uz pomoć unapred programiranih algoritamskih okvira, sistem mašinskog učenja može automatski da uči o interpretaciji vizuelnih podataka. Model može naučiti da razlikuje slične slike ako mu je dat dovoljno veliki skup podataka. Algoritmi omogućavaju sistemu da uči sam, tako da može da zameni ljudski rad u zadacima kao što je prepoznavanje slika.

Konvolucione neuronske mreže pomažu mašinskom učenju i modelima dubokog učenja u razumevanju tako što dele vizuelne elemente na manje delove koji mogu biti označeni. Uz pomoć oznaka, on vrši konvolucije, a zatim koristi tercijarnu funkciju da bi dao preporuke o sceni koju posmatra.

Sa svakim ciklusom, neuronska mreža vrši konvolucije i procenjuje istinitost svojih preporuka i kao rezultat ovih poteza počinje da percipira i identifikuje slike kao čovek. Kompjuterski vid je sličan rešavanju slagalice u stvarnom svetu. Upravo tako rade neuronske mreže unutar kompjuterskog vida. Kroz niz filtriranja računari mogu spojiti sve delove slike i onda sami razmišljati. Kompjuter se često napaja hiljadama slika koje ga obučavaju da prepozna određene objekte. Ovo omogućava računaru da razume različite karakteristike koje čine određenu sliku i da je odmah prepozna.⁷²

⁷² Simplilearn (2023). What Is Computer Vision: Applications, Benefits and How to Learn It. <https://www.simplilearn.com/computer-vision-article>

TREĆI DEO: UTICAJ VEŠTAČKE INTELIGENCIJE NA DRUŠTVO I POSLOVNI SVET

3.1. Potencijal veštačke inteligencije

Zbog sve veće količine dostupnih podataka i promenljivih interesovanja i složenosti klijenata, moderna preduzeća se više ne oslanjaju na vremenske metode vođenja poslovanja kako bi podstakla ekspanziju. Ove dramatične promene, koje je omogućila veštačka inteligencija, otvorile su novi horizont potencijalnih rezultata koji se mogu iskoristiti za pokretanje korporativnog rasta putem neizbrisivih iskustava zasnovanih na podacima o klijentima.⁷³

Upotreba veštačke inteligencije u poslovanju podrazumeva upotrebu lažnog ljudskog inteligentnog računara da bi se podržao prihod, poboljšalo zadovoljstvo kupaca, povećala profitabilnost i podstakla poslovni rast i promene. AI ima pozitivan uticaj na poslovanje zbog automatizacije poslovnih procesa. Poslovni procesi koji se mogu automatizovati mogu se naći u menadžmentu, poslovnim fazama, lancu snabdevanja, ljudskim resursima i marketingu, između ostalog.⁷⁴ Na primer, u odeljenjima za ljudske resurse, mašinsko učenje podržava analizu velikih količina podataka, predviđanje ishoda i prepoznavanje obrazaca. Rekrutovanje zaposlenih u odeljenju za ljudske resurse je olakšano pregledom biografija, zakazivanjem intervjua i praćenjem aplikacija. Pored toga, chatbotovi imaju glavnu ulogu u procesu zapošljavanja tako što analiziraju odgovore kako bi utvrdili da li kandidat odgovara kompaniji ili ne.⁷⁵

AI omogućava poboljšanje korisničkog iskustva za preduzeća. Za mnoge kompanije, dosledna interakcija sa klijentima podstiče profit. Pošto je veštačka inteligencija svake godine sve pametnija, čet-botovi mogu da obavljaju sve više posla od ljudi. Oni mogu da komuniciraju sa klijentima i da se bave složenim problemima dok stavljaju manje tereta na IT resurse za upravljanje i održavanje kontakt centra, što rezultira poboljšanim korisničkim iskustvom.⁷⁶

⁷³ Virginia Tech admin. (2021). The impact of artificial intelligence on business: VT India, Virginia Tech India. https://vtindia.in/the-impact-of-artificial-intelligence-on-business/#How_does_AI_Impact_your_Business

⁷⁴ Lawton, G. & Tucci, L. (2022). What is business process automation, CIO? TechTarget. <https://www.techtarget.com/searchcio/definition/business-process-automation>

⁷⁵ Virginia Tech admin. (2021). The impact of artificial intelligence on business: VT India, Virginia Tech India. https://vtindia.in/the-impact-of-artificial-intelligence-on-business/#How_does_AI_Impact_your_Business

⁷⁶ MacAraig, M. 2019. 11 pros and cons of AI for Businesses, <https://www.insightsforprofessionals.com>

AI pomaže preduzećima da smanje operativne troškove. AI može uštedeti vreme, ukloniti jednostavne zadatke sa liste obaveza i učiniti zaposlene produktivnijim i vrednijim. Stoga će primena automatizacije veštačke inteligencije u kompaniji uštedeti vreme, smanjiti ljudsku grešku (koja košta) i poboljšati korisničko iskustvo. Međutim, i dalje su prisutni nedostaci prilikom pokretanja AI u preduzećima. Najznačajniji nedostatak je to što je implementacija AI skupa investicija, jer preduzeća moraju da investiraju u AI okvire, uključujući najnoviji hardver i softver. Timovi za obuku o tome kako da koriste sisteme veštačke inteligencije snose dodatne troškove. Sve ovo čini implementaciju i održavanje AI sistema skupim i dostupnim samo bogatim poslovnim subjektima.

Pošto je veštačka inteligencija postala popularnija nego ikada ranije, svi shvatamo veliki potencijal veštačke inteligencije, koja nam čini život lakšim i praktičnijim, međutim, činjenica je da ovaj sistem ima i nedostatke.

Pozitivni aspekti korišćenja veštačke inteligencije su:

Smanjenje ljudskih grešaka: Ljudi su skloni greškama kada obavljaju dosadne i ponavljajuće zadatke, dok pravilno programirani računari mogu da izbegnu takve greške. Modeli veštačke inteligencije prave prognoze primenom algoritama za prikupljanje podataka, čime se smanjuju greške i povećava preciznost.

Lako upravlja ogromnim podacima: U veoma kratkom vremenskom periodu, AI može da radi sa ogromnim količinama podataka. Može brzo da uhvati i izvuče odgovarajuće podatke za analizu. Pored toga, AI može tumačiti i transformisati ove podatke za dodatnu obradu.

Ublažava rizike: Primena veštačke inteligencije u opasnim okruženjima predstavlja značajnu prednost. Sistemi veštačke inteligencije mogu da smanje rizike povezane sa takvim zadacima za ljude. Na primer, roboti koji podržavaju veštačku inteligenciju mogu da obavljaju opasne zadatke kao što su iskopavanje uglja, istraživanje morskog života, pomoć u operacijama spasavanja od prirodnih katastrofa i tako dalje.

Otvaranje naprednijih radnih mesta: Prema „Izveštaju o budućnosti radnih mesta 2020“ koji je objavio Svetski ekonomski forum, procenjuje se da će do 2025. godine biti otvoreno 97 miliona novih radnih mesta u 26 zemalja. Ipak, ova novootvorena radna mesta će zahtevati nove veštine, što iziskuje značajna ulaganja u „dokvalifikaciju“ i

„prekvalifikaciju“ mladih i odraslih.⁷⁷ Uprkos raznim prednostima veštačke inteligencije, postoji nekoliko nedostataka koje treba imati na umu.

Nedostaci veštačke inteligencije

Nezaposlenost: Iako će veštačka inteligencija stvarati radna mesta u budućnosti, ona takođe eliminiše mnoge od tradicionalnih poslova koje trenutno obavljaju ljudi. Procena je da će AI eliminisati 85 miliona radnih mesta 2025.⁷⁸

Veća je verovatnoća da će povećati ljudsku lenjost: Kako se sve više zadataka automatizuje i digitalni asistenti postaju široko dostupni, verovatno će doći do povećanja ljudske lenjosti kao rezultat našeg sve većeg oslanjanja na ove alate. Sposobnost ljudi da obavljaju svakodnevne aktivnosti koje zahtevaju pamćenje ili analizu može biti negativno pogođena ako se na veštačku inteligenciju previše oslanja za jednostavne zadatke poput malih proračuna ili pamćenja brojeva ili adresa.⁷⁹

Međutim, ne može se poreći činjenica da je veštačka inteligencija poboljšala kvalitet našeg života, ali treba da budemo svesni i prednosti i mana ove tehnologije.

3.2. Poslovni menadžment

Uticaj veštačke inteligencije u poslovnom menadžmentu je značajan. On pokreće transformaciju i inovacije, što dovodi do poboljšane efikasnosti i efektivnosti za preduzeća. Uz algoritme za mašinsko učenje i prediktivnu analitiku, AI pruža vredne uvide koji omogućavaju menadžerima da donose odluke zasnovane na podacima.⁸⁰ Najbolji primer je Amazon, koji koristi veštačku inteligenciju za analizu ponašanja i preferencija kupaca, omogućavajući im da u skladu sa tim prilagode svoje proizvode i usluge. AI takođe može pomoći u automatizaciji

⁷⁷ Kanade, V. (2022). What is Artificial Intelligence (AI)? definition, types, goals, challenges, and trends in 2022, Spiceworks. https://www.spiceworks.com/tech/artificial-intelligence/articles/what-is-ai/#_002

⁷⁸ Ibid

⁷⁹ Madhugiri, D. (2023). Advantages and disadvantages of Artificial Intelligence (AI), KnowledgeHut. Knowledgehut.

<https://www.knowledgehut.com/blog/data-science/advantages-and-disadvantages-of-artificial-intelligence>

⁸⁰ Leadershiptribe (2023). Artificial Intelligence in Business Management: A Revolution in the Making. <https://leadershiptribe.com/blog/artificial-intelligence-in-business-management-a-revolution-in-the-making>

svakodnevnih zadataka, dajući menadžerima više vremena da se koncentrišu na strateške projekte. Korišćenjem alata za upravljanje projektima zasnovanih na veštačkoj inteligenciji, menadžeri mogu da automatizuju delegiranje zadataka i praćenje napretka. Ovo povećava produktivnost tima i omogućava efikasniju alokaciju resursa. Upotreba veštačke inteligencije transformiše način na koji menadžeri rade i donose odluke, što dovodi do poboljšane efikasnosti, produktivnosti i strateških rezultata.⁸¹

AI omogućava menadžerima da predvide tržišne trendove, analiziraju velike količine podataka u realnom vremenu i optimizuju operacije. Prema studiji koju je sproveo Accenture, AI ima potencijal da poveća produktivnost do 40% kroz automatizaciju rutinskih zadataka. Uticaj AI na prakse upravljanja posebno je očigledan u analizi podataka. Sa sposobnošću da analiziraju ogromne količine podataka, AI algoritmi mogu otkriti obrasce i korelacije koje mogu izmicati ljudskim analitičarima.⁸² Ova novootkrivena sposobnost omogućava menadžerima da donose svrsishodnije odluke i odrede oblasti za poboljšanje u okviru svojih operacija. Na primer, trgovci na malo mogu da iskoriste AI tehnologiju da bi se udubili u istoriju kupovine i preferencije kupaca, što im na kraju omogućava da daju personalizovane preporuke za proizvode.

Četbotovi sa AI su još jedan primer uticaja AI na prakse upravljanja. Ovi virtuelni pomoćnici mogu da obrađuju upite klijenata, da odmah odgovaraju i obavljaju transakcije. Ovo poboljšava korisničku uslugu i oslobađa ljudske resurse da se fokusiraju na složenije zadatke. Oblast veštačke inteligencije je zabeležila izuzetan napredak poslednjih godina, posebno u oblastima kao što su mašinsko učenje i obrada prirodnog jezika. Ova poboljšanja su imala transformativni uticaj na organizacije, povećavajući njihov učinak i produktivnost. Na primer, Google-ov DeepMind je korišćen za optimizaciju potrošnje energije u svojim centrima podataka, što je dovelo do impresivnog smanjenja troškova energije od 40%. U upravljanju lancem snabdevanja, veštačka inteligencija se koristi za optimizaciju celog procesa, od predviđanja potražnje do upravljanja zalihama. Kompanije kao što su Valmart i Amazon koriste AI algoritme za analizu istorijskih podataka o prodaji i eksternih faktora kako bi precizno predvideli buduću

⁸¹ Takyar, A. (2023). AI in business management: Unveiling the next wave of operational excellence. <https://www.leewayhertz.com/ai-in-business-management/>

⁸² Ibid

potražnju. Ovo im omogućava da optimizuju nivoe zaliha, smanje otpad i poboljšaju zadovoljstvo kupaca.⁸³

AI takođe ostavlja svoj trag u oblasti finansijskog upravljanja. Robo-savetnici, pokretani AI algoritmima, pružaju automatizirane savete o ulaganjima i usluge upravljanja portfoliom. Ove platforme analiziraju tržišne trendove i individualne preferencije investitora kako bi ponudile personalizovane strategije ulaganja. Ovo ne samo da demokratizuje pristup finansijskim savetima, već i pruža isplativa rešenja. Integracija AI u poslovne procese dolazi sa mnoštvom prednosti. To uključuje inovativne cene, prilagođene preporuke, automatizovano zapošljavanje, poboljšanu korisničku podršku, poboljšanu sajber bezbednost, analitiku u realnom vremenu i prediktivnu analitiku.

Preduzeća poput Ubera koriste veštačku inteligenciju za dinamičko određivanje cena, prilagođavajući cene na osnovu ponude i potražnje. Ovo omogućava optimalne strategije određivanja cena i povećanja prihoda. Netflix koristi AI algoritme da preporuči filmove ili serije na osnovu istorije gledanja korisnika. Ovo poboljšava korisničko iskustvo i povećava angažovanje korisnika. AI može da pojednostavi zapošljavanje pregledavanjem biografija kandidata i uzm izborom kandidata na osnovu unapred definisanih kriterijuma. Ovo štedi vreme i poboljšava efikasnost procesa zapošljavanja. Poboljšana korisnička podrška: četboti sa veštačkom inteligencijom mogu da obrađuju upite klijenata 24/7, pružajući trenutne odgovore i rešenja. Ovo dovodi do poboljšanog zadovoljstva kupaca i smanjuje opterećenje osoblja za podršku.

AI može da otkrije anomalije i potencijalne bezbednosne pretnje, omogućavajući preduzećima da brzo reaguju. Ovo pomaže u zaštiti osetljivih podataka i zaštiti od sajber napada. AI omogućava kompanijama da analiziraju podatke u realnom vremenu, osnažujući ih da donose pravovremene odluke. Ova sposobnost olakšava agilno donošenje odluka i sposobnost brzog reagovanja na dinamične tržišne uslove. AI može analizirati istorijske podatke i predvideti buduće trendove ili ishode. Ovo pomaže preduzećima da proaktivno identifikuju mogućnosti i rizike i donose odluke na osnovu informacija.⁸⁴

⁸³ Leadershiptribe (2023). Artificial Intelligence in Business Management: A Revolution in the Making. <https://leadershiptribe.com/blog/artificial-intelligence-in-business-management-a-revolution-in-the-making>

⁸⁴ Ibid

Uprkos brojnim prednostima, primena veštačke inteligencije u poslovnom upravljanju dolazi sa izazovima. To uključuje visoke početne troškove ulaganja, zavisnost od mašina, nedostatak veština i rizik od pomeranja posla. Implementacija AI tehnologija često zahteva značajna ulaganja u infrastrukturu, softver i talente. Malim i srednjim preduzećima će možda biti potrebna pomoć da dodele resurse za integraciju veštačke inteligencije.⁸⁵ Veliko oslanjanje na AI može stvoriti zavisnost od mašina, zbog čega je neophodno imati rezervne planove u slučaju kvarova sistema ili tehničkih problema. Organizacije moraju osigurati da imaju planove za vanredne situacije za održavanje operacija tokom zastoja. Potražnja za profesionalcima za veštačku inteligenciju, kao što su naučnici podataka i stručnjaci za veštačku inteligenciju, raste. Međutim, postoji potreba za više pojedinaca sa potrebnim veštinama i stručnošću. Organizacijama će biti potrebna pomoć u zapošljavanju i zadržavanju talenata za veštačku inteligenciju. Pomeranje posla: Integrisanje tehnologija veštačke inteligencije može dovesti do izmeštanja posla pošto se specifični rutinski zadaci automatizuju.⁸⁶

Organizacije moraju da planiraju i efikasno komuniciraju kako bi ublažile potencijalne gubitke posla i pružile mogućnosti za obuku zaposlenih za prelazak na nove uloge. Uspon veštačke inteligencije doneo je promene u poslovima i zahtevima za veštinama. Iako je istina da neki poslovi mogu biti automatizovani, AI takođe može da poveća ljudske sposobnosti, što će dovesti do stvaranja novih funkcija. Na primer, sve je veća potražnja za stručnjacima za veštačku inteligenciju da dizajniraju i održavaju AI sisteme. Da bi uspeali na radnom mestu pod uticajem AI tehnologije, radnici moraju biti prilagodljivi i da stalno stiču nove veštine. Iako je tehnička stručnost važna, važne su i meke veštine kao što su kritičko razmišljanje, kreativnost i emocionalna inteligencija. Mogućnost kontinuiranog učenja i usavršavanja biće od suštinskog značaja da bi stručnjaci bili relevantni i da bi mogli da iskoriste mogućnosti koje proizilaze iz napretka u veštačkoj inteligenciji.⁸⁷

Ovo ukazuje na potencijal AI u transformaciji tržišta rada i preoblikovanju poslovanja. Iako određene uloge mogu postati nevažne kako tehnologija napreduje, pojaviće se nove pozicije

⁸⁵ Takyar, A. (2023). AI in business management: Unveiling the next wave of operational excellence. <https://www.leewayhertz.com/ai-in-business-management/>

⁸⁶ Whitney, R. (2023). 4 popular uses of artificial intelligence in business. How to maximize the potential of AI in your business? <https://firmbee.com/artificial-intelligence-in-business>

⁸⁷ Ibid

koje će podržati implementaciju i održavanje AI sistema. Organizacije i pojedinci treba da prihvate promene koje donosi veštačka inteligencija i proaktivno se prilagode tržištu rada koji se razvija. Doživotno učenje i način razmišljanja o rastu biće ključni za uspeh u ekonomiji vođenoj veštačkom inteligencijom.⁸⁸ Gledajući unapred, budućnost AI u poslovnom menadžmentu je obećavajuća. Novi trendovi kao što su automatizacija zasnovana na veštačkoj inteligenciji, prediktivna analitika i personalizovano korisničko iskustvo nastaviće da transformišu poslovanje preduzeća. Kako se tehnologija veštačke inteligencije razvija, njen potencijal za postizanje poslovnog uspeha će se samo povećavati. Očekujemo da ćemo videti dalji napredak u obradi prirodnog jezika, prepoznavanju slika i robotici, omogućavajući AI da ima još širi uticaj u svim industrijama. Pored toga, etička razmatranja u vezi sa veštačkom inteligencijom, kao što su transparentnost, pravičnost i odgovornost, igraće sve važniju ulogu u njenom usvajanju i primeni.

3.2.1. Podrška menadžmenta i otpor promenama

Podrška menadžmentu je suštinski element u određivanju poslovne namere da usvoji novu tehnologiju. Menadžment će podsticati usvajanje nove tehnologije ako to vodi ka boljim performansama korisnika, pozitivno utiče na percepciju korisnika i poboljšava usvajanje drugih tehnologija.⁸⁹ Istraživanja su pokazala da je efikasan viši menadžment snažan pokretač za implementaciju nove tehnologije. Podrška menadžmentu je interni olakšavajući uslov za usvajanje nove tehnologije. Podrška menadžmenta utiče na percepciju nove tehnologije od strane zaposlenih.⁹⁰ Menadžment će odrediti tempo implementacije nove tehnologije u organizaciji. Bez obzira na pogled zaposlenih na novi sistem, menadžment može da inkorporira novi sistem u organizaciju i polako odbacuje stari sistem efektivno utičući na usvajanje tehnologije.⁹¹ Ovakav

⁸⁸ Takyar, A. (2023). AI in business management: Unveiling the next wave of operational excellence. <https://www.leewayhertz.com/ai-in-business-management/>

⁸⁹ Sargent, K., Hyland, P., & Sawang, S. (2012). Factors Influencing the Adoption of Information Technology in a Construction Business. *Australasian Journal of Construction Economics and Building*, 12(2), 72-86.

⁹⁰ Gambatese, J. A., & Hallowell, M. (2011). Enabling and measuring innovation in the construction industry. *Construction Management and Economics*, 29(6), 553-567.

⁹¹ Sargent, K., Hyland, P., & Sawang, S. (2012). Factors Influencing the Adoption of Information Technology in a Construction Business. *Australasian Journal of Construction Economics and Building*, 12(2), 72-86.

uticaj na digitalizaciju preduzeća poznat je kao direktna kontrola. Pored toga, menadžment može uticati na uticaj tehnologije indirektno kroz manipulaciju strukturama institucija.⁹²

Takođe je evidentno da većina organizacionih promena nailazi na veliki otpor. Uvođenje nove tehnologije može se suočiti sa velikim otporom ako će prelazak na ovaj novi sistem biti dugotrajan. Zaposleni će takođe biti protiv promena koje će uticati na njihova radna mesta.⁹³ Otpor prema promenama se obično razvija u ranim fazama projekta i ponekad može proći nezapaženo. Prosečna dužina radnog staža zaposlenih će takođe odrediti stopu otpora promenama. Međutim, u organizaciji u kojoj je radna snaga kvalifikovana, otpor promenama će biti minimalan i zaposleni će stalno doživljavati promene bez suštinskih odstupanja od poslovanja.⁹⁴

Stabilnost i sigurnost baze kupaca je značajan faktor koji utiče na usvajanje nove tehnologije u organizaciji. Ulaganje u novu tehnologiju zahteva sigurnost da će postojati prihod koji će omogućiti otplatu ove investicije. Obaveza korisnika je garancija prvenstveno u smanjenju rizika koji su inherentni u odluci o usvajanju. Digitalizacija organizacionog procesa će smanjiti operativne troškove, povećati tržišnu moć organizacije i stabilnost odnosa firme sa kupcima. Odnos sa kupcem će garantovati prisustvo buduće potražnje. Potražnja kupaca je značajnija u nekim industrijama nego u drugim, na primer, da bi proizvođač automobila digitalizovao, postoji potreba da se osigura stalni kupci. Usvajanje nove tehnologije u organizaciji često je skupo zbog kupovine novih sistema i opreme. Neke od ovih tehnologija zahtevaju obuku zaposlenih, što dodatno povećava troškove implementacije nove tehnologije. U slučaju da postoje mrežni efekti, onda će komplementarni sistemi zahtevati ažuriranje ili zamenu. Neke tehnologije bi zahtevale gašenje poslovnih procesa, pa će zbog izgubljenog rezultata nastati dodatni troškovi. Dakle, kada je potražnja neizvesna, onda će poslovanje biti skeptično prema digitalizaciji proizvodnog procesa. Neizvesnost u potražnji će produžiti vreme oporavka, pa će

⁹² Gambatese, J. A., & Hallowell, M. (2011). Enabling and measuring innovation in the construction industry. *Construction Management and Economics*, 29(6), 553-567.

⁹³ Fan, Q. (2016). Factors Affecting Adoption of Digital Business: Evidence from Australia. *Global Journal of Business Research*, 10(3), 79-84.

⁹⁴ Sargent, K., Hyland, P., & Sawang, S. (2012). Factors Influencing the Adoption of Information Technology in a Construction Business. *Australasian Journal of Construction Economics and Building*, 12(2), 72-86.

ulaganje u digitalizaciju učiniti neisplativim. Međutim, kada je potražnja izvesna, onda organizacija može tačno da predvidi potražnju koja je podsticaj za usvajanje nove tehnologije.⁹⁵

Većina organizacija prihvata digitalizaciju koja je predstavljala veliki izazov za rukovodstvo ovih organizacija. Poslovni lideri i viši stručnjaci postali su svesni potrebe za digitalizacijom poslovnih procesa. Trenutna mantra poslovnog sveta je „prilagodi se ili umri“ prvenstveno kada se razmatra tehnologija. Bilo da se radi o novom softveru, automatizaciji, poslovnoj inteligenciji ili veštačkoj inteligenciji, brza konvergencija tehnologije i poslovanja uvek je bila povezana sa konkurentnošću, ponekad u ekstremnim slučajevima, opstankom. U preduzećima, opstanak podrazumeva razvoj okretnog, automatizovanog, prilagodljivog procesa vođenog podacima i komunikacijom. Nepoštovanje ovih evolucionih koraka će se pokazati štetnim za bilo koju kompaniju. Istraživanja su pokazala da bi do kraja naredne decenije oko 40 procenata sadašnjih S&P 500 kompanija nestalo.⁹⁶

Ovo izumiranje će biti posledica poremećaja prirode tehnologije. Čak i sa tehnologijom koja je odgovorna za izazivanje mnoštva ovih promena, imperativ je razumeti ko pokreće tehnologiju. Koliko god se tehnologija smatrala agentom poslovne promene, ona se ne razvija u vakuumu, odnosno pokreću je ljudi, pa će zanemarivanje ljudskog elementa tehnologije značiti propast biznisu. Usvajanje stava prihvatanja od strane bilo koje organizacije prema promenama i tehnološkim inovacijama, rukovodstvo u organizaciji treba da istinski prihvati tehnologiju i praktikuje sve za šta će se zalagati. Tehnološki razvoj organizacije treba da bude više od investicije, već potpuna integracija. Lideri treba da shvate da je tehnologija postala suštinski element organizacije, tako da je neophodan simbiotski odnos sa zaposlenima. Efikasna tehnologija je ona koju zaposleni lako prihvataju umesto da je posmatraju kao nerazumljivu ili zastrašujuću. Lideri moraju da shvate da se to na kraju svodi na ono što će tim dobiti od tehnologije. Izabrana tehnologija treba da učini zaposlene informisanijim i agilnijim u sprovođenju organizacionih strategija. Digitalizacija procesa u organizaciji je timski napor, stoga liderstvo igra vitalnu ulogu u tome da se ovo dovrši.⁹⁷

⁹⁵ Ibid

⁹⁶ Ioannou, L. (2014). A decade to mass extinction event in S&P 500. <https://www.cnbc.com/2014/06/04/15-years-to-extinction-sp-500-companies.html>

⁹⁷ Uzialko, A. C. (2017). Good Leadership Is the Key to Adopting New Technology. <https://www.businessnewsdaily.com/10067-business-technology-innovation-human-user.htm>

3.2.2. Strateško usklađivanje organizacije sa digitalizacijom

Strateška povezanost digitalizacije i poslovanja predviđa da su organizaciona strategija i ciljevi usklađeni kako bi se osiguralo da digitalizacija primeni. Usklađivanje ciljeva i digitalizacija je osnovna briga menadžmenta. Usvajanje nove tehnologije u većini kompanija donosi obim kontinuiranog digitalnog poboljšanja za razliku od poslovnih tema. Većina tehnoloških kompanija ne uspeva da unapredi svoju korporativnu agendu u poređenju sa svojim konkurentima. Primarni cilj digitalizacije je da omogući kompanijama da povećaju svoje strateške ciljeve ostvarivanja uticaja u korporativnom okruženju. Naterati digitalne kompanije da steknu svoju konkurentsku prednost na tržištu je nešto što zaslužuje nadogradnju.⁹⁸

Povezivanje moderacije poslovanja i digitalizacije je nešto što je zakasnilo. Kada digitalne kompanije postanu strateški unapređene u poslovnoj areni, one postaju strateški i relativno povoljne. Konkurentska prednost dolazi kao rezultat potrebe za generisanjem više prihoda i mapiranja celog društva na bolji način⁹⁹ Digitalizacija sa vremenom nastavlja da se mapira u svakodnevnom poslovanju svakog preduzeća. Vlasnici poslovnih i digitalnih platformi dele informacije kako bi poboljšali svoj doprinos razvoju. Prema Kvonu i Parku, ideja poslovnog plana je savršen teren u profesionalnoj poslovnoj areni, posebno u fazama zahteva za ulaganje. Zahvaljujući digitalizaciji, mnoge stvari su postale savremeno usvojene u poslovnom okruženju. Koncepti se pretvaraju u novac sa odgovarajućom brigom i energijom. Veza između tehnološkog plana i poslovnog plana ukazuje na savršen odnos između tehnologije i poslovnih ideologija.¹⁰⁰

S druge strane, prema Kvonu i Parku, veza između poslovnog plana i tehnološkog plana ostaje okosnica međusobne veze između njih. Koliko god se ova dva i dalje oslanjala na različite osnove, koncepti su jedinstveni u odnosu na informacije koje se razmatraju. Cela ideja planiranja, sofisticiranosti, neizvesnosti, uticaja, veličine organizacije se oslanja na efikasnost

⁹⁸ Daub, M., & Wiesinger, A. (2015). Acquiring the capabilities, you need to go digital. <https://www.mckinsey.com/business-functions/digital-mckinsey/our-insights/acquiring-the-capabilities-you-need-to-go-digital>

⁹⁹ Fan, Q. (2016). Factors Affecting Adoption of Digital Business: Evidence from Australia. *Global Journal of Business Research*, 10(3), 79-84.

¹⁰⁰ Kwon, E. H., & Park, M. J. (2017). Critical Factors on Firm's Digital Transformation Capacity: Empirical Evidence from Korea. *International Journal of Applied Engineering Research*, 12(22), 12585-12596.

između razvoja poslovanja i napretka digitalizacije.¹⁰¹ Pristup liderstva na platformama digitalnih preduzeća je model liderstva od vrha do dole. Cela ideja je da se zaposleni dovoljno trude da dođu do ideja koje ih održavaju relevantnim u igri. Izvršni direktor prati napredak na osnovu produktivnosti mlađih zaposlenih. Uobičajeno je da su zaposleni nižeg ranga, mlađe generacije kompanije. Cela industrija leži u grupi pojedinaca koji dolaze sa novim idejama u godinama koje se vremenom smatraju mlađim. Međutim, cela ideja je da se stvori tačna percepcija koja daje operativne prednosti operativnim tržištima. Ideja pravljenja strategije od dna ka vrhu je da se pruži konkurentska prednost u percepciji i liderstvu na tržištu.¹⁰² Harmonizacija digitalizacije i poslovanja je kapija ka uspehu i danima koji dolaze. Tehnologija brzo raste i daje ljudima priliku da ulažu u interpersonalne i rekreativne aktivnosti koje im nude opcije da steknu svoje motivacione prednosti na svojim interpersonalnim i profesionalnim platformama.¹⁰³

3.2.3. Izvršno digitalno liderstvo i organizacioni kapacitet za digitalizaciju

Pristup od vrha prema dole vođen menadžmentom je efikasan način pokretanja digitalizacije u organizacijama. Zaposleni će se angažovati na različite načine guranja i sprovođenja promena u organizaciji pod jakim vođstvom. Najviše rukovodstvo će sprovesti promene u kompaniji na osnovu odluke zasnovane na brzim promenama tržišnih uslova. Za percepciju i odlučnost zaposlenih biće odgovoran najviši menadžment. Rukovodstvo u organizaciji biće odgovorno za stvaranje snažne vizije za digitalizaciju organizacije. Najviši menadžment će biti u boljoj poziciji da pokreće promene izvan granica podela. Lideri će se angažovati i praviti promene kroz različite alate upravljanja kako bi omogućili digitalizaciju, a zatim saopštili promene zaposlenima. Izvršna vlast može rezultirati korporativnim upravljanjem i transformacijom sistema koja će omogućiti digitalnu transformaciju organizacije. Izvršna vlast će podsticati odgovarajuću organizacionu kulturu i razvoj talenata, tehničko vođstvo i neophodna ulaganja. Veil i Ros su se složili da je uloga generalnog direktora i CIO značajna, posebno

¹⁰¹ Ibid

¹⁰² Thomas, K. D., Boring, R. L., Hugo, J. V., & Hallbert, B. P. (2017, June 12). Human factors for main control room modernization. Nuclear News, pp. 46-50.

¹⁰³ Kwon, E. H., & Park, M. J. (2017). Critical Factors on Firm's Digital Transformation Capacity: Empirical Evidence from Korea. International Journal of Applied Engineering Research, 12(22), 12585-12596.

imajući u vidu činjenicu da obezbeđuju harmonizaciju IT i poslovne strategije u organizaciji.¹⁰⁴

Liderstvo organizacije je od vitalnog značaja jer stvara razumevanje važnosti digitalizacije organizacije.¹⁰⁵ Liderstvo u organizaciji će ispuniti odgovornost maksimiziranja vrednosti digitalizacije poslovnog procesa. Izvršni direktor će povezati tehnološku strategiju sa poslovnom strategijom organizacije. Tehničko vođstvo će voditi poslovne rukovodioce u generisanju najveće vrednosti digitalnih promena. Podela između rukovodstva će imati pozitivan uticaj na vezu između digitalizacije i poslovanja. Ova razmena znanja će takođe osigurati da su strateški ciljevi kompanije dobro povezani sa tehnologijom koja se koristi. Postoje dva različita stila koji se mogu koristiti u upravljanju digitalizacijom. Transakciono vođstvo će biti od vitalnog značaja za upravljanje međusobnim stavom koji postoji između lidera i zaposlenih. Transformaciono vođstvo će definisati kako lideri utiču na zaposlene. U transformacionom liderstvu, lideri će proširiti razumevanje ciljeva organizacije i pohvaliti se učinkom zaposlenih. Liderstvo transformacije će uticati na organizacionu kulturu i poboljšati radno okruženje. Unapređenje kulture i motivacije utiče na organizaciju u celini.

Transformacija u poslovnom kontekstu povezana je sa fundamentalnom promenom poslovnog procesa. Digitalna transformacija podrazumeva promene u digitalnoj strategiji, poslovnom modelu i kulturi i poslovnom modelu. Digitalizacija podrazumeva kombinaciju pristupa promenama i inovacijama u organizaciji usled implementacije različitih digitalnih procesa. Digitalizacija u organizaciji zavisi od kompanije, industrije i lidera. Većina istraživanja je insistirala da je tehnologija sredstvo, a ne svrha. Digitalizaciju ne treba suziti na samo digitalne inovacije kao što su mobilna tehnologija, poslovna inteligencija ili analitička rešenja, već je treba posmatrati u širem smislu.

Da bismo izmerili kapacitet organizacije za digitalizaciju, potrebno je da procenimo digitalnu transformaciju iz različitih uglova. Digitalne mogućnosti kompanije su preduslov za dugoročno takmičenje. Većina kompanija koje pokušavaju da se digitalizuju nisu postavile jasan put za postizanje digitalizacije. Ove organizacije još uvek nisu jasne o najboljem načinu da postave svoje digitalne alate. Većina ovih organizacija mora da razvije alate i talente potrebne za

¹⁰⁴Više o tome: Weill, P., & Ross, J. W. (2004). IT governance: How top performers manage IT decision rights for superior results. Boston: Harvard Business Press.

¹⁰⁵ Kwon, E. H., & Park, M. J. (2017). Critical Factors on Firm's Digital Transformation Capacity: Empirical Evidence from Korea. International Journal of Applied Engineering Research, 12(22), 12585-12596.

digitalizaciju. Većini organizacija nedostaje stručnost i kritični resursi potrebni za olakšavanje tranzicije. Posebna ekspertiza potrebna za uspešnu digitalizaciju. Digitalizacija organizacionog procesa podrazumeva uključivanje menadžera proizvođa koji su tehnički potkovani i pismeni u najsavremenijim tehnologijama koje će oblikovati put donošenja odluka potrošača. Digitalizacija kompanije može da podrazumeva korišćenje analitičara za deo organizacije za poslovnu inteligenciju, menadžera sadržaja kako bi se osiguralo da se podaci dopadaju ciljnoj publici.¹⁰⁶ Kompanije za tehnološke usluge su u boljem položaju da imaju više profesionalaca potrebnih za digitalizaciju jer imaju raznovrstan put karijere i put ličnog razvoja. Digitalizacija je od suštinske važnosti u svakoj organizaciji, međutim, razvoj stručnosti je kompetencija koja je potrebna lokalno u organizaciji. U cilju povoljnog nadmetanja, organizacija zahteva da usvoji dinamičan pristup u pristupu digitalnoj ekspertizi izvan organizacije.¹⁰⁷

Ključni cilj digitalizacije u transportnoj kompaniji biće brzo povećanje. Lideri u ovoj industriji nastoje da premeste jaz koji postoji u različitim tehnološkim oblastima, posebno u upravljanju uslugama. U transportu, IT odeljenje bi trebalo da bude dobro opremljeno da upravlja projektima digitalizacije i takođe nadgleda inicijativu pune digitalizacije. Postoji potreba za kontinuiranim naporima u zapošljavanju, razvoju i zadržavanju potrebnih talenata. Da bi se kompanija najbolje nadmetala, postoji potreba da se prihvati dinamičan pristup za procenu digitalnih mogućnosti, posebno iz spoljašnjeg okruženja. Uglavnom, kompanija će zahtevati da bude u stanju da izbalansira brzinu rada IT odeljenja. Značajno je da kompanije procenjuju svoje digitalne ciljeve, sposobnost da praktično izvrše promene i potrebne operativne modele. Radeći ovo, kompanija će moći da drastično i održivo poveća nove digitalne inicijative. Ovi napori će ubrzati upotrebu novih tehnologija uz usklađivanje fragmentiranih aktivnosti koje su od interesa za organizaciju. Promene treba da budu u skladu sa promenama pojedinačnih poslovnih grupa prvenstveno u razvoju odnosa sa dobavljačima koji se lako mogu razvijati bez varijacija u potrebama kupaca. Kompanije koje prolaze kroz konvencionalne transformacije, posebno digitalizaciju, imaju tendenciju da imaju sekvencijalni pristup implementaciji promena kroz inkorporaciju veština i alata. Talenat i tehnologija su međusobno isprepleteni prvenstveno u

¹⁰⁶ Daub, M., & Wiesinger, A. (2015). Acquiring the capabilities, you need to go digital. <https://www.mckinsey.com/business-functions/digital-mckinsey/our-insights/acquiring-the-capabilities-you-need-to-go-digital>

¹⁰⁷ Ibid

slučaju kada je tehnologija kruta i ne može se preneti kupcima pre nego što se napori digitalizacije završe. U digitalizaciji, ovi napori su obično nedovoljni, pa će proces implementacije digitalnih projekata zahtevati stalno pojašnjenje ciljeva i ažuriranje internih zahteva. Ključna prednost digitalizacije je u tome što preduzeća mogu da iskoriste šanse koje se pružaju za inovacije koje su od kraja do kraja usmerene na kupca. Ciljevi projekta će se stalno usavršavati i ažurirati.¹⁰⁸

3.3. Uticaj novih tehnologija na korporativnu bezbednost

Razvoj novih tehnologija u 21. veku donosi brojne prednosti korporacijama, uključujući povećanje efikasnosti, automatizaciju procesa, i unapređenje komunikacije. Međutim, paralelno s tim dolazi i do porasta bezbednosnih izazova. Nove tehnologije menjaju prirodu pretnji i zahtevaju redefinisane strategije korporativne bezbednosti. U ovom poglavlju analiziraju se ključni aspekti uticaja novih tehnologija na korporativnu bezbednost, identifikuju se rizici, i predlažu se strateške smernice za adaptaciju bezbednosnih politika u savremenom poslovnem okruženju.

Pod "novim tehnologijama" u kontekstu ovog rada podrazumevamo:

Internet stvari (IoT)

Veštačka inteligencija (AI) i mašinsko učenje (ML)

Cloud computing (računarstvo u oblaku)

Big Data analitika

Blockchain tehnologije

5G mreže i buduće komunikacione tehnologije

Automatizovani informacioni sistemi

¹⁰⁸ Gambatese, J. A., & Hallowell, M. (2011). Enabling and measuring innovation in the construction industry. *Construction Management and Economics*, 29(6), 553-567.

Svaka od ovih tehnologija ima direktan i indirektan uticaj na bezbednosnu arhitekturu korporacije. Na primer, IoT proširuje napadnu površinu, dok cloud okruženja zahtevaju novu paradigmu upravljanja pristupom i zaštitom podataka.

Uvođenje novih tehnologija dovodi do evolucije pretnji koje se više ne odnose samo na spoljne aktere, već i na kompleksne interne izazove. Ključni aspekti promena uključuju:

Kibernetičke pretnje visokog nivoa sofisticiranosti (APT - Advanced Persistent Threats)

Insajderske pretnje u digitalizovanom okruženju

Rastuća ranjivost kroz dobavljačke lance (Supply Chain Attacks)

Automatizovani napadi vođeni AI tehnologijama

Napadi na podatke u realnom vremenu i cloud infrastrukture

Zbog dinamičnog tehnološkog okruženja, tradicionalni modeli bezbednosti, zasnovani na perimetarskoj zaštiti, više nisu dovoljni. Potrebna je tranzicija ka Zero Trust modelima i integraciji bezbednosti u sve faze poslovanja.

Zero Trust arhitektura odnosi se na implementaciju principa "nikome se ne veruje po defaultu, čak ni unutar mreže", postaje ključna. To podrazumeva višefaktorsku autentifikaciju, mikrosegmentaciju mreže i kontinuiranu verifikaciju identiteta.

Razvijaju se novi modeli procene rizika, koji uključuju prediktivnu analitiku i simulacije scenarija upotrebom veštačke inteligencije. Upravljanje rizicima postaje proaktivan proces, a ne reaktivan.

Pitanje zaštite podataka u cloud okruženju, kao i poštovanje regulativa poput GDPR i lokalnih zakona o zaštiti podataka, stavlja dodatni pritisak na bezbednosne timove.

AI donosi velike prednosti u detekciji anomalija, predikciji napada i automatizaciji odgovora. Međutim, takođe se koristi i za sofisticirane napade, kao što su: Deepfake manipulacije, Automatizovani phishing napadi, Generisanje malvera koji izbegavaju tradicionalne alate za detekciju.

Zbog toga je neophodno razviti AI-sposobne bezbednosne alate koji mogu da odgovore u realnom vremenu.

Bez obzira na tehnološki napredak, ljudski faktor ostaje ključna karika u bezbednosnom lancu. U savremenom okruženju, izazovi uključuju:

Nedostatak edukacije i svesti o bezbednosnim pretnjama, loše upravljanje pristupnim privilegijama, rastući fenomen "digitalnog zamora" među zaposlenima.

Potrebno je kontinuirano ulagati u edukaciju i izgradnju kulture bezbednosti na svim nivoima korporacije.

Nove tehnologije postavljaju i pitanja pravne odgovornosti i etičkih dilema. Na primer: ko je odgovoran kada AI sistem donese pogrešnu odluku koja izazove bezbednosni incident? Kako uskladiti poslovne interese sa zaštitom privatnosti korisnika? Koje zakone treba primeniti u slučaju prekograničnih cyber napada?

Regulatorni okvir često kasni za tehnološkim razvojem, što stvara prostor za zloupotrebe.

Nove tehnologije redefinišu paradigmu korporativne bezbednosti. Sa jedne strane, one omogućavaju napredne metode zaštite, dok sa druge strane otvaraju nove vektore napada. Ključ uspešne adaptacije leži u balansiranju tehnoloških, ljudskih i organizacionih aspekata bezbednosti. Potrebna je kontinuirana procena rizika, integracija savremenih bezbednosnih praksi i razvoj fleksibilnih strategija koje mogu da odgovore na ubrzane promene u digitalnom okruženju.

GLAVA ČETVRTA: IMPLEMENTACIJA AI I PRISTUP PRAVNOJ REGULATIVI

4.1. Najava regulative

Krajem 2019. godine, tada novoizabrana predsjednica Evropske komisije Ursula fon der Lajen najavila je nameru da regulisanje upotrebe veštačke inteligencije (AI) bude glavni prioritet za Evropu. Obećala je da će „predložiti zakone za koordiniran evropski pristup o ljudskim i etičkim implikacijama veštačke inteligencije“ u toku svojih prvih 100 dana na funkciji.¹⁰⁹ Nakon ove objave, Evropska Unija (EU) je pokrenula zakonodavne napore za rešavanje potencijalnih rizika i šteta sa sistemima veštačke inteligencije. Međutim, dizajniranje vladine politike i regulative za sisteme veštačke inteligencije i njihovo korišćenje nije mali podvig; na kraju tih 100 dana Evropska komisija je izradila samo belu knjigu koja istražuje opcije politike.¹¹⁰ Komisija je 2021. godine iznela nacrt zakona o veštačkoj inteligenciji o kojoj tela i zainteresovane strane EU danas aktivno raspravljaju. Stvaranje politike veštačke inteligencije je teško; zahteva pažljivo razmatranje mnogih perspektiva koje uključuju ovu složenu tehnologiju i tekuće razgovore kako bi se pomoglo u budućim odlukama i kompromisnim rešenjima.

AI je postao centralni deo svakodnevnog života, transformišući društvo na mnogo načina. Koristi se u sistemima za preporuke proizvoda za onlajn tržišta, tehnologiji za prepoznavanje glasa i alatima koji pomažu u usmeravanju uličnog saobraćaja. Mnoge tehnologije za spasavanje života oslanjaju se na veštačku inteligenciju, od informisanja o odlukama o zdravstvenoj zaštiti do pružanja pomoći naporima za spasavanje. Dok su primena i koristi od veštačke inteligencije brojni, tu su i mnogi izazovi i rizici. Na primer, štetna pristrasnost u sistemima veštačke inteligencije može produžiti ili pogoršati istorijske nejednakosti u oblastima kao što su zapošljavanje, zdravstvena zaštita, finansije i stanovanje, a zlonamerni akteri mogu da koriste AI

¹⁰⁹ Ursula von der Leyen, *Political Guidelines for the Next European Commission 2019-2024*. Available at: https://ec.europa.eu/info/sites/default/files/political-guidelines-next-commission_en_0.pdf.

¹¹⁰ European Commission, *On Artificial Intelligence – A European approach to excellence and trust*, February 2020. Available at: https://ec.europa.eu/info/sites/default/files/commission-white-paper-artificial-intelligence-feb2020_en.pdf.

alate da manipulišu ljudima. Evropska komisija je postavila kao prioritet stvaranje regulatornog okvira koji bi sprečio i minimizirao negativne efekte veštačke inteligencije. Nasuprot tome, Sjedinjene Države su se manje fokusirale na regulaciju, a više na pristupe „mekog prava“, kao što je npr smernice, standardi i okviri, dopunjeni deliktom i drugim postojećim zakonima (kao što su zakoni o građanskim pravima) da bi pomogli da se ljudi budu odgovorni kada dođe do štete. Stručnjaci i kreatori politike se, međutim, ne slažu jedni sa drugima oko toga kako osmisliti politiku veštačke inteligencije i kakav balans ona treba da postigne između različitih vrednosti, kao što je proaktivan rad koji bi trebalo da spreči štetu kroz vladinu intervenciju i zauzimanje popustljivijeg pristupa za podsticanje eksperimentisanja i inovacija. Dvopartijski centar za politiku posvećen je pronalaženju zajedničkih rešenja za uspostavljanje poverenja, smanjenje štete i promovisanje inovacija u oblasti veštačke inteligencije. U izradi ovog dokumenta, BPC je radio sa zainteresovanim stranama iz akademske zajednice, industrije, civilnog društva i drugima iz Evropske unije i Sjedinjenih Država. Široko zasnovana pitanja relevantna za politiku veštačke inteligencije EU BPC je testirao putem obavljanja niza intervjuja sa predstavnicima akademske zajednice, industrije i civilnog društva u EU i Sjedinjenim Državama koji imaju stručnost u naporima EU u vezi sa AI.

4.2. Razmatranje regulatornog pristupa

I dalje traje debata o tome da li i koliko kreatori politike treba da imaju predostrožnosti u pristupu tehnologiji veštačke inteligencije. Pristup iz predostrožnosti nastoji da postigne ek-ante zaštitu, ublažavajući potencijalne štete čak i pre nego što se realizuju.¹¹¹ Na primer, vladino telo može da postavi zahteve ili smernice za sisteme veštačke inteligencije kako bi ispunili specifične standarde pre primene, koristeći bilo instrumenti tvrdog ili mekog prava. Imajući u vidu potencijal da veštačka inteligencija može da ozbiljno poremeti društvo na potencijalno negativne načine, neki stručnjaci preporučuju pristup iz predostrožnosti u odnosu na dozvoljavajući pristup kako bi obeshrabrili preterano preuzimanje rizika od strane programera prilikom izgradnje

¹¹¹ World Commission on the Ethics of Scientific Knowledge and Technology, UNESCO, *The Precautionary Principle*, 2005. Available at: <https://unesdoc.unesco.org/ark:/48223/pf0000139578>.

sistema. Dozvoljavajući pristup može dati inovatorima više slobode da eksperimentišu i primenjuju sisteme veštačke inteligencije bez eksplicitnog odobrenja vlade, ali ih i dalje može ostaviti otvorenim za tužbe i velike novčane kazne kako bi se obeshrabilo bezobzirno preuzimanje rizika. Međutim, zbog rizika od gušenja inovacija kroz regulatorne barijere, neki stručnjaci savetuju da se ne koristi pristup iz predostrožnosti i da se da prednost popustljivom pristupu. Dozvoljavajući pristup, ako je dobro osmišljen, i dalje može ublažiti i sprečiti mnoge neželjene ishode, istovremeno omogućavajući veće eksperimentisanje i inovacije.¹¹² EU često koristi princip predostrožnosti kada kreira regulativu i zakonodavstvo, kao što je Zakon o veštačkoj inteligenciji.¹¹³ Princip predostrožnosti se najviše koristi u zakonu EU o životnoj sredini; EU je utvrdila da su propisi predostrožnosti neophodni da bi se ljudi zaštitili od naučno dokazanih rizika po životnu sredinu i nepredvidivih posledica, kao što su aerosolni sprejevi koji uništavaju ozonski omotač ili neodrživo korišćenje ribljih resursa. Neki veruju da mnogi rizici i neizvesnosti AI opravdavaju slične propise predostrožnosti u digitalnom okruženju.

Sjedinjene Države generalno prihvataju pojam inovacija bez dozvole, kao što je prikazano u 1990-im kada su razvile pravila oko interneta i digitalne ekonomije. Početkom 1990-ih, Kongres je proširio federalno finansiranje istraživanja i ovlastio Nacionalnu naučnu fondaciju da otvori NSFNET (ono što je danas evoluiralo u internet) za komercijalne aktivnosti. Sjedinjene Države su zadržale ovu strategiju, kada je Kancelarija za upravljanje i budžet izdala smernice u skladu sa izvršnom naredbom o održavanju američkog liderstva u oblasti veštačke inteligencije. U naredbi je stajalo: „Agencije moraju izbegavati pristup iz predostrožnosti koji drži sisteme veštačke inteligencije na tako neverovatno visokom standardu da društvo ne može da uživa u njihovim prednostima.” Iako je važno razmotriti istorijske pristupe regulisanju veštačke inteligencije, današnji tehnološki pejzaž se značajno razlikuje od onog od pre nekoliko decenija, pa čak i pre samo nekoliko godina. Kreatori politike treba da odvagaju ono što su naučili iz istorije i da stečeno znanje uporede sa informacijama koje uče danas.

¹¹² Adam Thierer, “Embracing a Culture of Permissionless Innovation,” Cato Institute, November 17, 2014. Available at: <https://www.cato.org/cato-online-forum/embracing-culture-permissionless-innovation>.

¹¹³ European Parliamentary Research Service, *The Precautionary Principle: Definitions, Applications, and Governance*, December 2015. Available at: [https://www.europarl.europa.eu/thinktank/en/document/EPRS_IDA\(2015\)573876](https://www.europarl.europa.eu/thinktank/en/document/EPRS_IDA(2015)573876).

Kao deo svoje digitalne strategije, EU želi da reguliše veštačku inteligenciju (AI) kako bi obezbedila bolje uslove za razvoj i upotrebu ove inovativne tehnologije. AI može stvoriti mnoge prednosti, kao što je bolja zdravstvena zaštita; bezbedniji i čistiji transport; efikasnija proizvodnja; i jeftinija i održivija energija. U aprilu 2021. Evropska komisija je predložila prvi regulatorni okvir EU za AI: sistemi veštačke inteligencije se mogu koristiti u različitim aplikacijama analizirani i klasifikovani prema riziku koji predstavljaju za korisnike. Različiti nivoi rizika će značiti više ili manje regulacije. Prioritet parlamenta je da osigura da sistemi veštačke inteligencije, koji se koriste u EU, budu bezbedni, transparentni, sledljivi, nediskriminatorni i ekološki prihvatljivi. Sisteme veštačke inteligencije treba da nadgledaju ljudi, a ne automatizacija, kako bi se sprečili štetni ishodi. Parlament takođe želi da uspostavi tehnološki neutralnu, jednoobraznu definiciju za veštačku inteligenciju koja bi se mogla primeniti na buduće sisteme veštačke inteligencije.

4.2.1. Čvrsto pravo vs meko pravo

Čvrsto pravo je termin koji se koristi za regulatorne mehanizme koje sprovodi vlada, kao što su zakoni i propisi, dok se meki zakon odnosi na politike, kao što su standardi, najbolje prakse i smernice koje vlada manje sprovodi, ali su fleksibilnije. Svaki pristup ima prednosti i nedostatke. Čvrsti zakon može nametnuti veće terete poštovanja i biti sporiji za prilagođavanje, ali ima prednost vladine moći za sprovođenje. Meki zakon je fleksibilniji i manje opterećujući, ali ga je teže sprovesti. Zakon EU o veštačkoj inteligenciji koristi čvrst zakon za rizičnije primene veštačke inteligencije, a meki zakon za manje rizične. Kombinacija ova dva pristupa je zbog razlike u zakonodavstvu između različitih nivoa rizika i različite moći sprovođenja za svaki od njih. Zabranjena je upotreba veštačke inteligencije koja stvara neprihvatljive rizike. Sistemi veštačke inteligencije visokog rizika su obavezni da se usaglašavaju sa zakonskim zahtevima, uključujući ocenjivanje usaglašenosti, što je analiza rizika analogna procesu koji koristi Uprava za hranu i lekove za odobravanje medicinskih uređaja na tržištu. Sistemi veštačke inteligencije sa ograničenim rizikom podležu ograničenim obavezama transparentnosti, kao što je otkrivanje video snimka koji sadrže veštački generisanu sličnost sa stvarnom osobom (ili „duboko lažno“). Sistemi veštačke inteligencije koji predstavljaju nizak ili minimalan rizik ne podležu zakonskim

obavezama, ali mogu da izvrše procene uticaja koji se sami proveravaju i dobrovoljne kodekse ponašanja. Iako je razlika između tvrdog i mekog zakona često nejasna, sistemi veštačke inteligencije sa većim rizikom su očigledno podložni većim prinudnim obavezama u okviru Zakona o veštačkoj inteligenciji od veštačke inteligencije nižeg rizika.

Još uvek je u toku debata o tome kada koristiti čvrsti, a kada meki zakon za upravljanje AI. Eksperti su često isticali različite zabrinutosti kada su raspravljali o ovom pitanju. Neki su naglasili zabrinutost da bi strogi zakon mogao blokirati regulatorne zahteve za postojeće tehnologije, usporavajući tehnološki napredak. Mnogi drugi su naglasili potrebu za preskriptivnim propisima za kontrolu sistema veštačke inteligencije koji predstavljaju značajne pretnje po zdravlje ili bezbednost korisnika, kao što su one koje se koriste za medicinske uređaje. Takođe, još uvek je neusaglašen stav stručnjaka o tome da li i kada politika treba da primeni čvrsti ili meki zakon u slučajevima kada je nivo rizika nepoznat ili neizvestan.

Različita pravila za različite nivoe rizika

Nova pravila utvrđuju obaveze za provajdere i korisnike u zavisnosti od nivoa rizika od veštačke inteligencije. Iako mnogi sistemi veštačke inteligencije predstavljaju minimalan rizik, potrebno ih je proceniti.

Sistemi veštačke inteligencije neprihvatljivog rizika su sistemi koji se smatraju pretnjom po ljude i biće zabranjeni. To uključuje:

- Kognitivna biheviorna manipulacija ljudi ili specifičnih ranjivih grupa: na primer igračke koje se aktiviraju glasom koje podstiču opasno ponašanje kod dece;

- Socijalno bodovanje: klasifikovanje ljudi na osnovu ponašanja, socio-ekonomskog statusa ili ličnih karakteristika;

- Biometrijska identifikacija i kategorizacija ljudi;

- Sistemi biometrijske identifikacije u realnom vremenu i daljinski, kao što je prepoznavanje lica.

Neki izuzeci mogu biti dozvoljeni za potrebe sprovođenja zakona. Sistemi daljinske biometrijske identifikacije „u realnom vremenu“ biće dozvoljeni u ograničenom broju ozbiljnih slučajeva, dok će sistemi za daljinsku biometrijsku identifikaciju „post“ biti dozvoljeni za krivično gonjenje za teška krivična dela i to tek nakon odobrenja suda.

Sistemi veštačke inteligencije koji negativno utiču na bezbednost ili osnovna prava smatraće se visokim rizikom i biće podeljeni u dve kategorije:

1) Sistemi veštačke inteligencije koji se koriste u proizvodima koji potpadaju pod zakonodavstvo EU o bezbednosti proizvoda. Ovo uključuje igračke, avijaciju, automobile, medicinske uređaje i liftove.

2) Sistemi veštačke inteligencije koji spadaju u određene oblasti koje će morati da budu registrovane u bazi podataka EU:

Upravljanje i rad kritične infrastrukture

Obrazovanje i stručna obuka

Zapošljavanje, upravljanje radnicima i pristup samozapošljavanju

Pristup i uživanje osnovnih privatnih usluga i javnih usluga i beneficija

Sprovođenje zakona

Upravljanje migracijama, azilom i graničnom kontrolom

Pomoć u pravnom tumačenju i primeni zakona.

Svi visokorizični sistemi veštačke inteligencije biće procenjeni pre nego što budu stavljeni na tržište, kao i tokom njihovog životnog ciklusa. Ljudi će imati pravo da podnose žalbe na sisteme veštačke inteligencije određenim nacionalnim vlastima.

Generativna AI, kao što je ChatGPT, neće biti klasifikovana kao visokorizična, ali će morati da bude u skladu sa zahtevima transparentnosti i zakonom o autorskim pravima EU što podrazumeva:

Otkrivanje da je sadržaj generisao AI

Dizajniranje modela kako bi se sprečilo stvaranje nelegalnog sadržaja

Objavljivanje sažetaka podataka zaštićenih autorskim pravima koji se koriste za obuku

Modeli veštačke inteligencije opšte namene sa velikim uticajem koji bi mogli predstavljati sistemski rizik, kao što je napredniji AI model GPT-4, morali bi da prođu temeljne procene i svi ozbiljni incidenti bi morali da budu prijavljeni Evropskoj komisiji. Sadržaj koji se generiše ili menja uz pomoć veštačke inteligencije – slike, audio ili video datoteke (na primer deepfakes) – mora biti jasno označen kao AI generisan kako bi korisnici bili svesni kada naiđu na takav sadržaj.

Zakon ima za cilj da početnicima i malim i srednjim preduzećima ponudi mogućnosti da razvijaju i obučavaju AI modele pre njihovog puštanja u javnost. Zbog toga se zahteva da nacionalne vlasti obezbede kompanijama okruženje za testiranje koje simulira uslove bliske stvarnom svetu. On će se u potpunosti primenjivati 24 meseca nakon stupanja na snagu, ali neki delovi će se primenjivati ranije:

Zabrana AI sistema koji predstavljaju neprihvatljive rizike primenjivaće se šest meseci nakon stupanja na snagu.

Kodeksi prakse će se primenjivati devet meseci nakon stupanja na snagu.

Pravila o sistemima veštačke inteligencije opšte namene koji moraju da budu u skladu sa zahtevima transparentnosti primenjivaće se 12 meseci nakon stupanja na snagu. Sistemi visokog rizika će imati više vremena da ispune zahteve pošto će obaveze koje se na njih odnose postati primenjive 36 meseci nakon stupanja na snagu.

Mere EU koje se odnose na digitalizaciju nastoje da regulišu sledeće pojave:

Opasnosti od kriptovaluta i prednosti zakonodavstva EU

Borba protiv sajber kriminala

Podsticanje razmene podataka u EU

Digitalna tržišta EU i digitalne usluge

Nastojanje Evropskog parlamenta da zaštiti onlajn igrače.

4.3. Procena Opšte regulative

Opšta regulativa utvrđuje obaveze i politike koje se primenjuju na sve industrije ili sektore, dok se sektorski pristup fokusira na određeni sektor. Na primer, propis koji se primenjuje samo na farmaceutske lekove prati pristup specifičan za sektor. Nasuprot tome, zakon o deliktu koji se primenjuje u svakom sektoru privrede sledi opšti pristup. Zakon o veštačkoj inteligenciji EU ima opšti pristup regulisanju veštačke inteligencije, ali neki tvrde da ga uokvirivanje pristupa zasnovanog na riziku oblikuje u pristup specifičan za sektor.

Zakon o veštačkoj inteligenciji se generalno primenjuje na razvoj, proizvodnju i upotrebu AI sistema u svim sektorima. Uprkos svojoj nameni, neki stručnjaci oklevaju da kažu da Zakon o veštačkoj inteligenciji zaista koristi opšti pristup, sugerišući da neki jezici ciljaju na specifičnu

sektorsku upotrebu. Na primer, neke klasifikacije visokog rizika u nacrtu zakona mogu propisati pravila koja su relevantnija za neke industrije od drugih. EU i Sjedinjene Države imaju za cilj borbu protiv štetne pristrasnosti u algoritamskim odlukama. Pristrasni skupovi podataka ili dizajn algoritama mogu produžiti ili pogoršati istorijske nejednakosti i mogu naneti značajnu štetu ranjivoj populaciji i marginalizovanim grupama. Rešenja za rešavanje ovog izazova bila su glavni fokus mnogih zainteresovanih strana iz vlade, civilnog društva, akademske zajednice i industrije. Iskorištavanje ljudi sa različitim gledištima i pozadinom prilikom izrade i primene AI alata ili izrade politike veštačke inteligencije može pomoći da se pokriju slepe tačke i identifikuju načini za suzbijanje štetnih pristrasnosti. Rešavanje štetne pristrasnosti veštačke inteligencije zahteva višestruku strategiju koja koristi nekoliko pristupa. Međutim, detalji ove strategije su još uvek za raspravu. Omogućavanje industriji da prikuplja demografske podatke, kao što su starost, rasa ili pol, kako bi se pomoglo u otkrivanju, proceni i rešavanju algoritamskih izazova predrasuda je pristup koji su neki predložili za razmatranje. Na primer, demografski podaci se mogu koristiti za razumevanje i testiranje pristrasnosti u algoritmima za zapošljavanje organizacije. Međutim, veoma lična priroda ovih podataka zahteva dalje razmatranje prilikom evaluacije ovog pristupa. Zainteresovane strane i stručnjaci nisu se složili oko toga da li i kako prikupljati demografske podatke da bi se otkrila i ublažila štetna pristrasnost. Oni su generalno izražavali značajne rezerve, ali su mnogi bili otvoreniji za to ideja da li bi kreatori politike mogli da reše određene važne probleme.

Dve glavne brige odnosile su se na zaštitu privatnosti podataka i obezbeđivanje korišćenja ovih podataka za ublažavanje štetne pristrasnosti. Neki načini za ublažavanje ovih zabrinutosti mogu uključivati bolju sigurnost podataka, minimiziranje prikupljanja podataka na samo ono što je neophodno, sporazume o privatnosti kako bi se korisnici informisali o prikupljenim informacijama i njihovoj nameravanoj upotrebi, i mehanizmi za pozivanje kompanija na odgovornost kada naruše poverenje korisnika.

Pokrenuto je i diskutovano o nekoliko drugih pristupa. Jedna ideja je bila stvaranje regulatornih „peščanika“ za proučavanje tehnika i praksi ublažavanja pristrasnosti za određene AI sisteme u ograničenijem stvarnom okruženju. Regulatorni sandbokovi bi omogućili zainteresovanim stranama da procene izvodljivost prakse u kontrolisanom okruženju u stvarnom svetu. Druge tehnike o kojima se razgovaralo uključivale su testiranje sistema veštačke inteligencije pre nego što se primene i istraživanje upotrebe anonimizovanih ili sintetičkih

podataka za suzbijanje štetne pristrasnosti. Pregled zakonskih obaveza u sadašnjem zakonodavstvu je takođe važan kako bi se osiguralo da su oni najbolji pristup za rešavanje štetne pristrasnosti.

Zakon o veštačkoj inteligenciji prati niz drugih politika o tehnologiji i podacima koje je Evropska komisija predložila i usvojila u protekloj deceniji. EU je osmislila ove zakone da kolektivno regulišu i usmeravaju neke operacije tehnoloških kompanija, a svi deluju na nivou cele Evropske unije. Mnogi akti utiču na razvoj AI indirektno – AI zavisi od velikog obima podataka da bi efikasno funkcionisala, tako da će bilo kakva ograničenja u prikupljanju i korišćenju podataka uvek uticati na to kako AI sistemi funkcionišu. Štaviše, različiti zakoni i propisi međusobno deluju, tako da kreatori politike treba da imaju na umu dupliranja, nedoslednosti i praznine u ovim zakonima i propisima.

4.4. Regulativa veštačke inteligencije i poređenja između SAD i EU

Veštačka inteligencija i poređenja između SAD i EU u oblasti AI i mašinskog učenja nisu novi, ali postaju moćniji i uticajniji sa sve većom dostupnošću velikih skupova podataka, poboljšanom računarskom snagom i većim kapacitetom skladištenja. AI i podaci su neraskidivo povezani; potrebna je velika količina podataka za obuku i testiranje AI sistema, koji onda mogu analizirati velike količine dodatnih podataka. Sistemi veštačke inteligencije zahtevaju intenzivnu računarsku snagu, a nedavna eksponencijalna poboljšanja računarske moći omogućila su mnogim modernim AI sistemima da se razvijaju i dograđuju.¹¹⁴ Neki primeri AI aplikacija uključuju autonomna vozila i generisanje teksta nalik čoveku (kao što je GPT-3). Poboljšanja u AI mogu dovesti do veće upotrebe i većeg uticaja na to kako radimo, živimo i doživljavamo svet. Pored toga, kako programeri čine značajne korake unapređujući AI, važno je razmotriti implikacije politika koje se danas donose na budući napredak i primenu tehnologije. Pošto politike veštačke inteligencije u inostranstvu utiču na razvoj i upotrebu veštačke inteligencije u Sjedinjenim Državama, neophodna je promišljena međunarodna diskusija. Propisi EU često

¹¹⁴ Jaime Sevilla, Lennart Heim, et al., “Compute Trends Across Three Eras of Machine Learning,” arXiv:2202.05924, March 9, 2022. Available at: <http://arxiv.org/abs/2202.05924>.

efikasno postavljaju ili oblikuju globalne standarde – fenomen koji se naziva „Briselski efekat“.¹¹⁵ Na primer, mnogi veb-sajtovi poštuju Opštu uredbu o zaštiti podataka (GDPR) na međunarodnom nivou, iako su to propisi specifični za EU. I nacionalni interes i određeni nivo međunarodne saradnje verovatno će pokretati tehnologiju politike u 21. veku. Sjedinjene Države i EU, na primer, imaju zajedničku posvećenost sprovođenju odgovornog upravljanja pouzdanom veštačkom inteligencijom,¹¹⁶ koja je usredsređena na principe OECD o veštačkoj inteligenciji uključivanja, pravičnosti, transparentnosti, robusnosti i odgovornosti.¹¹⁷

Sjedinjene Države i EU dale su prioritet razvoju i korišćenju verodostojne veštačke inteligencije, ali se njihove domaće političke strategije razlikuju. Sjedinjene Države se fokusiraju na usklađivanje standarda preko Nacionalnog instituta za standarde i tehnologiju (NIST), savezne agencije, i objavila su svoj prvi nacrt Okvira za upravljanje rizikom od veštačke inteligencije kako bi se obezbedile smernice za razvoj AI.¹¹⁸ U EU, Evropska komisija je predložila širi regulatorni okvir za harmonizaciju politike veštačke inteligencije u državama članicama, poznat kao Zakon o veštačkoj inteligenciji (AI Act), koji prevazilazi samo standarde za stvaranje zakonskih i tehničkih obaveza za potencijalno štetnu veštačku inteligenciju. Mnoge američke države žele da donesu sopstvene zakone i propise oko veštačke inteligencije umesto saveznog zakona.

Razlike između pristupa SAD i EU u regulisanju veštačke inteligencije delimično potiču iz ključnih strukturnih razlika između ova dva regiona. Na primer, prema mišljenju mnogih stručnjaka, američki pravni sistem je evoluirao tako da ima stroži režim odgovornosti za delikte u odnosu na EU, što može delimično biti vođeno sistemom običajnog prava SAD. Stroži sistem odgovornosti ima troškove i koristi, ali verovatno smanjuje potrebu za strožom regulativom. Dalje, razlike između EU i Američke strukture upravljanja mogu uticati na regulatorni režim koji je najprikladniji za svaki region: EU je unija od 27 država članica, dok Sjedinjene Države imaju

¹¹⁵ Anu Bradford, *The Brussels Effect: How the European Union Rules the World* (New York: Oxford University Press, 2020).

¹¹⁶ U.S. Department of Commerce, “U.S.-EU Joint Statement of the Trade and Technology Council,” May 16, 2022. Available at: <https://www.commerce.gov/news/press-releases/2022/05/us-eu-joint-statement-trade-and-technology-council>.

¹¹⁷ OECD AI Policy Observatory, “OECD AI Principles overview,” accessed July 14, 2022. Available at: <https://oecd.ai/en/ai-principles>.

¹¹⁸ National Institute of Science and Technology, *AI Risk Management Framework*, Draft, 2022. Available at: <https://www.nist.gov/itl/ai-risk-management-framework>.

federalni sistem koji se reguliše prema Ustavu SAD. Sjedinjene Države i Evropa se globalno takmiče za liderstvo u AI, kako međusobno tako i sa drugim zemljama kao što je Kina. Sjedinjene Države su mnogo investirale u istraživanje i razvoj veštačke inteligencije: prema Globalnoj AI Vibranci Univerziteta Stanford Alat, 12 nacija je globalno vodeća u ukupnim privatnim investicijama u veštačku inteligenciju, grantovima za patente za veštačku inteligenciju i citatima iz AI repozitorijuma. Nedavna ulaganja EU u veštačku inteligenciju deo su njenog Koordinisanog plana da postane svetski lider u razvoju i primeni najsavremenije tehnologije kroz investicije od milijardu evra i planovi za privlačenje drugih javnih i privatnih doprinosa.¹¹⁹ AI izveštava da su globalne privatne investicije u veštačku inteligenciju u 2021. godini iznosile 93,5 milijardi dolara, što je više nego duplo u odnosu na prethodnu godinu.¹²⁰ Sjedinjene Države su prednjačile po iznosu investicija, a slede ih Kina i Evropska unija. Nedavna istraživanja i razvoj u oblasti veštačke inteligencije u velikoj meri su uticali na globalnu ekonomiju i društvo i očekuje se da će nastaviti da utiču na rast širom sveta. Očekuje se da će razvoj veštačke inteligencije značajno uticati na nekoliko industrija, od maloprodaje preko proizvodnje do zdravstvene zaštite.¹²¹ AI može poboljšati odluke u lancu snabdevanja, povećati efikasnost proizvodnih linija i pomoći zdravstvenim radnicima da dijagnostikuju bolesti. Tehnologija bi se mogla implementirati na različite načine, fokusirajući se na upravljanje rizikom, prediktivnu analitiku, korisničke usluge ili podršku za istraživanje i razvoj, između ostalog.

Široka svestranost AI aplikacija naglašava finansijske implikacije koje ih prate. Uprkos mnogim pozitivnim doprinosima, veštačka inteligencija je privukla negativnu pažnju javnosti i izazvala mnoge važne zabrinutosti poslednjih godina. AI je donela izazove algoritamske pristrasnosti, koja može da podstakne nejednakost i naškodi mnogim nedovoljno zastupljenim grupama i marginalizovanim zajednicama. Na primer, Scientific American je izvestio: „Algoritmi obučeni sa rodno neuravnoteženim podacima imaju lošije rezultate u čitanju

¹¹⁹ European Commission, “Excellence and trust in artificial intelligence,” accessed July 14, 2022. Available at: https://ec.europa.eu/info/strategy/priorities-2019-2024/europe-fit-digital-age/excellence-trust-artificial-intelligence_en.

¹²⁰ Stanford University Human-Centered Artificial Intelligence, “The AI Index Report – Artificial Intelligence Index,” accessed July 14, 2022. Available at: <https://aiindex.stanford.edu/report/>.

¹²¹ The Economist Intelligence Unit, “Intelligent Economies: AI’s transformation of industries and society,” 2018. Available at: <https://impact.economist.com/perspectives/technology-innovation/intelligent-economies-ais-transformation-industries-and-society/>

rendgenskih snimaka grudnog koša za nedovoljno zastupljen pol, a istraživači su već zabrinuti da algoritmi za otkrivanje raka kože, od kojih su mnogi obučeni prvenstveno na ljudima sa svetlom kožom, rade lošije u otkrivanju raka kože koji utiče na tamniju kožu.”¹²² Kombinacija nadzora na radnom mestu i prikupljanja podataka sa povećanom upotrebom veštačke inteligencije dovodi u prvi plan pitanja privatnosti i asimetrije moći. U članku BPC-a se navodi: „Netradicionalni alati za praćenje zaposlenih koji analiziraju e-poštu ili prikupljaju biometrijske podatke, postali su popularniji na radnom mestu tokom godina – Studija Gartnera pokazala je da je više od polovine kompanija ispitanih 2018. primenilo neki oblik alata za praćenje zaposlenih, što je povećanje od 30% u odnosu na 2015.”¹²³ Sve ovo je dovelo do veće svesti kreatora politike o veštačkoj inteligenciji i većeg osećaja hitnosti da se pozabave etičkim pitanjima veštačke inteligencije.

Istraživanja pokazuju porast zakona koji se odnose na veštačku inteligenciju na globalnom nivou, zajedno sa povećanim pominjanjem AI u zakonodavnim postupcima i političkim dokumentima. Na primer, Sjedinjene Države, Španija, Belgija, Francuska, Italija i Nemačka donele su zakone u vezi sa veštačkom inteligencijom 2021. Uzlazni trend bi se mogao nastaviti jer vlade i zakonodavna tela širom sveta razmatraju kako da se pozabave etičkim i društvenim izazovima veštačke inteligencije uz promovisanje odgovornih inovacija. Ovde ukazujemo na to da je za donošenje odluka o politikama za regulaciju AI prioritet u obrazovanju kreatora politike i zainteresovanih strana o važnim konceptima u upravljanju veštačkom inteligencijom. Kao pokretači razvoja veštačke inteligencije, ekonomskog rasta i zakonodavne akcije, ova dva upravljačka tela će biti veoma uticajna u budućnosti AI.

¹²² Amit Kaushal, Russ Altman, and Curt Langlotz, “Health Care AI Systems Are Biased,” *Scientific American*, November 17, 2020. Available at: <https://www.scientificamerican.com/article/health-care-ai-systems-are-biased/>.

¹²³ Babu Jackson, “Prevalence of Biometric Data and Security Concerns,” Bipartisan Policy Center, September 20, 2021. Available at: <https://bipartisanpolicy.org/blog/prevalence-of-biometric-data-and-security-concerns>.

GLAVA PETA: PRAVNA REGULATIVA U DOMENU SAJBER BEZBEDNOSTI

5.1. Pravna rešenja i strategije sajber bezbednosti

Dok održiva politika sajber bezbednosti uključuje širok spektar razmatranja, zakonske mere igraju ključnu ulogu u prevenciji i borbi protiv sajber kriminala. Donošenje zakona je potrebno u svim oblastima, uključujući kriminalizaciju, proceduralna ovlašćenja, nadležnost, međunarodnu saradnju i odgovornost provajdera internet usluga. Konkretno, na nacionalnom nivou, zakoni o sajber kriminalu najčešće se tiču kriminalizacije – utvrđivanja specijalizovanih krivičnih dela za osnovna dela sajber kriminala. Međutim, zemlje sve više prepoznaju potrebu za zakonodavstvom u drugim oblastima.¹²⁴ Tehnološki razvoj povezan sa sajber bezbednošću i sajber kriminalom znači da, dok se tradicionalni zakoni mogu donekle primeniti, zakonodavstvo takođe mora da se nosi sa novim konceptima i objektima, kao što su nematerijalni kompjuterski podaci, koji se tradicionalno ne obrađuju zakonom.

U mnogim državama, zakoni o tehničkom razvoju datiraju iz 19. veka. Ovi zakoni su bili, i u velikoj meri fokusirani na fizičke objekte, oko kojih se vrteo svakodnevni život industrijskog društva. Iz tog razloga, mnogi tradicionalni opšti zakoni ne uzimaju u obzir specifičnosti informacionih tehnologija koje su povezane sa sajber kriminalom i zločinima koji stvaraju elektronske dokaze. Ova dela u velikoj meri karakterišu novi nematerijalni objekti, kao što su podaci ili informacije. Iako se krivično pravo često smatra najrelevantnijim kada je u pitanju sajber kriminal, pravni odgovori na širu zabrinutost u vezi sa sajber bezbednošću takođe uključuju upotrebu drugih grana prava, kao što su građansko pravo i upravno pravo.

Dalje podele unutar ovih pravnih režima obuhvataju materijalno i procesno pravo, kao i regulatorne i ustavne, ili zakone zasnovane na pravima. U mnogim pravnim sistemima, svaki od ovih režima karakterišu specifični ciljevi, institucije i zaštitne mere. Zakoni o sajber kriminalu se najčešće nalazi u oblastima materijalnog i procesnog krivičnog prava. Međutim, brojne druge oblasti prava su takođe važne.¹²⁵

¹²⁴ Više o tome: Sandage, John et al. eds. (2013). *Comprehensive Study on Cybercrime* (United Nations Office on Drugs and Crime).

¹²⁵ Ibid

Pitanje kriminalizacije nepoželjnog ponašanja na internetu ima dvostruki efekat:

(a) stvaranje pravne osnove za retributivno suzbijanje ponašanja, i (b) stvaranje klime društvene neprihvatljivosti sajber kriminala, deromantizovanje i stigmatizovanje takvog ponašanja. Oni koji koriste internet za činjenje zločina odrasli su i bili su socijalizovani u klimi u kojoj je preovlađujući vid nezakonitih aktivnosti bio kriminal u stvarnom svetu, u svojim tradicionalnim obličjima,¹²⁶ dok hakovanje, na primer, ne izaziva osećaj društvene neprihvatljivosti. Trenutno, hakersko ponašanje karakteriše „laissez-faire“ stav prema odgovornosti i zakonitosti u mnogim jurisdikcijama na globalnom nivou. U tom smislu, konceptualizacija sajber kriminala kao krivičnog dela unapređuje i odmazdu i odvracanje.¹²⁷

5.2. Haktivizam i kriminalizacija

Neki komentatori predlažu da će ciljanje haktivizma (za razliku od hakovanja) u naporima za kriminalizaciju najočiglednije minimizirati problem u široj debati o sajber bezbednosti.¹²⁸ Možda je problem u tome što je kriminalizacija haktivizma previše složeno pravno pitanje koje uključuje specifičnu nameru da se promovišu, inter alia, politički ciljevi. To bi takođe moglo biti izazovno iz perspektive društvene prihvatljivosti takvog zločina. Primeri sajber napada pod vođstvom „Anonimusa“ (Anonymous) su ilustrativni upravo zbog „Robin Hud“ arome namera Anonimusa. Baš kao što je mafija nekada bila istaknuta kao lice organizovanog kriminala, vlade treba da iskoriste vidljivost Anonimusa, kada razgovaraju o sajber bezbednosti sa širom javnošću. Ovo izdvajanje Anonimusa ne bi bilo neopravdano. Prema izveštaju iz 2012. koji je objavio Verizon (Verizon), haktivisti su pretekli sajber kriminalce kao grupa odgovorna za najveću štetu nastalu od sajber napada. Anonymous posebno ima veliko javno priznanje zbog svog oslanjanja na društvene medije (Tviter fidove, JuTjub stranice i veb stranice), mehanizme brendiranja (ikonične maske Gaj Foksa (Guy Fawkes) i prakse imenovanja) i politički nabijene stavove tokom izvođenja sajber napada na visokopofilne žrtve.

¹²⁶ Brenner Susan (2004). *Toward a Criminal Law for Cyberspace: Distributed Security*, BOSTON UNIVERSITY JOURNAL OF SCIENCE & TECHNOLOGY LAW. Vol. 10, Issue 1, pp. 1-109

¹²⁷ Kelly Brian (2012). *Investing In a Centralized Cybersecurity Infrastructure: Why "Hacktivism" Can And Should Influence Cybersecurity Reform* BOSTON UNIVERSITY LAW REVIEW, Vol. 92. 1671-1673

¹²⁸ Ibid

Bez sumnje, veliki sektori američke javnosti su bili pogođeni sajber napadima (PayPal/Visa/MasterCard), nestankom kompanije Soni (Sony) i protestima pokreta „Okupaj“ (Occupy).¹²⁹

S druge strane, haktivizam je krovni pojam koji pokriva niz radnji koje su jedinstvene aktivnosti. Ovo bar važi za objektivne elemente. Neka od ovih dela su već kriminalizovana u različitim jurisdikcijama i predstavljaju krivična dela kao što je, na primer, nezakonit pristup, presretanje i mešanje u kompjuterske podatke. Haktivizam se, međutim, može razlikovati po svom subjektivnom elementu, odnosno specifičnoj nameri da se postigne određeni cilj ili rezultat.¹³⁰

5.3. Postupak, dokazi i harmonizacija zakona

Prekogranična priroda sajber kriminala i činjenje sajber kriminala u elektronskom okruženju su glavne poteškoće sa kojima se suočavaju organi za sprovođenje zakona. Tradicionalne pretpostavke o tome da je počinitelj posmatran kako se priprema da počinji delo ili beži od krivičnog dela više ne važe.¹³¹ Izazovi u istrazi sajber kriminala proizilaze iz kriminalnih inovacija počinitelja, poteškoća u pristupu elektronskim dokazima i iz internih resursa, kapaciteta i logističkih ograničenja. Osumnjičeni često koriste tehnologije anonimizacije i prikrivanja, a nove tehnike brzo prolaze do široke kriminalne publike preko onlajn tržišta kriminala.¹³²

Identifikacija počinioca, istraga i prikupljanje dokaza o zločinu mogu biti zahtevni iz različitih razloga. Pored izazova anonimizacije i zamagljivanja, zemlja, u kojoj se nalazi sajber kriminalac i njegove aktivnosti, možda neće definisati šta je učinjeno kao nezakonito i stoga možda neće moći da ga krivično goni ili sarađuje u njegovom izručenju radi krivičnog gonjenja na drugom mestu; država domaćin možda nema važeće sporazume sa državom žrtvom, koji je obavezuju da pomaže u prikupljanju dokaza koji se mogu koristiti protiv počinioca; ili izuzetno

¹²⁹ Ibid

¹³⁰ Brenner Susan (2004). Toward a Criminal Law for Cyberspace: Distributed Security, BOSTON UNIVERSITY JOURNAL OF SCIENCE & TECHNOLOGY LAW. Vol. 10, Issue 1, pp. 1-109

¹³¹ Ibid

¹³² Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

promenljivi elektronski dokazi su možda uništeni, namerno ili zato što su to rutinski podaci o transakcijama koje nije zadržao provajder internet usluga, a koje je počinitelj koristio da počinio zločin. Sajber prostor čini fizički prostor irelevantnim. Postaje lako viktimizirati nekoga ko je na kraju sveta kao i suseda.¹³³ Dakle, na proceduralnom nivou, glavni problem za nacionalnu primenu zakona je pomirenje istorijske činjenice da je policija operativno i organizaciono lokalizovana unutar granica jurisdikcije (inherentno povezana sa državnim suverenitetom), dok su sajber bezbednost i sajber kriminal globalizovani i poremećene jurisdikcije.¹³⁴

Istrage kibernetičkog kriminala, od strane organa za sprovođenje zakona, zahtevaju spajanje tradicionalnih i novih policijskih tehnika. Dok se neke istražne radnje mogu postići tradicionalnim ovlašćenjima, mnoge proceduralne odredbe se ne prevode dobro sa prostornog, objektno orijentisanog pristupa na pristup koji uključuje elektronsko skladištenje podataka i tokove podataka u realnom vremenu.¹³⁵

Dokaz je sredstvo kojim se utvrđuju činjenice relevantne za krivicu ili nevinost pojedinca na suđenju. Elektronski dokazi su sav takav materijal koji postoji u elektronskom, ili digitalnom, obliku, koji može biti uskladišten ili prolazan i može postojati u obliku računarskih datoteka, prenosa, evidencije, metapodataka ili mrežnih podataka. Digitalna forenzika se bavi vraćanjem, često nestabilnih i lako kontaminiranih, informacija koje mogu imati dokaznu vrednost. Forenzičke tehnike uključuju kreiranje „bit-za-bit“ (bit-for-bit) kopija uskladištenih i izbrisanih informacija kako bi se osiguralo da se originalne informacije ne menjaju i „heševe“ kriptografskih datoteka ili digitalnih potpisa koji mogu da demonstriraju promene u informacijama. To znači da je od strane organa za sprovođenje zakona potreban dovoljan broj forenzičkih istražitelja i dostupnost forenzičkih alata kako ne bi došlo do zaostatka zbog ogromne količine podataka za analizu. Osumnjičeni koriste šifrovanje, što čini pristup ovoj vrsti dokaza teškom za pristupanje. U većini zemalja zadatak analize elektronskih dokaza leži na organima za sprovođenje zakona.¹³⁶

¹³³ Brenner Susan (2004). *Toward a Criminal Law for Cyberspace: Distributed Security*, BOSTON UNIVERSITY JOURNAL OF SCIENCE & TECHNOLOGY LAW. Vol. 10, Issue 1, pp. 1-109

¹³⁴ Više o tome: WALL DAVID (2007). *CYBERCRIME: THE TRANSFORMATION OF CRIME IN THE INFORMATION AGE*. Polity Press.

¹³⁵ Više o tome: Sandage, John et al. eds. (2013). *Comprehensive Study on Cybercrime* (United Nations Office on Drugs and Crime).

¹³⁶ Ibid

Dodatni izazov je upotreba tehnologije od strane organa za sprovođenje zakona – čini se da u ovom trenutku sajber prestupnici bolje koriste tehnološke mogućnosti koje imaju. Prisustvo relevantnog tela sa posebnim znanjem i stručnošću u okviru policijskih snaga je ključni element u efikasnom regulisanju sajber prostora.¹³⁷ Tužioci moraju da vide i razumeju elektronske dokaze kako bi izgradili slučaj na suđenju. Mnoge zemlje u razvoju, na globalnom nivou, nemaju dovoljno resursa kako bi tužioci to i učinili. Kompjuterske veštine tužilaštva su obično niže nego kod istražitelja. Isto važi i za sudije koje se bave visoko specijalizovanim predmetima sajber kriminala. Obuka pravosuđa o zakonu o sajber kriminalu, prikupljanju dokaza i osnovnom i naprednom poznavanju računara predstavlja poseban prioritet.¹³⁸ Mnoge jurisdikcije ne prave pravnu razliku između elektronskih dokaza i fizičkih dokaza.¹³⁹ Iako se pristupi razlikuju, mnoge zemlje smatraju ovo dobrom praksom, jer obezbeđuje pravičnu prihvatljivost pored svih drugih vrsta dokaza. Brojne zemlje van Evrope uopšte ne prihvataju elektronske dokaze, zbog čega je krivično gonjenje sajber kriminala, kao i bilo kog drugog zločina dokazanog elektronskim informacijama, neizvodljivo. Iako zemlje generalno nemaju posebna pravila o elektronskim dokazima, brojne zemlje su se pozvale na principe kao što su: pravilo o najboljim dokazima, relevantnost dokaza, pravilo iz druge ruke, autentičnost i integritet, od kojih svi mogu imati posebnu primenu na elektronske dokaze.¹⁴⁰

Mnoge zemlje imaju elemente pravnog okruženja koji se bavi sajber bezbednošću i sajber kriminalom, ali se ovi nacionalni pravni okviri uveliko razlikuju u pogledu načina na koji se ova pitanja rešavaju.¹⁴¹ U današnjem globalizovanom svetu, zakon se sastoji od mnoštva nacionalnih, regionalnih i međunarodnih pravnih sistema. Interakcije između ovih sistema se javljaju na više nivoa. Kao rezultat toga, odredbe su ponekad u suprotnosti jedna sa drugom, što dovodi do

¹³⁷ Više o tome: WALL DAVID (2007). CYBERCRIME: THE TRANSFORMATION OF CRIME IN THE INFORMATION AGE. Polity Press.

¹³⁸ Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

¹³⁹ Brenner Susan (2004). Toward a Criminal Law for Cyberspace: Distributed Security, BOSTON UNIVERSITY JOURNAL OF SCIENCE & TECHNOLOGY LAW. Vol. 10, Issue 1, pp. 1-109

¹⁴⁰ Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

¹⁴¹ Više o tome: Satola David & Judy Henry (2011). Towards a Dynamic Approach to Enhancing International Cooperation and Collaboration in Cybersecurity Legal Frameworks: Reflections on the Proceedings of the Workshop on Cybersecurity Legal Issues at the 2010 United Nations Internet Governance Forum 37 WILLIAM MITCHELL LAW REVIEW.

sukoba zakona, ili se ne preklapaju u dovoljnoj meri, ostavljajući praznine u nadležnostima. Ove razlike između nacionalnih zakona dovode do pitanja da li, i ako jesu, koliko daleko se nacionalne pravne razlike u zakonima o sajber kriminalu mogu i kako ih smanjiti? Drugim rečima, koliko je važno uskladiti zakone o sajber kriminalu? Ovo se može preduzeti na više načina, uključujući i obavezujuće i neobavezujuće međunarodne ili regionalne inicijative. Osnova harmonizacije može biti jedinstveni nacionalni pristup.¹⁴²

Jedan od glavnih argumenata u korist ujedinjenja zakona u svim jurisdikcijama je izbegavanje sigurnih utočišta za počinioce. Dakle, ako su štetna dela koja uključuju internet kriminalizovana, na primer, u državi A, ali ne i u državi B, počinilac u državi B može biti slobodan da cilja žrtve u državi A putem interneta. U takvim slučajevima, država A ne može, sama, efikasno da zaštiti od efekata takvih transnacionalnih aktivnosti. Čak i tamo gde njen krivični zakon dozvoljava utvrđivanje nadležnosti nad počinioce u državi B, i dalje će biti potreban pristanak ili pomoć države B – bilo u vezi sa prikupljanjem dokaza ili sa izručenjem identifikovanog počinioce. Da bi zaštitila osobe pod svojom jurisdikcijom, malo je verovatno da će država B pomoći tamo gde ponašanje nije inkriminisano u sopstvenoj zemlji.¹⁴³

Harmonizacija takođe može omogućiti globalno prikupljanje dokaza. Usklađivanje procesnog prava je drugi neophodan uslov za efikasnu međunarodnu saradnju. U gornjem primeru, ako država B nema neophodna proceduralna ovlašćenja za ubrzano čuvanje kompjuterskih podataka, na primer, onda država A neće moći da zatraži ovu mogućnost putem uzajamne pravne pomoći. Drugim rečima, zamoljena država može pružiti pomoć samo na svojoj teritoriji, u meri u kojoj bi to mogla da učini za ekvivalentnu nacionalnu istragu.¹⁴⁴

Svetska zdravstvena organizacija (WHO) objavila je novu publikaciju u kojoj su navedena ključna regulatorna razmatranja o veštačkoj inteligenciji (AI) za zdravlje.¹⁴⁵ Publikacija naglašava važnost uspostavljanja bezbednosti i efikasnosti sistema veštačke inteligencije, brzog stavljanja odgovarajućih sistema na raspolaganje onima kojima su potrebni i

¹⁴² Više o tome: Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).

¹⁴³ Ibid

¹⁴⁴ Ibid

¹⁴⁵ <https://www.who.int/news/item/19-10-2023-who-outlines-considerations-for-regulation-of-artificial-intelligence-for-health>

podsticanja dijaloga među zainteresovanim stranama, uključujući programere, regulatore, proizvođače, zdravstvene radnike i pacijente. Sa sve većom dostupnošću podataka o zdravstvenoj zaštiti i brzim napretkom u analitičkim tehnikama, bilo da su mašinsko učenje, zasnovane na logici ili statistički, AI alati bi mogli da transformišu zdravstveni sektor. SZO prepoznaje potencijal AI u poboljšanju zdravstvenih ishoda jačanjem kliničkih ispitivanja; poboljšanje medicinske dijagnoze, lečenja, samopomoći i nege usmerene na osobu; i dopunjavanje znanja, veština i kompetencija zdravstvenih radnika. Na primer, veštačka inteligencija bi mogla da bude korisna u okruženjima sa nedostatkom medicinskih specijalista, npr. u tumačenju skeniranja mrežnjače i radioloških slika između mnogih drugih. Međutim, AI tehnologije, uključujući velike jezičke modele, se brzo primenjuju, ponekad bez potpunog razumevanja kako mogu da rade, što bi moglo da koristi ili naškodi krajnjim korisnicima, uključujući zdravstvene radnike i pacijente.

Kada koriste zdravstvene podatke, sistemi veštačke inteligencije bi mogli da imaju pristup osetljivim ličnim informacijama, što zahteva čvrste zakonske i regulatorne okvire za zaštitu privatnosti, bezbednosti i integriteta, kojima ova publikacija želi da pomogne u postavljanju i održavanju. „Veštačka inteligencija obećava veliko zdravlje, ali takođe dolazi sa ozbiljnim izazovima, uključujući neetičko prikupljanje podataka, pretnje sajber bezbednosti i sve veće pristrasnosti ili dezinformacije“, rekao je dr Tedros Adhanom Gebrejesus, generalni direktor SZO. „Ove nove smernice će podržati zemlje da efikasno regulišu veštačku inteligenciju, da iskoriste njen potencijal, bilo da se radi o lečenju raka ili otkrivanju tuberkuloze, uz minimiziranje rizika. Kao odgovor na rastuću potrebu zemlje da odgovorno upravlja brzim porastom zdravstvenih tehnologija AI, publikacija navodi šest oblasti za regulisanje AI za zdravlje. Da bi podstakla poverenje, publikacija naglašava važnost transparentnosti i dokumentacije, kao što je dokumentovanje celog životnog ciklusa proizvoda i praćenje razvojnih procesa. Za upravljanje rizikom, pitanja kao što su 'namenska upotreba', 'kontinuirano učenje', ljudske intervencije, modeli obuke i pretnje sajber bezbednosti moraju se sveobuhvatno baviti, sa modelima koji su jednostavni.

Eksterna validacija podataka i jasnost o nameravanoj upotrebi veštačke inteligencije pomaže da se obezbedi bezbednost i olakša regulacija. Posvećenost kvalitetu podataka, kao što je rigorozno procenjivanje sistema pre objavljivanja, je od vitalnog značaja da se obezbedi da sistemi ne pojačavaju pristrasnosti i greške. Izazovi koje postavljaju važni, složeni propisi, kao

što su Opšta uredba o zaštiti podataka (GDPR) u Evropi i Zakon o prenosivosti i odgovornosti zdravstvenog osiguranja (HIPAA) u Sjedinjenim Američkim Državama, rešavaju se sa naglaskom na razumevanju delokruga nadležnosti i zahtevima za saglasnost, u cilju zaštite privatnosti i zaštite podataka. Podsticanje saradnje između regulatornih tela, pacijenata, zdravstvenih radnika, predstavnika industrije i vladinih partnera može pomoći da se osigura da proizvodi i usluge ostanu u skladu sa propisima tokom njihovog životnog ciklusa. Sistemi veštačke inteligencije su složeni i zavise ne samo od koda sa kojim su izgrađeni, već i od podataka na kojima su obučeni, a koji potiču iz kliničkih okruženja i interakcija korisnika. Bolja regulativa može pomoći u upravljanju rizicima da bi se izbegla pristrasnost u podacima o obuci AI. Na primer, AI modelima može biti teško da precizno predstavljaju raznolikost populacija, što dovodi do pristrasnosti, netačnosti ili čak neuspeha. Da bi se ublažili ovi rizici, propisi se mogu koristiti kako bi se osiguralo da su atributi kao što su pol, rasa i etnička pripadnost ljudi koji su predstavljeni u podacima o obuci prijavljeni i da skupovi podataka namerno budu reprezentativni.

GLAVA ŠESTA: AI U PRAVNIČKOJ STRUCI - DEFINISANI I TRENUTNO KORIŠTENI PROGRAMI

6.1. Programi veštačke inteligencije kreirani za pravnu industriju

AI je definisan na nekoliko načina tokom svog postojanja, ali generalno AI je „sposobnost mašine da imitira inteligentno ljudsko ponašanje.“¹⁴⁶ Ovo uključuje „prepoznavanje govora i objekata, donošenje odluka na osnovu podataka i prevođenje jezici.“¹⁴⁷ AI može da oponaša ljudsku inteligenciju na dva načina: prvo, AI programi mogu biti obučeni unosom podataka, što je istorijski bio jedini način na koji su AI programi mogli da oponašaju ljudsku inteligenciju; drugo, napredniji AI programi mogu sami da uče putem pokušaja i grešaka. Iako koncept veštačke inteligencije postoji već duže vreme, upotreba naprednijih programa u pravnoj oblasti beleži tek neznatan razvoj. Početak regulisanja AI u pravnoj oblasti još uvek je skroman.¹⁴⁸ Bilo je samo nekoliko kompanija koje su razvile AI programe prilagođene pravnim poslovima, a programske funkcije su uglavnom bile ograničene na eDiscoveri, pregled ugovora i dužnu pažnju. Štaviše, prvobitno jedini subjekti koji su koristili AI u pravnom polju su bile najveće advokatske firme koje su radile na najvećim poslovima. Vremenom se proširila uloga programera i korisnika AI u pravnom polju. Danas sve više kompanija razmatra kako se AI može koristiti. Za upotrebu izvan eDiscoveri-a, to je bio uski fokus: pregled ugovora i pravna pažnja koju koriste najveće advokatske firme za najveće poslove. Danas, AI više nije dostupna samo najbogatijim advokatskim kancelarijama. Od februara 2018. National Law Journal je identifikovao preko pedeset kompanija koje nude programe veštačke inteligencije kreirane za

¹⁴⁶ *Artificial intelligence*, MERRIAM-WEBSTER, <https://www.merriam-webster.com/dictionary/artificial%20intelligence> [https://perma.cc/3ZBP-PLAJ]

¹⁴⁷ Lauri Donahue, *A Primer on Using Artificial Intelligence in the Legal Profession*, HARV. J. L. & TECH. (Jan. 3, 2018), <http://jolt.law.harvard.edu/digest/a-primer-on-using-artificial-intelligence-in-the-legal-profession> [https://perma.cc/H65H-6A5A]

¹⁴⁸ Keith Mullen, *Artificial Intelligence: Shiny Object? Speeding Train?*, AM. BAR ASS'N RPTE EREPORT 2018 FALL ISSUE, https://www.americanbar.org/groups/real_property_trust_estate/publications/ereport/rpte-2018/artificial-intelligence/ [https://perma.cc/HCE7-WNVV]

pravnu industriju.¹⁴⁹ AI sada može da se koristi za sastavljanje ugovora, pregled ugovora, digitalni potpis, pravna pitanja i upravljanje pitanjima, automatizaciju ekspertize, pravnu analitiku, zadatke menadžmenta i apstraktima zakupa. Manje firme takođe počinju da se oslanjaju na AI tehnologiju, sa namerom da im ova tehnologija pomogne da se takmiče na legalnom tržištu.

Trenutni podaci pokazuju da oko „85 posto advokata u manjim advokatskim kancelarijama koristi veštačku inteligenciju, smanjujući ili eliminišući ono što su nekada bile prednosti resursa i osoblja u većim advokatskim firmama.”¹⁵⁰ Kako tehnologija veštačke inteligencije bude naprednija, funkcije, proizvođači i korisnička baza AI će verovatno nastaviti da rastu.

6.1.1. Napredne platforme za pravno istraživanje

Pravna istraživanja su jedna oblast u kojoj je veštačka inteligencija napravila značajan napredak. Pravna istraživanja su prešla dug put od vremena kada su studenti prava i saradnici morali da čitaju teške knjige slučajeva kako bi pronašli relevantan presedan. Danas većina advokata koristi online platforme za pravna istraživanja poput LexisNexis ili Westlaw koje koriste AI tehnologiju. Poslednjih godina razvijene su još naprednije platforme za pravna istraživanja koje uključuju modernije AI tehnologije. Jedan primer je ROSS Intelligence. ROSS Intelligence je pokrenut 2018. godine i oglašava se kao „prvi veštački inteligentni advokat na svetu“.¹⁵¹ Program košta 69 dolara mesečno po godišnjem planu, i koristi ga nekoliko velikih advokatskih firmi, uključujući Baker Hostetler, Latham & Watkins, Jackson Levis, i Dentons.

¹⁴⁹ Mullen, *supra* note 7 (“As recent as 2014, only a handful of companies pointed artificial intelligence (“AI”) at legal documents. For uses outside of eDiscovery, it was a narrow focus: contract reviews and legal due diligence – used by the largest of law firms on the grandest of deals.”)

¹⁵⁰ Jake Heller, *Is AI the Great Equalizer For Small Law?*, ABOVE THE LAW (Aug. 15, 2018), <https://abovethelaw.com/2018/08/is-a-i-the-great-equalizer-for-small-law/> [https://perma.cc/7UPA-CK4W].

¹⁵¹ Matthew Griffin, *Meet Ross, The World's First AI Lawyer*, 311 INST. (Jul. 11, 2016), <https://www.311institute.com/meet-ross-the-worlds-first-ai-lawyer/> [https://perma.cc/89RH-C9M9].

6.1.2. Aplikacije za pravnu samopomoć

Pored toga što su postali napredniji, AI programi su takođe postali dostupniji, posebno preko legalnog tržišta aplikacija za samopomoć. Jedan primer pravne aplikacije za samopomoć je DoNotPai. DoNotPai je 2016. godine pokrenuo Džošua Brauder, koji je u to vreme bio devetnaestogodišnji student koji je kreirao aplikaciju da „pomogne svojoj porodici i prijateljima da ospore kazne za parkiranje.“¹⁵² Aplikacija je stekla popularnost, imala je mogućnosti koje su se proširile i sada se proglašava kao „Prvi robot advokat na svetu“. Aplikacija je besplatna za preuzimanje i trenutno omogućava korisnicima da podnesu tužbu bilo kom sudu za sporove male vrednosti u zemlji, dobiju zelene karte i vize, bore se protiv naknada za kreditne kartice, tuže tehnološke kompanije za kršenje podataka i, naravno, da se bore protiv kazni za parkiranje. Da bi koristili aplikaciju, korisnici moraju jednostavno odgovoriti na nekoliko relevantnih pitanja. Aplikacija zatim uzima njihove odgovore da automatski popuni pravne formulare koji se mogu poslati direktno potrebnom primaocu. Od svog pokretanja, aplikacija je uspešno poništila 160.000 kazni za parking. Aplikacije kao što je DoNotPai zaobilaze potrebu za ljudskim advokatima i povećavaju pristup pravdi.

6.1.3. Pregled ugovora

Pregled ugovora je još jedna oblast u kojoj AI počinje da ostvaruje veliki napredak. Nekoliko kompanija, kao što su Salesforce, Home Depot i eBay, koriste AI tehnologiju za pregled ugovora u svojim svakodnevnim operacijama. Kako sve više advokata počinje da koristi AI za pregled ugovora, praksa postaje uobičajena. Pre doba veštačke inteligencije, pregled ugovora je bio posao ljudskih advokata, kao jedan od dosadnijih zadataka advokata, pa nije ni čudo što je to jedan od prvih zadataka koji se prenosi na AI programe. AI programi koji mogu „precizno da čitaju ugovore u bilo kom formatu, pružaju analitiku o podacima izvučenim iz

¹⁵² Julie Fishbach, *Coder, 19, Builds Chatbot That Fights Parking Tickets*, NBC NEWS (Jul. 21, 2016), <https://www.nbcnews.com/feature/college-game-plan/coder-19-builds-chatbot-fights-parking-tickets-n612326> [https://perma.cc/UEN7-DBAM].

ugovora i izvlače podatke o ugovorima mnogo brže nego što bi to bilo moguće sa timom advokata“ već postoje.

Nekoliko startapa uključujući Lavgeek, Klariti, Clearlav i LekCheck je prilagodio ovu tehnologiju pravnom polju. Ovi startapi nastoje da kreiraju programe koji „automatski unose predložene ugovore, analiziraju ih u potpunosti koristeći tehnologiju obrade prirodnog jezika (NLP) i određuju koji delovi ugovora su prihvatljivi, a koji problematični.” Međutim, ovi programi u nastajanju nisu potpuno lišeni ljudskog elementa. Advokati tek treba da donesu konačne suštinske odluke o tačnom sadržaju i jeziku koji se koristi u ugovoru nakon što pregledaju sugestije iz programa AI. Advokati takođe treba da budu ti koji će pregovarati o ugovoru. Bez obzira na to, „kako NLP mogućnosti napreduju, nije teško zamisliti budućnost u kojoj se ceo proces sprovodi od strane AI programa koji su ovlašćeni, u okviru unapred programiranih parametara, da sklapaju sporazume.“ Ugovaranje ima pravo. potencijal da postane sve automatizovaniji proces.

6.2. Korišćenje AI za unapređenje kvaliteta pravnih usluga

AI programi imaju potencijal da poboljšaju kvalitet pravnih usluga povećanjem tačnosti i efikasnosti advokata. Napredne platforme za pravna istraživanja opremljene AI su učinile pravna istraživanja bržim i lakšim, dajući advokatima kapacitet da urade više za kraće vreme. Ove napredne platforme za pravna istraživanja takođe omogućavaju advokatima da sa lakoćom provere svoj rad, povećavajući svoju tačnost. U kontekstu revizije ugovora, AI programi su već pokazali sposobnost da rade brže i sa višom stopom tačnosti od ljudskih pravnika.¹⁵³ U takmičenju za reviziju ugovora između iskusnih korporativnih advokata i AI, AI program je „postigao nivo tačnosti uočavanja rizika u ugovorima od 94%“ za 26 sekundi. S druge strane, advokati su u proseku „potrošili 92 minuta da bi postigli nivo tačnosti od 85%“. Povećana tačnost i efikasnost takođe mogu uštedeti novac klijentima i povećati profit advokatima. Brži rad

¹⁵³ Audrey Herrington, *How AI Can Make Legal Services More Affordable*, SMARTLAWYER (Jul. 23, 2019), <http://www.nationaljurist.com/smartlawyer/how-ai-can-make-legal-services-more-affordable> [https://perma.cc/TJ97-CA4H].

sa višim stopama tačnosti verovatno znači manje naplativih sati koje naplaćuju advokati, što znači više novca uštedenog za klijente.

Iako ovo može u početku izgledati kao gubitak profita za advokate, to bi zapravo moglo da proizvede veće profitne marže. To je zato što brži rad sa višom stopom tačnosti može podstaći klijente da se vrate sa novim poslovima i omogućiti advokatskim firmama da prihvate više klijenata. Iz ovih razloga, zaista bi mogla biti izuzetno skupa odluka da se tehnologija veštačke inteligencije ne koristi u svojoj pravnoj praksi. U poređenju sa kompanijama kao što su Google i Adobe, koje imaju bruto marže od šezdeset do devedeset procenata, advokatske firme moraju da se bave utvrđenom strukturom troškova i bore se da njihove marže pređu preko četrdeset procenata. AI može pomoći advokatskim firmama da se oslobode postojeće strukture troškova kako bi povećali marže. Dok su povećana tačnost i efikasnost i niži troškovi od ogromne koristi za pravnu industriju, upotreba AI takođe implicira zabrinutost u oblasti pravne etike. Modelpravila profesionalnog ponašanja, koja služe kao etičke smernice za pravnike, usvojena su u svom originalnom obliku 1983. godine.¹⁵⁴ Dakle, Model pravila nisu kreirana da bi se razmatrali napredniji AI programi kao što su ROSS Intelligence i DoNotPai. Ipak, Pravila su napisana sa namerom da budu prilagodljiva modernim vremenima. U uvodu se navodi da je pri donošenju Pravila „želja pisaca bila da sačuvaju sve što je vredno i trajno u vezi sa postojećim Modelom pravila, a da ih istovremeno prilagode realnostima savremene pravne prakse i granicama profesionalne discipline.”¹⁵⁵ Dakle, Model pravila profesionalnog ponašanja su u stanju da regulišu upotrebu naprednijih AI programa u pravnoj oblasti. Jedna etička implikacija koja proizilazi iz upotrebe AI u pravnom polju je kompetencija advokata. Model pravila 1.1 zahteva od advokata da kompetentno zastupaju klijente. Pravilo kaže: „Advokat će obezbediti kompetentno zastupanje klijenta. Kompetentno zastupanje zahteva pravno znanje, veštinu, temeljnost i pripremu koja je razumno neophodna za zastupanje.” Pored toga, komentar 8 na

¹⁵⁴ MODEL RULES OF PROF'L CONDUCT: PREFACE (2019), https://www.americanbar.org/groups/professional_responsibility/publications/model_ruler_of_professional_conduct/model_rules_of_professional_conduct_preface/ [https://perma.cc/E8K6-ZYY9].

¹⁵⁵ MODEL RULES OF PROF'L CONDUCT: PREFACE, ETHICS 2000 CHAIR'S INTRODUCTION (2002), https://www.americanbar.org/groups/professional_responsibility/publications/model_ruler_of_professional_conduct/model_rules_of_professional_conduct_preface/ethics_2000_cchai_introduction/ [https://perma.cc/7F3K-N4XG].

pravilo 1.1, koji je dodat 2012. godine, proširuje koncept kompetentnog zastupanja u svetlu tehnoloških napretka u pravnoj oblasti.

Čitanje pravila 1.1 i komentara 8 zajedno ukazuje na to da advokati imaju etičku obavezu da budu u toku sa tehnologijom koja se koristi u pravnoj oblasti kako bi obezbedili kompetentno zastupanje klijenata. Tačnije, čini se da pravilo ukazuje da u doba veštačke inteligencije advokati imaju dve etičke dužnosti. Prvo, pravnici moraju imati osnovno razumevanje programa AI koje odlučuje da koriste u svojoj praksi. Pošto je AI grana računarske nauke i često uključuje tehničko znanje van stručnosti većine advokata, razumevanje načina na koji AI programi funkcionišu može biti teško za advokate. Bez obzira na to, advokati i dalje moraju da održavaju osnovnu liniju znanja o AI programima koje koriste, uključujući: (1) zašto AI program daje svoje rezultate i (2) za šta je AI program sposoban, a za šta nije.¹⁵⁶ Bez ovog osnovnog znanja, advokati neće moći da koriste programe veštačke inteligencije sa punom kompetencijom, čime će ugroziti njihovu sposobnost da obezbede kompetentno zastupanje svojim klijentima. Brojne države su to shvatile i uspostavile sopstvena pravila koja regulišu nadležnost advokata u pogledu upotrebe tehnologije. Ukupno trideset šest država je već primenilo svoja sopstvena pravila koja regulišu upotrebu tehnologije.

Pravilo Floride „sugerše da bi kontinuirano obrazovanje moglo biti neophodno da bi se razumeli rizici povezani sa korišćenjem tehnologije.“¹⁵⁷ Štaviše, „Njujork je proglasio pravilo da advokati mora koristiti razumnu pažnju i ostati u toku sa tehnološkim napretkom.“¹⁵⁸ Razumevanje tehnologije koju neko koristi u svojoj praksi je imperativ za pružanje kompetentne pravne usluge. Drugi ključni aspekt koji pravnici treba da shvate u ispunjavanju dužnosti kompetentnosti je da rezultati veštačke inteligencije ne bi trebalo automatski da budu prihvaćeni kao istiniti. Iako su mnogi noviji AI programi tehnički ispravni, oni su i dalje nesavršeni. Advokati i dalje moraju biti oprezni kada koriste ove programe. Stoga, advokati moraju (1)

¹⁵⁶ Lat, *supra* note 1; Tashea & Economou, *supra* note 52; Roy D. Simon, *Artificial Intelligence, Real Ethics*, 90-APR N.Y. ST. B.J. 34, 34 (2018).

¹⁵⁷ Katherine Medianik, *Artificially Intelligent Lawyers: Updating the Model Rules of Professional Conduct in Accordance with the New Technological Era*, 39 CARDOZO L. REV. 1497, 1515 (2018) (citing Fla. Bar Prof'l Ethics Comm., Op. 06-2 (2006)).

¹⁵⁸ Medianik, *supra* note 55 (quoting N.Y. State Bar Ass'n Comm. on Prof'l Ethics, Formal Op. 782 (2004)).

redovno proveravati da bi bili sigurni da AI program koji koriste ispravno radi i (2) pregledati rezultate programa kako bi obezbedili kompetentno pravno zastupanje. Obaveza advokata da pregledaju programe veštačke inteligencije i rezultate koje oni proizvode dodatno je potkrepljena Modelom pravila 5.3, koje utvrđuje obavezu advokata da nadgledaju neadvokate. Pravilo 5.3(b), najrelevantnija odredba, kaže: „advokat koji ima ulogu direktnog nadzornog organa nad neadvokatima će uložiti razumne napore da osigura da je ponašanje te osobe u skladu sa profesionalnim obavezama advokata.” Iako AI program nije osoba, to je mašina, on će oponašati ljudsku inteligenciju.

Prema tome, prema pravilu 5.3, AI bi se mogao smatrati neadvokatom kome advokat delegira posao, što pokreće dužnost advokata da osigura da je rad koji proizvodi AI program kompetentan. Čitanje Modela pravila u modernim vremenima ukazuje na to da advokati moraju da imaju osnovno razumevanje o tome kako deluju AI programi koje koriste, a ne da automatski prihvate njihova rešenja, da bi mogli da obezbede kompetentno pravno zastupanje. Trenutno, ovo može izgledati samo po sebi razumljivo. U praksi prava, slepo prihvatanje rezultata programa ne samo da bi bilo neetično, već i nepromišljeno. U budućnosti, međutim, ovi osnovni pojmovi o tome kako koristiti AI na etički način mogu postati mnogo manje očigledni.

Sve više oslanjanje na AI u pravnom polju može izazvati dodatnu zabrinutost u pogledu etike u pogledu kompetentnog pravnog zastupanja. Iako gore pomenuta rasprava upozorava na etičke implikacije korišćenja AI, mogu postojati i etičke implikacije odbijanja upotrebe AI. Kako se tehnologija veštačke inteligencije poboljšava i postaje sve raširenija u pravnom polju, odbijanje upotrebe veštačke inteligencije u nečijoj pravnoj praksi može značajno da ugrozi sposobnost da obezbedi kompetentno pravno zastupanje. Ovo je posebno tačno zato što ukoliko se više koristi veštačka inteligencija, program postaje korisniji: „alati veštačke inteligencije u narednih nekoliko godina će iskoristiti privatne podatke advokatskih firmi kako bi stvorili jedinstvene uvide koristeći kolektivno iskustvo advokata i njihov radni proizvod. Dakle, odbijanje upotrebe tehnologije koja čini pravni rad preciznijim i efikasnijim može se smatrati odbijanjem da se klijentima obezbedi kompetentno pravno zastupanje.

6.3. Korišćenje AI za poboljšanje pristupa pravdi

Još jedan pozitivan efekat korišćenja veštačke inteligencije u pravnom polju je mogućnost povećanja individualnog pristupa pravdi. Ujedinjene nacije definišu pristup pravdi kao „osnovni princip vladavine prava“. Ujedinjene nacije takođe objašnjavaju šta pristup pravdi podrazumeva: U nedostatku pristupa pravdi, ljudi ne mogu da ostvaruju svoja prava, osporavaju diskriminaciju ili pozivaju donosioce odluka na odgovornost. Preambula Modela pravila predviđa da advokati treba da traže poboljšanje zakona i pristup pravnom sistemu za sve one koji zbog ekonomskih ili društvenih prepreka ne mogu da priušte ili obezbede adekvatan pravni savet.¹⁵⁹ Sjedinjene Države trenutno doživljavaju krizu pristupa pravdi jer nema dovoljno ljudi da priušte ili dobiju pravne usluge kada su im potrebne.

6.4. Pravna etika

Iako zakon još uvek nije jasno definisao kako koristiti programe veštačke inteligencije u skladu sa principima pravne etike, smernice bi ipak trebalo da budu izvedene iz Modela pravila. Model pravila zahteva ljudski element u radu advokata. Advokati ne smeju da koriste AI programe da zamene svoj rad bez kršenja svoje dužnosti da obezbede kompetentno zastupanje, navedeno je u Pravilu 1.1.90. Advokati mogu da koriste AI programe, međutim, samo da bi poboljšali svoj rad. Međutim, granica između zamene i povećanja nije uvek jasna. Kako bi se osiguralo da je obaveza kompetentnosti ispunjena kada se koristi AI, advokati treba da se pridržavaju sledećih sugestija. Da bi osigurali kompetentno zastupanje, advokati treba da imaju osnovno razumevanje programa AI koje odluče da koriste u svojoj praksi i da se uzdrže od automatskog prihvatanja rezultata AI programa koje koriste. Ovo se odnosi na sve programe

¹⁵⁹ MODEL RULES OF PROF'L CONDUCT: PREFACE, ETHICS 2000 CHAIR'S INTRODUCTION (2002), https://www.americanbar.org/groups/professional_responsibility/publications/model_ruler_of_professional_conduct/model_rules_of_professional_conduct_preamble/ethics_2000_cchai_introduction/ [https://perma.cc/7F3K-N4XG].

veštačke inteligencije, bilo da se radi o naprednoj platformi za pravna istraživanja ili programu koji vrši pregled ugovora.

Osnovno razumevanje AI programa koji se koristi podrazumeva razumevanje zašto AI program proizvodi svoje rezultate i za šta je program sposoban, a za šta nije. Ne prihvatanje automatskih rezultata AI programa koji se koristi kao istinito podrazumeva redovnu proveru programa kako bi se uverio da radi ispravno, a to podrazumeva i pregled rezultata programa. Programi veštačke inteligencije koji ne uključuju ljudske advokate ne bi trebalo da pružaju pravne savete jer bi to bila neovlašćena pravna praksa prema pravilu 5.5.91. Kada se koriste programi kao što su aplikacije za pravnu samopomoć i robo-obraci, na primer, nije dozvoljeno da AI program daje suštinske pravne savete, jer bi to bila neovlašćena pravna praksa i stoga kršenje Modela pravila profesionalnog ponašanja. Dakle, nezaobilazni ljudski element u advokaturi u doba veštačke inteligencije funkcioniše u oba smera. Ljudski advokati ne bi trebalo da se u potpunosti oslanjaju na programe veštačke inteligencije da bi davali pravne savete, a AI programi ne bi trebalo da daju pravne savete osim ako je ljudski advokat uključen.

6.5. Pravci buduće regulative

Uprkos medijskom ludilu oko objavljivanja OpenAI-ja ChatGPT i potonja eksplozija generativne AI proizvoda, AI je već decenijama među nama. Gugl svakodnevno koristi mašinsko učenje da napaja svoje algoritme da obezbede relevantnije informacije preko Google pretraživača. Drugi primer je gde je AI omogućio računarima da komuniciraju sa ljudima putem uređaja za glasovnu kontrolu, kao što je Amazon Echo. AI je takođe korišćen u medicini da prouči šta ljude čini zdravim. Kompanije kao što su Microsoft i Alphabet su koristile veštačku inteligenciju za analizu svetlosnih signala za proizvodnju slike crnih rupa. Vlade takođe imaju implementiran AI, na primer država Viktorija, gde Vlada koristi AI senzore za uključivanje bezbednosnih metrika na saobraćajne signale u realnom vremenu, a viktorijanska policija koristi tehnologije automatskog prepoznavanja registarskih tablica u svom nadzoru javnosti. Međutim, AI tehnologije predstavljaju niz jedinstvenih izazova koji povećavaju rizik od štete ljudima, organizacijama i društvu. Ovi uključuju:

Hiperskalabilnost automatizacije omogućene AI, što rezultira široko rasprostranjenim negativnim uticajima u slučaju grešaka u načinima na koje sistemi veštačke inteligencije donose predviđanja ili odluke.

Poteškoće u ispitivanju i objašnjavanju izlaza izuzetno složenih algoritamskih modela, koje ljudski programeri nisu eksplicitno programirali. Nedostatak jasnih principa za odgovarajuću dodelu odgovornosti, vlasništva i odgovornosti za rezultate AI modela, posebno kada su sistemi veštačke inteligencije kreirani kroz složene i međuzavisne lance vrednosti. Prilagodljivost i samopromena (tj. „učenje“) za AI modele znače da pristup „postavi i zaboravi“ poput tradicionalnih softverskih proizvoda ne funkcioniše. Nema sumnje da moć veštačke inteligencije otključava značajnu efikasnost i mogućnosti, sa skoro neograničenim obimom primene. Međutim, sa ovim mogućnostima dolazi mnoštvo novih rizika i potencijalnih šteta koje AI predstavlja organizacijama, ljudima i društvu koje se moraju uzeti u obzir. Neki od ovih primera su:

Jedna od prvih briga sa kojima se moraju pozabaviti AI programeri, korisnici i operateri jesu pitanja privatnosti koja su pokrenuta razvojem i upotrebom AI, posebno u oblasti pristanka i uključivanja ličnih podataka u AI modeliranje i unose. Pristrasnost i diskriminacija AI može biti pod uticajem slučajne pristrasnosti i diskriminacije. Postoje inherentni problemi koji su ugrađeni u skupove ulaznih podataka koji se koriste za obuku AI algoritama koji mogu dovesti do neželjenih ishoda.

Realnost načina na koji su sistemi veštačke inteligencije dizajnirani (u suštini to je „crna kutija“) i implementirani znači da se postojeći standardi odgovornosti moraju promeniti.

AI će nesumnjivo imati transformativne efekte na ekonomiju, jer su čak su i računari, kada su prvi put predstavljeni, značajno promenili mnoge uslove zapošljavanja. Postoje ekonomski i društveni aspekti kojima se mora upravljati odgovarajućom regulativom AI.

Ovi rizici i štete, između ostalog, razlog su zašto se vlade širom sveta bore da uspostave odgovarajuće kontrole i smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju. Otvoreno pismo od 22. marta 2023. godine, zahtev za trenutnim moratorijumom na obuku određenih sistema veštačke inteligencije zbog potencijalnih opasnosti od nekontrolisanog razvoja veštačke inteligencije, potpisalo je više od 1.100 ljudi u jednoj nedelji, uključujući suosnivača Apple-a Stiva Voznijaka i tehnološkog milijardera Elona Maska. U pismu se sugerise da programeri veštačke inteligencije treba da rade sa kreatorima politike i

vladama kako bi zajednički razvili robusne sisteme upravljanja veštačkom inteligencijom i protokole za napredni dizajn i razvoj AI.

Regulisanje veštačke inteligencije je teško zbog potrebe da se uspostavi ravnoteža između omogućavanja inovacije i evolucije, dok se AI koristi na bezbedan i odgovoran način. Jedna stvar je jasna a to je da omogućavanje nesputanog razvoja veštačke inteligencije može biti neodrživo.

Postoji niz pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI. Vrlo malo je rečeno o tome koja organizacija u lancu razvoja AI treba da popravi takva pitanja kada se pojave, otvarajući jaz u odgovornosti. Životni ciklus veštačke inteligencije je složen i uključuje mnogo različitih operatera, od kojih svi imaju potencijal da utiču na kvalitet tehnologije.

Postoje, takođe, problemi sa samim konceptom vlasništva nad podacima. Zakonski ne postoji jedinstvena definicija ili okvir za vlasništvo nad podacima, a različite jurisdikcije mogu usvojiti različite pristupe u zavisnosti od prirode, izvora i upotrebe podataka. Dopisivanje prava vlasništva podataka podacima koje koristi ili generiše AI predstavlja brojne pravne izazove, uključujući: 1. Podaci koje obrađuje, generiše ili koristi AI često uključuju preklapanje interesa različitih strana. To takođe nije jasna linearna podela prava, već zavisi od prirode podataka, doprinosa strana koje su obezbedile, kreirale ili kontrolisale podatke i ugovornih ili pravnih aranžmana između tih strana. Na primer, ko poseduje podatke ili sadržaj koji generiše sistem veštačke inteligencije koji je sakupio podatke iz javnog izvora? Ko poseduje podatke ili sadržaj koji se generiše korišćenjem nejavnih podataka, npr. platforma zasnovana na pretplati? Ovo će se nesumnjivo onda okrenuti razmatranjima licenciranja, kao i pravima na privatnost i poverljivost. Kao i kod većine pitanja veštačke inteligencije, još uvek nema utvrđenog stava po ovom pitanju. 2. Kada se utvrdi da podaci koje koristi ili kreira AI zapravo mogu biti predmet vlasničkih prava, onda se postavlja pitanje kako se ta prava mogu sprovesti. Uopšteno govoreći, ugovorni podaci i poverljivost i pravične obaveze poverljivosti i prava intelektualne svojine pružaju osnovu i okvir za uspostavljanje prava vlasnika podataka da ograniči upotrebu i otkrivanje podataka (osim ličnih podataka). Ali čak i tada, ostaje otvoreno pitanje kako ovi vlasnici podataka mogu da kontrolišu pristup, otkrivanje, prenos, brisanje ili modifikaciju podataka koji se nalaze u „crnoj kutiji“?

6.6. Sajber bezbednost na nacionalnom nivou

Bil Clinton (Bill Clinton) je 1997. osnovao predsedničku komisiju za zaštitu kritične infrastrukture (PCCIP). Njegov konačni izveštaj „Kritične osnove: zaštita američke infrastrukture“, identifikovao je osam sektora američke privrede koji imaju stratešku bezbednosnu vrednost: telekomunikacije, električna energija, gas i nafta, voda, transport, bankarstvo i finansije, hitne službe i kontinuitet vlade. PCCIP je prepoznao ne samo zavisnost nacije od njene kritične infrastrukture (CI), već i zavisnost savremenih KI od IT sistema. Takođe se navode „prožimajuće“ ranjivosti koje su bile otvorene za napad „širokog spektra“ potencijalnih protivnika. Pošto je pretnja sajber napadima na KI strateškog obima, nacionalni odgovor mora biti jednak zadatku, a to podrazumeva podizanje svesti javnosti, ulaganje u obrazovanje, naučna istraživanja, razvoj sajber prava i međunarodne saradnje. Stvorene su i nove agencije i ekonomski „sektorski koordinatori“, ali PCCIP je naglasio da nijedan pojedinac ili organizacija ne može biti odgovoran za CIP jer su kritične infrastrukture kolektivna imovina kojom vlada i privatni sektor moraju upravljati zajedno.¹⁶⁰ Rusija je 4. decembra 1998. sponzorisa Rezoluciju Generalne skupštine Ujedinjenih nacija (UN) 53/70, „Razvoj u oblasti informacija i telekomunikacija u kontekstu međunarodne bezbednosti“. U njemu se navodi da nauka i tehnologija igraju važnu ulogu u međunarodnoj bezbednosti i da, iako moderna informaciona i komunikaciona tehnologija (IKT) nudi civilizaciji „najšire pozitivne mogućnosti“, IKT je ipak podložna zloupotrebi od strane kriminalaca i terorista. Od država članica UN-a je stoga zatraženo da obaveste generalnog sekretara o svojoj zabrinutosti u vezi sa zloupotrebom IKT-a i ponude predloge za unapređenje njihove bezbednosti u budućnosti. UN su usvojile Rezoluciju 53/70 4. januara 1999.¹⁶¹ Najuspešniji međunarodni sporazum o sajber bezbednosti do sada jeste Konvencija Saveta Evrope o sajber kriminalu, otvoren je za potpisivanje 2001. godine. Ovaj sporazum ima za cilj regulisanje kršenja autorskih prava, prevara, dečije pornografije i kršenje politike bezbednosti mreže. Nudi smernice za sprovođenje zakona u vezi presretanja podataka i pretraživanja računarskih mreža. Njegov krajnji cilj je zajednička politika o sajber kriminalu, širom sveta, putem nacionalnog zakonodavstva i

¹⁶⁰ Neumann, P.G. (1998). "Protecting the Infrastructures," *Communications of the ACM* 41(1) 128.

¹⁶¹ Ibid

međunarodne saradnje. Trenutno, Konvencija ima četrdeset sedam nacionalnih potpisnika, trideset ratifikacija i predstavlja primarni pravni instrument koji je dostupan nacionalnim državama u pogledu sajber bezbednosti.¹⁶² Na nivou nacionalne bezbednosti, desile su se najvidljivije promene u strategiji sajber odbrane. Nastupila je prekretnica 2008. godine, kada su nepoznati hakeri uspešno kompromitovali poverljive vojne sisteme u „najznačajnijem proboju” američkih vojnih kompjutera ikada. Zamenik američkog sekretara za odbranu Vilijam Lin (William Lynn) napisao je za „Forin Afers“ (Foreign Affairs) da protivnici sada imaju moć da ometaju kritičnu američku informacionu infrastrukturu i da asimetrična priroda hakovanja znači da „desetak” kompjuterskih programera može predstavljati pretnju po nacionalnu bezbednost. Dugoročno, Lin je verovao da hakovanje računara može dovesti do gubitka intelektualne imovine koja može da liši naciju ekonomske vitalnosti. Najopipljiviji odgovor SAD-a bilo je stvaranje posebnog ogranka njihove vojske „Sajber Komanda“ (Cyber Command - USCYBERCOM) u 2010. godini. USCYBERCOM ima tri osnovne misije: odbranu računarske mreže, koordinaciju nacionalnih resursa za sajber ratovanje i vezu sajber-bezbednosti sa domaćim i stranim partnerima. Njegova prva misija, odbrana, oslanja se na kombinaciju tradicionalnih najboljih praksi u kompjuterskoj bezbednosti i poverljivim obaveštajnim pretnjama, kako bi se stvorio jedinstven, „aktivan” položaj američke sajber odbrane.¹⁶³ Potraga za strateškom sajber odbranom je svoj poslednji korak napred napravila u Lisabonu, u Portugaliji, u novembru. 2010. godine, kada je dvadeset osam nacionalnih lidera iz najvećeg svetskog političkog i vojnog saveza objavilo novi „strateški koncept“ za Severnoatlantski savez (NATO). Ovaj dokument jasno ilustruje brz porast uočene veze između računarske bezbednosti i nacionalne bezbednosti. Prethodni Strateški koncept, napisan 1999. godine, nije sadržao ni jedan pomen kompjutera, mreža, ili hakera. Novi dokument pod nazivom „Aktivno angažovanje, moderna odbrana“, opisuje sajber napade kao „česte, organizovane i skupe“ i da su sada dostigli prag na kojem prete „nacionalnom i evroatlantskom prosperitetu, bezbednosti i stabilnosti“. Ako se NATO nada da će odbraniti sajber domen, mora poboljšati svoju sposobnost sprečavanja, otkrivanja, suprotstavljanja i oporavka od sajber napada. U najmanju ruku, ovo zahteva

¹⁶² The Council of Europe Convention on Cybercrime. <http://conventions.coe.int/>.

¹⁶³ Lynn, W.J. (2010) “Defending a New Domain: The Pentagon’s Cyberstrategy,” *Foreign Affairs* 89(5) 97-108.

stavljanje svih NATO tela pod centralizovanu sajber zaštitu, unapređenje kapaciteta za sajber odbranu država članica i koordinaciju napora nacionalnih i organizacionih resursa sajber bezbednosti. U stvari, NATO je možda najbolje mesto za početak odgovaranja na pitanja nacionalne sajber bezbednosti. Internet je sada međunarodno dobro i kao takav, pretnje na njemu zahtevaju međunarodni odgovor. Kao velika grupa bogatih nacija sa zajedničkom političkom i vojnom agendom, moguće je da bi NATO danas mogao zadati značajan udarac jednoj od najvećih prednosti hakera – anonimnosti.¹⁶⁴

6.8. Bezbednosni aspekti unutar Evropske unije

Prve nacionalne strategije sajber bezbednosti razvijene su u prvoj deceniji 21. veka. Na čelu ovog pokreta bile su Sjedinjene Američke Države koje su 2003. objavile „Nacionalnu strategiju za bezbednost sajber prostora”.¹⁶⁵ Ona je bila deo ukupne Nacionalne strategije za unutrašnju bezbednost, koja je razvijena kao odgovor na terorističke napade 11. septembra 2001. godine. Akcioni planovi i strategije sa ograničenim fokusom na državne akcije u sajber prostoru razvijeni su za slične inicijative širom Evrope u narednim godinama. Nemačka je 2005. objavila „Nacionalni plan za zaštitu informacione infrastrukture (NPSI)“. Sledeće godine, Švedska je usvojila „Strategiju za poboljšanje internet bezbednosti u Švedskoj“. Nakon široko objavljenog sajber napada na Estoniju 2007. godine, ova zemlja je bila prva članica Evropske unije (EU) koja je objavila široku nacionalnu strategiju sajber bezbednosti 2008. godine. Od tada je značajan broj država objavio nacionalnu strategiju sajber bezbednosti. Iako je nekoliko država članica EU razvilo strategije sajber bezbednosti, njihova područja fokusa nisu uvek ista.¹⁶⁶

Estonija naglašava neophodnost bezbednog nacionalnog sajber prostora i fokusira se na zaštitu informacionih sistema. Sve njihove preporučene mere su građanske prirode i koncentrisane su na obrazovanje, regulisanju i saradnji.¹⁶⁷

¹⁶⁴ Ibid

¹⁶⁵ CISA (2003). The National Strategy to Secure Cyberspace. <https://www.cisa.gov/national-strategy-secure-cyberspace>.

¹⁶⁶ Ncsss (2016). "National Cyber Security Strategies (Ncsss) Map — ENISA". *Enisa.europa.eu*. N.p.

¹⁶⁷ Ibid

Francuska se fokusirala na stvaranje robusnih informacionih sistema kako bi se oduprla događajima u sajber prostoru, koji bi mogli da ugroze integritet, dostupnost ili poverljivost podataka. Francuska je naglasak stavila na tehnička sredstva koja se odnose na bezbednost digitalne infrastrukture i borbu protiv sajber kriminala, ali i uspostavljanje kapaciteta sajber odbrane.¹⁶⁸

Nemačka se fokusirala na prevenciju i krivično gonjenje sajber napadača i na sprečavanje poremećaja u informacionim uslugama, posebno kada su u pitanju kritične infrastrukture. Strategija je postavila osnovu za zaštitu kritične digitalne infrastrukture.¹⁶⁹

Holandska strategija sajber bezbednosti imala je za cilj obezbeđivanje bezbedne i pouzdane digitalne infrastrukture uz istovremenu zaštitu otvorenosti i slobode interneta. Holandija daje definiciju sajber bezbednosti u svojoj strategiji: „Sajber bezbednost je oslobođenje od opasnosti ili štete uzrokovane prekidom ili opadanjem IKT-a ili zloupotrebom IKT-a. Opasnost ili šteta usled zloupotrebe, poremećaja ili ispadanja može se sastojati od ograničenja dostupnosti i pouzdanosti IKT-a, kršenja poverljivosti informacija pohranjenih u IKT-u ili oštećenja integriteta tih informacija.”¹⁷⁰

Pristup Ujedinjenog Kraljevstva je spoj nacionalnih ciljeva i oblasti sajber bezbednosti u razvoju. Ujedinjeno Kraljevstvo nastoji da postane najveća ekonomija inovacija, investicija i kvaliteta u oblasti informacionih tehnologija, i na taj način da bude u mogućnosti da u potpunosti iskoristi potencijal i prednosti sajber prostora. Ova strategija se takođe bavi pretnjama iz sajber prostora kao što su sajber napadi od kriminalaca, terorista i država, kako bi ga učinili sigurnim utočištem za korporacije i građane.¹⁷¹

6.8.1. Odgovornost za štetnu upotrebu AI

Ko je na kraju odgovoran za AI i štetu ili štetu koju ona uzrokuje? Zakon o ovome je nejasan bez jasnog uputstva zakonodavaca, ostavljajući strankama da se oslanjaju na postojeće

¹⁶⁸ Ncsss (2016). "National Cyber Security Strategies (Ncsss) Map — ENISA". *Enisa.europa.eu*. N.p.

¹⁶⁹ Ibid

¹⁷⁰ Ibid

¹⁷¹ Ibid

zakone u oblastima kao što su nemar, zakon o potrošačima, ugovori i zakon o korporacijama. Postoji veliki broj strana uključenih u razvoj i korišćenje AI sistema: dobavljač podataka, dizajner, proizvođač, programer, korisnik i sam sistem veštačke inteligencije. Svaki slučaj kvara može zahtevati jedinstvenu analizu okolnih algoritama i okolnosti kako bi se utvrdilo ko je kriv. S obzirom na ovo, vlade će morati da razmotre da li je možda potreban pojednostavljen pristup odgovornosti. Na kraju krajeva, ovo će zavisiti od toga u kojoj oblasti prava se podnosi zahtev. Na primer, u slučajevima kada je došlo do kršenja ugovora, posledice će se verovatno odlučiti alokacijom rizika unutar ugovora. Ovo pitanje odgovornosti biće sve više postavljano kako propagiranje AI povećava štetu. U nastavku su navedene neke ključne oblasti prava u kojima odgovornost za rad AI može doći u igru.

Prekršaj nemara može pomoći u utvrđivanju odgovornosti AI koja je pošla naopako. Primer koji se često pominje je slučaj testiranja samovozećeg automobila za Uber koji je udario i ubio ženu. Postojali su različiti faktori koji su doprineli incidentu, od algoritma koji je nije identifikovao dovoljno brzo da prepozna potrebu za prekidom, do toga da operater bezbednosti nije obraćao pažnju. Ovaj incident je doveo do toga da je vozač optužen za krivični nemar, ali nije imalo pravnih posledica na Uber, jer je tužilaštvo odbilo da krivično goni kompaniju uprkos tome što je glavni uzrok bio kvar sistema.

Razvoj običajnog prava kroz ovakve slučajeve će odrediti šta će predstavljati povredu dužnosti brige u korišćenju AI. Najverovatnije, odgovornost će zavisiti od toga kada je neko dužan da vodi računa o korišćenju i razvoju veštačke inteligencije, kao i od predvidljivosti da se ona koristi onako kako je bila definisana. Važno je napomenuti da postoje uspostavljene kategorije odnosa, kao što su lekari i pacijenti, gde će automatski već postojati obaveza nege, tako da će ove strane morati da uvedu dodatnu pažnju ako odluče da koriste AI. Kako budemo pratili dalje razvojni lanac veštačke inteligencije, biće više problema sa blizinom i predvidljivošću. Sankcije za nemar i pomoćnu odgovornost mogu pomoći u balansiranju odgovornosti između različitih uključenih strana. Šteta koja se može potraživati takođe ima mogućnost uključivanja psihičkih povreda i pogrešno datih saveta, što je važno za razmatranje kada se koristi AI.

6.8.2. Pristrasnost u korišćenju AI i Anti-diskriminacija

Pristrasnost u sistemima veštačke inteligencije predstavlja stvaran rizik da dovede do potencijalnog kršenja zakona protiv diskriminacije kada organizacije koriste ove sisteme za donošenje odluka u vezi sa pojedincima. Ona proizilazi iz alata za donošenje odluka AI koji generišu nepravedne i diskriminatorne rezultate, često kao rezultat nekog oblika statističke pristrasnosti u osnovnim podacima o obuci. AI sve češće koriste organizacije za donošenje važnih odluka, uključujući i vladu. Ovi sistemi veštačke inteligencije su generalno obučeni za primenu pristupa zasnovanog na modelu verovatnoće na upite. Ovo često uključuje korišćenje istorijskih podataka da bi se odredio tretman novih podataka koji su im predstavljeni.

Pitanje se javlja kada su podaci koji se koriste za obuku AI sami po sebi pristrasni ili diskriminatorni, a to je obrazac koji će AI koristiti za buduće odluke. Na primer, ako se od veštačke inteligencije zatraži da odluči o očekivanom finansijskom uspehu nekoga na osnovu toga što je obučen na osnovu podataka gde je postojao rodni jaz u platama, onda će AI verovatno zaključiti da bi odgovarajući iznos za plaćanje muškarca trebalo da bude veći od plaćanja zaposlene osobe ženskog pola. Ovo se na sličan način može primeniti u krivičnom kontekstu gde mnoge zemlje već koriste alate za predviđanje.

Već više od 20 godina, Velika Britanija koristi Sistem za procenu prestupnika (Oasis) u svom sistemu krivičnog pravosuđa da bi predviđala pitanja kao što su odobravanje kaucije, vrsta izrečene kazne, klasifikacija zatvorske sigurnosti, upućivanje na rehabilitaciju, pa čak i ishodi imigracionih slučajeva. Za sve ovo vreme, nijednom naučniku nije bio dozvoljen pristup podacima koji se koriste za nezavisnu analizu njihovog rada ili tačnosti. Nedavna studija analizirala je učinak predviđanja takvih alata za procenu rizika od kriminala i otkrila da je prediktivni učinak bio kontroverzan, u rasponu od lošeg do umerenog. Dalje se otkrilo da većina studija validacije ima visok rizik od pristrasnosti, delimično zbog neprikladnog analitičkog pristupa koji se koristi. Jedan takav slučaj u kojem je to bilo očigledno je Džordan Svini, osuđeni ubica koji je pušten iz zatvora 2022. godine nakon što je klasifikovan kao „srednji rizik“, da bi nakon nekoliko dana brutalno napao i ubio mladu ženu koja je bila sama u kući. Pregledom je utvrđeno da je trebalo da bude klasifikovan kao „visoko rizičan“.

Poslodavci se takođe sve više oslanjaju na AI tehnologije i alate za proveru u procesu zapošljavanja kako bi izabrali „najboljeg kandidata“. Ovi poslodavci treba da osiguraju da budu

u toku sa zakonima protiv diskriminacije. Čak i tamo gde nema zlobe ili loših namera poslodavca, problem je u tome što korisnici veštačke inteligencije ne mogu u potpunosti da razumeju „crnu kutiju“, pa stoga ne mogu da obezbede da se ona ne koristi na diskriminatorski način. Kao takvi, oni koji žele da koriste sisteme veštačke inteligencije na ovaj način moraju da obrate veliku pažnju da aktivno prate rezultate u pogledu indikacija nepravednosti ili pristrasnosti.

SEDMI DEO: REZULTATI EMPIRIJSKOG ISTRAŽIVANJA

U anketiranju je učestvovalo ukupno $N = 253$ ispitanika, od toga $n = 139$ (54,9%) ispitanika muškog pola i $n = 114$ (45,1%) ispitanika ženskog pola. Srednja vrednost zastupljenosti po polu je 1.45 a standardna devijacija .499.

U istraživanju je učestvovalo najviše ispitanika starosne grupe 26-35 godina $n = 77$ (30,4%), a najmanje iz starosne grupe 56-65 godina $n = 24$ (9,5%). Srednja vrednost zastupljenosti po starosnoj strukturi je 2,68 dok je standardna devijacija 1.283.

Istraživanjem bilo obuhvaćeno najviše ispitanika sa višom/visokom školskom spremom $n = 103$ (40,7%), a najmanje sa doktoratom $n = 18$ (7,1%) i osnovnom školom $n = 3$ (1,2%). Srednja vrednost zastupljenosti po stepenu obrazovanja je 3,05 dok je standardna devijacija .918.

Istraživanjem su obezbeđeni stavovi učesnika na anketiranju koji se odnose na primenu veštačke inteligencije i oblasti regulisanja na međunarodnom nivou, s obzirom na to da nacionalna zakonodavstva još uvek nemaju preduzete radnje za regulaciju ove oblasti, kao i implikacijama primene AI. Empirijsko istraživanje baziralo se na utvrđivanju i analizi dobijenih stavova ispitanika u vezi sa različitim aspektima ispitivane pojave.

Za testiranje postavljenih hipoteza utvrđivane su uzročno-posledične veze odabranih varijabli. Cilj istraživanja bio je skoncentrisan na prikaz dobijenih odgovora ispitanika, prikazivanje njihovih stavova i utvrđivanje parametara koji su važni za dokazivanje ili opovrgavanje postavljenih hipoteza.

Deskriptivno longitudinalno istraživanje, kao dizajn istraživanja korišćenog u ovoj disertaciji, bio je najbolji izbor za opisivanje karakteristika fenomena koji se proučavao, jer ne odgovara na pitanja o tome kako, kada i zašto su se te karakteristike pojavile, već se bavi karakteristikama fenomena koji se proučava.

Alat za prikupljanje podataka koji se koristio imao je dva dela, prvi je fokusiran na demografski profil ispitanika, a drugi deo je sastavljen od iznetih tvrdnji koje je ispitanici trebalo da potvrde ili da opovrgnu. Prikupljeni podaci su zbrojeni i predstavljeni u broju učestalosti, procentima i ponderisanim srednjim vrednostima da bi se omogućila jednostavna deskriptivna analiza.

Ispitanicima su ponuđeni odgovori rangirani na osnovu Likertove skale od 1. Apsolutno se ne slažem do 5. Apsolutno se slažem. Za obradu rezultata anketnog istraživanja korišćen je

IBM SPSS softver za obradu podataka u oblasti društvenih nauka. Odgovori su obrađeni primenom deskriptivne statistike, komparativne metode i korelacione analize.

Za potrebe ovog istraživanja pripremljen je poseban upitnik koji je prosleđen elektronskim putem na 280 mejl adresa. Za finalnu obradu odabrano je ukupno 253 ispravna upitnika.

Dobijeni rezultati prikazani su tekstualno, tabelarno, a analiza postavljenih hipoteza uključuje i grafički prikaz dobijenih rezultata.

Demografski podaci ispitanika:

Tabela 1. Polna struktura ispitanika-statistički podaci

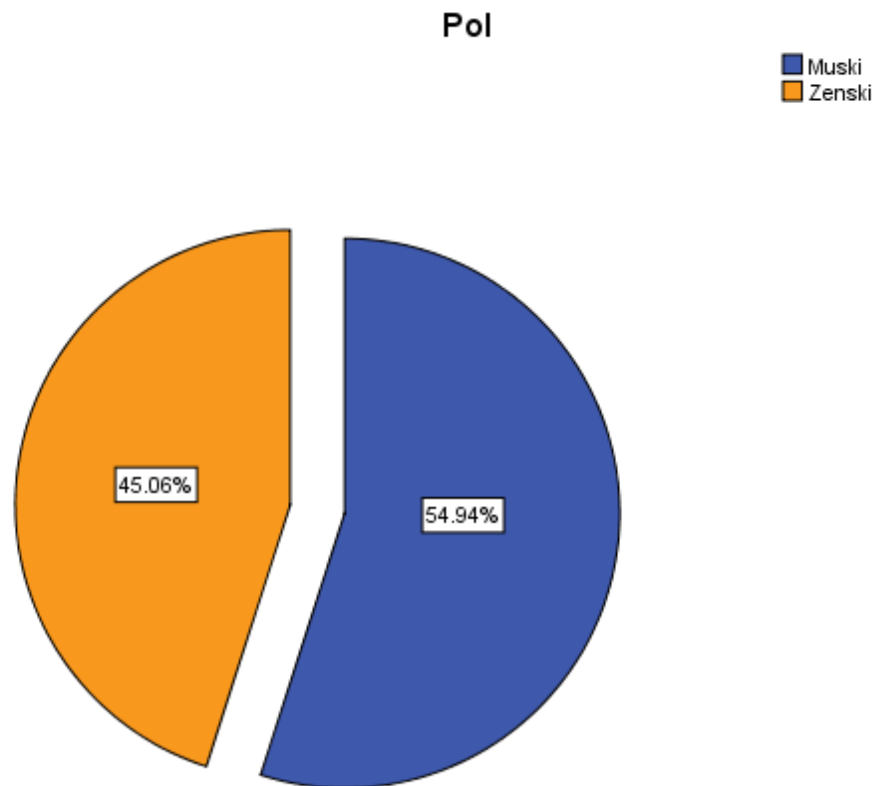
Statistics		
Pol		
N	Valid	253
	Missing	0
Mean		1.45
Median		1.00
Mode		1
Std. Deviation		.499
Sum		367

Tabela 2. Polna struktura ispitanika-frekvencija i procenti

Pol

	Frequency	Percent	Valid Percent	Cumulative Percent
Muski	139	54.9	54.9	54.9
Zenski	114	45.1	45.1	100.0
Total	253	100.0	100.0	

U anketiranju je učestvovalo $n = 139$ (54,9%) ispitanika muškog pola i $n = 114$ (45,1%) ispitanika ženskog pola. Srednja vrednost zastupljenosti po polu je 1.45 a standardna devijacija .499.



Grafički prikaz 1. Prikaz statističkih pokazatelja vezanih za polnu strukturu ispitanika

Tabela 3. Starosna struktura ispitanika-statistički podaci

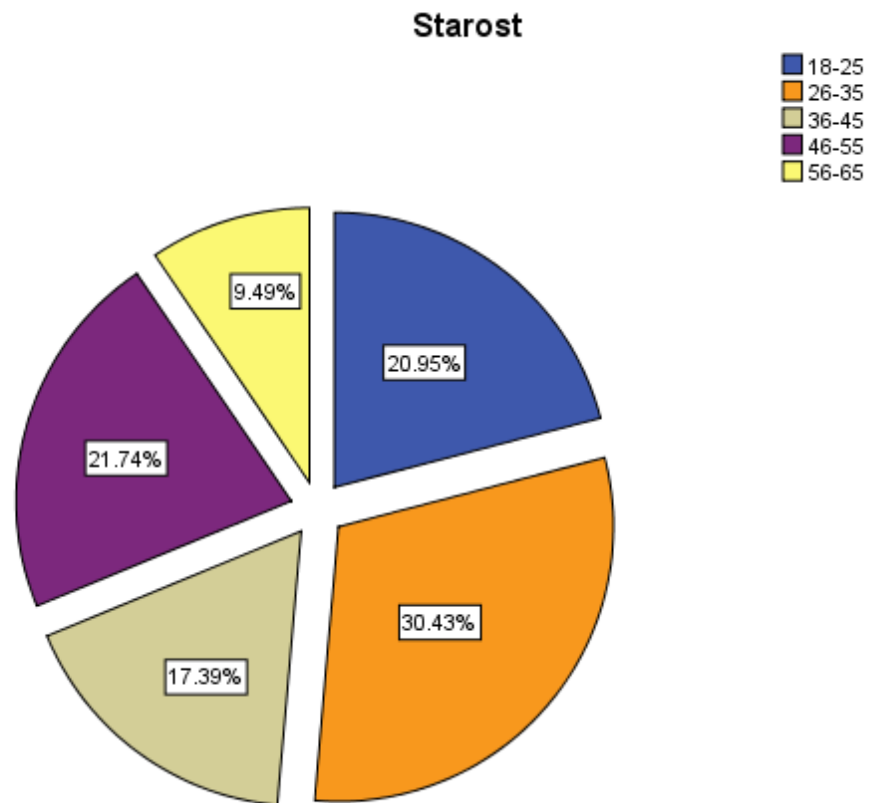
Statistics		
Starost		
N	Valid	253
	Missing	0
Mean		2.68
Median		2.00
Mode		2
Std. Deviation		1.283
Sum		679

Tabela 4. Starosna struktura ispitanika -frekvencija i procenti

Starost

	Frequency	Percent	Valid Percent	Cumulative Percent
18-25	53	20.9	20.9	20.9
26-35	77	30.4	30.4	51.4
36-45	44	17.4	17.4	68.8
46-55	55	21.7	21.7	90.5
56-65	24	9.5	9.5	100.0
Total	253	100.0	100.0	

Istraživanje je obezbedilo prikaz starosne grupe na sledeći način: najviše je ispitanika starosti 26-35 godina $n=77$ (30,4%), a najmanje 56-65 godina $n=24$ (9,5%). M starosne strukture je 2,68 dok je SD 1.283.



Grafički prikaz 2. Prikaz starosnih grupa ispitanika

Tabela 5. Stepen obrazovanja-statistički podaci

Statistics

Obrazovanje

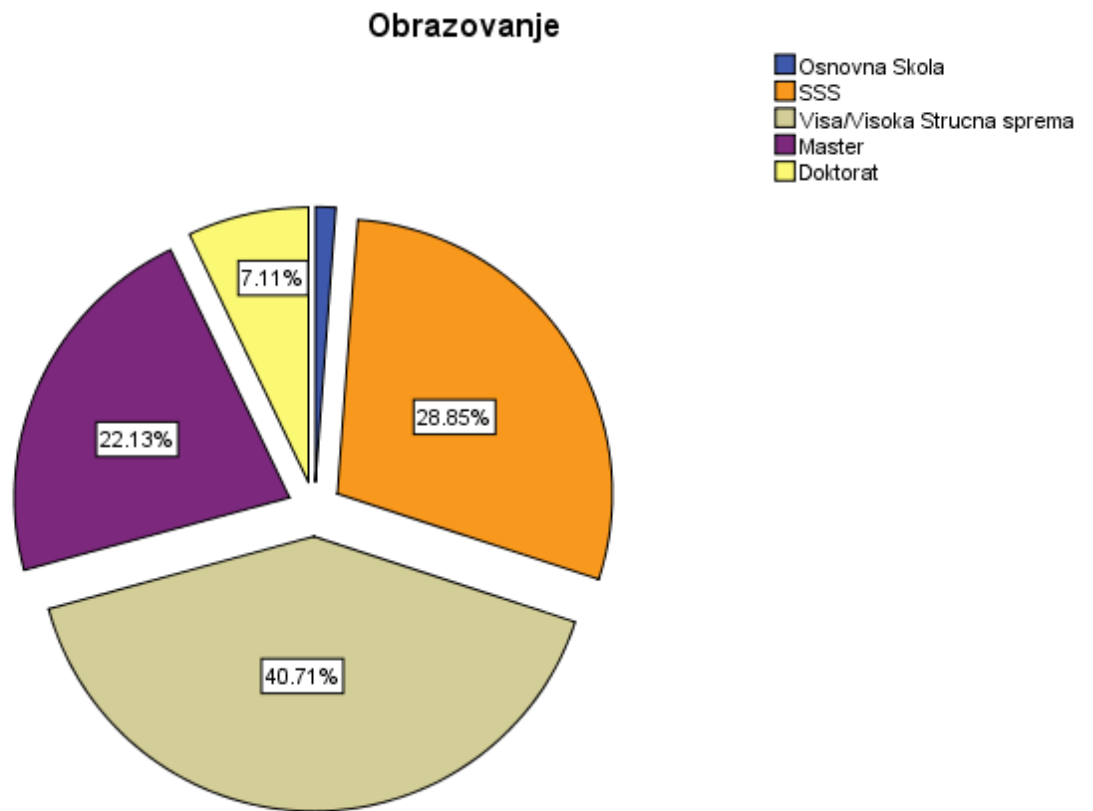
N	Valid	253
	Missing	0
Mean		3.05
Median		3.00
Mode		3
Std. Deviation		.918
Sum		772

Tabela 6. Stepen obrazovanja - frekvencija i procenti

Obrazovanje

	Frequency	Percent	Valid Percent	Cumulative Percent
Osnovna Skola	3	1.2	1.2	1.2
SSS	73	28.9	28.9	30.0
Visa/Visoka Strucna sprema	103	40.7	40.7	70.8
Master	56	22.1	22.1	92.9
Doktorat	18	7.1	7.1	100.0
Total	253	100.0	100.0	

Istraživanje je obezbedilo prikaz školske sprema na sledeći način: najviše ispitanika je sa višom/visokom školskom spremom $n=103$ (40,7%), a najmanje sa doktoratom $n=18$ (7,1%) i osnovnom školom $n=3$ (1,2%). M je 3,05 dok je SD .918.



Grafički prikaz 3. Prikaz stepena obrazovanja ispitanika

Ispitivanje stavova ispitanika

P1. Da li podržavate stav da zakonodavstvo mora da uspostavi odgovarajuću kontrolu i da jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju?

Tabela 7. Mišljenje ispitanika u vezi pitanja P1

Statistics

P1

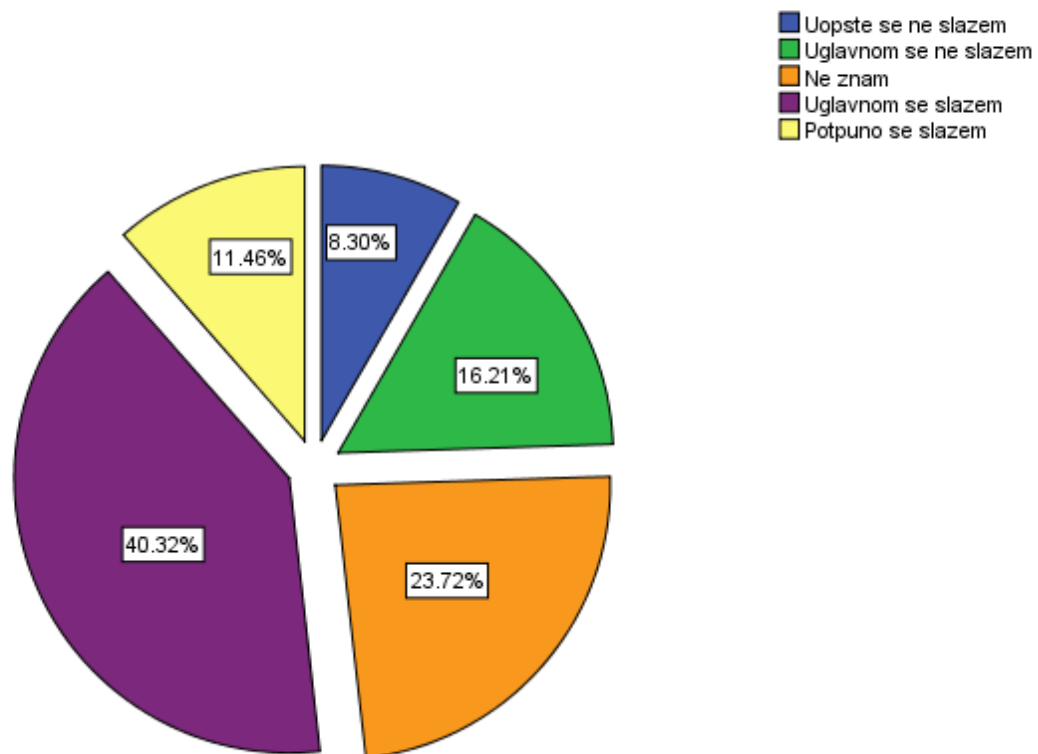
N	Valid	253
	Missing	0
Mean		3.30
Median		4.00
Mode		4
Std. Deviation		1.126
Sum		836

P1

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	21	8.3	8.3	8.3
Uglavnom se ne slazem	41	16.2	16.2	24.5
Ne znam	60	23.7	23.7	48.2
Uglavnom se slazem	102	40.3	40.3	88.5
Potpuno se slazem	29	11.5	11.5	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika izjasnio se da se uglavnom slaže sa iznetom tvrdnjom n=102 (40,3%) i da se potpuno slaže n=29 (11,5%). M je 3,30 dok je SD 1.126.

P1



Grafički prikaz 4. Grafički prikaz stavova ispitanika u odnosu na pitanje P1

P2. Da li podržavate stavda upotreba AI može dovede do štetnih posledica?

Tabela 8. Mišljenje ispitanika u vezi pitanja P2

Statistics

P2

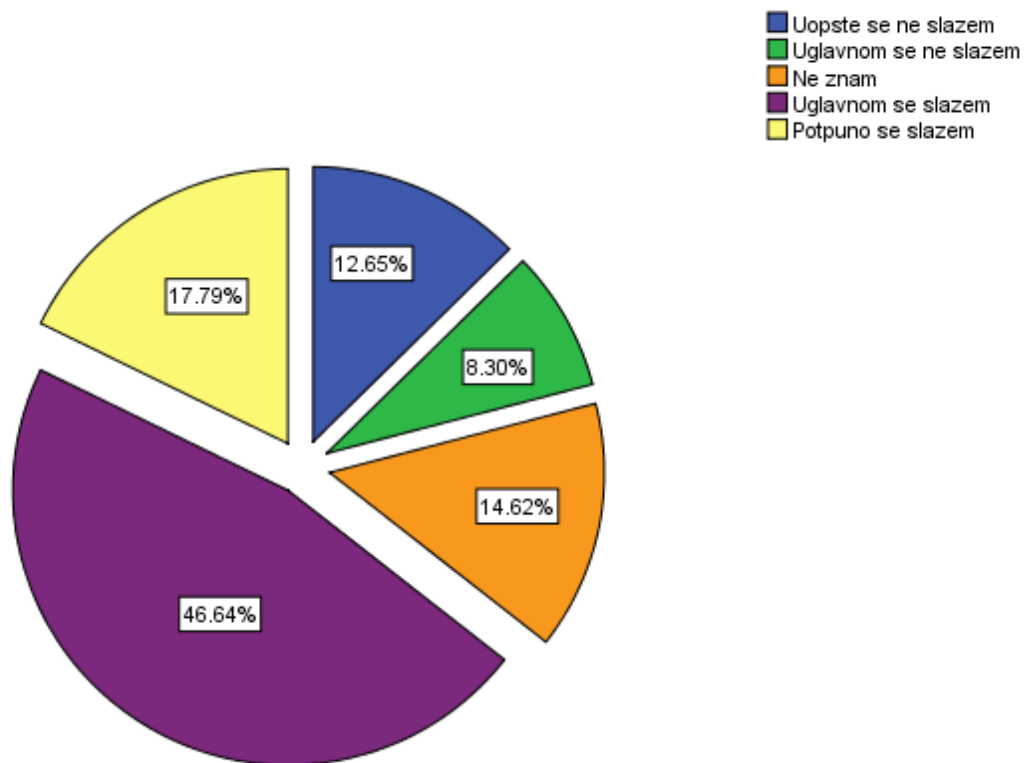
N	Valid	253
	Missing	0
Mean		3.49
Median		4.00
Mode		4
Std. Deviation		1.240
Sum		882

P2

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	32	12.6	12.6	12.6
Uglavnom se ne slazem	21	8.3	8.3	20.9
Ne znam	37	14.6	14.6	35.6
Uglavnom se slazem	118	46.6	46.6	82.2
Potpuno se slazem	45	17.8	17.8	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika se uglavnom slaže sa iznetom tvrdnjom $n=118$ (46,6%) i potpuno se slaže $n=45$ (17,8%). M je 3,49 a SD 1.240.

P2



Grafički prikaz 5. Grafički prikaz stavova ispitanika u odnosu na pitanje P2

P3. Da li podržavate stavda regulisanje veštačke inteligencije mora da postigne ravnotežu između regulisanja kontrole AI i omogućavanja inovacije i evolucije?

Tabela 9. Mišljenje ispitanika u vezi pitanja P3

Statistics

P3

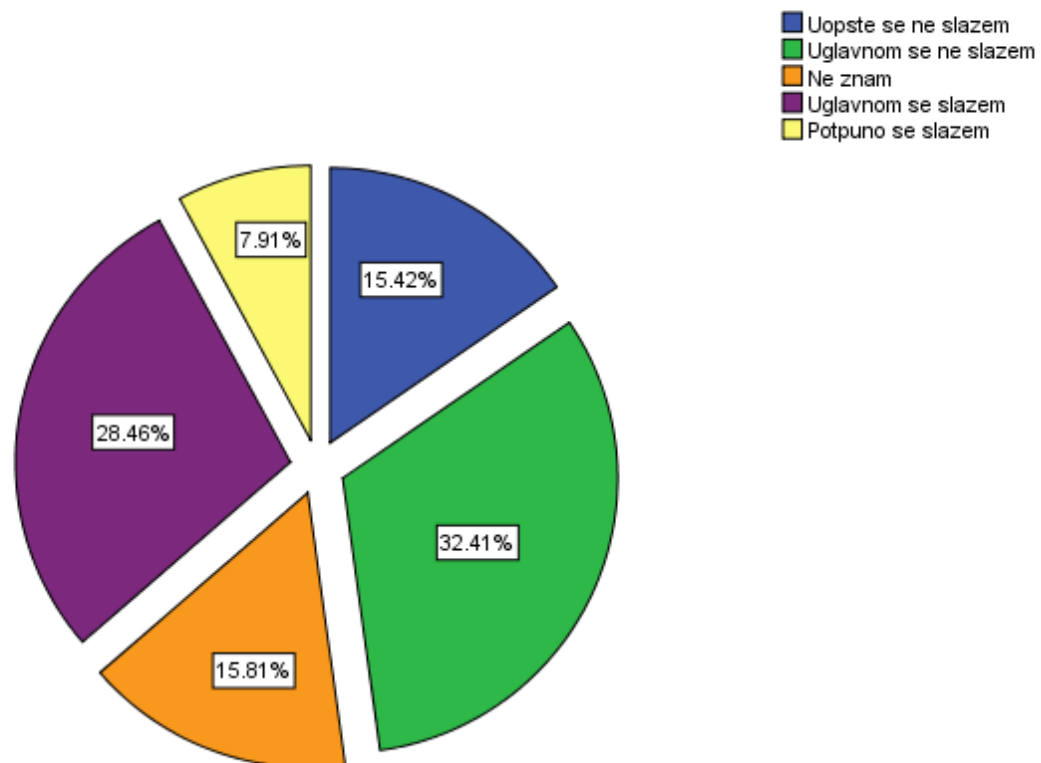
N	Valid	253
	Missing	0
Mean		2.81
Median		3.00
Mode		2
Std. Deviation		1.229
Sum		711

P3

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	39	15.4	15.4	15.4
Uglavnom se ne slazem	82	32.4	32.4	47.8
Ne znam	40	15.8	15.8	63.6
Uglavnom se slazem	72	28.5	28.5	92.1
Potpuno se slazem	20	7.9	7.9	100.0
Total	253	100.0	100.0	

Ispitanici se uglavnom ne slažu sa iznetom tvrdnjom n=82 (32,4%), uglavnom se slažu n=72 (28,5%) i potpuno se slaže n=20 (7,9%). M je 2,81 a SD 1.229.

P3



Grafički prikaz 6. Grafički prikaz stavova ispitanika u odnosu na pitanje P3

P4. Da li podržavate stavda AI mora da se koristi na bezbedan i odgovoran način?

Tabela 10. Mišljenje ispitanika u vezi pitanja P4

Statistics

P4

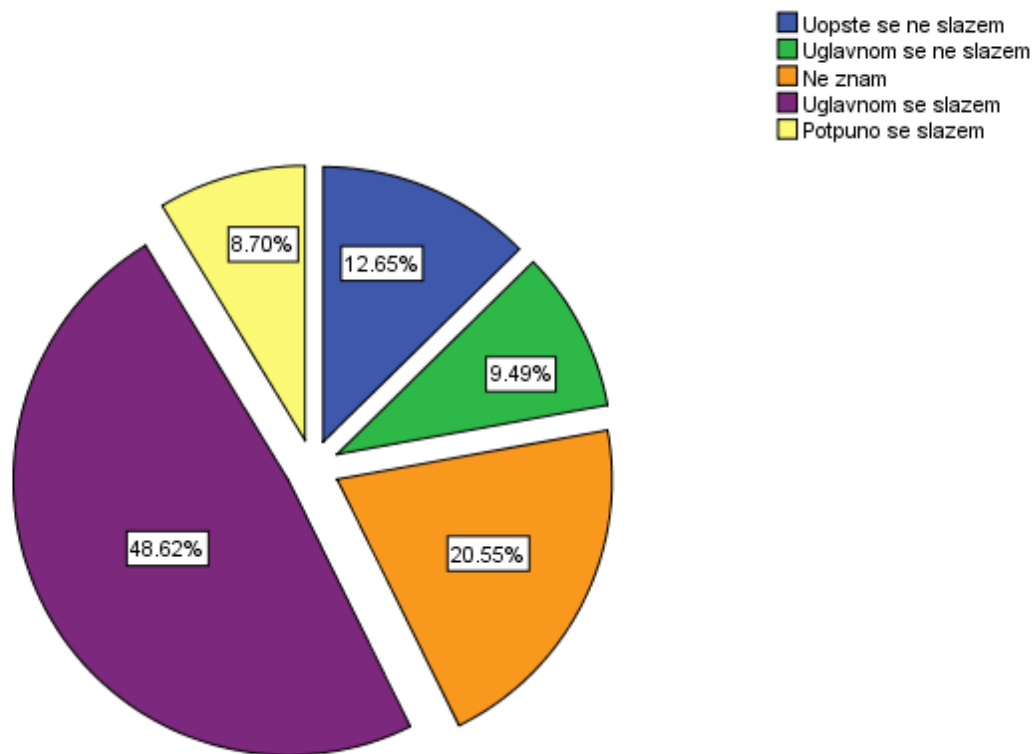
N	Valid	253
	Missing	0
Mean		3.31
Median		4.00
Mode		4
Std. Deviation		1.159
Sum		838

P4

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	32	12.6	12.6	12.6
Uglavnom se ne slazem	24	9.5	9.5	22.1
Ne znam	52	20.6	20.6	42.7
Uglavnom se slazem	123	48.6	48.6	91.3
Potpuno se slazem	22	8.7	8.7	100.0
Total	253	100.0	100.0	

Ispitanici se uglavnom slažu sa iznetom tvrdnjom $n=123$ (48,6%) i $n=22$ (8,7%) ispitanika se potpuno slaže. M je 3,31 a SD 1.159.

P4



Grafički prikaz 7. Grafički prikaz stavova ispitanika u odnosu na pitanje P4

P5. Da li podržavate stavda podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana?

Tabela 11. Mišljenje ispitanika u vezi pitanja P5

Statistics

P5

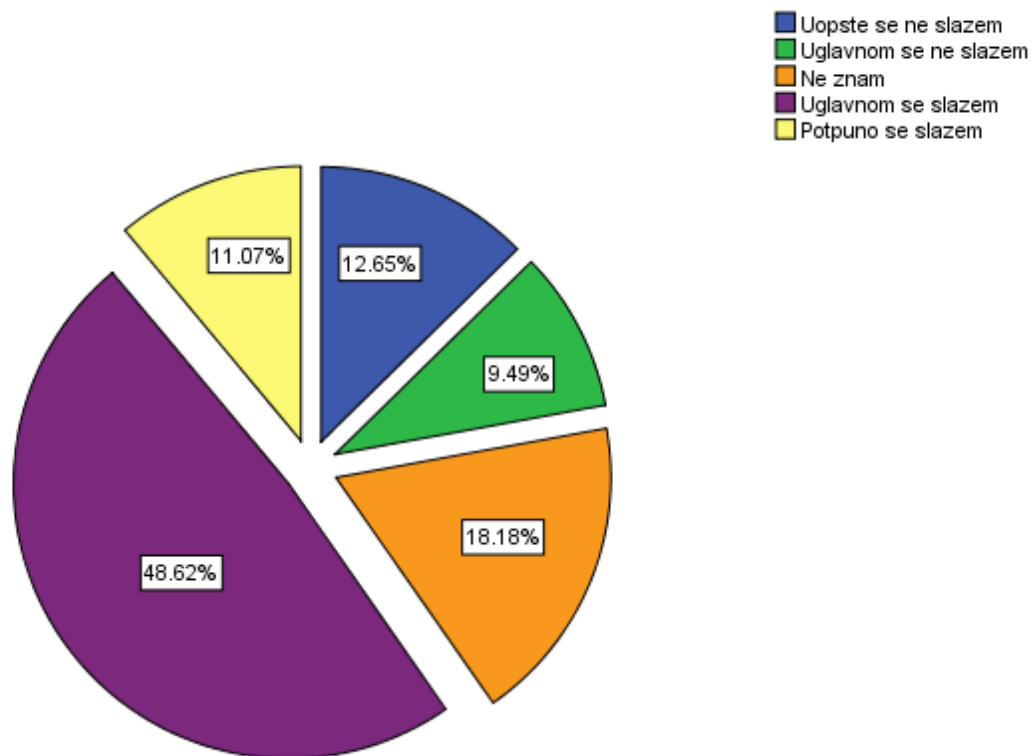
N	Valid	253
	Missing	0
Mean		3.36
Median		4.00
Mode		4
Std. Deviation		1.186
Sum		850

P5

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	32	12.6	12.6	12.6
Uglavnom se ne slazem	24	9.5	9.5	22.1
Ne znam	46	18.2	18.2	40.3
Uglavnom se slazem	123	48.6	48.6	88.9
Potpuno se slazem	28	11.1	11.1	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=123 (48,6%) uglavnom slaže sa iznetom tvrdnjom, dok se n=28 (11,1%) potpuno slaže. M je 3,36 a SD 1.186.

P5



Grafički prikaz 8. Grafički prikaz stavova ispitanika u odnosu na pitanje P5

P6. Da li podržavate stavda upotreba AI otvara pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI?

Tabela 12. Mišljenje ispitanika u vezi pitanja P6

Statistics

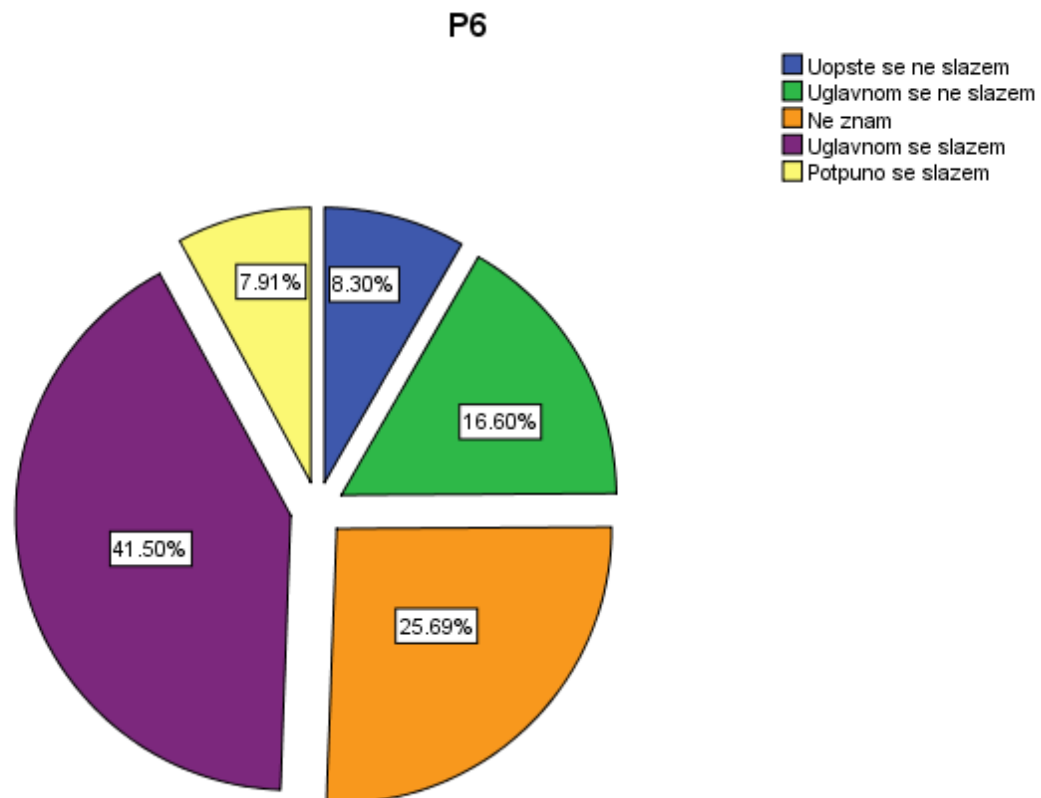
P6

N	Valid	253
	Missing	0
Mean		3.24
Median		3.00
Mode		4
Std. Deviation		1.084
Sum		820

P6

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	21	8.3	8.3	8.3
Uglavnom se ne slazem	42	16.6	16.6	24.9
Ne znam	65	25.7	25.7	50.6
Uglavnom se slazem	105	41.5	41.5	92.1
Potpuno se slazem	20	7.9	7.9	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=105 (41,5%) se uglavnom slaže sa iznetom tvrdnjom, n=20 (7,9%) se potpuno slaže. M je 3,24 a SD 1.084.



Grafički prikaz 9. Grafički prikaz stavova ispitanika u odnosu na pitanje P6

P7. Da li podržavate stavda upotreba AI sistema mora da se klasifikuje u odnosu na stepen rizika do koga dolazi prilikom primene AI?

Tabela 13. Mišljenje ispitanika u vezi pitanja P7

Statistics

P7

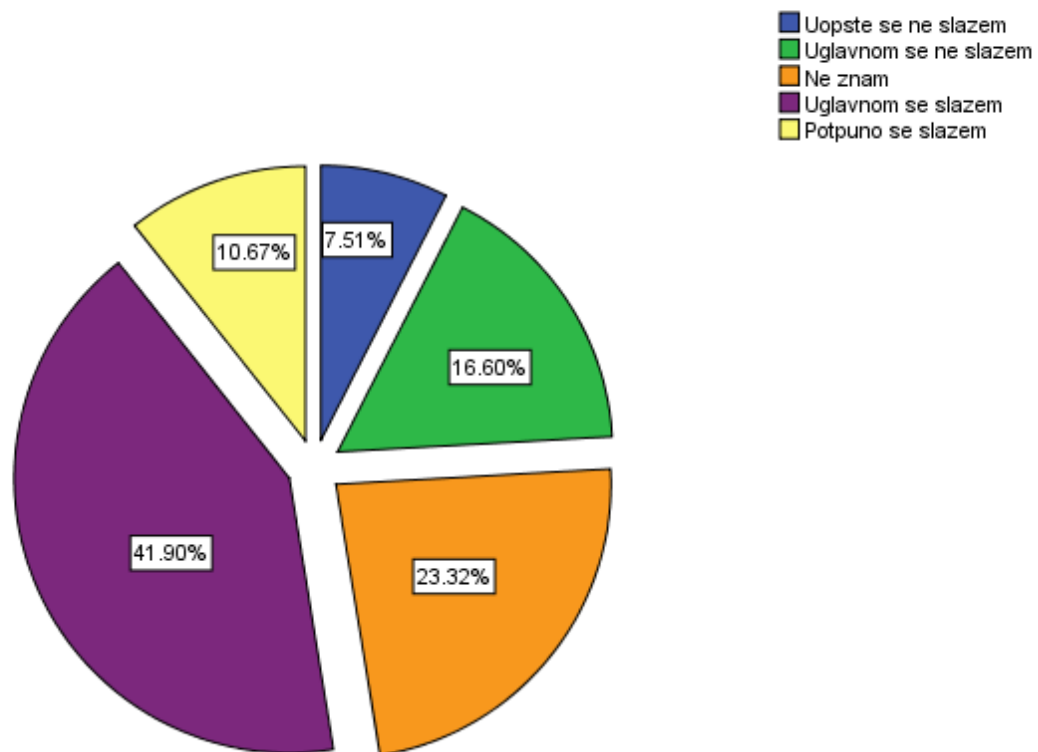
N	Valid	253
	Missing	0
Mean		3.32
Median		4.00
Mode		4
Std. Deviation		1.103
Sum		839

P7

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	19	7.5	7.5	7.5
Uglavnom se ne slazem	42	16.6	16.6	24.1
Ne znam	59	23.3	23.3	47.4
Uglavnom se slazem	106	41.9	41.9	89.3
Potpuno se slazem	27	10.7	10.7	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=106 (41,9%) se uglavnom slaže sa iznetom tvrdnjom, n=27 (10,7%) se potpuno slaže. M je 3,32 a SD 1.103.

P7



Grafički prikaz 10. Grafički prikaz stavova ispitanika u odnosu na pitanje P7

P8. Da li podržavate stavda AI sistemi visokog rizika moraju da budi obuhvaćeni strogo definisanim načinima delovanja?

Tabela 14. Mišljenje ispitanika u vezi pitanja P8

Statistics

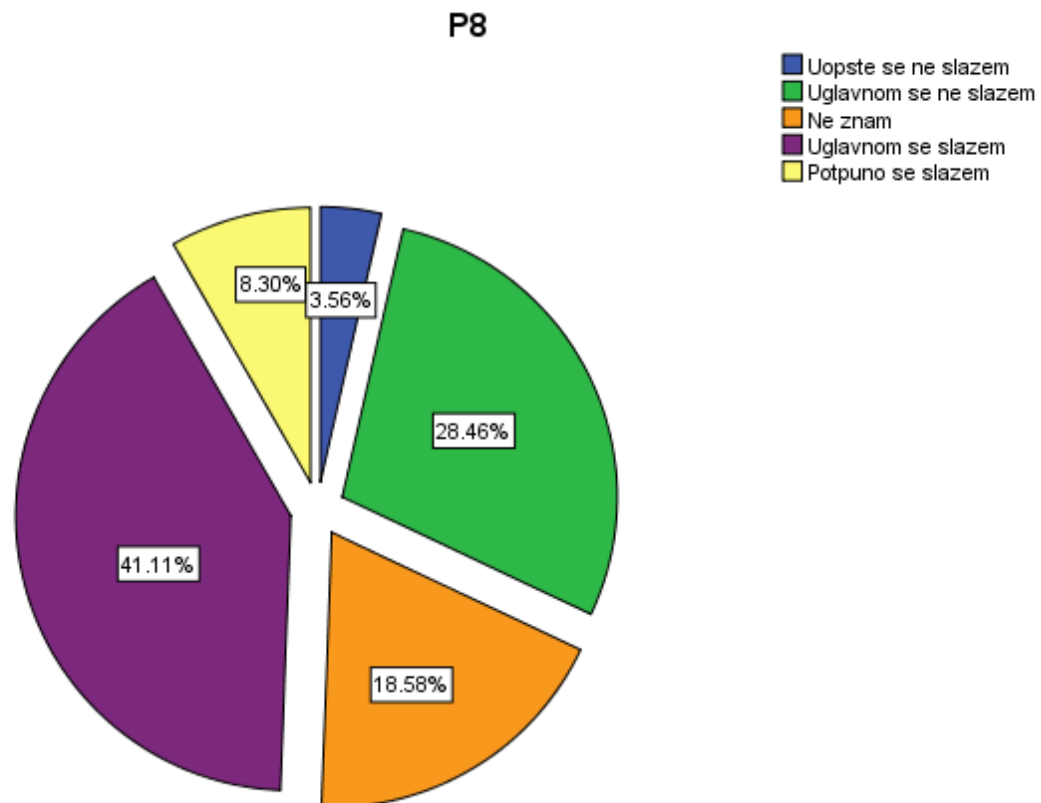
P8

N	Valid	253
	Missing	0
Mean		3.22
Median		3.00
Mode		4
Std. Deviation		1.061
Sum		815

P8

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	9	3.6	3.6	3.6
Uglavnom se ne slazem	72	28.5	28.5	32.0
Ne znam	47	18.6	18.6	50.6
Uglavnom se slazem	104	41.1	41.1	91.7
Potpuno se slazem	21	8.3	8.3	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika $n=104$ (41,4%) se uglavnom slaže sa iznetom tvrdnjom, dok se $n=21$ (8,3%) potpuno slaže. M je 3,22 a SD 1.061.



Grafički prikaz 11. Grafički prikaz stavova ispitanika u odnosu na pitanje P8

P9. Da li podržavate stavda će bilo kakvo definisanje AI vrlo brzo biti prevaziđeno razvojem ove tehnologije?

Tabela 15. Mišljenje ispitanika u vezi pitanja P9

Statistics

P9

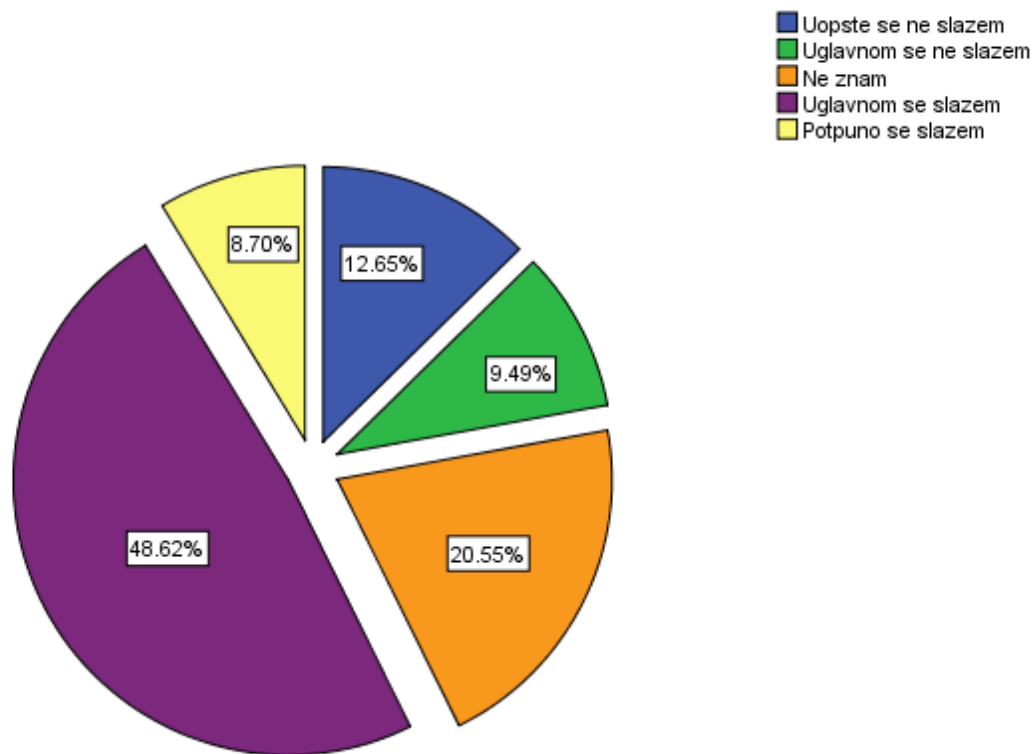
N	Valid	253
	Missing	0
Mean		3.31
Median		4.00
Mode		4
Std. Deviation		1.159
Sum		838

P9

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	32	12.6	12.6	12.6
Uglavnom se ne slazem	24	9.5	9.5	22.1
Ne znam	52	20.6	20.6	42.7
Uglavnom se slazem	123	48.6	48.6	91.3
Potpuno se slazem	22	8.7	8.7	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=123 (48,6%) se uglavnom slaže sa iznetom tvrdnjom, dok se n=22 (8,7%) potpuno slaže. M je 3,31 a SD 1.159.

P9



Grafički prikaz 12. Grafički prikaz stavova ispitanika u odnosu na pitanje P9

P10. Da li podržavate stavda problem za funkcionisanje AI modela i njegovo efikasno funkcionisanje nastaje ako je skup podataka za testiranje nekompletan i ne pokriva dovoljan opseg mogućih situacija?

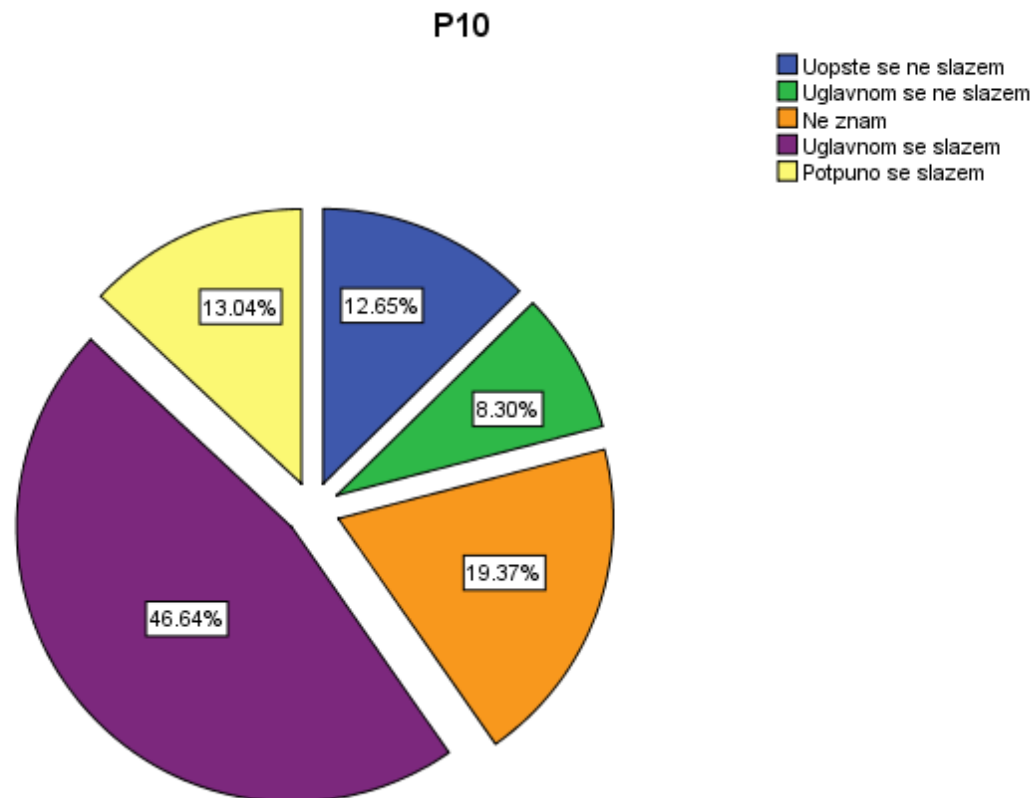
Tabela 16. Mišljenje ispitanika u vezi pitanja P10

Statistics		
P10		
N	Valid	253
	Missing	0
Mean		3.39
Median		4.00
Mode		4
Std. Deviation		1.196
Sum		858

P10

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	32	12.6	12.6	12.6
Uglavnom se ne slazem	21	8.3	8.3	20.9
Ne znam	49	19.4	19.4	40.3
Uglavnom se slazem	118	46.6	46.6	87.0
Potpuno se slazem	33	13.0	13.0	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=118 (46,6%) se uglavnom slaže sa iznetom tvrdnjom, dok se n=33 (13,0%) potpuno slaže. M je 3,39 a SD 1.196.



Grafički prikaz 13. Grafički prikaz stavova ispitanika u odnosu na pitanje P10

P11. Da li podržavate stavda integracija veštačke inteligencije (AI) i mašinskog učenja (ML) u različitim industrijama se ubrzano povećava poslednjih godina i da se od njih čekuje da automatizuju mnoge zadatke koje trenutno obavljaju ljudi?

Tabela 17. Mišljenje ispitanika u vezi pitanja P11

Statistics

P11

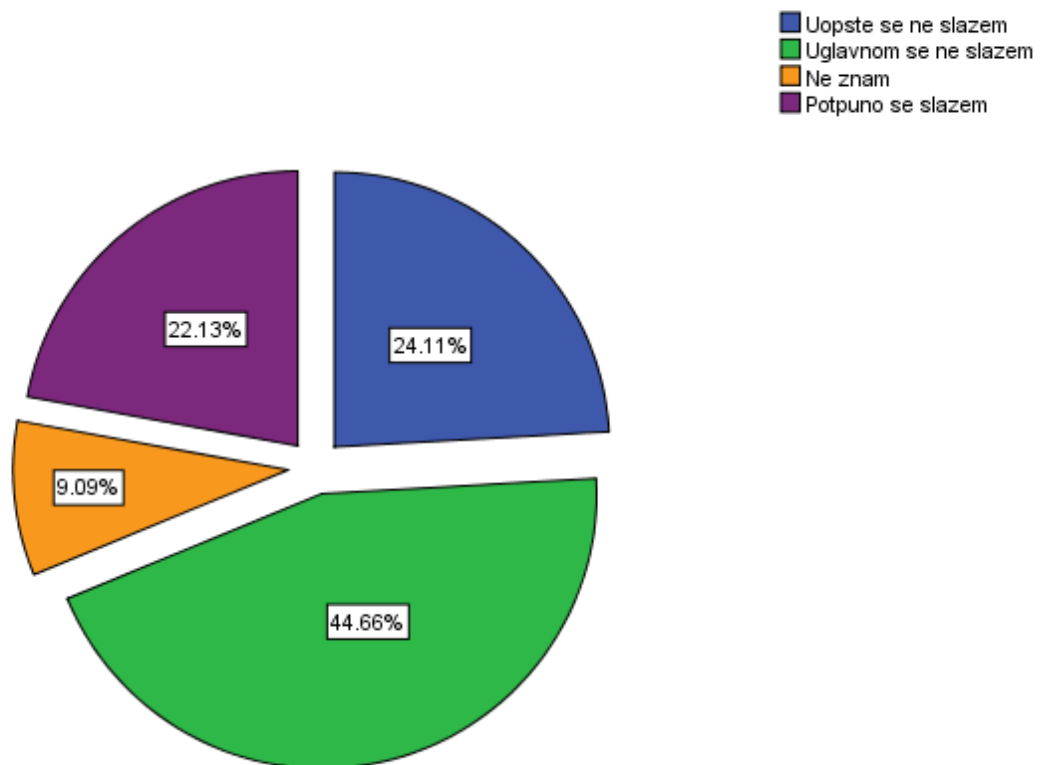
N	Valid	253
	Missing	0
Mean		2.51
Median		2.00
Mode		2
Std. Deviation		1.438
Sum		636

P11

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	61	24.1	24.1	24.1
Uglavnom se ne slazem	113	44.7	44.7	68.8
Ne znam	23	9.1	9.1	77.9
Potpuno se slazem	56	22.1	22.1	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=113 (44,7%) se uglavnom ne slaže u vezi sa iznetom tvrdnjom, dok se n=56 (22,1%) potpuno slaže. M je 2,51 a SD 1.438.

P11



Grafički prikaz 14. Grafički prikaz stavova ispitanika u odnosu na pitanje P11

P12. Da li podržavate stavda upotreba veštačke inteligencije u poslovanju podrazumeva upotrebu lažnog ljudskog inteligentnog računara da bi se podržao prihod, poboljšalo zadovoljstvo kupaca, povećala profitabilnost i podstakla poslovni rast i promene?

Tabela 18. Mišljenje ispitanika u vezi pitanja P12

Statistics

P12

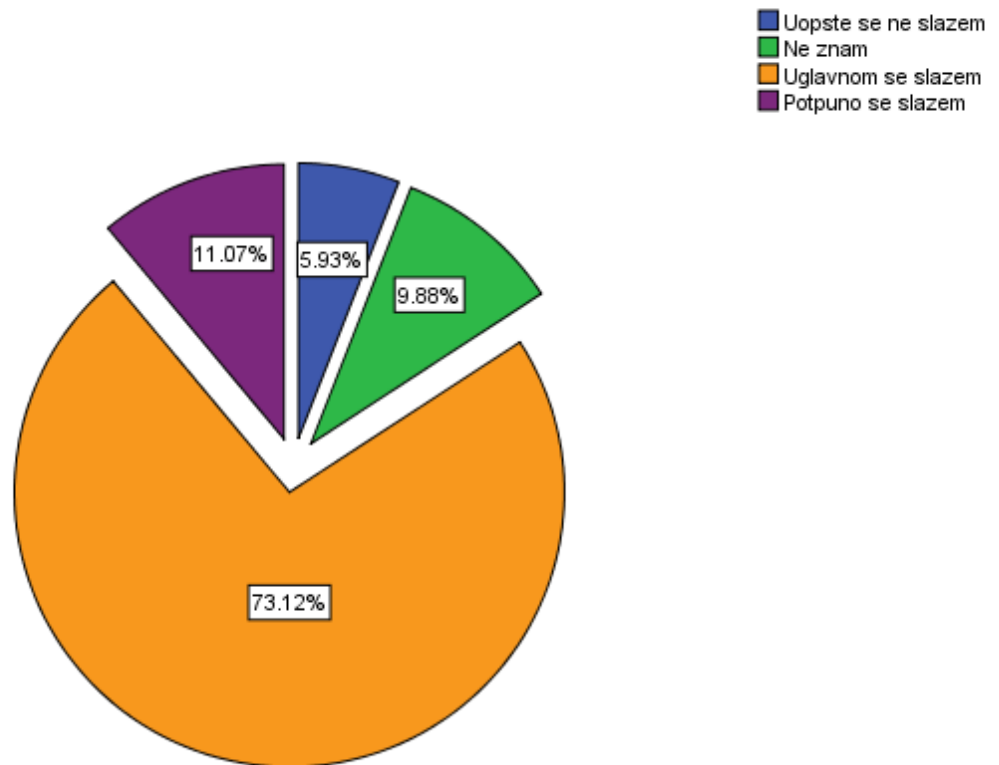
N	Valid	253
	Missing	0
Mean		3.83
Median		4.00
Mode		4
Std. Deviation		.848
Sum		970

P12

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	15	5.9	5.9	5.9
Ne znam	25	9.9	9.9	15.8
Uglavnom se slazem	185	73.1	73.1	88.9
Potpuno se slazem	28	11.1	11.1	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=185 (73,1%) se uglavnom slaže sa iznetom tvrdnjom, dok se n=28 (11,1%) potpuno slaže. M je 3,83 a SD .848.

P12



Grafički prikaz 15. Grafički prikaz stavova ispitanika u odnosu na pitanje P12

P13. Da li podržavate stavda štetna pristrasnost u sistemima veštačke inteligencije može produžiti ili pogoršati istorijske nejednakosti u oblastima kao što su zapošljavanje, zdravstvena zaštita, finansije i stanovanje, a zlonamerni akteri mogu da koriste AI alate da manipulišu ljudima?

Tabela 19. Mišljenje ispitanika u vezi pitanja P13

Statistics

P13

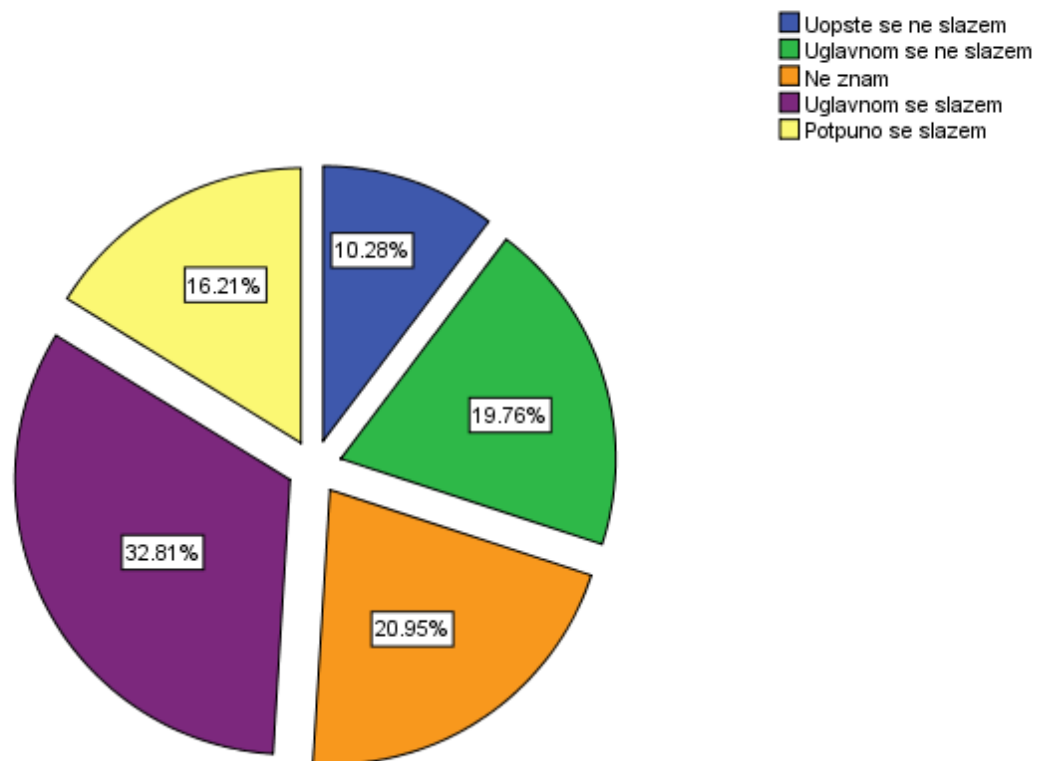
N	Valid	253
	Missing	0
Mean		3.25
Median		3.00
Mode		4
Std. Deviation		1.237
Sum		822

P13

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	26	10.3	10.3	10.3
Uglavnom se ne slazem	50	19.8	19.8	30.0
Ne znam	53	20.9	20.9	51.0
Uglavnom se slazem	83	32.8	32.8	83.8
Potpuno se slazem	41	16.2	16.2	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=83 (32,8%) se uglavnom slaže sa iznetom tvrdnjom, dok se n=41 (16,2%) potpuno slaže. M je 3,25 a SD 1.237.

P13



Grafički prikaz 16. Grafički prikaz stavova ispitanika u odnosu na pitanje P13

P14. Da li podržavate stavda je prisutna dilema oko toga kako osmisлити politiku veštačke inteligencije i kakav balans ona treba da postigne između različitih vrednosti, kao što je proaktivan rad koji bi trebalo da spreči štetu kroz vladinu intervenciju i zauzimanje popustljivijeg pristupa za podsticanje eksperimentisanja i inovacija?

Tabela 20. Mišljenje ispitanika u vezi pitanja P14

Statistics

P14

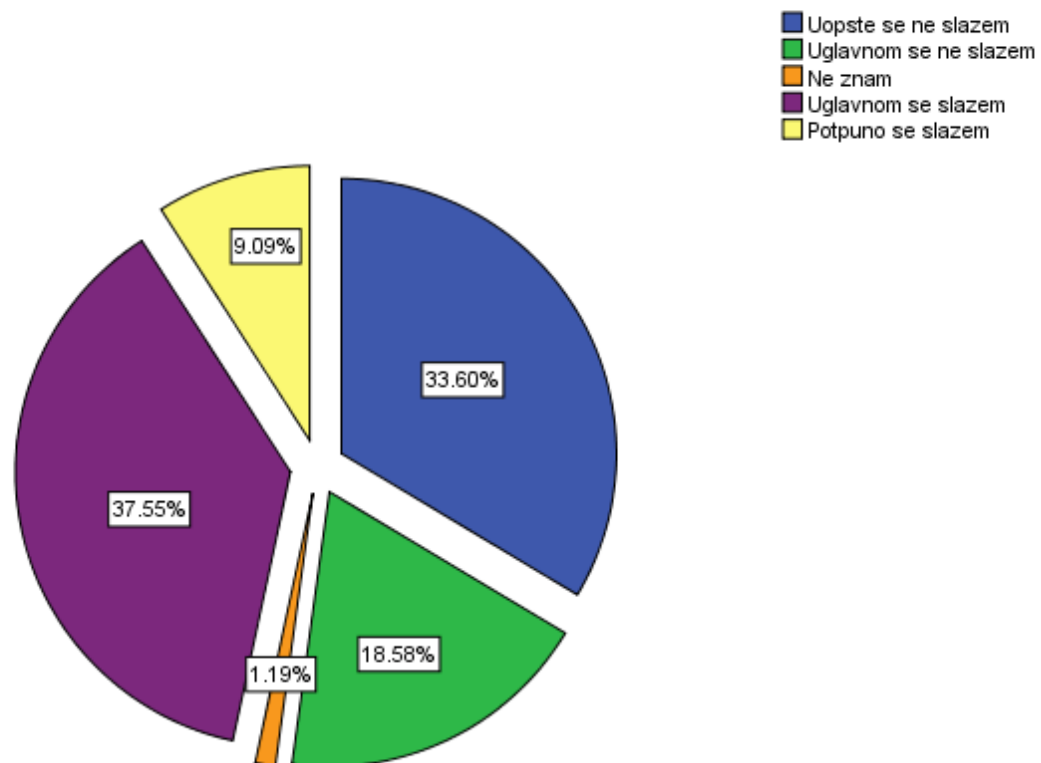
N	Valid	253
	Missing	0
Mean		2.70
Median		2.00
Mode		4
Std. Deviation		1.479
Sum		683

P14

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	85	33.6	33.6	33.6
Uglavnom se ne slazem	47	18.6	18.6	52.2
Ne znam	3	1.2	1.2	53.4
Uglavnom se slazem	95	37.5	37.5	90.9
Potpuno se slazem	23	9.1	9.1	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=95 (37,5%) se uglavnom slaže sa iznetom tvrdnjom, dok se n=23 (9,1%) potpuno slaže. M je 2,70 a SD 1.479.

P14



Grafički prikaz 17. Grafički prikaz stavova ispitanika u odnosu na pitanje P14

P15. Da li podržavate stavda kao deo svoje digitalne strategije, EU želi da reguliše veštačku inteligenciju (AI) kako bi obezbedila bolje uslove za razvoj i upotrebu ove inovativne tehnologije?

Tabela 21. Mišljenje ispitanika u vezi pitanja P15

Statistics

P15

N	Valid	253
	Missing	0
Mean		4.14
Median		4.00
Mode		4
Std. Deviation		.475
Sum		1048

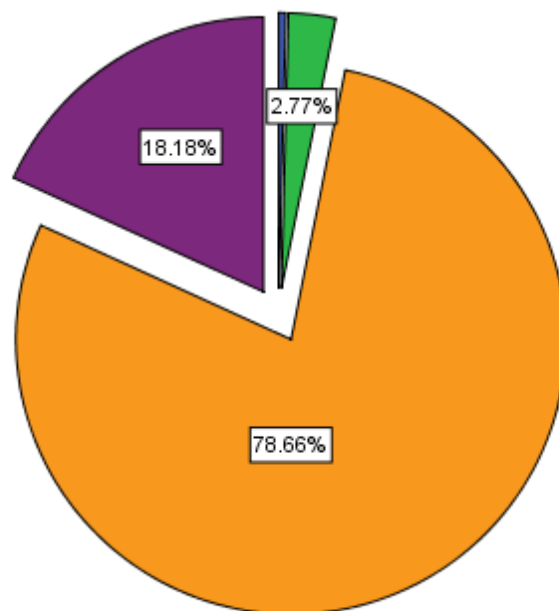
P15

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	1	.4	.4	.4
Ne znam	7	2.8	2.8	3.2
Uglavnom se slazem	199	78.7	78.7	81.8
Potpuno se slazem	46	18.2	18.2	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=199 (78,7%) se uglavnom slaže sa iznetom tvrdnjom, dok se n=46 (18,2%) potpuno slaže. M je 4,14 a SD .475.

P15

- Uopšte se ne slažem
- Ne znam
- Uglavnom se slažem
- Potpuno se slažem



Grafički prikaz 18. Grafički prikaz stavova ispitanika u odnosu na pitanje P15

P16. Da li podržavate stavda je opravdano da Zakon EU o veštačkoj inteligenciji koristi čvrst zakon za rizičnije primene veštačke inteligencije, a meki zakon za manje rizične?

Tabela 22. Izjašnjavanje ispitanika u odnosu na pitanje P16

Statistics

P16

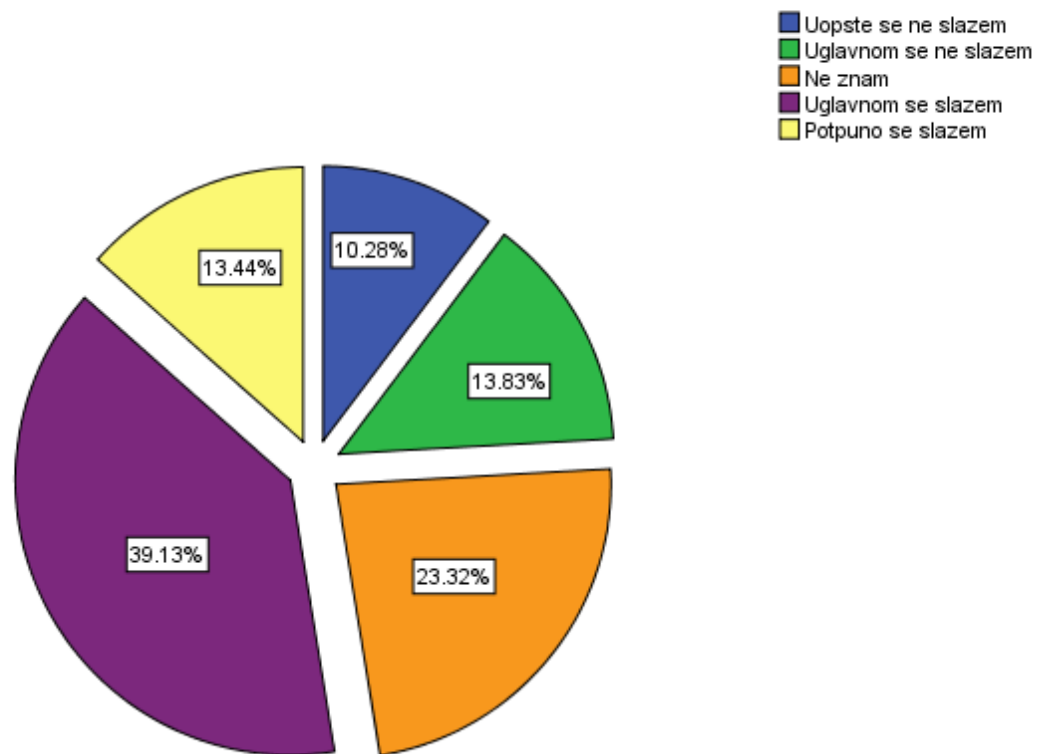
N	Valid	253
	Missing	0
Mean		3.32
Median		4.00
Mode		4
Std. Deviation		1.176
Sum		839

P16

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	26	10.3	10.3	10.3
Uglavnom se ne slazem	35	13.8	13.8	24.1
Ne znam	59	23.3	23.3	47.4
Uglavnom se slazem	99	39.1	39.1	86.6
Potpuno se slazem	34	13.4	13.4	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=99 (39,1%) se uglavnom slaže sa iznetom tvrdnjom, dok se n=34 (13,4%) potpuno slaže. M je 3,32 a SD 1.176.

P16



Grafički prikaz 19. Grafički prikaz stavova ispitanika u odnosu na pitanje P16

P17. Da li podržavate stavda je prilikom korišćenja AI neophodna kontrola zbog zaštite privatnosti podataka i obezbeđivanja korišćenja ovih podataka za ublažavanje štetne pristrasnosti?

Tabela 23. Izjašnjavanje ispitanikau odnosu na pitanje P17

Statistics

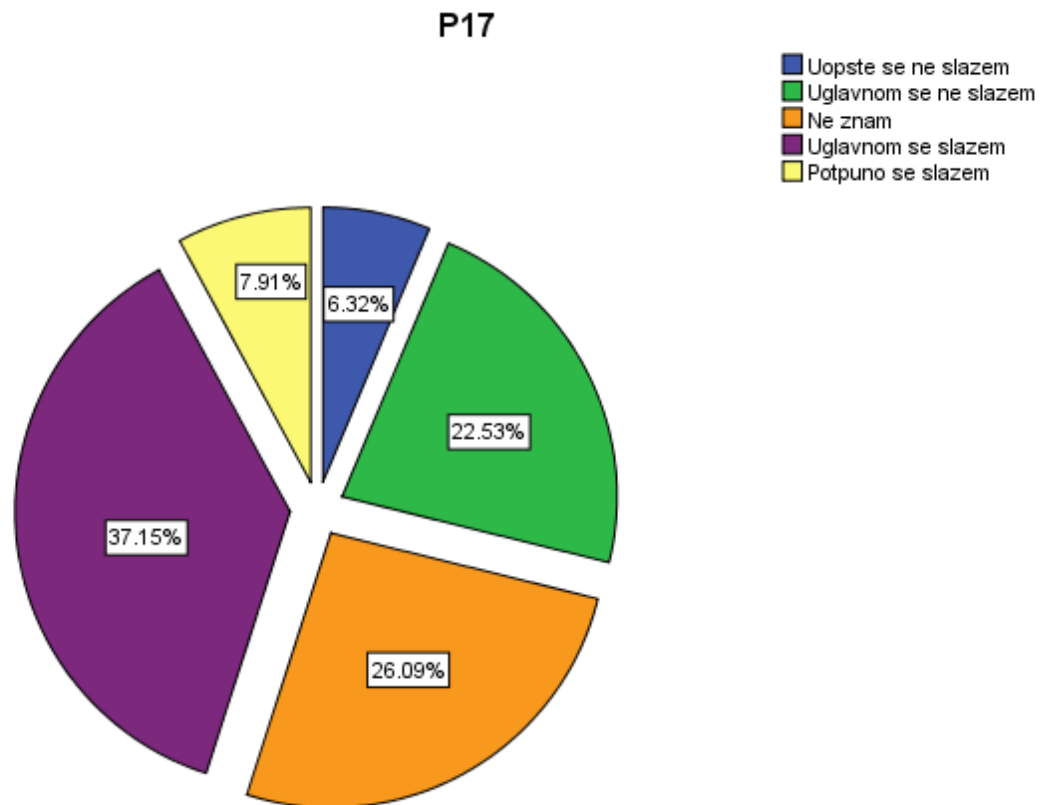
P17

N	Valid	253
	Missing	0
Mean		3.18
Median		3.00
Mode		4
Std. Deviation		1.067
Sum		804

P17

	Frequency	Percent	Valid Percent	Cumulative Percent
Uopste se ne slazem	16	6.3	6.3	6.3
Uglavnom se ne slazem	57	22.5	22.5	28.9
Ne znam	66	26.1	26.1	54.9
Uglavnom se slazem	94	37.2	37.2	92.1
Potpuno se slazem	20	7.9	7.9	100.0
Total	253	100.0	100.0	

Najveći broj ispitanika n=94 (37,2%) se uglavnom slaže sa iznetom tvrdnjom, dok se n=20 (7,9%) potpuno slaže. M je 3,18 a SD 1.067.



Grafički prikaz 20. Grafički prikaz stavova ispitanika u odnosu na pitanje P17

Testiranje postavljenih hipoteza

Glavna hipoteza

H0: Ukoliko je zakonodavstvo uspostavilo odgovarajuću kontrolu i dalo jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju utoliko je manja mogućnost da upotreba AI dovede do štetnih posledica.

Pomoćne hipoteze:

H1: Ukoliko je regulisanje veštačke inteligencije uspostavljeno tako da je postignuta ravnoteža između regulisanja kontrole AI i omogućavanja inovacije i evolucije, utoliko je veća mogućnost da se AI koristi na bezbedan i odgovoran način.

H2: Ukoliko podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana utoliko je veća mogućnost da se otvore pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI.

H3: Ukoliko je zakonodavni pristup definisan oslanjanjem na procenu stepena rizika od upotrebe AI utoliko je veća mogućnost da se AI sistemi visokog rizika suoče sa strogo definisanim načinima delovanja.

Pre nego što smo pristupili testiranju postavljenih hipoteza, poštujući metodološka pravila obavili smo proveru skale. Provera skale koja je sastavljena od 17 pitanja kojom smo ispitivali materiju kojom se bavimo u ovoj doktorskoj disertaciji. Proverili smo unutrašnju saglasnost za skalu sastavljenu da bismo dobili potvrdu da je skala ispravno postavljena i da meri pojavu koju istražujemo.

Tabela 24. Rezime obrade za skalu

Case Processing Summary		
	N	%
Valid	253	100.0
Cases Excluded ^a	0	.0
Total	253	100.0

a. Listwise deletion based on all variables in the procedure.

Tabela 25. Koeficijent Cronbach's Alpha za skalu

Reliability Statistics		
Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.888	.893	17

Tabela 26. Pojedinačna statistika (M; SD) za skalu

Item Statistics

	Mean	Std. Deviation	N
P1	3.30	1.126	253
P2	3.49	1.240	253
P3	2.81	1.229	253
P4	3.31	1.159	253
P5	3.36	1.186	253
P6	3.24	1.084	253
P7	3.32	1.103	253
P8	3.22	1.061	253
P9	3.31	1.159	253
P10	3.39	1.196	253
P11	2.51	1.438	253
P12	3.83	.848	253

P13	3.25	1.237	253
P14	2.70	1.479	253
P15	4.14	.475	253
P16	3.32	1.176	253
P17	3.18	1.067	253

Tabela 27. Rezime statistike stavki za skalu

Summary Item Statistics

	Mean	Minimum	Maximum	Range	Maximum Minimum	Variance	N of Items
Item Means	3.276	2.514	4.142	1.628	1.648	.142	17

Za unutrašnju saglasnost skale koristili smo Cronbach alphu. Dobijeni Cronbachov koeficijent alfa $\alpha=.888$ pokazuje visoku vrednost unutrašnje saglasnosti skale što potvrđuje da je istraživački zadatak dobro postavljen i da se varijable odnose na materiju koju smo ispitali. Sledeći korak u sprovedenom istraživanju odnosio se na proveru subskale koju smo sastavili od varijabli izabranih za testiranje glavne hipoteze i pomoćnih hipoteza. Te varijable su:

P1. Da li podržavate stavda zakonodavstvo mora da uspostavi odgovarajuću kontrolu i da jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju?

P2. Da li podržavate stavda upotreba AI može dovesti do štetnih posledica?

P3. Da li podržavate stavda regulisanje veštačke inteligencije mora da postigne ravnotežu između regulisanja kontrole AI i omogućavanja inovacije i evolucije?

P4. Da li podržavate stavda AI mora da se koristi na bezbedan i odgovoran način

P5. Da li podržavate stavda podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana?

P6. Da li podržavate stavda upotreba AI otvara pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI?

P7. Da li podržavate stavda upotreba AI sistema mora da se klasifikuje u odnosu na stepen rizika do koga dolazi prilikom primene AI?

P8. Da li podržavate stavda AI sistemi visokog rizika moraju da budi obuhvaćeni strogo definisanim načinima delovanja?

Prvo smo pristupili testiranju normalnosti raspodele za varijable koje čine subskalu.

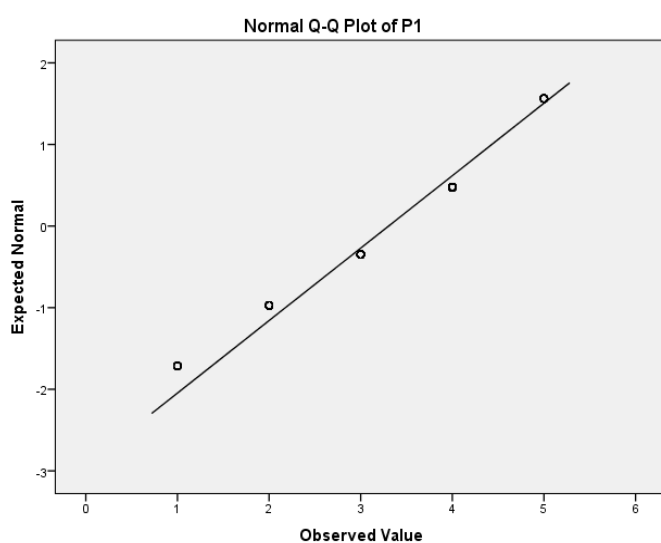
Tabela 28. Procena normalnosti raspodele

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
P1	.249	253	.000	.888	253	.000
P2	.305	253	.000	.832	253	.000
P3	.223	253	.000	.891	253	.000
P4	.297	253	.000	.832	253	.000
P5	.302	253	.000	.834	253	.000
P6	.252	253	.000	.881	253	.000
P7	.258	253	.000	.884	253	.000

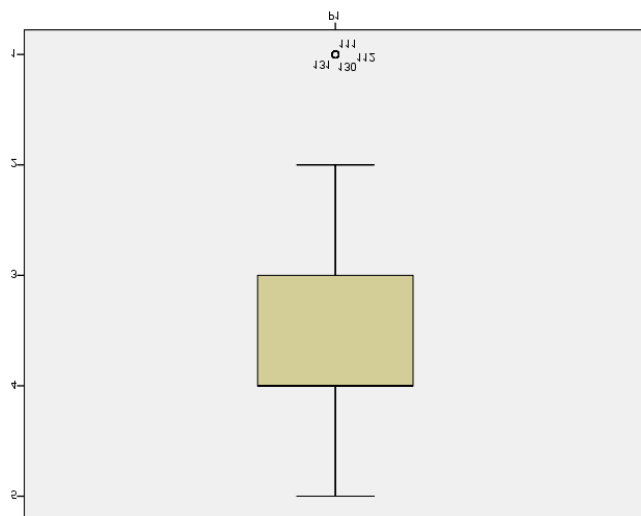
P8	.263	253	.000	.870	253	.000
----	------	-----	------	------	-----	------

a. Lilliefors Significance Correction

Normalnost raspodele se pokazuje statističkim odstupanjem od normalnosti odnosno Sig većim od 0,05. Iz tabele vidimo da je Sig .000, zbog čega zaključujemo da primenom ovih testova pretpostavka o normalnosti raspodele nije potvrđena. Međutim, ovakav rezultat je čest kada su u pitanju veliki uzorci ispitanika >50. Zbog toga smo pristupili analizi histograma, kao pouzdanijoj metodi provere normalnosti.



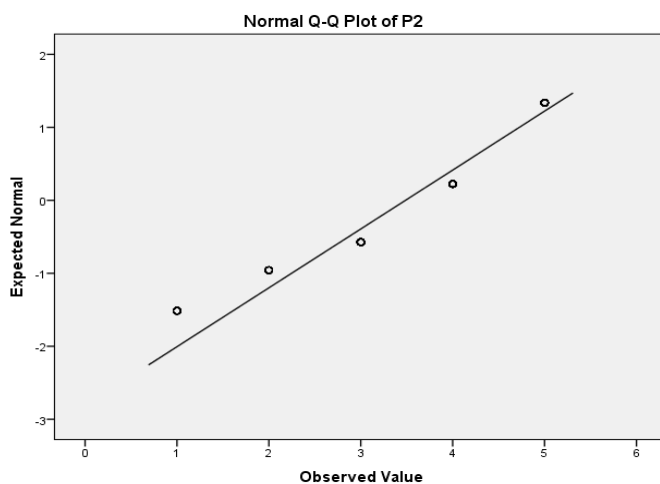
Grafički prikaz 21. Testiranje normalnosti raspodele P1



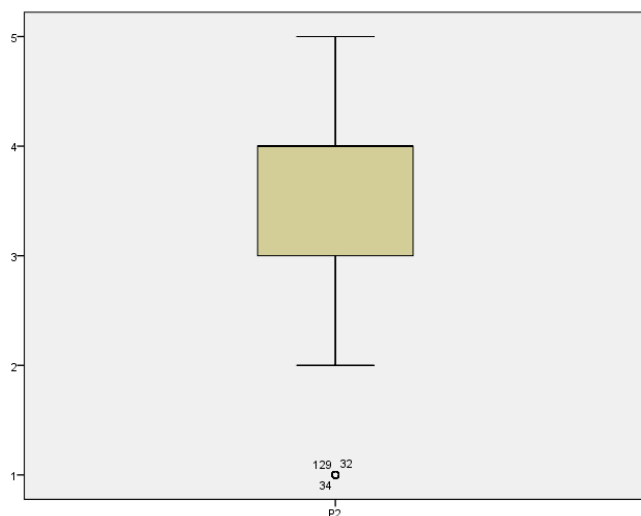
Grafički prikaz 22. Netipične tačke za P1

Rezultati su normalno raspoređeni što potvrđuje izgled krive normalne verovatnoće, s obzirom na to da su rezultati blizu pravoj liniji.

Pravougaoni dijagram pokazuje četiri netipične tačke koje su od ivice pravougaonika udaljene više od 1,5 njegovih dužina. S obzirom na to da nema ekstremnih vrednosti koje se označavaju * prihvatamo pretpostavku normalnosti raspodele.



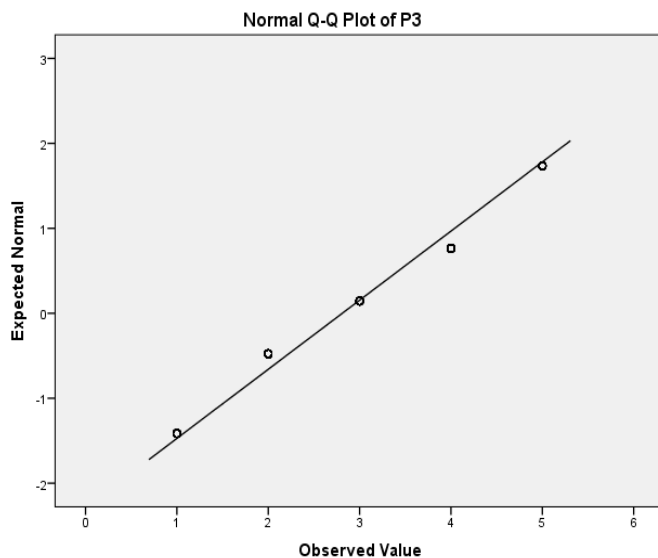
Grafički prikaz 23. Testiranje normalnosti raspodele P2



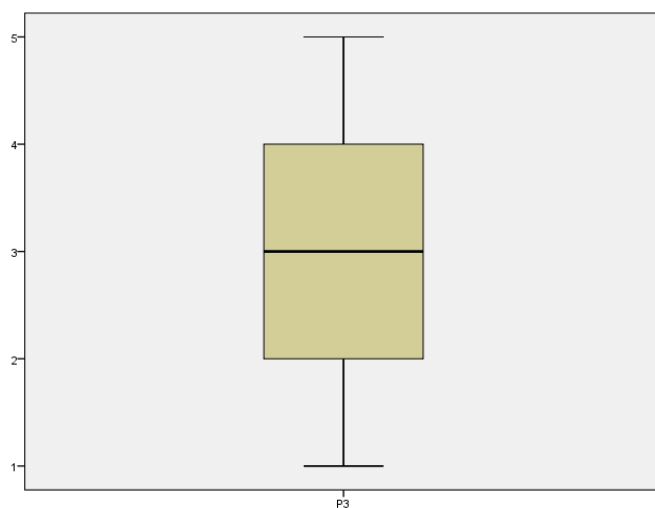
Grafički prikaz 24. Netipične tačke za P2

Rezultati su normalno raspoređeni što potvrđuje izgled krive normalne verovatnoće, s obzirom na to da su rezultati blizu pravoj liniji.

Pravougaoni dijagram pokazuje tri netipične tačke koje su od ivice pravougaonika udaljene više od 1,5 njegovih dužina. S obzirom na to da nema ekstremnih vrednosti koje se označavaju * prihvatamo pretpostavku normalnosti raspodele.



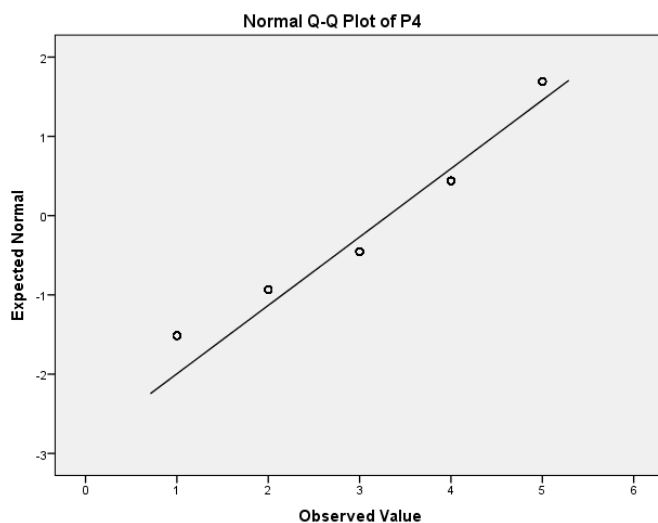
Grafički prikaz 25. Testiranje normalnosti raspodele P3



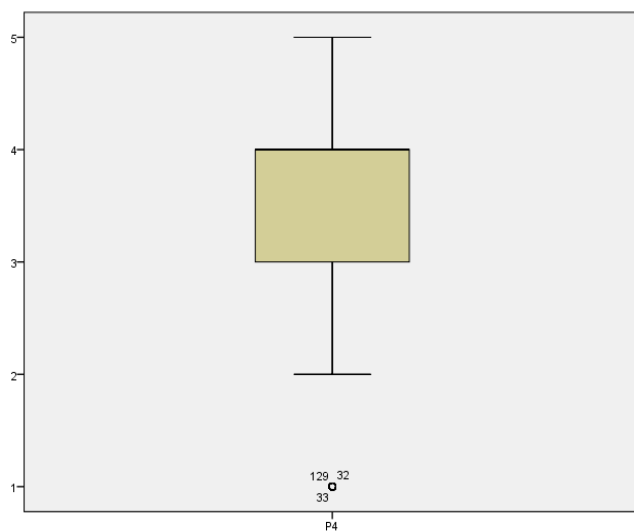
Grafički prikaz 26. Netipične tačke za P3

Rezultati su normalno raspoređeni što potvrđuje izgled krive normalne verovatnoće, s obzirom na to da su rezultati blizu pravoj liniji.

Pravougaoni dijagram pokazuje da nema netipičnih tačkaka koje su od ivice pravougaonika udaljene više od 1,5 njegovih dužina. S obzirom na to da nema ekstremnih vrednosti koje se označavaju * prihvatamo pretpostavku normalnosti raspodele.



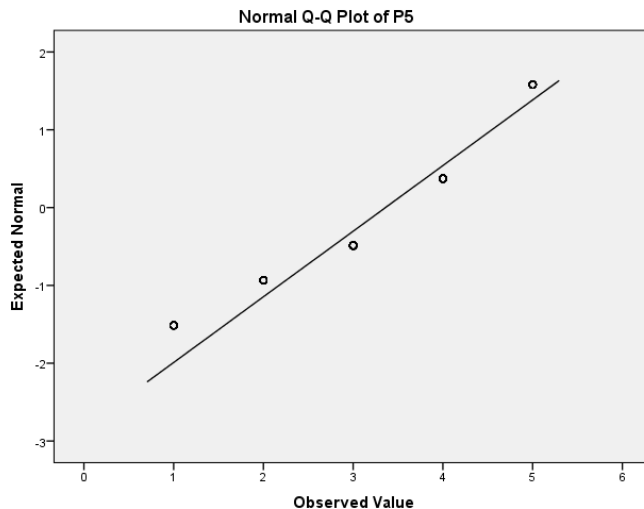
Grafički prikaz 27. Testiranje normalnosti raspodele P4



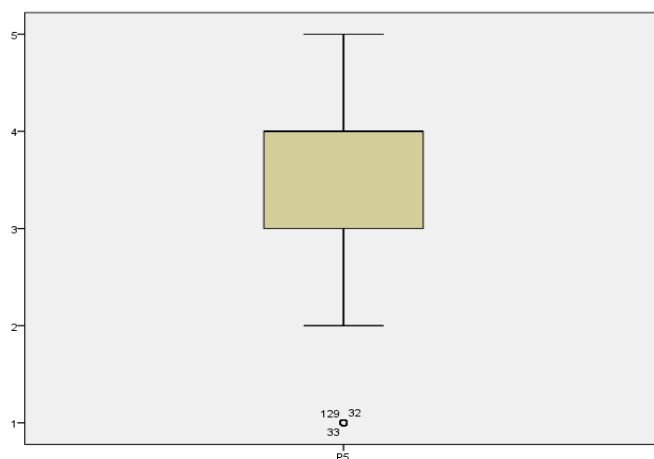
Grafički prikaz 28. Netipične tačke za P4

Rezultati su normalno raspoređeni što potvrđuje izgled krive normalne verovatnoće, s obzirom na to da su rezultati blizu pravoj liniji.

Pravougaoni dijagram pokazuje tri netipične tačke koje su od ivice pravougaonika udaljene više od 1,5 njegovih dužina. S obzirom na to da nema ekstremnih vrednosti koje se označavaju * prihvatamo pretpostavku normalnosti raspodele.



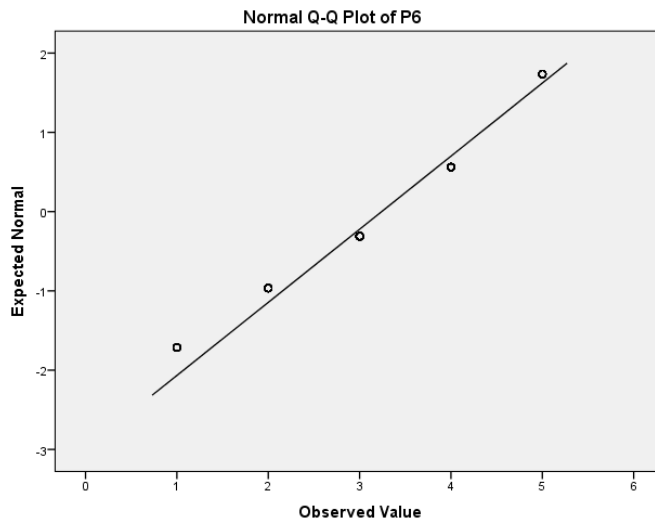
Grafički prikaz 29. Testiranje normalnosti raspodele P5



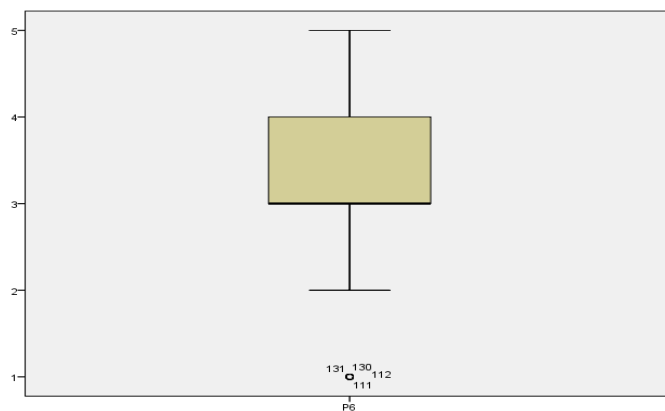
Grafički prikaz 30. Netipične tačke za P5

Rezultati su normalno raspoređeni što potvrđuje izgled krive normalne verovatnoće, s obzirom na to da su rezultati blizu pravoj liniji.

Pravougaoni dijagram pokazuje tri netipične tačke koje su od ivice pravougaonika udaljene više od 1,5 njegovih dužina. S obzirom na to da nema ekstremnih vrednosti koje se označavaju * prihvatamo pretpostavku normalnosti raspodele.



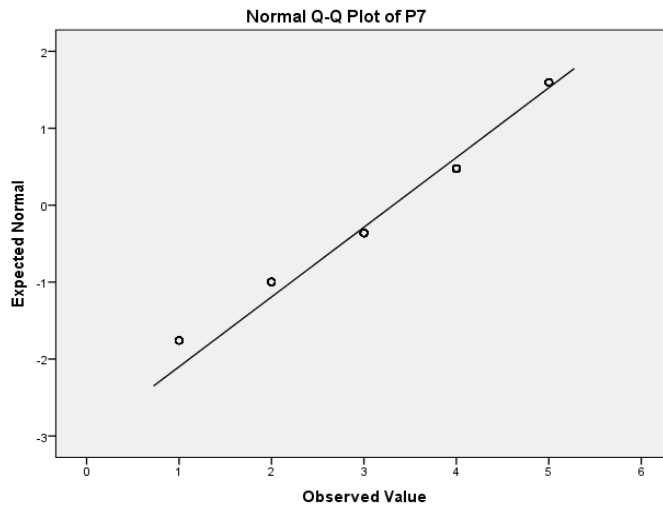
Grafički prikaz 31. Testiranje normalnosti raspodele P6



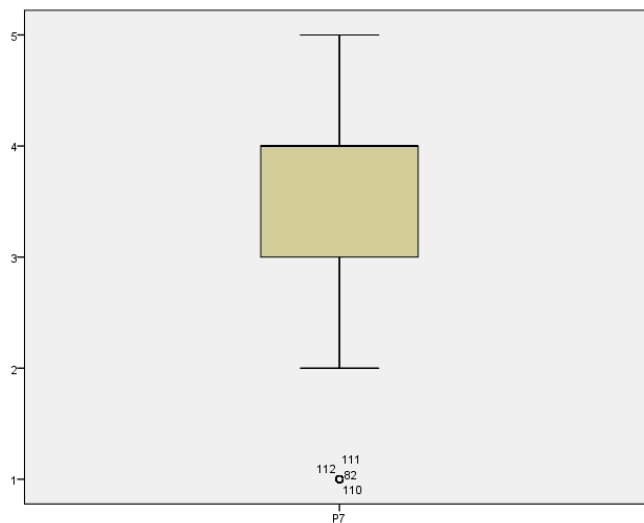
Grafički prikaz 32. Netipične tačke za P6

Rezultati su normalno raspoređeni što potvrđuje izgled krive normalne verovatnoće, s obzirom na to da su rezultati blizu pravoj liniji.

Pravougaoni dijagram pokazuje četiri netipične tačke koje su od ivice pravougaonika udaljene više od 1,5 njegovih dužina. S obzirom na to da nema ekstremnih vrednosti koje se označavaju * prihvatamo pretpostavku normalnosti raspodele.



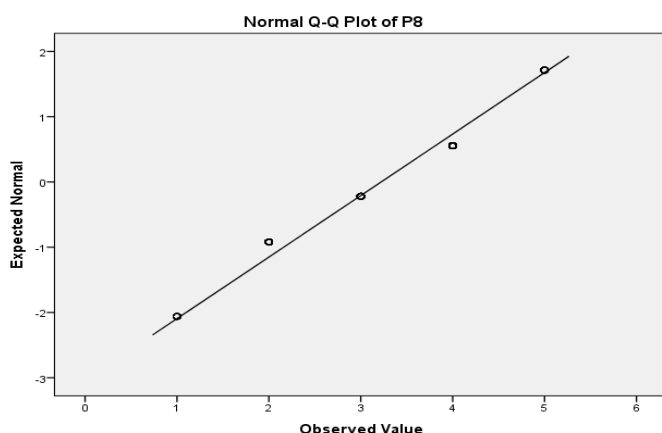
Grafički prikaz 33. Testiranje normalnosti raspodele P7



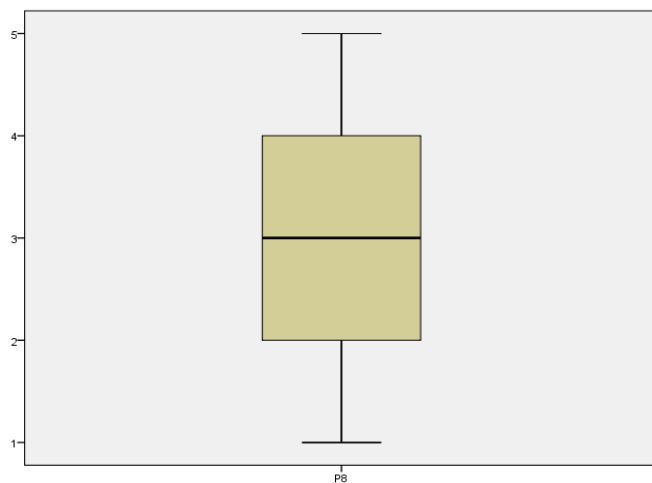
Grafički prikaz 34. Netipične tačke za P7

Rezultati su normalno raspoređeni što potvrđuje izgled krive normalne verovatnoće, s obzirom na to da su rezultati blizu pravoj liniji.

Pravougaoni dijagram pokazuje četiri netipične tačke koje su od ivice pravougaonika udaljene više od 1,5 njegovih dužina. S obzirom na to da nema ekstremnih vrednosti koje se označavaju * prihvatamo pretpostavku normalnosti raspodele.



Grafički prikaz 35. Testiranje normalnosti raspodele P8



Grafički prikaz 36. Netipične tačke za P8

Rezultati su normalno raspoređeni što potvrđuje izgled krive normalne verovatnoće, s obzirom na to da su rezultati blizu pravoj liniji.

Pravougaoni dijagram pokazuje da nema netipičnih tačaka koje su od ivice pravougaonika udaljene više od 1,5 njegovih dužina. S obzirom na to da nema ekstremnih vrednosti koje se označavaju * prihvatamo pretpostavku normalnosti raspodele.

Dalji postupak uključivao je testiranje unutrašnje saglasnosti subskale koju smo sastavili od varijabli izabranih za testiranje opšte hipoteze i posebnih hipoteza.

Tabela 29. Rezime obrade za subskalu

Case Processing Summary

		N	%
Cases	Valid	253	100.0
	Excluded ^a	0	.0
	Total	253	100.0

a. Listwise deletion based on all variables in the procedure.

Tabela 30. Koeficijent Cronbach's Alpha za subskalu

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
.911	.912	8

Tabela 31. Pojedinačna statistika (M; SD) za subskalu

Item Statistics

	Mean	Std. Deviation	N
P1	3.30	1.126	253
P2	3.49	1.240	253
P3	2.81	1.229	253
P4	3.31	1.159	253
P5	3.36	1.186	253
P6	3.24	1.084	253
P7	3.32	1.103	253
P8	3.22	1.061	253

Tabela 32. Korelaciona matrica za subskalu

Inter-Item Correlation Matrix

	P1	P2	P3	P4	P5	P6	P7	P8
1	1.000	.542	.128	.587	.563	.889	.871	.638
2	.542	1.000	.535	.900	.874	.488	.581	.518
3	.128	.535	1.000	.585	.589	.198	.208	.151
4	.587	.900	.585	1.000	.966	.512	.509	.518
5	.563	.874	.589	.966	1.000	.491	.486	.498
6	.889	.488	.198	.512	.491	1.000	.891	.498
7	.871	.581	.208	.509	.486	.891	1.000	.591
8	.638	.518	.151	.518	.498	.498	.591	1.000

Tabela 33. Rezime statistike stavki za subskalu

Summary Item Statistics							
	Mean	Minimum	Maximum	Range	Maximum Minimum	Variance	N of Items
Item Means	3.256	2.810	3.486	.676	1.241	.039	8

Unutrašnju saglasnost subskale potvrdili smo uz pomoć Crombach alphe. Kronbahov koeficijent alfa $\alpha=.911$ ukazao je na izuzetno visoku vrednost unutrašnje saglasnosti subskale i na taj način smo dobili potvrdu da su istraživačka pitanja adekvatna za testiranje postavljenih hipoteza.

U sumiranju rezultata dobijanih testiranjem subskale rukovali smo se time da Kronbahov koeficijent alfa daje rezultat koji se odnosi na prosečnu korelaciju između vrednosti na skali. Mogući rezultati su na lestvici od 0 do 1, na osnovu čega se gradi stepen pouzdanosti. Karakteristično za očekivane rezultate je to da vrednost manja od 0,7 ukazuje da skala nije dobro odabrana. Naša skala je pokazala vrednost .911.

S obzirom na to da smo u testiranje uključili 8 varijabli, osvrnuli smo se na stav istraživača koji iskazuju bojazan da testiranje skala sa brojem varijabli ispod 10 može da se dobije neadekvatan rezultat, izvršili smo i dodatnu proveru skalu tako što smo izračunali srednje vrednosti korelacije između svakog para vrednosti. Rezultat ove provere je pokazao da se pojedinačna statistika (M; SD) za subskalu kreće u rasponu od 2.81 do 3.49, čime smo dobili pozitivan odgovor s obzirom na to da se kao optimalna srednja vrednost korelacije između parova vrednosti na skali prihvata rezultat od 0.2 do 0.4.

Testiranje glavne hipoteze:

Za proveru opravdanosti glavne hipoteze: H0: Ukoliko je zakonodavstvo uspostavilo odgovarajuću kontrolu i dalo jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju utoliko je manja mogućnost da upotreba AI dovede do štetnih posledica, testirali smo odgovore ispitanika na iznete tvrdnje:

P1. Da li podržavate stavda zakonodavstvo mora da uspostavi odgovarajuću kontrolu i da jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju?

P2. Da li podržavate stavda upotreba AI može dovesti do štetnih posledica?

Tabela 34. Vrednosti Chi-Square Testa za generalnu hipotezu

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	328.768 ^a	16	.000
Likelihood Ratio	280.068	16	.000
Linear-by-Linear Association	73.959	1	.000
N of Valid Cases	253		

a. 9 cells (36.0%) have expected count less than 5. The minimum expected count is 1.74.

Tabela 35. Simetrične mere za generalnu hipotezu

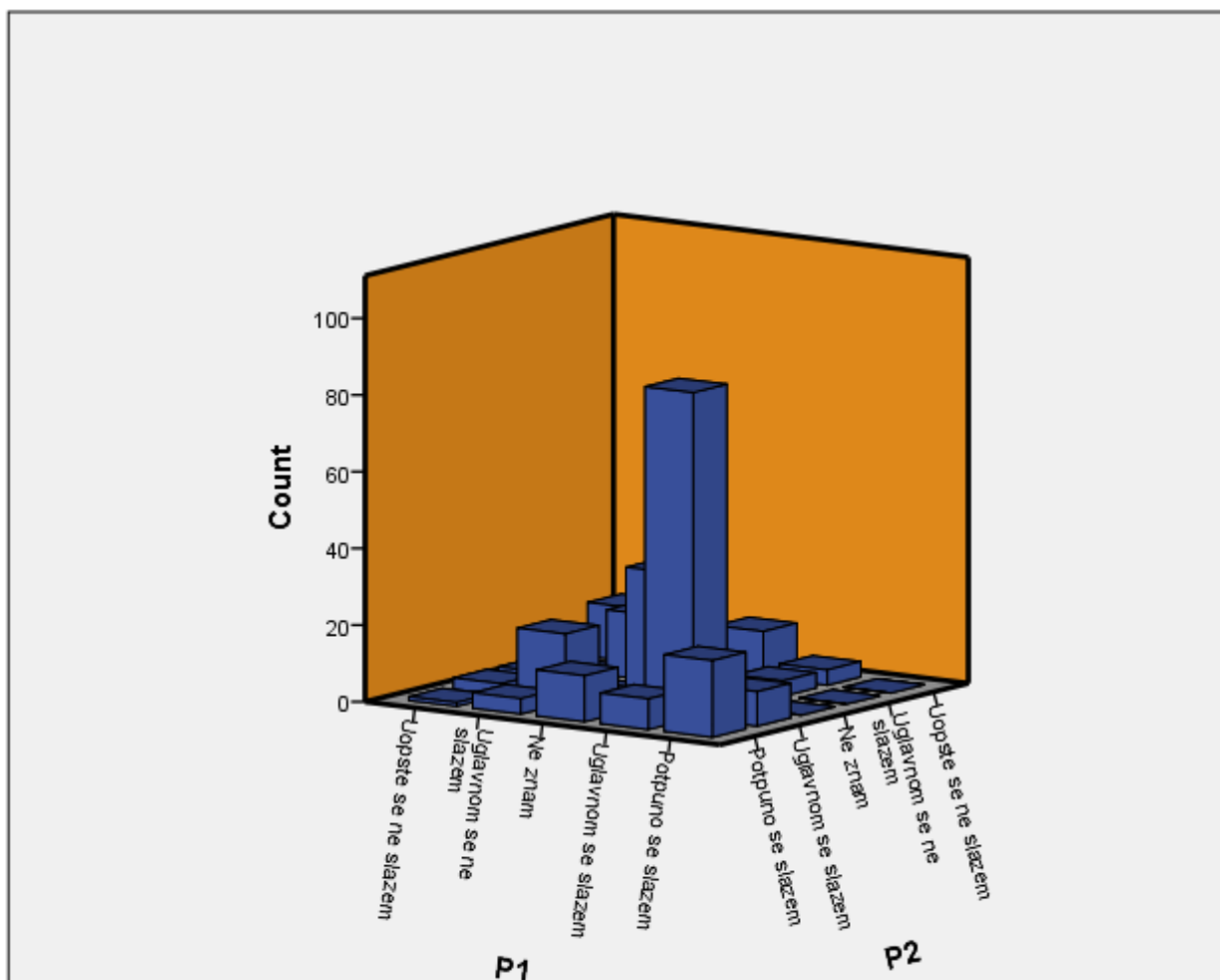
Symmetric Measures					
		Value	Asymp. Std. Error ^a	Approx. T ^b	Approx. Sig.
Ordinal by Ordinal	Gamma	.585	.054	9.384	.000
	Spearman Correlation	.540	.053	10.166	.000 ^c
	Pearson's R	.542	.053	10.211	.000 ^c
Interval by Interval					
N of Valid Cases		253			

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

c. Based on normal approximation.

Rezultat prikazan u tabelama pokazuje da je $X^2(16,1) = 328.768^a$, $p < 0.01$, što ukazuje na postojanje statističke značajnosti dobijenih rezultata. Koeficijent Pirsonove korelacije $r = .542$ potvrđuje da je između ispitivanih varijabli prisutan visok stepen pozitivne korelacije. Vrednost Spearman Correlation rangova je $Rho = .540$. Dobijen Gamma rezultat .585 usmerava nas na to da stepen saglasnosti sa prvom tvrdnjom omogućava predviđanje ishoda odgovora na drugu tvrdnju za 58,5%.



Grafički prikaz 37. Grafički prikaz testiranja generalne hipoteze.

Na osnovu statističke obrade testiranih varijabli zaključujemo da je generalna hipoteza: H0: Ukoliko je zakonodavstvo uspostavilo odgovarajuću kontrolu i dalo jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju utoliko je manja mogućnost da upotreba AI dovede do štetnih posledica, potvrđena.

Posebne hipoteze:

Za proveru opravdanosti prve posebne hipoteze H1 koja glasi: Ukoliko je regulisanje veštačke inteligencije uspostavljeno tako da je postignuta ravnoteža između regulisanja kontrole

AI i omogućavanja inovacije i evolucije, utoliko je veća mogućnost da se AI koristi na bezbedan i odgovoran način, testirali smo odgovore ispitanika na iznete tvrdnje:

P3. Da li podržavate stavda regulisanje veštačke inteligencije mora da postigne ravnotežu između regulisanja kontrole AI i omogućavanja inovacije i evolucije?

P4. Da li podržavate stavda AI mora da se koristi na bezbedan i odgovoran način?

Tabela 36. Vrednosti Chi-Square Testa za prvu posebnu hipotezu

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	355.013 ^a	16	.000
Likelihood Ratio	240.060	16	.000
Linear-by-Linear Association	86.227	1	.000
N of Valid Cases	253		

a. 9 cells (36.0%) have expected count less than 5. The minimum expected count is 1.74.

Tabela 37. Simetrične mere za prvu posebnu hipotezu

Symmetric Measures

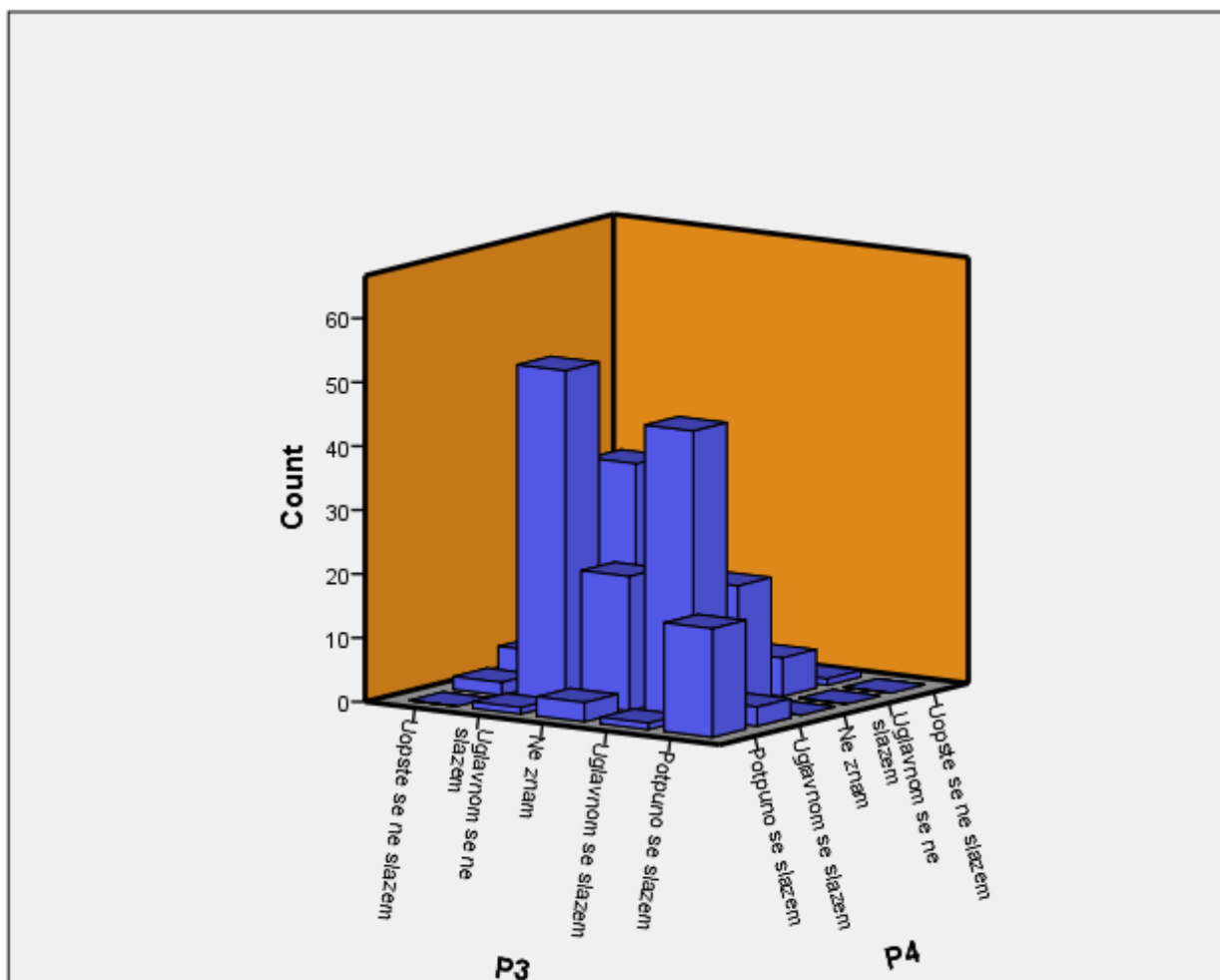
		Value	Asymp Std. Error ^a	Ap prox. T ^b	Appr ox. Sig.
Ordinal by Ordinal	Gamma	.603	.057	8.529	.000
	Spearman Correlation	.539	.055	10.149	.000 ^c
	Interval by Interval	.585	.045	11.426	.000 ^c
N of Valid Cases		253			

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

c. Based on normal approximation.

Rezultat prikazan u tabelama pokazuje da je $X^2(16,1) = 355.013^a$, $p < 0.01$, što ukazuje na postojanje statističke značajnosti dobijenih rezultata. Koeficijent Pirsonove korelacije $r = .585$ potvrđuje da je između ispitivanih varijabli prisutan visok stepen pozitivne korelacije. Vrednost Spearman Correlation rangova je $Rho = .539$. Dobijen Gamma rezultat .603 usmerava nas na to da stepen saglasnosti sa prvom tvrdnjom omogućava predviđanje ishoda odgovora na drugu tvrdnju za 60,3%.



Grafički prikaz 38. Grafički prikaz korelacije testiranja prve posebne hipoteze

Na osnovu statističke obrade testiranih varijabli zaključujemo da je prva posebna hipoteza: H1: Ukoliko je regulisanje veštačke inteligencije uspostavljeno tako da je postignuta ravnoteža između regulisanja kontrole AI i omogućavanja inovacije i evolucije, utoliko je veća mogućnost da se AI koristi na bezbedan i odgovoran način, potvrđena.

Za proveru opravdanosti druge posebne hipoteze H2 koja glasi: Ukoliko podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana utoliko je veća mogućnost da se otvore pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI, testirali smo odgovore ispitanika na iznete tvrdnje:

P5. Da li podržavate stavda podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana?

P6. Da li podržavate stavda upotreba AI otvara pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI?

Tabela 38. Vrednosti Chi-Square Testa za drugu posebnu hipotezu

Chi-Square Tests			
	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	132.950 ^a	16	.000
Likelihood Ratio	108.011	16	.000
Linear-by-Linear Association	60.739	1	.000
N of Valid Cases	253		

a. 10 cells (40.0%) have expected count less than 5. The minimum expected count is 1.90.

Tabela 39. Simetrične mere za drugu posebnu hipotezu

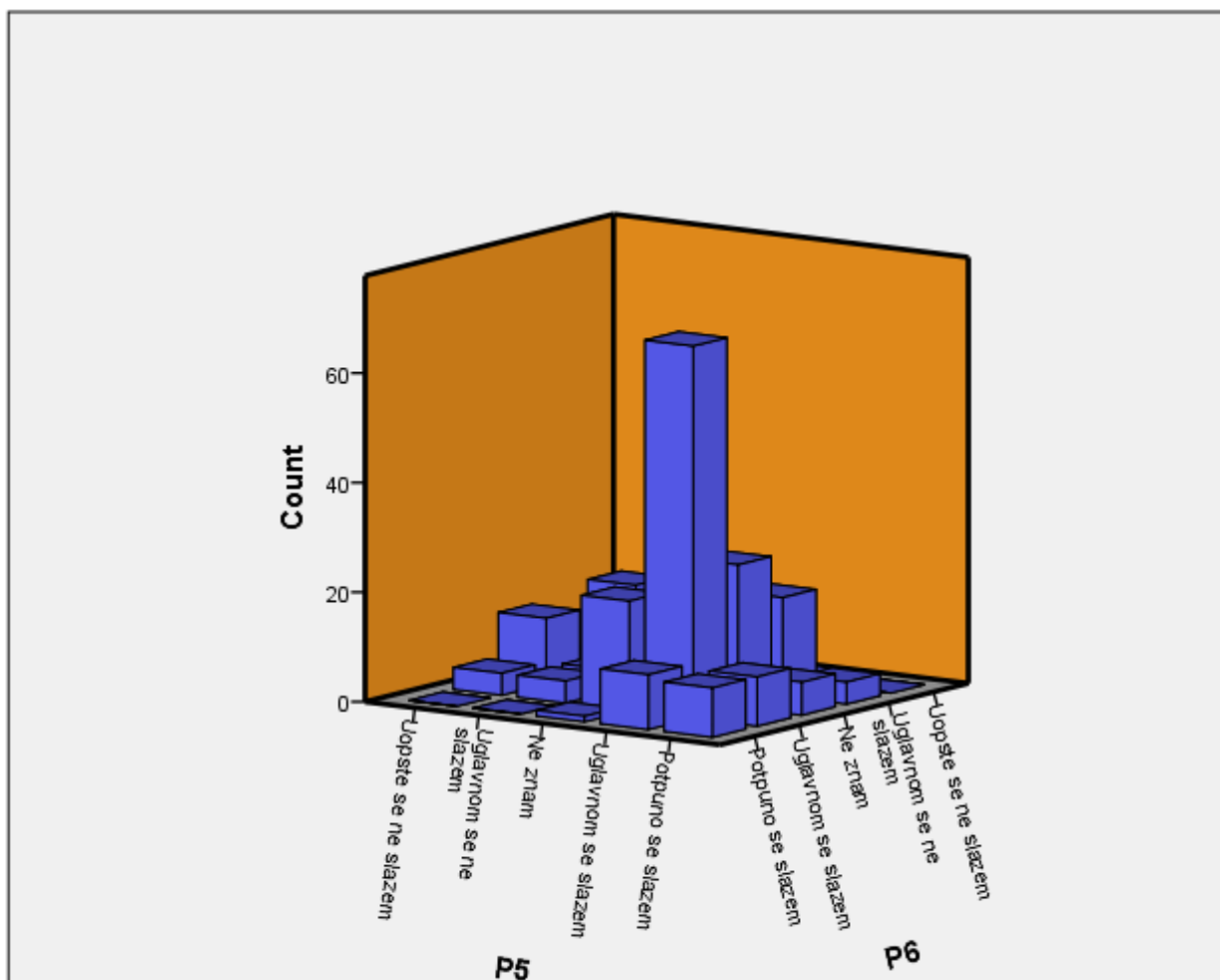
Symmetric Measures					
		Value	Asymp. Std. Error ^a	Approx. T ^b	Approx. Sig.
Ordinal by Ordinal	Gamma	.517	.060	7.437	.000
	Spearman Correlation	.445	.055	7.864	.000 ^c
	Pearson's R	.491	.053	8.928	.000 ^c
Interval by Interval					
N of Valid Cases		253			

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

c. Based on normal approximation.

Rezultat prikazan u tabelama pokazuje da je $X^2(16,1) = 132.950^a$, $p < 0.01$, što ukazuje na postojanje statističke značajnosti dobijenih rezultata. Koeficijent Pirsonove korelacije $r = .491$ potvrđuje da je između ispitivanih varijabli prisutan srednje visok stepen pozitivne korelacije. Vrednost Spearman Correlation rangova je $Rho = .445$. Dobijen Gamma rezultat .571 usmerava nas na to da stepen saglasnosti sa prvom tvrdnjom omogućava predviđanje ishoda odgovora na drugu tvrdnju za 57,1%.



Grafički prikaz 39. Grafički prikaz korelacije testiranja druge posebne hipoteze

Na osnovu statističke obrade testiranih varijabli zaključujemo da je druga posebna hipoteza: H2: Ukoliko podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana utoliko je veća mogućnost da se otvore pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI, potvrđena.

Za proveru opravdanosti treće posebne hipoteze H3 koja glasi: Ukoliko je zakonodavni pristup definisan oslanjanjem na procenu stepena rizika od upotrebe AI utoliko je veća mogućnost da se AI sistemi visokog rizika suoče sa strogo definisanim načinima delovanja, testirali smo odgovore ispitanika na iznete tvrdnje:

P7. Da li podržavate stavda upotreba AI sistema mora da se klasifikuje u odnosu na stepen rizika do koga dolazi prilikom primene AI?

P8. Da li podržavate stavda AI sistemi visokog rizika moraju da budi obuhvaćeni strogo definisanim načinima delovanja?

Tabela 40. Vrednosti Chi-Square Testa za treću posebnu hipotezu

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	222.183 ^a	16	.000
Likelihood Ratio	193.440	16	.000
Linear-by-Linear Association	88.009	1	.000
N of Valid Cases	253		

a. 10 cells (40.0%) have expected count less than 5. The minimum expected count is .68.

Tabela 41. Simetrične mere za treću posebnu hipotezu

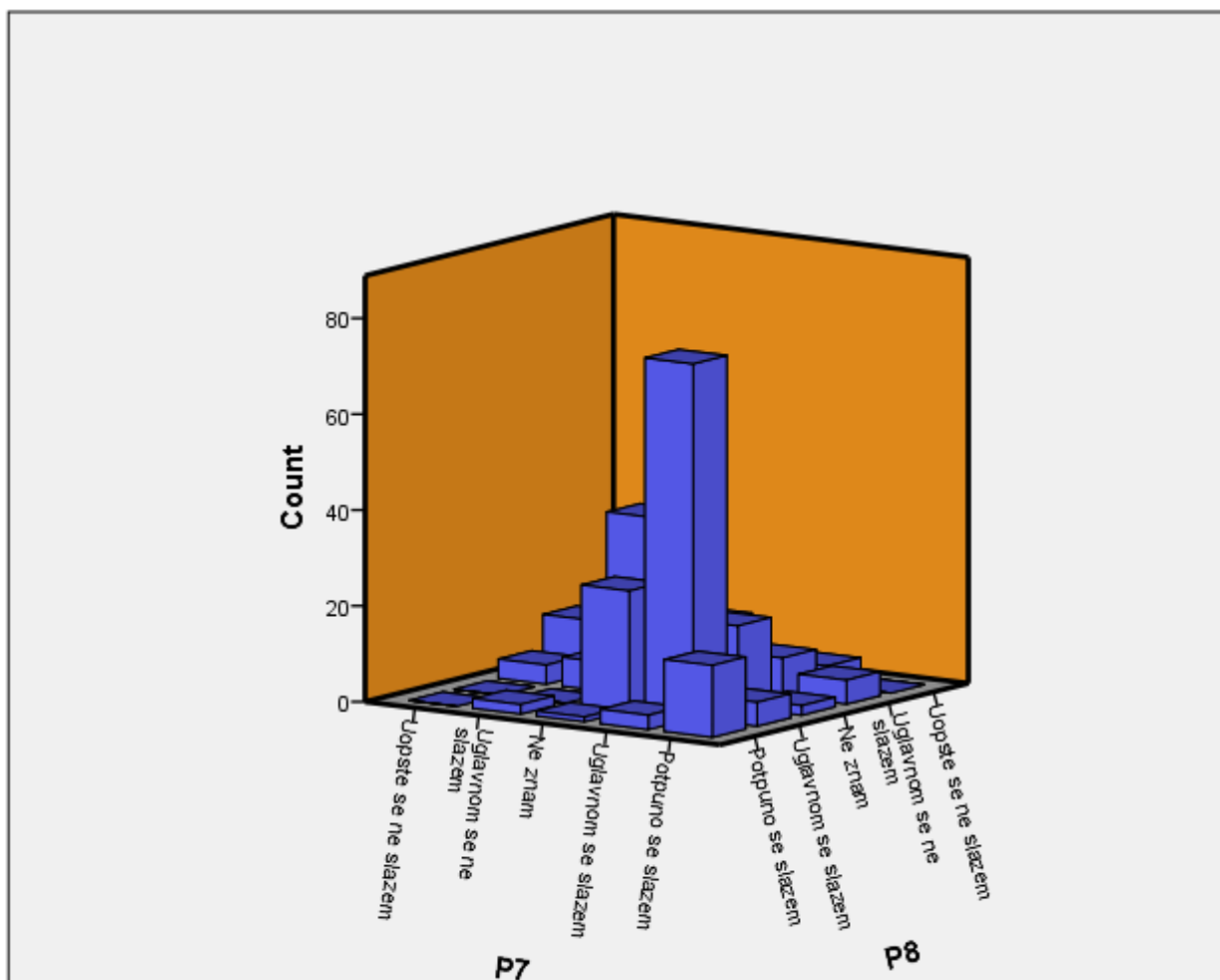
Symmetric Measures					
		Value	Asymp. Std. Error ^a	Approx. T ^b	Approx. Sig.
Ordinal by Ordinal	Gamma	.665	.056	10.467	.000
	Spearman Correlation	.587	.052	11.494	.000 ^c
	Pearson's R	.591	.051	11.606	.000 ^c
Interval by Interval					
N of Valid Cases		253			

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

c. Based on normal approximation.

Rezultat prikazan u tabelama pokazuje da je $X^2(16,1) = 222.183^a$, $p < 0.01$, što ukazuje na postojanje statističke značajnosti dobijenih rezultata. Koeficijent Pirsonove korelacije $r = .591$ potvrđuje da je između ispitivanih varijabli prisutan visok stepen pozitivne korelacije. Vrednost Spearman Correlation rangova je $Rho = .587$. Dobijen Gamma rezultat .665 usmerava nas na to da stepen saglasnosti sa prvom tvrdnjom omogućava predviđanje ishoda odgovora na drugu tvrdnju za 66,5%.



Grafički prikaz 40. Grafički prikaz korelacije testiranja treće posebne hipoteze

Na osnovu statističke obrade testiranih varijabli zaključujemo da je treća posebna hipoteza: H3: Ukoliko je zakonodavni pristup definisan oslanjanjem na procenu stepena rizika od upotrebe AI utoliko je veća mogućnost da se AI sistemi visokog rizika suoče sa strogo definisanim načinima delovanja, potvrđena.

ZAKLJUČNA RAZMATRANJA

Regulisanje veštačke inteligencije je teško zbog potrebe da se uspostavi ravnoteža između omogućavanja inovacije i evolucije, dok se AI koristi na bezbedan i odgovoran način. Jedna stvar je jasna a to je da omogućavanje nesputanog razvoja veštačke inteligencije može biti neodrţivo.

Nema sumnje da moć veštačke inteligencije otključava značajnu efikasnost i mogućnosti, sa skoro neograničenim obimom primene. Međutim, sa ovim mogućnostima dolazi mnoštvo novih rizika i potencijalnih šteta koje AI predstavlja organizacijama, ljudima i društvu koje se moraju uzeti u obzir. Neki od ovih primera su:

Jedna od prvih briga sa kojima se moraju pozabaviti AI programeri, korisnici i operateri jesu pitanja privatnosti koja su pokrenuta razvojem i upotrebom AI, posebno u oblasti pristanka i uključivanja ličnih podataka u AI modeliranje i unose. Pristrasnost i diskriminacija AI može biti pod uticajem slučajne pristrasnosti i diskriminacije. Postoje inherentni problemi koji su ugrađeni u skupove ulaznih podataka koji se koriste za obuku AI algoritama koji mogu dovesti do neţeljenih ishoda.

Postoji niz pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI. Vrlo malo je rečeno o tome koja organizacija u lancu razvoja AI treba da popravi takva pitanja kada se pojave, otvarajući jaz u odgovornosti. Źivotni ciklus veštačke inteligencije je sloţen i uključuje mnogo različitih operatera, od kojih svi imaju potencijal da utiču na kvalitet tehnologije.

Tehnološki razvoj povezan sa sajber bezbednošću i sajber kriminalom znači da, dok se tradicionalni zakoni mogu donekle primeniti, zakonodavstvo takođe mora da se nosi sa novim konceptima i objektima, kao što su nematerijalni kompjuterski podaci, koji se tradicionalno ne obrađuju zakonom.

Ova dela u velikoj meri karakterišu novi nematerijalni objekti, kao što su podaci ili informacije. Iako se krivično pravo često smatra najrelevantnijim kada je u pitanju sajber kriminal, pravni odgovori na širu zabrinutost u vezi sa sajber bezbednošću takođe uključuju upotrebu drugih grana prava, kao što su građansko pravo i upravno pravo.

Lansiranje ChatGPT-a u novembru 2022. fundamentalno je promenilo naš odnos sa veštačkom inteligencijom. Aplikacija je postala viralna, postavši najbrže rastuća aplikacija za potrošače u istoriji i upознаvši nas sa generativnom veštačkom inteligencijom, tehnologijom sa mogućnošću generisanja novog sadržaja koji se ne razlikuje – čak i superiorniji – od ljudskog rezultata. Od tog datuma, Generative AI se razvijao vrtoglavim tempom, šireći se od teksta do generisanja detaljnih slika, proizvodnje čitavih video sekvenci i stvaranja novih muzičkih kompozicija. Mnogi ljudi sada rutinski koriste tehnologiju da poboljšaju svoj svakodnevni život, za zadatke koji se kreću od administrativnih (uređenje rasporeda praznika) do duboko ličnih (trening kroz težak lični izazov). Uticaj ove transformativne tehnologije se takođe oseća u našim profesionalnim životima. Nedavno Deloitteovo istraživanje izvršnih direktora otkrilo je da velika većina (79%) očekuje da će generativna AI transformisati njihove organizacije u roku od tri godine. Poslovni lideri u svim industrijama prepoznaju da će tehnologija imati dubok strateški uticaj, a vodeće organizacije već reaguju. Prioriteti C-Suite-a odražavaju ovo, a AI i računarstvo u oblaku su rangirane kao dve tehnologije koje će najverovatnije podstaći rast tokom sledeće godine.

Uprkos transformativnom obećanju „pravne tehnologije“ tokom poslednjih godina, rad pravne funkcije ostao je uglavnom nepromenjen. Pravne tehnologije poput upravljanja dokumentima, e-naplate i upravljanja IP-om obično su imale usku primenu, dovodeći samo do inkrementalnog povećanja efikasnosti. Verujemo da je generativna veštačka inteligencija drugačija i da će ova tehnologija pokrenuti istinsku, održivu promenu u načinu pružanja pravnih usluga. Takođe verujemo da će uticaj na pravne timove biti dalekosežan, iz četiri ključna razloga. Prvo, pravna funkcija često zaostaje za drugim funkcijama omogućavanja u digitalnoj zrelosti, što znači da generativna AI može da pruži značajno i merljivo poboljšanje. Drugo, generativna AI prevazilazi mogućnosti tradicionalnih tehnologija, uključujući AI tehnologije, u analizi dokumenata i nestrukturiranih podataka. Ovo čini tehnologiju dobro prilagođenom legalnom radu, sa fokusom na nestrukturisane podatke sa teškim tekstom. Treće, generativna AI svestranost znači da ima primenu u nizu pravnih aktivnosti, otključavajući potencijal u svim oblastima pravnog rada. Četvrto, pored upotrebe veštačke inteligencije od strane advokata, pravno odeljenje ima ključnu ulogu u uvođenju generativne veštačke inteligencije u posao i osigurava da se ona koristi bezbedno, etički i odgovorno. Ovo gledište je snažno podržano nalazima naše ankete. 79% naših ispitanika veruje da generativna AI nije samo reklama, i da će

imati umeren do značajan dugoročni uticaj na to kako se obavlja pravni posao. Dalje, 49% očekuje više od same transformacije, pri čemu su neki zadaci potpuno zastareli zbog Generativne AI.

Osnovne mogućnosti generativne veštačke inteligencije (inteligentna pretraga/analiza, sumiranje i generisanje sadržaja) su veoma efikasne kada se primenjuju na velike količine dokumenata sa teškim tekstom: • Generativna AI može pomoći ugovornim i komercijalnim timovima da pregledaju i sumiraju ugovore, procene ugovorne rizike, uporede ugovore sa priručnicima i predlože kako da izmene ugovore kako bi se uskladili sa priručnikom. • Timovi za spajanje i pripajanje mogu da koriste dinamičke upite na prirodnom jeziku da pretražuju ogromne količine dokumenata na više jezika kako bi identifikovali potencijalne rizike i prednosti. • Timovi za usklađenost sa zakonima i propisima mogu brzo da uporede sadržaj sa novim ili postojećim pravilima i propisima, identifikujući oblasti potencijalne neusklađenosti sa prilagođenim predlozima za sanaciju.

U pravnim operacijama, zadaci kao što su „vađenje podataka, upravljanje i analiza“ imaju visoke rezultate, što smatramo posebno moćnim u omogućavanju efikasnijeg upravljanja znanjem. Isto važi i za pregled dokumenata i ekstrakciju podataka u kontekstu M&A. Primer mogućnosti velikog uticaja koja preseca više zadataka je prevođenje, koje, kada se kombinuje sa drugim generativnim AI mogućnostima (kao što je poređenje ugovora na nekoliko jezika sa regulatornim standardom, ili pretraživanje i analiza presedana u višejezičkom spremištu znanja), donosi korak promene u efikasnosti i kvalitetu. Efikasno prevođenje je moćna prednost za globalne organizacije, omogućavajući bolju saradnju i efikasnost preko jezičkih i kulturnih granica. Koristeći prevodilačke mogućnosti Generativne AI, organizacije neće biti vezane za regrutovanje u određenim jurisdikcijama za pristup jezičkim sposobnostima, što rezultira radnom snagom koja je fleksibilnija i agilnija.

Primenom generativne veštačke inteligencije na ceo pravni ili poslovni proces, koristeći je za poboljšanje svakog pojedinačnog koraka u lancu, prednosti tehnologije se mogu povećati.

Uspostavljanjem pravih zaštitnih mehanizama, GC-i zapravo mogu ubrzati vreme do vrednosti ohrabrujući bezbedno i odgovorno eksperimentisanje i korišćenje GenAI.

Verujemo da generativna veštačka inteligencija može da stvori značajnu i održivu vrednost za korporativno pravo odeljenja, kroz poboljšanje efikasnosti, iskustva i sposobnosti.

Dalje, očekujemo da će generativna AI pojednostaviti i poboljšati interakciju pravnog odeljenja sa poslovanjem. Dok advokati mogu instinktivno da se osećaju neprijatno da koriste generativnu veštačku inteligenciju kao prvu kontakt tačku za poslovna pitanja, realnost je da to predstavlja poboljšanje današnje pozicije. Putem generativnih AI omogućenih chat botova i legalnih ulaznih vrata, poslovni korisnici mogu postavljati pitanja na jednostavnom jeziku i dobijati brze, prilagođene i kontekstualizovane odgovore. Pored toga, ljudski razgovor će poboljšati korisničko iskustvo i podstaći korisnike da angažuju – umesto da zaobiđu – pravno odeljenje. Očekujemo da će ovo poboljšati opšte upravljanje rizikom. Sposobnost personalizacije odgovora na poslovanje i pojednostavljenja korisničkog iskustva podiže generativnu veštačku inteligenciju do efektivne prve tačke kontakta, postavljajući rutinske upite i dodatno oslobađajući vreme advokata.

Pošto je pretnja sajber napadima na KI strateškog obima, nacionalni odgovor mora biti jednak zadatku, a to podrazumeva podizanje svesti javnosti, ulaganje u obrazovanje, naučna istraživanja, razvoj sajber prava i međunarodne saradnje.

U njemu se navodi da nauka i tehnologija igraju važnu ulogu u međunarodnoj bezbednosti i da, iako moderna informaciona i komunikaciona tehnologija (IKT) nudi civilizaciji „najšire pozitivne mogućnosti“, IKT je ipak podložna zloupotrebi od strane kriminalaca i terorista.

Generativna AI takođe može kombinovati interno institucionalno znanje (o kupcu, tržištu ili konkurentu) sa eksternim izvorima podataka, pomažući pravnim odeljenjima da steknu bolje razumevanje određenog pitanja i oblikuju strateški odgovor. Da biste izvukli najveću vrednost koristeći generativnu AI, treba je primeniti na podatke koji su centralizovani, očišćeni i odgovarajuće strukturirani.

Pun potencijal generativne veštačke inteligencije biće ostvaren samo ako je uspešno integrisan u pravne tokove rada. Pravni timovi imaju značajan jaz koji treba zatvoriti ako žele da ispune ove ambicije usvajanja. Ovo uključuje više od dve trećine preduzimanja internih inicijativa u vezi sa obukom i podizanjem svesti, kao i eksternih inicijativa za partnerstvo sa provajderima u slučajevima ciljane upotrebe.

Za ispunjenje zahteva postavljenih u metodološkom delu ove doktorske disertacije urađeno je kvalitativno ispitivanje stavova ispitanika. U anketiranju je učestvovalo ukupno N=

253 ispitanika, od toga $n = 139$ (54,9%) ispitanika muškog pola i $n = 114$ (45,1%) ispitanika ženskog pola. Srednja vrednost zastupljenosti po polu je 1.45 a SD .499.

U istraživanju je učestvovalo najviše ispitanika starosne grupe 26-35 godina $n=77$ (30,4%), a najmanje iz starosne grupe 56-65 godina $n=24$ (9,5%). Srednja vrednost zastupljenosti po starosnoj strukturi je 2,68 dok je SD 1.283.

Istraživanjem bilo obuhvaćeno najviše ispitanika sa višom/visokom školskom spremom $n=103$ (40,7%), a najmanje sa doktoratom $n=18$ (7,1%) i osnovnom školom $n=3$ (1,2%). Srednja vrednost zastupljenosti po stepenu obrazovanja je 3,05 dok je SD .918.

Empirijsko istraživanje baziralo se na utvrđivanju i analizi dobijenih stavova ispitanika u vezi sa različitim aspektima ispitivane pojave.

Za testiranje postavljenih hipoteza utvrđivane su uzročno-posledične veze odabranih varijabli. Cilj istraživanja bio je skoncentrisan na prikaz dobijenih odgovora ispitanika, prikazivanje njihovih stavova i utvrđivanje parametara koji su važni za dokazivanje ili opovrgavanje postavljenih hipoteza.

Deskriptivno longitudinalno istraživanje je dizajn istraživanja korišćenog u ovoj studiji jer je to najbolji izbor za opisivanje karakteristika populacije ili fenomena koji se proučava i ne odgovara na pitanja o tome kako, kada i zašto su se te karakteristike pojavile, već bavi se karakteristikama populacije ili situacije koja se proučava.

Pregled rezultata dobijenih testiranjem postavljenih hipoteza

Za proveru opravdanosti glavne hipoteze: H_0 : Ukoliko je zakonodavstvo uspostavilo odgovarajuću kontrolu i dalo jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju utoliko je manja mogućnost da upotreba AI dovede do štetnih posledica, testirali smo odgovore ispitanika na iznete tvrdnje:

P1. Da li podržavate stav da zakonodavstvo mora da uspostavi odgovarajuću kontrolu i da jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju?

P2. Da li podržavate stav da upotreba AI može dovesti do štetnih posledica?

Rezultat prikazan u tabelama pokazuje da je $X^2(16,1) = 328.768^a$, $p < 0.01$, što ukazuje na postojanje statističke značajnosti dobijenih rezultata. Koeficijent Pirsonove korelacije $r = .542$ potvrđuje da je između ispitivanih varijabli prisutan visok stepen pozitivne korelacije. Vrednost

Spearman Correlation rangova je $Rho=.540$. Dobijen Gamma rezultat .585 usmerava nas na to da stepen saglasnosti sa prvom tvrdnjom omogućava predviđanje ishoda odgovora na drugu tvrdnju za 58,5%.

Na osnovu statističke obrade testiranih varijabli zaključujemo da je generalna hipoteza: H_0 : Ukoliko je zakonodavstvo uspostavilo odgovarajuću kontrolu i dalo jasne smernice u vezi sa upotrebom i razvojem alata koji omogućavaju veštačku inteligenciju utoliko je manja mogućnost da upotreba AI dovede do štetnih posledica, potvrđena.

Testiranje pomoćnih hipoteza

Za proveru opravdanosti prve omoćne hipoteze H_1 koja glasi: Ukoliko je regulisanje veštačke inteligencije uspostavljeno tako da je postignuta ravnoteža između regulisanja kontrole AI i omogućavanja inovacije i evolucije, utoliko je veća mogućnost da se AI koristi na bezbedan i odgovoran način, testirali smo odgovore ispitanika na iznete tvrdnje:

P3. Da li podržavate stav da regulisanje veštačke inteligencije mora da postigne ravnotežu između regulisanja kontrole AI i omogućavanja inovacije i evolucije?

P4. Da li podržavate stav da AI mora da se koristi na bezbedan i odgovoran način?

Rezultat prikazan u tabelama pokazuje da je $X^2(16,1)=355.013^a$, $p<0.01$, što ukazuje na postojanje statističke značajnosti dobijenih rezultata. Koeficijent Pirsonove korelacije $r=.585$ potvrđuje da je između ispitivanih varijabli prisutan visok stepen pozitivne korelacije. Vrednost Spearman Correlation rangova je $Rho=.539$. Dobijen Gamma rezultat .603 usmerava nas na to da stepen saglasnosti sa prvom tvrdnjom omogućava predviđanje ishoda odgovora na drugu tvrdnju za 60,3%.

Na osnovu statističke obrade testiranih varijabli zaključujemo da je prva posebna hipoteza: H_1 : Ukoliko je regulisanje veštačke inteligencije uspostavljeno tako da je postignuta ravnoteža između regulisanja kontrole AI i omogućavanja inovacije i evolucije, utoliko je veća mogućnost da se AI koristi na bezbedan i odgovoran način, potvrđena.

Za proveru opravdanosti druge posebne hipoteze H_2 koja glasi: Ukoliko podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana utoliko je veća mogućnost da se otvore pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI, testirali smo odgovore ispitanika na iznete tvrdnje:

P5. Da li podržavate stav da podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana?

P6. Da li podržavate stav da upotreba AI otvara pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI?

Rezultat prikazan u tabelama pokazuje da je $X^2 (16,1) = 132.950^a$, $p < 0.01$, što ukazuje na postojanje statističke značajnosti dobijenih rezultata. Koeficijent Pirsonove korelacije $r = .491$ potvrđuje da je između ispitivanih varijabli prisutan srednje visok stepen pozitivne korelacije. Vrednost Spearman Correlation rangova je $Rho = .445$. Dobijen Gamma rezultat $.571$ usmerava nas na to da stepen saglasnosti sa prvom tvrdnjom omogućava predviđanje ishoda odgovora na drugu tvrdnju za 57,1%.

Na osnovu statističke obrade testiranih varijabli zaključujemo da je druga posebna hipoteza: H2: Ukoliko podaci koje obrađuje, generiše ili koristi AI uključuju preklapanje interesa različitih strana utoliko je veća mogućnost da se otvore pitanja privatnosti i sajber bezbednosti koja su pokrenuta razvojem, upotrebom i primenom AI, potvrđena.

Za proveru opravdanosti treće posebne hipoteze H3 koja glasi: Ukoliko je zakonodavni pristup definisan oslanjanjem na procenu stepena rizika od upotrebe AI utoliko je veća mogućnost da se AI sistemi visokog rizika suoče sa strogo definisanim načinima delovanja testirali smo odgovore ispitanika na iznete tvrdnje:

P7. Da li podržavate stav da upotreba AI sistema mora da se klasifikuje u odnosu na stepen rizika do koga dolazi prilikom primene AI?

P8. Da li podržavate stav da AI sistemi visokog rizika moraju da budu obuhvaćeni strogo definisanim načinima delovanja?

Rezultat prikazan u tabelama pokazuje da je $X^2 (16,1) = 222.183^a$, $p < 0.01$, što ukazuje na postojanje statističke značajnosti dobijenih rezultata. Koeficijent Pirsonove korelacije $r = .591$ potvrđuje da je između ispitivanih varijabli prisutan visok stepen pozitivne korelacije. Vrednost Spearman Correlation rangova je $Rho = .587$. Dobijen Gamma rezultat $.665$ usmerava nas na to da stepen saglasnosti sa prvom tvrdnjom omogućava predviđanje ishoda odgovora na drugu tvrdnju za 66,5%.

Na osnovu statističke obrade testiranih varijabli zaključujemo da je treća posebna hipoteza: H3: Ukoliko je zakonodavni pristup definisan oslanjanjem na procenu stepena rizika od

upotrebe AI utoliko je veća mogućnost da se AI sistemi visokog rizika suoče sa strogo definisanim načinima delovanja, potvrđena.

LITERATURA

- 3Blue1Brown (2018). Neural Networks & Deep Learning. 3Blue1Brown.
- Acemoglu, D., and Restrepo, P. (2018). "Robots and Jobs: Evidence from US Labor Markets." NBER Working Paper, no. 24190, 2018.
- Adam Thierer (2015). "Embracing a Culture of Permissionless Innovation," Cato Institute, November 17, 2014. <https://www.cato.org/cato-online-forum/embracing-culture-permissionless-innovation>.
- AlgorithmWatch and Access Now (2021), "EU Policy Makers: Protect People's Rights, Don't Narrow Down the Scope of the AI Act!" November 23, 2021. <https://algorithmwatch.org/en/statement-scope-of-eu-ai-act/>
- Amit Kaushal, Russ Altman, and Curt Langlotz (2020). "Health Care AI Systems Are Biased," *Scientific American*,
- ANI (2020). "AI Should Never Be Allowed to Substitute Human Discretion in Judicial Functioning: CJI," ETCIO, 27 January 2020, <https://cio.economic-times.indiatimes.com/news/government-policy/ai-should-never-be-allowed-to-substitute-human-discretion-in-judicial-functioning-cji/73651703>
- Anu Bradford, *The Brussels Effect* (2020). How the European Union Rules the World (New York: Oxford University Press, 2020). November 17, 2020. Available at: <https://www.scientificamerican.com/article/health-care-ai-systems-are-biased/>.
- Artificial intelligence (2023). MERRIAM-WEBSTER, <https://www.merriam-webster.com/dictionary/artificial%20intelligence> [<https://perma.cc/3ZBP-PLAJ>]
- Audrey Herrington (2019). *How AI Can Make Legal Services More Affordable*, SMARTLAWYER (Jul. 23, 2019), <http://www.nationaljurist.com/smartlawyer/how-ai-can-make-legal-services-more-affordable> [<https://perma.cc/TJ97-CA4H>].
- Babu Jackson (2021). "Prevalence of Biometric Data and Security Concerns," Bipartisan Policy Center, September 20, 2021. <https://bipartisanpolicy.org/blog/prevalence-of-biometric-data-and-security-concerns>.

- Becker, M., & Buchkremer, R. (2018). Ranking of current information technologies by risk and regulatory compliance officers at financial institutions – a German perspective. *The Review of Finance and Banking*, 10(1) <https://search.proquest.com/docview/2208135738?accountid=27513>
- Big Data Value Association (2021). *BDVA Position Paper: Response to the European Commission's proposal for AI Regulation*, August 4, 2021. Available at: https://www.bdva.eu/sites/default/files/BDVA_DAIRO%20response-feedback%20AI%20Regulation_Final.pdf.
- Brenner Susan (2004). Toward a Criminal Law for Cyberspace: Distributed Security, *BOSTON UNIVERSITY JOURNAL OF SCIENCE & TECHNOLOGY LAW*. Vol. 10, Issue 1, pp. 1-109
- Chopra, K. N. (2018). Overview and application of artificial intelligence concepts and some important coevolving modern issues on management of organization and commerce. *Journal of Internet Banking and Commerce*, 23(3), 1-22. <https://search.proquest.com/docview/2216881012?accountid=27513>
- Chowdhury EK (eds.) (2021) *The essentials of machine learning in finance and accounting: prospects and challenges of using artificial intelligence in the audit process*. Taylor and Francis Inc, London.
- Centre for Data Ethics and Innovation, United Kingdom Government (2021). "The European Commission's Artificial Intelligence Act Highlights the Need for an Effective AI Assurance Ecosystem," May 11, 2021. <https://cdei.blog.gov.uk/2021/05/11/the-european-commissions-artificial-intelligence-act-highlights-the-need-for-an-effective-ai-assurance-ecosystem>.
- CISA (2003). *The National Strategy to Secure Cyberspace*. <https://www.cisa.gov/national-strategy-secure-cyberspace>.
- Daub, M., & Wiesinger, A. (2015). *Acquiring the capabilities, you need to go digital*. <https://www.mckinsey.com/business-functions/digital-mckinsey/our-insights/acquiring-the-capabilities-you-need-to-go-digital>

- Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard business review*, 96(1), 108-116.
- Dogru, A. K., & Keskin, B. B. (2020). AI in operations management: applications, challenges and opportunities. *Journal of Data, Information and Management*, 2(2), 67-74.
- Ergen M (2019) What is artificial intelligence? Technical considerations and future perception. *Anatol J Cardiol* 22(2):5–7. <https://doi.org/10.14744/AnatolJCardiol.2019.79091>
- European Parliamentary Research Service (2015). *The Precautionary Principle: Definitions, Applications, and Governance*, December 2015. [https://www.europarl.europa.eu/thinktank/en/document/EPRS_IDA\(2015\)573876](https://www.europarl.europa.eu/thinktank/en/document/EPRS_IDA(2015)573876).
- European Commission (2020). *On Artificial Intelligence – A European approach to excellence and trust*, February 2020. https://ec.europa.eu/info/sites/default/files/commission-white-paper-artificial-intelligence-feb2020_en.pdf.
- European Commission (2021). *Annexes to the Proposal for Regulation of the European Parliament and Council*, April 21, 2021. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52021PC0206>.
- European Commission (2022). “Excellence and trust in artificial intelligence,” accessed July 14, 2022. https://ec.europa.eu/info/strategy/priorities-2019-2024/europe-fit-digital-age/excellence-trust-artificial-intelligence_en.
- Fan, Q. (2016). Factors Affecting Adoption of Digital Business: Evidence from Australia. *Global Journal of Business Research*, 10(3), 79-84.
- Frey, C. B., and Osborne, M. A. (2017). "The Future of Employment: How Susceptible are Jobs to Computerisation?" *Technological Forecasting and Social Change*, vol. 114, 2017, pp. 254-280.
- Gambatese, J. A., & Hallowell, M. (2011). Enabling and measuring innovation in the construction industry. *Construction Management and Economics*, 29(6), 553-567.
- Ghanoum S, Alaba FM (2020). Integration of artificial intelligence in auditing: the effect on auditing. <https://www.diva-portal.org>.

- Goodfellow, I., Bengio, Y., & Courville, A. (2017). Deep learning. Cambridge, MA: MIT Press.
- Hastie, T., et al. (2009). The elements of statistical learning: data mining, inference, and prediction, Springer Science & Business Media.
- Hathaway, I., Maxim, R., Muro, M., & Whiton, J. (2019). Automation and artificial intelligence: How machines are affecting people and places. *Metropolitan Policy Program at Brookings*.
- Hermann, J. M. Del Balso (2017). “Meet Michelangelo: Uber’s Machine Learning Platform,” Uber Engineering, 5 September 2017, <https://eng.uber.com/michelangelo-machine-learning-platform/>
- Hitachi (2017). “Take on This Unpredictable Business Age Together With Hitachi AI Technology/H,” 2017, https://social-innovation.hitachi/-/media/project/hitachi/sib/en/solutions/ai/pdf/ai_en_170310.pdf
- Holzinger, A. (2018). From machine learning to explainable AI. 2018 world symposium on digital intelligence for systems and machines (DISA), IEEE.
- Information Commissioner’s Office (ICO) (2019). “Fully Automated Decision Making AI Systems: The Right to Human Intervention and Other Safeguards,” United Kingdom, 5 August 2019, <https://ico.org.uk/about-the-ico/news-and-events/ai-blog-fully-automated-decision-making-ai-systems-the-right-to-human-intervention-and-other-safeguards/>
- Institute of Electrical and Electronics Engineers (2019). IEEE Position Statement of Artificial Intelligence, June 24, 2019.: <https://globalpolicy.ieee.org/wp-content/uploads/2019/06/IEEE18029.pdf>.
- Ioannou, L. (2014). A decade to mass extinction event in S&P 500. <https://www.cnbc.com/2014/06/04/15-years-to-extinction-sp-500-companies.html>
- Jaime Sevilla, Lennart Heim, et al., (2022). “Compute Trends Across Three Eras of Machine Learning,” arXiv:2202.05924, March 9, 2022. <http://arxiv.org/abs/2202.05924>.
- Jake Heller (2018). *Is AI the Great Equalizer For Small Law?*, ABOVE THE LAW (Aug. 15, 2018), <https://abovethelaw.com/2018/08/is-a-i-the-great-equalizer-for-small-law/> [<https://perma.cc/7UPA-CK4W>].

- Jordan, M. I., & Mitchell, T. M. (2015). Machine learning: Trends, perspectives, and prospects. *Science*, 349(6245), 255-260.
- Julie Fishbach (2016). *Coder, 19, Builds Chatbot That Fights Parking Tickets*, NBC NEWS (Jul. 21, 2016), <https://www.nbcnews.com/feature/college-game-plan/coder-19-builds-chatbot-fights-parking-tickets-n612326>[<https://perma.cc/UEN7-DBAM>].
- Kanade, V. (2022). What is Artificial Intelligence (AI)? definition, types, goals, challenges, and trends in 2022, Spiceworks. https://www.spiceworks.com/tech/artificial-intelligence/articles/what-is-ai/#_002
- Katherine Medianik (2018). Artificially Intelligent Lawyers: Updating the Model Rules of Professional Conduct in Accordance with the New Technological Era, 39 CARDOZO L. REV. 1497, 1515 (2018) (citing Fla. Bar Prof'l Ethics Comm., Op. 06-2 (2006)).
- Kelly Brian (2012). Investing In a Centralized Cybersecurity Infrastructure: Why "Hacktivism" Can And Should Influence Cybersecurity Reform BOSTON UNIVERSITY LAW REVIEW, Vol. 92. 1671-1673
- Keith Mullen (2018). *Artificial Intelligence: Shiny Object? Speeding Train?*, AM. BAR ASS'N RPTE EREPORT 2018 FALL ISSUE, https://www.americanbar.org/groups/real_property_trust_estate/publications/ereport/rpte-ereport-fall-2018/artificial-intelligence/ [<https://perma.cc/HCE7-WNWV>]
- Kolbjørnsrud, V., Amico, R., Thomas, R. J. (2015). The promise of artificial intelligence: Redefining management in the workforce of the future. Accenture.
- Kozlowski S (2018). An audit ecosystem to support blockchain-based accounting and assurance. Emerald Publishing Limited, Bingley. <https://doi.org/10.1108/978-1-78743-413-420181015>
- Kremer Jan-Frederik & Müller Benedikt (2014). *Cyberspace and International Relations: Theory, Prospects and Challenges*. Springer.
- Krishna, R. (2017). *Computer Vision: Foundations and Applications*. Stanford University.

- Kwon, E. H., & Park, M. J. (2017). Critical Factors on Firm's Digital Transformation Capacity: Empirical Evidence from Korea. *International Journal of Applied Engineering Research*, 12(22), 12585-12596.
- Lat, *supra* note 1 (2018). Tashea & Economou, *supra* note 52; Roy D. Simon, *Artificial Intelligence, Real Ethics*, 90-APR N.Y. ST. B.J. 34, 34.
- Lauri Donahue, *A Primer on Using Artificial Intelligence in the Legal Profession*, HARV. J. L. & TECH. (Jan. 3, 2018), <http://jolt.law.harvard.edu/digest/a-primer-on-using-artificial-intelligence-in-the-legal-profession> [<https://perma.cc/H65H-6A5A>]
- Lawton, G. & Tucci, L. (2022). What is business process automation, CIO? TechTarget. <https://www.techtarget.com/searchcio/definition/business-process-automation>
- Leadershiptribe (2023). Artificial Intelligence in Business Management: A Revolution in the Making. <https://leadershiptribe.com/blog/artificial-intelligence-in-business-management-a-revolution-in-the-making>
- Lee, D. (2016). "Tay: Microsoft Issues Apology Over Racist Chatbot Fiasco," BBC, 25 March 2016, <https://www.bbc.com/news/technology-35902104>
- Lynn, W.J. (2010) "Defending a New Domain: The Pentagon's Cyberstrategy," *Foreign Affairs* 89(5) 97-108.
- MacAraig, M. (2019). 11 pros and cons of AI for Businesses, <https://www.insightsforprofessionals.com>.
- Madhugiri, D. (2023). Advantages and disadvantages of Artificial Intelligence (AI), KnowledgeHut. Knowledgehut.
- <https://www.knowledgehut.com/blog/data-science/advantages-and-disadvantages-of-artificial-intelligence>
- Matthew Griffin (2016). *Meet Ross, The World's First AI Lawyer*, 311 INST. (Jul. 11, 2016), <https://www.311institute.com/meet-ross-the-worlds-first-ai-lawyer/> [<https://perma.cc/89RH-C9M9>].
- McKinsey Global Institute (2018). "Jobs Lost, Jobs Gained: What the Future of Work Will Mean for Jobs, Skills, and Wages." 2018.

- Medianik (2004). *supra* note 55 (quoting N.Y. State Bar Ass’n Comm. on Prof’l Ethics, Formal Op. 782 (2004)).
- Merriam-Webster. (n.d.). *Automation*. Merriam-Webster: <https://www.merriam-webster.com/dictionary/automation>
- MODEL RULES OF PROF’L CONDUCT: PREFACE, ETHICS 2000 CHAIR’S INTRODUCTION (2002). https://www.americanbar.org/groups/professional_responsibility/publications/model_ruler_of_professional_conduct/model_rules_of_professional_conduct_preface/ethics_2000_cchai_introduction/ [https://perma.cc/7F3K-N4XG].
- Mullen, *supra* note 7 (“As recent as 2014, only a handful of companies pointed artificial intelligence (“AI”) at legal documents. For uses outside of eDiscovery, it was a narrow focus: contract reviews and legal due diligence – used by the largest of law firms on the grandest of deals.”)
- National Institute of Science and Technology (2022). *AI Risk Management Framework*, 2022. <https://www.nist.gov/itl/ai-risk-management-framework>
- Neumann, P.G. (1998). “Protecting the Infrastructures,” *Communications of the ACM* 41(1) 128. Ncsss (2016). "National Cyber Security Strategies (Ncsss) Map — ENISA". *Enisa.europa.eu*. N.p.
- Nick Gondek (2021). “Learning about Machine Learning,” Bipartisan Policy Center, August 4, 2021. <https://bipartisanpolicy.org/explainer/learning-about-machine-learning>.
- OECD AI Policy Observatory (2022). “OECD AI Principles overview,” accessed July 14, 2022. <https://oecd.ai/en/ai-principles>.
- Panda B, Leepsa NM (2017) Agency theory: review of theory and evidence on problems and perspectives. *Indian J Corp Gov* 10(1):74–95. <https://doi.org/10.1177/0974686217701467>
- PricewaterhouseCoopers (PwC) (2017). “PwC Releases Report on Global Impact and Adoption of AI,” 25 April 2017, <https://www.pwc.com/us/en/press-releases/2017/report-on-global-impact-and-adoption-of-ai.html>

- Sajja, P. S. (2021). Introduction to Artificial Intelligence. Illustrated Computational Intelligence, Springer: 1-25.
- Sandage, John et al. eds. (2013). Comprehensive Study on Cybercrime (United Nations Office on Drugs and Crime).
- Sargent, K., Hyland, P., & Sawang, S. (2012). Factors Influencing the Adoption of Information Technology in a Construction Business. Australasian Journal of Construction Economics and Building, 12(2), 72-86.
- Satola David & Judy Henry (2011). Towards a Dynamic Approach to Enhancing International Cooperation and Collaboration in Cybersecurity Legal Frameworks: Reflections on the Proceedings of the Workshop on Cybersecurity Legal Issues at the 2010 United Nations Internet Governance Forum 37 WILLIAM MITCHELL LAW REVIEW.
- Schmitt Michael (2013). Tallinn Manual on the International Law Applicable to Cyber Warfare. Cambridge University Press.
- Shaw, J., et al. (2019). "Artificial intelligence and the implementation challenge." Journal of medical Internet research 21(7): e13659.
- Shim, J. K., & Rice, J. S. (1988). Expert systems applications to managerial accounting. Journal of Systems Management, 39(6), 6.
<https://search.proquest.com/docview/199850554?accountid=27513>
- Simplilearn (2023). What Is Computer Vision: Applications, Benefits and How to Learn It. <https://www.simplilearn.com/computer-vision-article>
- Spremić, M., Ph.D. (2017). Enterprise information systems in digital economy. Zagreb.
- Stahl, William M. (2011). "The Uncharted Waters of Cyberspace: Applying the Principles of International Maritime Law to the Problem of Cybersecurity." Georgia Journal of International and Comparative Law. Vol. 40. (2011): 247-274.
- Stanford University Human-Centered Artificial Intelligence (na), "The AI Index Report – Artificial Intelligence Index," Available at: <https://aiindex.stanford.edu/report/>.
- Stevenson, W. J. (2020). *Operations Management*. McGraw Hill.

- Takyar, A. (2023). AI in business management: Unveiling the next wave of operational excellence. <https://www.leewayhertz.com/ai-in-business-management/>
- The Council of Europe Convention on Cybercrime (na). <http://conventions.coe.int/>.
- The Economist Intelligence Unit (2018). “Intelligent Economies: AI’s transformation of industries and society,”. <https://impact.economist.com/perspectives/technology-innovation/intelligent-economies-ais-transformation-industries-and-society/>
- Thomas, K. D., Boring, R. L., Hugo, J. V., & Hallbert, B. P. (2017). Human factors for main control room modernization. *Nuclear News*, pp. 46-50.
- Urbas Gregor (2012). Cybercrime, Jurisdiction and Extradition: The Extended Reach of Cross-Border Law Enforcement, *JOURNAL OF INTERNET LAW*, Vol. 16, Issue 1: 7-17.
- Ursula von der Leyen (na). *Political Guidelines for the Next European Commission 2019-2024*. https://ec.europa.eu/info/sites/default/files/political-guidelines-next-commission_en_0.pdf.
- U.S. Congress (2019). *Algorithmic Accountability Act of 2019*, H.R.223, 116th Congress, introduced in House April 11, 2019. <https://www.congress.gov/bill/116th-congress/house-bill/2231/text>.
- U.S. Department of Commerce (2022). “U.S.-EU Joint Statement of the Trade and Technology Council,” May 16, 2022. <https://www.commerce.gov/news/press-releases/2022/05/us-eu-joint-statement-trade-and-technology-council>.
- Uzialko, A. C. (2017). Good Leadership Is the Key to Adopting New Technology. <https://www.businessnewsdaily.com/10067-business-technology-innovation-human-user.htm>
- Virginia Tech admin. (2021). The impact of artificial intelligence on business: VT India, Virginia Tech India. https://vtindia.in/the-impact-of-artificial-intelligence-on-business/#How_does_AI_Impact_your_Business
- WALL DAVID (2007). *CYBERCRIME: THE TRANSFORMATION OF CRIME IN THE INFORMATION AGE*. Polity Press.

- Whitney, R. (2023). 4 popular uses of artificial intelligence in business. How to maximize the potential of AI in your business? <https://firmbee.com/artificial-intelligence-in-business>
- WHO (2013). <https://www.who.int/news/item/19-10-2023-who-outlines-considerations-for-regulation-of-artificial-intelligence-for-health>
- Weill, P., & Ross, J. W. (2004). IT governance: How top performers manage IT decision rights for superior results. Boston: Harvard Business Press.
- Wood, T. (2024). Business uses of natural language processing - the 2024 capabilities of NLP. <https://fastdatascience.com/business-uses-natural-language-processing/>
- World Commission on the Ethics of Scientific Knowledge and Technology (2005). UNESCO, *The Precautionary Principle*, 2005 <https://unesdoc.unesco.org/ark:/48223/pf0000139578>.
- Zhong, B. X. Pan, P.E. Love, J. Sun, C. Tao, (2020). Hazard analysis: a deep learning and text mining framework for accident prevention, *Adv. Eng. Inform.* 46, 101152, <https://doi.org/10.1016/j.aei.2020.101152>.