

УНИВЕРЗИТЕТ У БЕОГРАДУ

Факултет организационих наука

НАСТАВНО-НАУЧНОМ ВЕЋУ

Предмет: Подобност теме и кандидата Небојше Аврамовића за израду докторске дисертације

Одлуком Наставно-научног већа Факултета организационих наука - Универзитета у Београду 05-01 бр. 3/101-6 од 02.10.2025. године именовани смо за чланове Комисије за преглед и одбрану приступног рада и оцену научне заснованости теме докторске дисертације кандидата **Небојше Аврамовића** под насловом:

„Примена вештачке интелигенције у алатима за пословно одлучивање и процену ризика у реалном времену у дигиталном пословању”

На основу материјала приложеног уз Захтев кандидата, Комисија подноси следећи

РЕФЕРАТ

1. Подаци о кандидату

1.1. Биографски подаци

Небојша Аврамовић је рођен у Београду, Земун, 8.4.1981. године. Отац је двоје деце и живи и ради у Земуну. Завршио је основну школу Светозар Милетић у Земуну, а потом и средњу електротехничку школу Никола Тесла у Београду. Дипломирао је на Војнотехничкој академији (ВТА) у Београду 2005. године, са просечном оценом током студија 8,07 и дипломским радом на тему „Перспективе развоја мобилне телефоније“.

Докторске академске студије уписао је на Факултету организационих наука 2020. године, смер Информациони системи и квантитативни менаџмент.

Небојша Аврамовић је власник, оснивач и генерални директор компаније Consulting IT DNA у Београду која се успешно бави консултантским услугама у области пословних и дигиталних трансформација коришћењем најсавременијих алата и технологија као што су, између осталог, вештачка интелигенција и машинско учење. Компанија успешно послује са више од 10 запослених.

Пре оснивања своје компаније, Небојша је био запослен у компанији Oracle од 2017. године, на позицији Директора за консултантске услуге за југоисточну Европу и водио тим од 50 консултаната. Од доласка на ту позицију, Небојша је успео да оствари раст од више од 250% промета и значајно прошири тим и саме капацитете тима у смислу знања, вештина и потребне експертизе начинивши тај тим најзначајнијим партнером клијентима из финансијског сектора на тржишту са фокусом на пројекте које се тичу трансформације и оптимизације пословања. Пре ове позиције, Небојша је био директор продајног тима где је такође остварио раст од 350% реализујући, између осталог, неколико значајних трансформационих пројеката са кључним клијентима.

Пре Oracle-а, Небојша је радио као Директор тима за развој пословања за територију југоисточне Европе, у компанији Ingram Micro. У свом тиму, Небојша је имао преко 20 запослених и био је одговоран за пословање неколико реномираних произвођача технологија. Небојша је у компанији Ingram Micro радио од 2015. до 2017. године.

IBM је био Небојшин послодавац од 2008. године до 2015. године на позицијама регионалног Директора продаје за безбедносна решења специфично везана за финансијски сектор за територију централне и источне Европе, Директора одељења за хардверске производе за Србију, Црну Гору, Северну Македонију и Албанију, Представника продаје за кључне кориснике, Представника продаје за финансијски сектор и Представника продаје за IBM технолошке сервисе. Небојша је био један од најистакнутијих чланова тима и редовни годишњи добитник награде "100% club" за најбоље представнике продаје у целој компанији.

Небојша је, пре IBM-а, радио у компанији МДС Информатички Инжењеринг на позицији Представника продаје, од 2006. до 2008. године.

Професионалну каријеру је започео у Војсци, радећи као официр на позицији Начелника центра за телекомуникације, водећи одељење са 10 запослених. На овој позицији је радио од 2005. до 2006. године.

1.2. Стечено научноистраживачко искуство

У наставку су приказане научноистраживачке активности кандидата. Оне обухватају позивна предавања и радове објављене на конференцијама у земљи и иностранству, од године уписа докторских студија, активности у извођењу наставе и преглед положених испита.

2025

- Infotech Konferencije i Izložbe, Tema: „AI i Cloud tehnologije u savremenom poslovanju“, Predavanje po pozivu
- Infotech Konferencije i Izložbe, Tema: „Tehnološke inovacije i IT bezbednost kao temelji buduće IT infrastrukture“, Učesnik na panelu

2024

- Infotech Konferencije i Izložbe, Tema: „Oracle Database 23ai AI“, Uvodno predavanje
- Infotech Konferencije i Izložbe, Tema: „Dalji koraci u digitalizaciji poslovnih procesa primenom veštačke inteligencije i bezbednosnih rešenja“, Učesnik na panelu

2023

- Infotech Konferencije i Izložbe, Tema: „Cloud – lokalno dostupan globalni resurs“, Predavanje po pozivu
- Zornić, N., Marković, A., & Avramović, N. (2023). A Bibliometric Review of Agent-Based Models Application in Economics. In *Central European Conference on Information and Intelligent Systems* (pp. 447-453). Faculty of Organization and Informatics, Varazdin.
- Zornić, N., Marković, A., & Avramović, N. (2023). *Agent-based models in economics: A review of applications* [Conference presentation]. The Fourteenth International Conference: Challenges of Europe: Design for the Next Generation, Bol, Island Brač, Croatia.

Следи списак положених предмета на докторским студијама са оценама и ЕСПБ бодовима:

Називи предмета	Оцена	ЕСПБ
Електронско пословање - одабрана поглавља	10 (десет)	10
Системи за управљање пословним процесима	10 (десет)	10
Квантитативни модели и методе у менаџменту	10 (десет)	10
Менаџмент електронског пословања - одабрана поглавља	10 (десет)	10
Пословна интелигенција - одабрана поглавља	10 (десет)	10
Статистика у менаџменту	10 (десет)	10
Наука о менаџменту	10 (десет)	10
Одлучивање - одабрана поглавља	10 (десет)	10
Методологија научно-истраживачког рада	10 (десет)	10

Кандидат је дана 07.11.2025. године успешно одбранио приступни рад за израду докторске дисертације под називом „Примена вештачке интелигенције у алатима за пословно одлучивање и процену ризика у реалном времену у дигиталном пословању”.

1.3. Оцена подобности кандидата за рад на предложеној теми

Узимајући у обзир:

- резултате остварене током досадашњег образовања;
- искуство у области информационих технологија и пословања кроз предавања и учешћа на панел дискусијама на значајним конференцијама које су обухватале теме као што су: вештачка интелигенција, Cloud технологије, IT безбедност, дигитализација пословних процеса, Oracle технологије и еУправа.
- научно-истраживачки допринос потврђен публикованим и презентованим радовима у области агентско-базираних модела у економији, што указује на способност академског истраживања и публиковања.

Закључује се да је Небојша Аврамовић у потпуности квалификован и припремљен да тему докторске дисертације самостално истражује и пружи научно-стручне доприносе у областима које покривају његове досадашње активности и објављени радови.

2. Предмет и циљ истраживања

2.1. Предмет истраживања

Дигитална трансформација је померила тежиште одлучивања са спорих, периодичних анализа на непрекидно, подацима вођено одлучивање у реалном времену. У таквом окружењу вештачка интелигенција (ВИ) више није помоћни алат, већ кључни механизам за предвиђање, прилагођавање и контролу ризика. Поред техничких користи (тачност, латенција, скалабилност), успех зависи од поверења врховног менаџмента (ТМТ), објашњивости и транспарентности модела и организационе зрелости (политике, процедуре, људски надзор). Овај рад полази од интегрисаног приступа: технолошке перформансе, објашњивост и управљање ВИ (*AI governance*) заједнички одређују намеру усвајања (спремност организације — ТМТ и кључних корисника — да уведе, финансира и користи ВИ решења у пракси) и квалитет пословних одлука.

Предмет истраживања докторске дисертације је улога ВИ у системима за пословно одлучивање и управљање ризиком у реалном времену (одлучивање у реалном времену – *Real Time Decision* – *RTD*), са фокусом на то како објашњивост, поверење, перцепција ризика, подршка врховног менаџмента и зрелост управљања ВИ обликују намеру усвајања и исходе одлука.

Одлучивање у реалном времену (енг. *Real Time Decision* - *RTD*) означава аутоматизовано или хибридно доношење одлука са ниском латенцијом, засновано на токовима података и ажурним моделима/правилима, тако да је исход одлуке доступан у року прикладном за оперативну акцију (типично од милисекунде до неколико секунди). *RTD* се разликује од периодичног, „ванмрежног“ („офлајн“) одлучивања по томе што користи континуиране улазе и механизме надзора (праћење перформанси, ревизорски траг, људска интервенција) ради безбедног и транспарентног извршавања.

Проблем је двострук: (1) организације неретко уводе *RTD* алате без адекватне ВИ подршке или са недовољно објашњивим моделима, што доводи до слабијих перформанси и погрешних менаџерских увида; (2) чак и када су техничке перформансе задовољене, недостатак поверења, нејасна објашњења и недовољно развијен оквир управљања ВИ (политике, процедуре, одговорности и надзор) спутавају усвајање ВИ у *RTD* системима и њихову делотворну употребу.

Илустративан пример (финансијски сектор): при стотинама хиљада до милиона дневних логовања на е-канале (е-банкинг и м-банкинг), *RTD* системи генеришу понуде финансијских производа према вероватноћи прихватања. Без адекватно уклопљене ВИ, системи не прате брзину промена, дају погрешне препоруке, стварају незадовољство клијената и пружају неразумљиве извештаје менаџменту. Супротно томе, објашњива ВИ (енг. *Explainable AI* - *XAI*), уз јасно дефинисане политике и људски надзор, обезбеђује одлуке које су тачније, непристрасније и подложне контроли и ревизији — што је нарочито важно у секторима са строгим регулативама (нпр. финансије, телекомуникације, здравство и сл.).

Актуелност истраживања проистиче из потребе да се формализује интегрисани модел који повезује техничке и организационе факторе (*XAI*, поверење, ризик, подршка ТМТ, управљање ВИ) и да се емпиријски провери како и под којим условима ВИ доприноси квалитетнијем одлучивању, нарочито у сценаријима реалног времена.

2.2. Циљеви истраживања

Општи циљ истраживања будуће докторске дисертације јесте да се, након критичког сагледавања ограничења традиционалних система за анализу ризика и подршку одлучивању, предложи унапређени (хибридни) оквир који омогућава делотворну, објашњиву и одговорну примену вештачке интелигенције у системима одлучивања у реалном времену у оквиру дигиталног пословања.

Овако дефинисани општи циљ подразумева разраду следећих специфичних циљева:

- **Развој интегрисаног модела** који повезује објашњивост и транспарентност ВИ, поверење врховног менаџмента, перцепцију ризика, управљање ВИ и намеру усвајања, уз јасно дефинисане медијације, модерације и нелинеарности.
- **Операционализација конструката и развој инструмента** (скеале за објашњивост, поверење, перципирани ризик, услове који олакшавају, перформансе одлука), уз проверу поузданости и конструктне валидности.
- **Емпиријска провера хипотеза** кроз валидацију мерног модела и тестирање структурних односа, уз испитивање посредничких, модерационих и евентуално нелинеарних ефеката.
- **Процена хетерогености** (разлике по индустријама и нивоу дигиталне зрелости) и **робусности налаза** (алтернативне спецификације модела).
- **Формулисање професионалних смерница** (политике, процедуре, људски надзор, ревизорски траг) за одговорну имплементацију ВИ у RTD системима.

Остварење ових циљева омогућиће да се теоријски увиди преведу у оперативне препоруке за организације које уводе ВИ у одлучивање у реалном времену.

Почетна листа библиографских извора који ће се користити приликом израде докторске дисертације:

1. Adlung, L., Cohen, Y., Mor, U., & Elinav, E. (2021). Machine learning in clinical decision making. *Med*, 2(6), 642–665.
2. Ali, I., Warraich, N. F., & Butt, K. (2025). Acceptance and use of artificial intelligence and AI-based applications in education: A meta-analysis and future direction. *Information Development*, 41(3), 859–874.
3. Alshammari, M. M., & Al-Mamary, Y. H. (2025). User acceptance of AI-powered training: extending the technology acceptance model (TAM). *Future Business Journal*, 11(1), 1–16.
4. Arena, S., Florian, E., Zennaro, I., Orrù, P. F., & Sgarbossa, F. (2022). A novel decision support system for managing predictive maintenance strategies based on machine learning approaches. *Safety science*, 146, 105529.
5. Azam, A., Boari, C., & Bertolotti, F. (2018). Top management team international experience and strategic decision-making. *Multinational Business Review*, 26(1), 50–70.
6. Aziz, S., & Dowling, M. (2019). Machine learning and AI for risk management. *Disrupting finance*, 33–50.
7. Balasubramanian, N., Ye, Y., & Xu, M. (2022). Substituting human decision-making with machine learning: Implications for organizational learning. *Academy of Management Review*, 47(3), 448–465.

8. Bandaragoda, T., Adikari, A., Nawaratne, R., Nallaperuma, D., Luhach, A. K., Kempitiya, T., ... & Chilamkurti, N. (2020). Artificial intelligence-based commuter behaviour profiling framework using Internet of things for real-time decision-making. *Neural Computing and Applications*, 32(20), 16057–16071.
9. Banh, L., & Strobel, G. (2023). Generative artificial intelligence. *Electronic Markets*, 33(1), 63.
10. Baquero, J. A., Burkhardt, R., Govindarajan, A., & Wallace, T. (2020). Derisking AI by design: How to build risk management into AI development. McKinsey & Company.
11. Baroni, I., Calegari, G. R., Scandolari, D., & Celino, I. (2022). AI-TAM: a model to investigate user acceptance and collaborative intention in human-in-the-loop AI applications. *Human Computation*, 9(1), 1–21.
12. Baryannis, G., Validi, S., Dani, S., & Antoniou, G. (2019). Supply chain risk management and artificial intelligence: state of the art and future research directions. *International journal of production research*, 57(7), 2179–2202.
13. Bastani, H., Bastani, O., & Sinchaisri, P. (2022). Improving human decision-making with machine learning. In *Academy of Management Proceedings* (Vol. 2022, No. 1, p. 17725). Briarcliff Manor, NY 10510: Academy of Management.
14. Bek-Pedersen, E., Lind, M., & Asheim, B. A. (2019, November). AI based real-time decision making. In *Abu Dhabi International Petroleum Exhibition and Conference* (p. D022S152R001). SPE.
15. Bevilacqua, S., Ferraris, A., Kozel, R., & Vincurova, Z. (2025). Exploring the artificial intelligence integration in top management team decision-making: an empirical analysis. *Business Process Management Journal*.
16. Bevilacqua, S., Masárová, J., Perotti, F. A., & Ferraris, A. (2025). Enhancing top managers' leadership with artificial intelligence: insights from a systematic literature review. *Review of Managerial Science*, 1–37.
17. Bhui, R., Lai, L., & Gershman, S. J. (2021). Resource-rational decision making. *Current Opinion in Behavioral Sciences*, 41, 15–21.
18. Brynjolfsson, E., & McAfee, A. (2017). The business of artificial intelligence. *Harvard Business Review*. <https://hbr.org/2017/07/the-business-of-artificial-intelligence>
19. ValidMind. (2025). AI in model risk management for financial services. Preuzeto sa: <https://validmind.com/blog/ai-in-model-risk-management-financial-services/>
20. Vasarhelyi, M. A., Kogan, A., & Tuttle, B. M. (2015). Big data in accounting: An overview. *Accounting Horizons*, 29(2), 381–396. <https://doi.org/10.2308/acch-51071>
21. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30.

22. Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478.
23. Vesna, B. A. (2021). Challenges of financial risk management: AI applications. *Management: Journal of Sustainable Business and Management Solutions in Emerging Economies*, 26(3), 27–34.
24. Vose, D. (2008). *Risk analysis: A quantitative guide* (2nd ed.). Wiley.
25. Vukmirović, D., & Jovanović-Milenković, M. (2025). Generativna veštačka inteligencija: Principi i primena u savremenom poslovanju. *Fakultet organizacionih nauka*.
26. Giudici, P., Centurelli, M., & Turchetta, S. (2024). Artificial intelligence risk measurement. *Expert Systems with Applications*, 235, 121220. <https://doi.org/10.1016/j.eswa.2023.121220>
27. Goll, I., & Rasheed, A. A. (2005). The relationships between top management demographic characteristics, rational decision making, environmental munificence, and firm performance. *Organization studies*, 26(7), 999–1023.
28. Guikema, S. (2020). Artificial intelligence for natural hazards risk analysis: Potential, challenges, and research needs. *Risk Analysis*, 40(6), 1117–1123.
29. Gupta, V., & Yang, H. (2024). Generative artificial intelligence (AI) technology adoption model for entrepreneurs: case of ChatGPT. *Internet Reference Services Quarterly*, 28(2), 223–242.
30. Gupta, V., & Yang, H. (2024). Study protocol for factors influencing the adoption of ChatGPT technology by startups: Perceptions and attitudes of entrepreneurs. *Plos one*, 19(2), e0298427.
31. Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
32. De Laat, P. B. (2018). Algorithmic decision-making based on machine learning from big data: can transparency restore accountability?. *Philosophy & technology*, 31(4), 525–541.
33. Deepa, R., Sekar, S., Malik, A., Kumar, J., & Attri, R. (2024). Impact of AI-focussed technologies on social and technical competencies for HR managers—A systematic review and research agenda. *Technological Forecasting and Social Change*, 202, 123301.
34. Delen, D., Zaim, H., Kuzey, C., & Zaim, S. (2013). A comparative analysis of machine learning systems for measuring the impact of knowledge management practices. *Decision support systems*, 54(2), 1150–1160.
35. Del Giudice, M., Scuotto, V., Papa, A., Tarba, S., & Warkentin, M. (2021). Internationalization, digitalization, and sustainability: Are SMEs ready? *Technological Forecasting and Social Change*, 166, 120650. <https://doi.org/10.1016/j.techfore.2021.120650>
36. Deloitte. (2024). Validating GenAI models for financial services. Deloitte. Preuzeto sa: <https://www.deloitte.com/uk/en/Industries/financial-services/blogs/validating-genai-models.html>

37. Devan, M., Prakash, S., & Jangoan, S. (2023). Predictive maintenance in banking: leveraging AI for real-time data analytics. *Journal of Knowledge Learning and Science Technology* ISSN: 2959-6386 (online), 2(2), 483–490.
38. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
39. Duan, Y., Edwards, J. S., & Dwivedi, Y. K. (2019). Artificial intelligence for decision making in the era of big data: Evolution, challenges and research agenda. *International Journal of Information Management*, 48, 63–71. <https://doi.org/10.1016/j.ijinfomgt.2019.01.021>
40. Dwivedi, Y. K., Hughes, L., Rana, N. P., & Kar, A. K. (2024). Artificial intelligence and the future of work: Implications for research, practice, and policy. *Information Systems Frontiers*, 26(3), 1011–1034. <https://doi.org/10.1007/s10796-023-10473-9>
41. Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., ... & Williams, M. D. (2021). Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International Journal of Information Management*, 57, 101994.
42. Erel, I., Stern, L. H., Tan, C., & Weisbach, M. S. (2021). Selecting directors using machine learning. *The Review of Financial Studies*, 34(7), 3226–3264.
43. European Commission. (2024). EU Artificial Intelligence Act: Regulation (EU) 2024/1689. *Official Journal of the European Union*. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32024R1689>
44. European Parliament. (2023). EU AI Act: first regulation on artificial intelligence. *European Parliament*.
45. Yazdi, M., Zarei, E., Adumene, S., & Beheshti, A. (2024). Navigating the power of artificial intelligence in risk management: A comparative analysis. *Safety*, 10(2), 42. <https://doi.org/10.3390/safety10020042>
46. Jayatilake, S. M. D. A. C., & Ganegoda, G. U. (2021). Involvement of machine learning tools in healthcare decision making. *Journal of healthcare engineering*, 2021(1), 6679512.
47. Jane, J. B., & Ganesh, E. N. (2019). A review on big data with machine learning and fuzzy logic for better decision making. *Int. J. Sci. Technol. Res*, 8(10), 1221–1225.
48. Ji, Y., Rossi, L., & Zheng, J. (2023). Open Domain Multi-document Summarization: A Comprehensive Study of Model Brittleness under Retrieval. In *Findings of the Association for Computational Linguistics: EMNLP 2023* (pp. 8177–8199). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2023.findings-emnlp.549>
49. Jorzik, P., Yigit, A., Kanbach, D. K., Kraus, S., & Dabić, M. (2023). Artificial intelligence-enabled business model innovation: Competencies and roles of top management. *IEEE transactions on engineering management*, 71, 7044–7056.

50. Kalodimos, J., & Leavitt, K. (2020). Experimental Shareholder Activism: A novel approach for studying top management decision making and employee career issues. *Journal of Vocational Behavior*, 120, 103429.
51. Kaminski, M. E. (2023). Regulating the Risks of AI. *BUL Rev.*, 103, 1347.
52. Kandasamy, U. C. (2024). Ethical leadership in the age of AI: Challenges, opportunities and framework for ethical leadership. *arXiv preprint arXiv:2410.18095v2*.
<https://doi.org/10.48550/arXiv.2410.18095>
53. Canals, J., & Heukamp, F. (2020). *The future of management in an AI world*. Palgrave Macmillan.
54. Kaplan, S., & Garrick, B. J. (1981). On the quantitative definition of risk. *Risk Analysis*, 1(1), 11–27.
55. Carpenter, M. A., Geletkanycz, M. A., & Sanders, W. G. (2004). Upper echelons research revisited: Antecedents, elements, and consequences of top management team composition. *Journal of Management*, 30(6), 749–778.
56. Cao, G., Duan, Y., Edwards, J. S., & Dwivedi, Y. K. (2021). Understanding managers' attitudes and behavioral intentions towards using artificial intelligence for organizational decision-making. *Technovation*, 106, 102312. <https://doi.org/10.1016/j.technovation.2021.102312>
57. Cao, M., Chychyla, R., & Stewart, T. (2015). Big data analytics in financial statement audits. *Accounting Horizons*, 29(2), 423–429. <https://doi.org/10.2308/acch-51068>
58. Celestin, M., & Vanitha, N. (2015). Artificial intelligence vs human intuition: Who wins in risk management. *International Journal of Multidisciplinary Research and Modern Education*, 1(1), 699–706.
59. Khan, F. A., Khan, N. A., & Aslam, A. (2024). Adoption of artificial intelligence in human resource management: an application of TOE-TAM model. *Research and review: human resource and labour management*, 5, 22–36.
60. Coito, T., Firme, B., Martins, M. S., Vieira, S. M., Figueiredo, J., & Sousa, J. M. (2021). Intelligent sensors for real-Time decision-making. *Automation*, 2(2), 62–82.
61. Kohli, J. K., Raj, R., Rawat, N., Gupta, A., & Kumar, V. (2024, March). Ai empowered moocs usage and its impact on service quality in higher education institute in india. In *2024 2nd International Conference on Device Intelligence, Computing and Communication Technologies (DICCT)* (pp. 559–563). IEEE.
62. Kolbjørnsrud, V., Amico, R., & Thomas, R. J. (2017). Partnering with AI: how organizations can win over skeptical managers. *Strategy & Leadership*, 45(1), 37–43.
63. Koponen, J., Julkunen, S., Laajalahti, A., Turunen, M., & Spitzberg, B. (2025). Work characteristics needed by middle managers when leading AI-integrated service teams. *Journal of Service Research*, 28(1), 168–185.

64. Korzyński, P., Silva, S. C. E., Górska, A. M., & Mazurek, G. (2024). Trust in AI and top management support in generative-AI adoption. *Journal of Computer Information Systems*, 1–15.
65. KPMG. (2024). Model Risk Management: Key considerations for AI/ML based models. Preuzeto sa: <https://assets.kpmg.com/content/dam/kpmgsites/in/pdf/2024/11/model-risk-management.pdf>
66. Cubric, M. (2020). Drivers, barriers and social considerations for AI adoption in business and management: A tertiary study. *Technology in Society*, 62, 101257.
67. Cunha, M. P., Rego, A., & Clegg, S. (2023). Fairness and accountability in AI-based management decisions. *Organization Studies*. Advance online publication. <https://doi.org/10.1177/01708406231154722>
68. Laney, D. (2001). 3D data management: Controlling data volume, velocity and variety. META Group Research Note, 6(70), 1.
69. Laudon, K. C., & Laudon, J. P. (2020). *Management information systems: Managing the digital firm* (16th ed.). Pearson.
70. Lauritzen, C., & Moksness, S. S. (2024). Navigation AI Adoption In Marketing: Does TAM Clarify The Journey? (Master's thesis, Handelshøyskolen BI).
71. Li, L., Chen, Y., & Zhang, J. (2018). Deep reinforcement learning for real-time optimization and control of complex systems. *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, 48(9), 1420–1432. <https://doi.org/10.1109/TSMC.2018.2819644>
72. Li, X., Wang, P., & Cui, Z. (2024). Data-driven interoperability frameworks for AI-based decision systems. *Information Systems Frontiers*, 26(1), 145–162. <https://doi.org/10.1007/s10796-023-10422-6>
73. Liao, H., He, Y., Wu, X., Wu, Z., & Bausys, R. (2023). Reimagining multi-criterion decision making by data-driven methods based on machine learning: A literature review. *Information Fusion*, 100, 101970.
74. Loftus, T. J., Tighe, P. J., Filiberto, A. C., Efron, P. A., Brakenridge, S. C., Mohr, A. M., ... & Bihorac, A. (2020). Artificial intelligence and surgical decision-making. *JAMA surgery*, 155(2), 148–158.
75. Lyell, D., & Coiera, E. (2017). Automation bias and verification complexity: A systematic review. *Journal of the American Medical Informatics Association*, 24(2), 423–431. <https://doi.org/10.1093/jamia/ocw105>
76. Manduva, V. (2021). Exploring the Role of Edge-AI in Autonomous Vehicle Decision-Making: A Case Study in Traffic Management. *International Journal Of Scientific Research and Managment*, 9(02), 588–607.
77. Marocco, M., Barbieri, C., & Talamo, A. (2024). Exploring facilitators and barriers to managers' adoption of AI-based systems in decision making. *AI & Society*, 39(1), 77–99. <https://doi.org/10.1007/s00146-023-01679-3>

78. McKinsey. (2023, 1. avgust). The state of AI in 2023: Generative AI's breakout year. McKinsey & Company. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai-in-2023-generative-ais-breakout-year>
79. Merkert, J., Mueller, M., & Hubl, M. (2015). A survey of the application of machine learning in decision support systems.
80. Meyer, G., Adomavicius, G., Johnson, P. E., Elidrisi, M., Rush, W. A., Sperl-Hillen, J. M., & O'Connor, P. J. (2014). A machine learning approach to improving dynamic decision making. *Information Systems Research*, 25(2), 239–263.
81. Mnih, V., Kavukcuoglu, K., Silver, D., Rusu, A. A., Veness, J., Bellemare, M. G., ... Hassabis, D. (2015). Human-level control through deep reinforcement learning. *Nature*, 518(7540), 529–533. <https://doi.org/10.1038/nature14236>
82. Mökander, J., Stix, J., & Widdershoven, L. (2024). Auditing generative AI systems: A three-layered approach. In J. M. S. Widdershoven (Ed.), *The Research Handbook on the Law of Artificial Intelligence* (pp. 300–320). Edward Elgar Publishing.
83. Munoko, I., Brown-Liburd, H., & Vasarhelyi, M. A. (2020). The ethical implications of using artificial intelligence in auditing. *Journal of Business Ethics*, 167(3), 553–574. <https://doi.org/10.1007/s10551-019-04407-1>
84. Nieto, Y., Gacía-Díaz, V., Montenegro, C., González, C. C., & Crespo, R. G. (2019). Usage of machine learning for strategic decision making at higher educational institutions. *Ieee Access*, 7, 75007–75017.
85. NIST. (2024). AI Test, Evaluation, Validation, and Verification (TEVV) Programs. Preuzeto sa: <https://www.nist.gov/ai-test-evaluation-validation-and-verification-tevv>
86. Novelli, C., Taddeo, M., & Floridi, L. (2024). Accountability in artificial intelligence: What it is and how it works. *AI & Society*, 39(4), 1871–1882. <https://doi.org/10.1007/s00146-023-01635-y>
87. Papadakis, V. M., & Barwise, P. (2002). How much do CEOs and top managers matter in strategic decision-making?. *British Journal of management*, 13(1), 83–95.
88. Pereira, A. C., & Romero, F. (2017). A review of the meanings and the implications of the Industry 4.0 concept. *Procedia Manufacturing*, 13, 1206–1214.
89. Pinski, M., von Kortzfleisch, H., & Hanelt, A. (2024). AI literacy for the top management: An upper echelons perspective on corporate AI orientation and implementation ability. *Electronic Markets*, 34(24), 1–18. <https://link.springer.com/article/10.1007/s12525-024-00707-1>
90. Pomerol, J. C. (1997). Artificial intelligence and human decision making. *European Journal of Operational Research*, 99(1), 3–25.
91. Power, D. J. (2002). *Decision support systems: Concepts and resources for managers*. Greenwood Publishing Group.

92. Provost, F., & Fawcett, T. (2013). *Data science for business: What you need to know about data mining and data-analytic thinking*. O'Reilly.
93. Punniyamoorthy, M., & Sridevi, P. (2016). Identification of a standard AI based technique for credit risk analysis. *Benchmarking: An International Journal*, 23(5), 1381–1390.
94. Raes, A. M., Bruch, H., & De Jong, S. B. (2013). How top management team behavioural integration can impact employee work outcomes: Theory development and first empirical tests. *Human Relations*, 66(2), 167–192.
95. Raji, I. D., Gebru, T., & Mitchell, M. (2020). Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. In *Proceedings of the 2020 ACM Conference on Fairness, Accountability, and Transparency* (pp. 33–44). ACM. <https://doi.org/10.1145/3351095.3372873>
96. Raji, I. D., Olanrewaju, J., & Sambasivan, N. (2022). Outsider oversight: Designing a third party audit ecosystem for AI governance. In *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 557–571). ACM.
97. Rane, N. L., Paramesha, M., Choudhary, S. P., & Rane, J. (2024). Artificial intelligence, machine learning, and deep learning for advanced business strategies: a review. *Partners Universal International Innovation Journal*, 2(3), 147–171.
98. Ransbotham, S., Kiron, D., LaFountain, B., & Khodabandeh, S. (2022). Achieving individual—and organizational—value with AI. *MIT Sloan Management Review* and Boston Consulting Group. <https://sloanreview.mit.edu/projects/achieving-individual-and-organizational-value-with-ai/>
99. Ren, Y. (2025). AI Agents and Agentic AI—Navigating a Plethora of Concepts, Architectures, and Applications. *arXiv preprint arXiv:2507.01376*.
100. Rockart, J. F., & Crescenzi, A. D. (1984). Engaging top management in information technology. *Sloan Management Review* (pre-1986), 25(4), 3.
101. Romero, M. L., & Suyama, R. (2025). Agentic AI for Intent-Based Industrial Automation. *arXiv preprint arXiv:2506.04980*.
102. Selbst, A. D., & Barocas, S. (2021). Fairness and explanation in AI-informed decision making. *AI and Ethics*, 1(2), 1–18. <https://doi.org/10.1007/s43681-021-00068-0>
103. Simsek, Z., Veiga, J. F., Lubatkin, M. H., & Dino, R. N. (2005). Modeling the multilevel determinants of top management team behavioral integration. *Academy of Management Journal*, 48(1), 69–84.
104. Shin, D. (2023). How do people judge the credibility of artificial intelligence? Focus on algorithmic transparency and human–AI collaboration. *Computers in Human Behavior*, 141, 107642. <https://doi.org/10.1016/j.chb.2022.107642>
105. Song, Y., Qiu, X., & Liu, J. (2025). The impact of artificial intelligence adoption on organizational decision-making: An empirical study based on the technology acceptance model in business management. *Systems*, 13(8), 683. <https://doi.org/10.3390/systems13080683>

106. Steimers, A., & Schneider, M. (2022). Sources of risk of AI systems. *International Journal of Environmental Research and Public Health*, 19(6), 3641.
107. Tabesh, P., & Vera, D. M. (2020). Top managers' improvisational decision-making in crisis: a paradox perspective. *Management Decision*, 58(10), 2235–2256.
108. Taddeo, M., & Floridi, L. (2021). How AI can be a force for good. *Science*, 372(6544), 26–29. <https://doi.org/10.1126/science.abg5991>
109. Tao, F., Zuo, Y., Xu, L. D., & Zhang, L. (2014). IoT-based intelligent perception and access of manufacturing resource toward cloud manufacturing. *IEEE Transactions on Industrial Informatics*, 10(2), 1547–1557.
110. Tien, J. M. (2017). Internet of things, real-time decision making, and artificial intelligence. *Annals of Data Science*, 4(2), 149–178.
111. Tulabandhula, T., & Rudin, C. (2014). On combining machine learning with decision making. *Machine learning*, 97(1), 33–64.
112. Turban, E., Sharda, R., & Delen, D. (2011). *Decision support and business intelligence systems* (9th ed.). Pearson.
113. Uzonwanne, F. C. (2018). Rational model of decision-making. In *Global encyclopedia of public administration, public policy, and governance* (pp. 5334–5339). Springer, Cham.
114. Faheem, M. A. (2021). AI-driven risk assessment models: Revolutionizing credit scoring and default prediction. *Iconic Research And Engineering Journals*, 5(3), 177–186.
115. Floridi, L., & Cows, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 5(1). <https://onlinelibrary.wiley.com/doi/abs/10.1002/9781119815075.ch45>
116. Hambrick, D. C., & Mason, P. A. (1984). Upper echelons: The organization as a reflection of its top managers. *Academy of Management Review*, 9(2), 193–206.
117. Harrison, J. S., Josefy, M. A., Kalm, M., & Krause, R. (2023). Using supervised machine learning to scale human-coded data: A method and dataset in the board leadership context. *Strategic Management Journal*, 44(7), 1780–1802.
118. Hassan, T. W. A. S. (2025). Large Language Models for Document Summarization: A Systematic Evaluation and Analysis of Efficiency and Factual Consistency. *International Journal of Advanced Computer Science and Applications*, 16(6), 356–371.
119. Herrmann, J. W. (2014). Rational decision making. *Wiley StatsRef: Statistics Reference Online*, 1–9.
120. Hickson, D. J. (1987). Decision-making at the top of organizations. *Annual review of sociology*, 13(1), 165–192.
121. Hough, P. K., & Duffy, N. M. (1987). Top management perspectives on decision support systems. *Information & Management*, 12(1), 21–30.

122. Chen, H., Chiang, R. H. L., & Storey, V. C. (2012), "Business Intelligence and Analytics: From Big Data to Big Impact", *MIS Quarterly*, 36(4), 1165–1188, <https://doi.org/10.2307/41703503>
123. Constantiou, I. D., & Kallinikos, J. (2015). New games, new rules: Big data and the changing context of strategy. *Journal of Information Technology*, 30(1), 44–57. <https://doi.org/10.1057/jit.2014.17>
124. Carmeli, A., & Halevi, M. Y. (2009). How top management team behavioral integration and behavioral complexity enable organizational ambidexterity: The moderating role of contextual ambidexterity. *The leadership quarterly*, 20(2), 207–218.
125. Shin, D. (2023). How do people judge the credibility of artificial intelligence? Focus on algorithmic transparency and human–AI collaboration. *Computers in Human Behavior*, 141, 107642. <https://doi.org/10.1016/j.chb.2022.107642>
126. Yörük, N. (2025). Enhancing SME Operations with Machine Learning and Business Intelligence: A Case Study of Kolay.ai. *Journal of Data Analytics and Artificial Intelligence Applications*, 2(3).

3. Полазне хипотезе

Полазећи од теоријског оквира и циљева истраживања, формулисане су хипотезе које оперативно повезују кључне конструкте: објашњивост/транспарентност ВИ (ХАИ), поверење врховног менаџмента (ТМТ), перципирани ризик, управљање ВИ, намера усвајања и перформансе одлука у реалном времену (RTD).

Ради јасноће и тестабилности, хипотезе су груписане на опште (главне) — које представљају темељне односе у моделу, посебне — које разјашњавају механизме дејства (медијације, модерације и нелинеарности) уз сваку главну хипотезу, и помоћне — које обухватају контролне, хетерогеност и робусност. Очекивани смерови ефеката назначени су уз сваку хипотезу, док су поступци емпиријске верификације дати у поглављу Методологија истраживања.

Опште (главне) хипотезе

- H1. Виши ниво објашњивости и транспарентности система вештачке интелигенције повећава поверење врховног менаџмента у алгоритамске одлуке.
- H2. Јача подршка врховног менаџмента повећава намеру организације да усвоји вештачку интелигенцију у одлучивању у реалном времену.
- H3. Виши ниво перципираног ризика у вези са применом вештачке интелигенције смањује намеру усвајања ових система.
- H4. Постоји нелинеарни однос између поверења у вештачку интелигенцију и перформанси одлука – оптималне перформансе се остварују при умереним нивоима поверења.

Посебне хипотезе (механизми уз опште хипотезе)

Уз Н1 (ХАИ → Поверење ТМТ):

Н1.1 (медијација). Виша ВИ писменост повећава јасноћу и разумљивост ХАИ објашњења, што посредно подиже поверење врховног менаџмента.

Н1.2 (медијација). Објашњивост ХАИ модела повећава перципирану корисност (PU) и лакоћу употребе (PEOU), које посредују утицај ХАИ на намеру усвајања (BI).

Н1.3 (модерација). ВИ писменост појачава позитиван ефекат објашњивости ХАИ модела на поверење ТМТ (интеракцијски ефекат).

Уз Н2 (ТМТ подршка → BI):

Н2.1 (модерација). Усклађеност и сарадња између ТМТ и ИТ менаџмента појачавају ефекат подршке ТМТ на намеру усвајања ВИ решења.

Уз Н3 (Ризик → BI):

Н3.1 (медијација). Виши ниво етичке усклађености и зрелости управљања ВИ смањује перципирани ризик и посредно повећава поверење и намеру усвајања.

Н3.2 (модерација). Зрео *AI governance* и доступност ХАИ функционалности ублажавају негативан ефекат перципираног ризика на намеру усвајања.

Уз Н4 (нелинеарност поверења):

Н4.1 (експлораторно). Ниво ВИ писмености и присуство људског надзора могу мењати оптималну тачку односа између поверења и перформанси одлука.

Помоћне хипотезе (контроле, хетерогеност, робусност)

С1 (контролне варијабле). Величина организације и дигитална зрелост позитивно утичу на намеру усвајања ВИ, независно од главних предиктора.

С2 (хетерогеност). Ефекти објашњивости ХАИ, перципираног ризика и *AI governance* зрелости израженији су у високо регулисаним индустријама (нпр. финансије, телекомуникације).

С3 (робусност). Налази остају статистички стабилни на алтернативним техничким спецификацијама (PLS-SEM и коваријациони SEM модели) и приликом редукције конструктора са високом колинеарношћу.

4. Научне методе истраживања

При изради овог истраживања примениће се више комплементарних научних метода.

Опште научне методе

- **Аналитичко-дедуктивна метода:** за анализу теоријских поставки и релевантних стандарда/оквира (нпр. модели прихватања технологије, оквири управљања ВИ), као и преглед постојеће стручне праксе у одлучивању и управљању ризиком у реалном времену.

- **Метода моделовања:** за концептуално осмишљавање и формализацију интегрисаног модела односа између објашњивости/транспарентности ВИ поверења врховног менаџмента перципираног ризика, управљања ВИ, намере усвајања и перформанси одлука у реалном времену.
- **Компаративна метода:** за упоредну анализу традиционалних приступа одлучивању и приступа заснованих на ВИ (са и без ХАИ и људског надзора), укључујући разлике по индустријама и нивоу дигиталне зрелости.
- **Синтетичка метода:** за интеграцију теоријских сазнања и емпиријских налаза у јединствен оквир и извођење закључака истраживања.

Посебне научне методе

- **Квалитативне методе:**
 - Полуструктурисани интервјуи са представницима ТМТ и ИТ/аналитичких тимова у организацијама у Србији, ради идентификовања мотива, баријера и стратегија усвајања ВИ у RTD процесима (уз акценат на објашњивост, поверење, ризик и управљање ВИ).
 - Квалитативна анализа садржаја документације/политика и интерних материјала (нпр. процедуре, ревизорски траг, смернице за ХАИ) у циљу издвајања образаца и критеријума добре праксе.
 - Студија случаја илустративних RTD решења (анонимизованих) ради приказа функционисања оквира у реалним условима и мапирања тачака контроле (људски надзор, прагови интервенције).
- **Квантитативне методе :**
 - Анкетно истраживање са стандардизованим Ликерт инструментом за мерење кључних конструката (објашњивост/транспарентност ВИ, поверење ТМТ, перципирани ризик, услови који олакшавају, намера усвајања, перформансе одлука у реалном времену).
 - Претест/пилот инструмента ради провере разумљивости и метријских својстава.
 - Статистичке методе:
 - CFA (конфирматорна факторска анализа) за валидацију мерног модела;
 - SEM (структурално моделовање једначина) за тестирање хипотеза, укључујући посредничке (медијационе) и модерационе (интеракционе) ефекте, као и нелинеарности
 - Провера инваријантности мерења по индустријама и нивоу дигиталне зрелости;
 - Процедуралне и статистичке контроле

Интердисциплинарност

Истраживање је изразито интердисциплинарно, јер обједињује увиде из рачунарства (вештачка интелигенција, објашњивост – ХАИ, обрада токова података), менаџмента и организационих наука (ТМТ, управљање ВИ), статистике (CFA/SEM) и етике/управљања ризиком (правичност, одговорност, транспарентност; ревизорски траг и људски надзор).

5. Очекивани научни допринос

Научни доприноси

- **Интегрисани теоријско-емпиријски модел:** развој и формализација модела који повезује објашњивост и транспарентност ВИ, поверење врховног менаџмента, перципирани ризик, управљање ВИ, и намеру усвајања у одлучивању у реалном времену — са јасно постављеним хипотезама и механизмима (медијације, модерације, нелинеарности).
- **Теоријска интеграција оквира:** консолидација налаза ТАМ/UTAUT са приступима управљања ВИ (*governance*) и моделима прихватања/избегавања — у јединствен, тестабилан оквир применљив на високо ризичне домене одлучивања.
- **Операционализација конструката и развој инструмента:** дизајн и валидација скала (CFA/SEM) за објашњивост, поверење, перципирани ризик, писменост о ВИ, услове који олакшавају и перформансе одлука у реалном времену; потврда поузданости и конструктне валидности на узорку организација.
- **Емпиријско доказивање нелинеарности:** тестирање и потврда инвертиране U-везе између поверења у ВИ и перформанси одлука, са идентификацијом „оптималне зоне” поверења и улоге људског надзора.
- **Модел механизма утицаја:** емпиријско потврђивање (i) посредничких ефеката објашњивости преко перципиране корисности/лакоће и (ii) ублажавајуће улоге управљања ВИ и писмености о ВИ на негативан утицај ризика.
- **Методолошки допринос:** предлог стандардизованог протокола мерења перформанси одлука у реалном времену (нпр. тачност/учинак кроз временске прозоре, прагови упозорења, ревизорски траг, *override* процедуре) који је преносив на различите индустрије.
- **Компаративни увид:** мултигрупна анализа између сектора (високо регулисани VS остали) и различитих нивоа дигиталне зрелости, ради мапирања граница и општости модела.

Стручни (професионални) доприноси

- **Оквир за одговорну имплементацију ВИ:** практичне смернице за ТМТ и ИТ за увођење објашњиве и контролисане ВИ (политике, процедуре, матрице одговорности, контролне тачке, ревизорски траг).
- **Водич за процену спремности и план усвајања:** *check*-листе и дијагностички упитници (писменост о ВИ, *governance*, подаци/инфраструктура, услови који олакшавају) са приоритетима улагања и мапом пута имплементације.
- **Контрола ризика и спречавање прекомерног ослањања:** протокол за уградњу људског надзора, прагова интервенције и метрика које балансирају ефикасност и безбедност одлучивања.
- **Модел зрелости (*maturity*) и KPI за одлучивање у реалном времену:** дефинисани нивои зрелости и показатељи успеха (нпр. време реакције, стабилност перформанси, стопа ескалација), за управљање и праћење напретка.
- **Шаблони за интерне политике:** предлошци политика и процедура (управљачки преглед, одговорност, транспарентност, заштита података) спремни за прилагођавање организацијама.

Друштвени доприноси

- **Повећање поверења у алгоритамско одлучивање:** препоруке које јачају транспарентност и одговорност система ВИ у организацијама од јавног интереса, што доприноси већој прихваћености технологије.
- **Подршка регулаторним телима и стандардизацији:** предлози елемената који могу информисати смернице и интерне правилнике у секторима са вишим ризиком (нпр. финансије, телекомуникације, јавне услуге).
- **Развој компетенција:** образовни модули и материјали за подизање писмености о ВИ код руководиоца и стручних тимова, што олакшава сигурно и етичко усвајање.
- **Економски и управљачки ефекти:** допринос ефикаснијем управљању ризицима и квалитетнијим одлукама, са потенцијалом за смањење трошкова, брже реаговање и боље усклађивање са прописима.
- **Етика и правичност:** операционални кораци који смањују ризик пристрасности и јачају правичност и заштиту корисника кроз објашњивост и “аудитибилност” система.

6. План истраживања и структура рада

План истраживања докторске дисертације приказан је у следећој табели:

ФАЗА	Кључни задатак	Методе / технике / стандарди
1. Анализа и синтеза литературе	• Преглед савремене литературе о ВИ, одлучивању у реалном времену и анализи ризика • Идентификација фактора усвајања ВИ на основу постојећих теоријских модела (ТАМ, UTAUT, IAAAM, <i>AI Readiness</i>, UET)	• Аналитичко-дедуктивна метода • Систематски преглед литературе • Компаративна метода – поређење традиционалних и ВИ приступа одлучивања
2. Дефинисање концептуалног модела	• Формулисање концептуалног оквира односа: XAI → Поверење (TMT) → Перцепција ризика → Намера усвајања → Перформансе RTD система • Дефинисање конструката и логичких веза у моделу	• Метода моделовања (дијаграми, логичке структуре) • Синтетичка метода – интеграција теорије и емпиријских увида
3. Планирање и развој методолошког приступа	• Дефинисање општих и посебних научних метода • Одабир приступа истраживања • Одређивање циљних група (TMT, ИТ менаџери)	• Комбинација општих и посебних научних метода (аналитичка, моделовање, компаративна, синтетичка) • Методолошко планирање засновано на квалитативним техникама
4. Прикупљање података (квалитативни део)	• Спровођење полуструктурираних интервјуа са TMT и ИТ стручњацима ради анализе мотива, баријера, ризика и фактора прихватања ВИ	• Полуструктурирани интервјуи • Квалитативна анализа садржаја документације и политика (XAI смернице, процедуре, ревизорски траг)
5. Анализа података	• Квалитативно тумачење ставова, образаца и фактора усвајања • Мапирање улоге <i>poverenja</i> TMT и XAI у одлучивању	• Квалитативна анализа садржаја • Синтетичка интеграција налаза

6. Формулисање интегрисаног модела поверења у ВИ	• Комбиновање теоријских основа, квалитативних увида и идентификованих фактора у јединствени модел	• Синтетичка метода – интеграција налаза у један теоријски оквир
7. Дискусија и научни допринос	• Повезивање налаза са постојећим теоријским моделима • Објашњење утицаја техничких (XAI), организационих (TMT) и етичких фактора (<i>AI governance</i>)	• Компаративна анализа са постојећим истраживањима • Аналитичко-синтетичка интерпретација
8. Закључци и препоруке	• Сумирање кључних налаза • Предлог смерница за одговорну и ефикасну примену ВИ у одлучивању у реалном времену	• Анализа и синтеза • Формулисање препорука за праксу и будућа истраживања

ФАЗА	МЕСЕЦ											
	I	II	III	IV	V	VI	VII	VIII	IX	X	XI	XII
1	*	*										
2		*	*									
3			*	*	*							
4				*	*	*						
5					*	*	*					
6						*	*	*	*			
7							*	*	*	*		
8								*	*	*	*	*

Оквирно, структуру докторске дисертације сачињаваће следећа поглавља:

Увод

- 1.1. Теоријски контекст и значај теме
- 1.2. Предмет, проблем и актуелност истраживања
- 1.3. Циљеви истраживања
- 1.4. Полазне хипотезе
- 1.5. Научне методе и допринос истраживања
- 1.6. Структура рада

2. Вештачка интелигенција у савременом пословању

3. Одлучивање у реалном времену и алати ВИ за анализу ризика

4. Теоријски оквир и преглед литературе

- 4.1. Усвајање вештачке интелигенције: покретачи и ограничавајући фактори
- 4.2. Поверење и интеракција човека и вештачке интелигенције у менаџерском одлучивању
- 4.3. Спремност организација и улога *Top Menadžmenta (TMT)* у имплементацији ВИ

5. Теоријски модели прихватања и поверења у вештачку интелигенцију

- 5.1. *Technology Acceptance Model (TAM)* – основни принципи и примена
- 5.2. *Unified Theory of Acceptance and Use of Technology (UTAUT)*
- 5.3. *Integrated AI Acceptance–Avoidance Model (IAAAM)*
- 5.4. *Technology Threat Avoidance Theory (TTAT)*
- 5.5. Модели организационе спремности за примену вештачке интелигенције
- 5.6. *Upper Echelons Theory (UET)* и улога лидерства у усвајању ВИ
- 5.7. Компаративна анализа модела и њихова применљивост у ВИ алатима за *RTD* и *risk analytics*

6. Улога менаџера у процесу усвајања ВИ алата

- 6.1. Перспектива *TMT* менаџмента: стратегија, лидерство и култура поверења
- 6.2. Перспектива ИТ менаџмента: технички аспекти, интероперабилност и одговорност
- 6.3. Интеракција *TMT* и ИТ менаџера у имплементацији ВИ алата
- 6.4. ВИ писменост (*AI literacy*) као фактор прихватања и поверења
- 6.5. Организационе политике и етичко управљање ВИ системима

7. Интегрисани теоријски модел прихватања и поверења у ВИ

- 7.1. Три димензије модела
 - 7.1.1. Технолошка димензија
 - 7.1.2. Организациона димензија
 - 7.1.3. Психолошко-етичка димензија
- 7.2. Механизми међусобног деловања фактора
- 7.3. Концептуални оквир истраживања (Емпиријски модел)
- 7.4. Емпиријски оквир истраживања
 - 7.4.1. Концептуална основа емпиријског оквира
 - 7.4.2. Операционализација димензија
 - 7.4.3. Истраживачке хипотезе
 - 7.4.4. Емпиријски модел истраживања

8. Емпиријски модел и методологија истраживања

- 8.1. Емпиријски модел: конструкти и хипотезе
- 8.2. Операционализација варијабли и развој истраживачких инструмената
- 8.3. Дизајн истраживања и узорак
 - 8.3.1. Циљна популација и критеријуми избора испитаника
 - 8.3.2. Анкетни инструменти:
 - Упитник за *TMT* менаџере
 - Упитник за ИТ менаџере
- 8.4. Поступак прикупљања података
- 8.5. Методе анализе података (дескриптивна, инференцијална, *SEM* анализа)
- 8.6. Валидација и поузданост мерних инструмената

9. Резултати емпиријског истраживања

- 9.1.** Опис узорка и основне карактеристике испитаника
- 9.2.** Дескриптивна статистика и корелациона анализа
- 9.3.** Тестирање хипотеза и анализа структуралних односа
- 9.4.** Поређење *TMT* и ИТ перспективе у перцепцији ВИ
- 9.5.** Дискусија о резултатима и интерпретација налаза

10. Дискусија и интерпретација резултата

- 10.1.** Поређење са претходним истраживањима и теоријским моделима
- 10.2.** Импликације резултата за менаџерску теорију и праксу
- 10.3.** Улога поверења, ВИ писмености и организационе културе у усвајању ВИ
- 10.4.** Ограничења истраживања

11. Закључак и препоруке

- 11.1.** Главни закључци дисертације
- 11.2.** Теоријски допринос
- 11.3.** Практични допринос и импликације за менаџмент и политику
- 11.4.** Смернице за будућа истраживања



7. Закључак и предлог

Из изложеног се може закључити да кандидат Небојша Аврамовић испуњава све услове предвиђене Законом о високом образовању за одобрење израде докторске дисертације под насловом „**Примена вештачке интелигенције у алатима за пословно одлучивање и процену ризика у реалном времену у дигиталном пословању**”.

Небојша Аврамовић поседује неопходна знања из електронског пословања, пословног одлучивања и информационих технологија, за успешан рад на докторској дисертацији.

Тема припада ужој научној области Електронско пословање, мултидисциплинарна је и актуелна. Добијени резултати ће допринети формализацији интегрисаног модела прихватања и управљања ВИ, те унапређењу квалитета пословног одлучивања и контроле ризика у организацијама које користе системе у реалном времену. Истраживање ће имати директну практичну примену кроз развој смерница и *maturity* модела за одговорну имплементацију ВИ.

На основу свега наведеног, комисија предлаже Наставно-научном већу да прихвати предложену тему и одобри израду пријављене докторске дисертације. За ментора докторске дисертације предлаже се др Драган Вукмировић, редовни професор Факултета организационих наука, Универзитета у Београду.

ЧЛАНОВИ КОМИСИЈЕ:

др Драган Вукмировић, редовни професор, Универзитет у
Београду- Факултет организационих наука

др Маријана Деспотовић-Зракић, редовни професор,
Универзитет у Београду Факултет организационих наука

др Борис Делибашић, редовни професор,
Универзитет у Београду-Факултет организационих наука

др Александар Ђоковић, ванредни професор,
Универзитет у Београду-Факултет организационих наука

др Саша Милић, начни саветник, Електротехнички
институт Никола Тесла у Београду

Београд, 14.11.2025.