

UMELÁ INTELIGENCIA A ROZHODOVANIE V RÁMCI DISKRECIÓNÁLNEHO OPRÁVNENIA VO VECIACH STAVEBNÝCH PRIESTUPKOV - EMPIRICKÁ ŠTÚDIA¹

ARTIFICIAL INTELLIGENCE AND DISCRETIONARY DECISION- MAKING IN CONSTRUCTION OFFENSE CASES – AN EMPIRICAL STUDY

Ján Škrobák²

Abstrakt: Príspevok sa zaoberá problematikou integrácie umelej inteligencie do rozhodovania v rámci diskrečionálneho oprávnenia. Zameriava sa pritom na rozhodovanie v rámci postihovania priestupkov v stavebných veciach, keďže práve na tomto úseku v ostatnom čase v rámci rekodifikačných procesov bol zrejmy silný trend elektronizácie rozhodovacích procesov. Zároveň však ide o rozhodovacie procesy, ktoré predstavujú osobitnú výzvu z hľadiska vstupných informácií i z hľadiska existujúcich súborov dát o aplikačnej praxi pri sankcionovaní. Vychádzajúc z vedeckej literatúry a priemetu takto získaných poznatkov do roviny slovenských reálií rozoberá prínosy a silné stránky a negatíva a riziká rozhodovania alebo účasti na rozhodovaní umelej inteligencie. V ďalšej časti príspevok prezentuje výsledky experimentu zameraného na testovanie, ako sa verejnosti bežne dostupné modely umelej inteligencie dokážu vysporiadať s rozhodovaním v sankčných veciach v rámci správnej úvahy. Testuje, ako – vychádzajúc z opisu skutkového stavu fiktívneho prípadu doplneného o základné informácie o relevantných ustanoveniach normatívnej úpravy – umelá inteligencia navrhuje rozhodnúť o uložení sankcie a nakoľko je pri tom konzistentná. Na základe teoretických i empirických východísk príspevok sumarizuje prínosy a riziká rozhodovania v sankčných veciach v rámci diskrečionálneho oprávnenia umelou inteligenciou.

Kľúčové slová: umelá inteligencia, automatizované rozhodovanie, konanie o priestupku, ukladanie sankcií, správne trestanie, nediskriminácia

Abstract: The paper deals with the issue of integrating artificial intelligence into discretionary decision-making. It focuses on decision-making within the

¹ Táto práca bola podporená Agentúrou na podporu výskumu a vývoja na základe Zmluvy č. APVV-24-0543.

² Univerzita Komenského v Bratislave, Právnická fakulta.

framework of sanctioning construction offences, since in this area a strong trend of electronic decision-making processes has recently been evident within the framework of recodification processes. However, these decision-making processes also represent a special challenge in terms of input information and in terms of existing data sets on application practice in sanctioning. Based on the scientific literature and the projection of the knowledge thus obtained into the Slovak reality, it discusses the benefits and strengths and the negatives and risks of decision-making or participation in decision-making by artificial intelligence. In the next part, the paper presents the results of an experiment aimed at testing how artificial intelligence models commonly available to the public can deal with discretionary decision-making in sanctioning cases. It tests how - based on a description of the facts of a fictitious case supplemented with basic information about the relevant legal regulation - artificial intelligence proposes a decision on the imposition of a sanction and how consistent it is in doing so. Based on theoretical and empirical foundations, the paper summarizes the benefits and risks of discretionary decision-making in sanction matters by artificial intelligence.

Key words: artificial intelligence, automated decision-making, misdemeanor proceedings, imposition of sanctions, administrative punishment, non-discrimination

Úvod

„Rozhodovanie v stavebných priestupkoch, kde má diskrečné oprávnenie kľúčovú úlohu, môže pri nasadení umelej inteligencie získať na efektívnosti, no zároveň v sebe nesie tiché riziko, že spravodlivosť ustúpi pred algoritmickou pohodlnosťou.“ Svoj článok začínam vetou v úvodzovkách, pretože ju napísal ChatGPT na základe požiadavky, aby vytvoril úvodnú vetu vedeckého článku, ktorý sa zaoberá možnosťami a limitmi používania umelej inteligencie pri rozhodovaní v rámci diskrecionálneho oprávnenia vo veciach stavebných priestupkov; táto veta mala byť odborná, ale mať trochu aj šarm bonmotu a vyjadrovať mierne znepokojenie.

Podarilo sa? Myslím, že celkom áno. Je to len drobná ilustrácia možností súčasných modelov generatívnej umelej inteligencie. Umelá inteligencia nezadržateľne preniká aj do sveta práva a stavia právnych teoretikov i praktikov pred množstvo otázok jej využitia a limitov jej využitia. Tieto otázky sú zrejme citlivejšie vo verejnom práve, keďže verejné právo je založené na rigoróznejšom prístupe k zákonnosti v tom zmysle, že pre konanie verejnoprávných entít sa vyžaduje zákonný základ a dominujú v ňom kogentné a striktné právne normy.

Práve využívanie diskrečionálneho oprávnenia je jednou zo sfér, kde sa uplatňuje výnimka z prevencie striktných noriem a používajú sa právne normy ekvitalne.

Cieľom výskumu, ktorého výsledky prezentujem v tomto príspevku, bolo nájsť odpoveď na otázku, do akej miery prichádza do úvahy zapojenie umelej inteligencie do rozhodovania vo veciach stavebných priestupkov, kde sa rozhoduje na základe diskrečionálneho oprávnenia, aké sú možné prínosy a aké sú možné riziká implementovania um v týchto prípadoch.

Osobitnú pozornosť som pritom zameral na otázku stability a kontinuity rozhodovania pri ukladaní sankcií, teda na to, či je dostatočne zachovaná požiadavka nediskriminácie a zachovania materiálnej rovnosti, ktorá správnym orgánom prikazuje rešpektovať kontinuitu rozhodovania.³

Uskutočniť vo vzťahu k týmto otázkam výskum práve vo sfére rozhodovania vo veciach stavebných priestupkov sa mi javilo ako obzvlášť vhodné z dvoch dôvodov:

Prvým je veľký trend k elektronizácii rozhodovacích procesov v stavebníctve a pri projektovaní, uskutočňovaní a užívaní stavieb⁴. Druhým je pnutie medzi determinantmi prínosov a rizík využívania umelej inteligencie, ktoré je obzvlášť citlivé práve v tejto oblasti: na jednej strane vidíme zložitosť zisťovania skutkového stavu stavebných priestupkov, ktorý nie je spravidla v úplnosti opísateľný len slovným naratívom, a častú potrebu zisťovania skutkového stavu in situ, na druhej strane možnosti vyplývajúce z možného etablovania informačných modelov stavby (BIM) (digitálne, interaktívne a aktualizovateľné vyjadrenie stavby s integrovanou databázou informácií), z ktorých vyplýva možnosť automatizovaného získavania podkladu pre rozhodnutie⁵.

³ KOŠIČIAROVÁ, S.: Správny poriadok. Komentár s novelou účinnou od 1. januára 2004. Šamorín : Heuréka, 2004, s. 25. ISBN 80-89122-14-0.

⁴ Ide o také prvky normatívnej úpravy, ako to, že dokumentácia stavby sa vypracuje spravidla v elektronickej podobe a ukladá sa v informačnom systéme, prerokovanie stavebného zámeru sa realizuje v informačnom systéme, ďalej príkladom je elektronické podávanie žiadosti na začatie konania, ako aj to, že doručovanie vo výstavbe sa uskutočňuje elektronicky podľa zákona o e-Governmente; ak sa doručuje elektronický dokument, ktorý je adresátovi dostupný prostredníctvom informačného systému, možno doručenie elektronického dokumentu vykonať dorúčením informácie s priamym odkazom na miesto, kde je prostredníctvom informačného systému adresátovi dostupný, a pod.

⁵ Ide najmä o najpokročilejšie dimenzie BIM, ako je gD a 10D, ktoré predstavujú formy pokročilého digitálneho modelovania stavby (BIM gD predstavuje používanie optických alebo laserových skenerov na zachytenie reálneho prostredia, komponentov alebo budov, pričom získané informácie sa prevádzajú do digitálnych modelov; BIM 10D predstavuje ponorenie prvkov virtuálnej reality do reálneho sveta). Pozri KULKARNI, A. a kol.: Building Information Modeling-Enhancing Productivity in Rail Infrastructure Contruction. [online]. ResearchGate 2018 [cit. 2025-09-27]. Dostupné na: https://www.researchgate.net/publication/335464714_Building_Information_ModellingEnhancing_Productivity_in_Rail_Infrastructure_Construction

Teoretický vstup do problematiky

Využívanie umelej inteligencie na rozhodovanie vo veciach viny zo spáchania verejnoprávnych deliktov a ukladania sankcií za ne je nepochybne citlivá otázka. V odbornej literatúre sa uvádzajú možné výhody a prínosy jej implementovania, ale upozorňuje sa aj na riziká, nevýhody či nedostatky. Ako zrejme preukáže táto časť príspevku, kvantitatívne zrejme prevažujú práve tieto riziká či nedostatky. Už samotný rešerš literatúry teda vyznieva z väčšej časti ako hlasný caveat.

Pri skúmaní teoretických východísk problematiky som vychádzal primárne zo zahraničnej literatúry a to v nemalom rozsahu aj z literatúry trestnoprávnej, ktorá sa týkala sa najmä rozhodovania o ukladaní sankcií v trestných veciach („criminal sentencing“). Tento prístup je možný a prínosný s ohľadom na:

- typovú príbuznosť rozhodovacích procesov (a etablované nazeranie na obvinenie z priestupkov ako na „trestné obvinenie“⁶),
- fakt, že AI systémy zamerané na podporu rozhodovania sa javia byť rozšírené najmä v trestných veciach (pozri napr. zoznam AI systémov zameraných na podporu rozhodovania, ktorý vedie Resource Centre Cyberjustice and AI Európskej komisie pre efektívnosť súdnictva /CEPEJ/)⁷, tu však treba upozorniť, že pri výskume som sa stretol najmä s informáciami o systémoch zameraných na posudzovanie rizika recidívy na účely rozhodovania o uložení trestu,⁸ čo je aspekt rozhodovania, ktorý v slovenskom priestupkovom konaní zjavne nemá ekvivalent,
- dostupnosť zdrojov skôr trestnoprávnej proveniencie.

Prv, než pristúpim k samotnému predstaveniu doktrinálnych pohľadov na pozitíva a negatíva využívania umelej inteligencie pri sankcionovaní v rámci správneho uváženia, je potrebné vyjasniť, čo presne sa pod takýmto využitím môže rozumieť.

Do úvahy prichádza široký diapazón možností siahajúcich od relatívne okrajového využitia umelej inteligencie pri pomocných úkonov až po plne automatizované rozhodnutie, kde nositeľom rozhodovacej činnosti nebude človek, ale umelá inteligencia. Je dôležité vnímať rozdiely medzi tým, keď sa umelá inteligencia do

⁶ Pozri HAVELKOVÁ, M.: Použitie ustanovení trestného práva v oblasti správneho trestania - 1. časť. [online]. Projjustice.sk. [cit. 2025-09-27]. Dostupné na: <https://www.projjustice.sk/spravne-pravo/pouzitie-stanoveni-trestneho-prava-v-oblasti-spravneho-trestania>

⁷ Pozri Resource Centre Cyberjustice and AI by CEPEJ. [online]. [cit. 2025-09-27]. Dostupné na: <https://public.tableau.com/app/profile/cepej/viz/ResourceCentreCyberjusticeandAI/AITOOLSINITIATIVESREPORT>

⁸ Pozri napr. FAZEL, S. a kol.: The predictive performance of criminal risk assessment tools used at sentencing: Systematic review of validation studies. In Journal of Criminal Justice. Zväzok 81, Júl – august 2022, nestránkované. <https://doi.org/10.1016/j.jcrimjus.2022.101902> Pozri aj dielo A. Završnika, ktorý uvádza aj príklady takýchto softvérov, predovšetkým algoritmus Arnold Foundation, ktorý sa používa v 21 štátoch USA – pozri ZAVRŠNIK, A.: Algorithmic justice: Algorithms and big data in criminal justice settings. In European Journal of Criminology, 2021, Zväzok 18, číslo 5, s. 623-642. DOI: 10.1177/1477370819876762.

rozhodovacieho procesu integruje ako podporný nástroj (napríklad na anonymizáciu dokumentov, na vyhľadávanie rozhodnutí v obdobných veciach, na posudzovanie kontinuity rozhodovania, na vytváranie návrhov rozhodnutí, na konzultovanie a oponovanie návrhov rozhodnutí), alebo je priamo aktérom automatizovaného rozhodovania.

V tejto súvislosti musím upozorniť, že vidím rozdiel medzi konceptom rozhodovania prostredníctvom umelej inteligencie a konceptom automatizovaného rozhodovania. Zatiaľ čo rozhodovanie s využitím AI má vždy prvok automatizovaného rozhodovania, automatizované rozhodovanie nemusí byť nevyhnutne založené na umelej inteligencii. Existujú aj systémy automatizovaného rozhodovania, ktoré nie sú založené na umelej inteligencii (napríklad ak elektronický systém zaznamená prekročenie najvyššej povolenej rýchlosti a automaticky vygeneruje rozhodnutie o pokute na základe algoritmu fixne definovanej sadzby s fixne preddefinovaným odôvodnením, nepôjde zrejme o rozhodovanie založené na použití umelej inteligencie). Toto rozlišovanie však neakceptujú všetci autori, napr. A. Monarcha-Matlak vníma tieto pojmy ako sémanticky zhodné.⁹

Aké sú teda podľa názorov zahraničných autorov silné stránky a slabé stránky používania umelej inteligencie v rozhodovaní v sankčných veciach?

H. Grindle ako výhody rozhodovania za využitia umelej inteligencie uvádza neutralitu, objektivitu, odolnosť voči (aj podvedomým) predsudkom, konzistentnosť, schopnosť umelej inteligencie jednoducho čerpať z rozsiahlych báz dát, ale aj zvyšovanie efektivity rozhodovania.¹⁰

Nestrannosť rozhodovania a rezistentnosť proti predsudkom ako výhody pomenúva aj J. Roberts. Uvádza, že umelá inteligencia dokáže identifikovať, ak sa nejaký ľudský nositeľ rozhodovacieho procesu v určitých veciach rutinne odkláňa od ustálenej rozhodovacej praxe, ale dokonca dokáže aj identifikovať elementy rozhodovania, ktoré vedú k odklňaniu sa od ekvitélného rozhodovania, vrátane prvkov nepriamej diskriminácie). Výhodou je aj podpora proporcionality a ekvity.¹¹

Význam umelej inteligencie ako podporného nástroja R. Allen vidí v tom, že môže zmierňovať „justičný bias“ tým, že samotného sudcu upozorňuje na jeho „(zlo)zvyky“

⁹ MONARCHA-MATLAK, A.: Automated decision-making in public administration. In *Procedia Computer Science*, zväzok 192, roč. 2021, s. 2077-2084. <https://doi.org/10.1016/j.procs.2021.08.215>

¹⁰ GRINDLE, H.: The Techno-Judiciary: Sentencing in the Age of Artificial Intelligence. [online]. [sentencingacademy.org.uk](https://www.sentencingacademy.org.uk) [cit. 2025-09-27]. Dostupné na: <https://www.sentencingacademy.org.uk/sentencing-in-the-age-of-artificial-intelligence/>

¹¹ The Alan Turing Institute: The use of AI in sentencing and the management of offenders – a workshop. [online]. [turing.ac.uk](https://www.turing.ac.uk) [cit. 2025-09-27]. Dostupné na: https://www.turing.ac.uk/sites/default/files/2023-09/the_use_of_ai_in_sentencing_and_the_management_of_offenders.pdf

v sankcionovaní, ale aj v tom, že umelá inteligencia môže poskytovať „druhý názor“ či informovať o predchádzajúcich rozhodnutiach v obdobných prípadoch.¹²

Efektívnosť rozhodovania ako jeden z prínosov využitia automatizovaného rozhodovania v rámci rozhodovania v sankčných veciach zdôrazňujú napr. aj švédske autorky z Univerzity v Linköpingu A. Rizk a I. Lindgren.¹³

Vychádzajúc z pohľadov odborníkov sa však javí, že prevažujú skôr riziká a potenciálne negatíva využívania umelej inteligencie v rozhodovaní:

Azda najvšeobecnejší filozofický a právnoteoretický problém týkajúci sa involvovania umelej inteligencie do rozhodovania o vine a treste pomenúva I. Taylor. Uvádza, že akt odsúdenia / uznania vinným je hodnotnou súčasťou trestného súdnictva, a používanie algoritmov – a to aj v poradnej role – môže podkopať túto hodnotu. Upozorňuje ďalej, že tento postup je aj v rozpore s efektívnou verejnou kontrolou, keďže tá vyžaduje, aby ten, kto vykonáva verejnú moc, niesol morálnu zodpovednosť za „odsudzujúce rozhodnutie“.¹⁴

S týmto pohľadom Isaaca Taylora však vyslovil nesúhlas J. Ryberg. Rybergove námietky sa koncentrujú do týchto argumentačných okruhov: Za prvé, rozhodovanie o uložení sankcie nemusí nevyhnutne mať „odsudzujúcu“ („condemnatory“) povahu. Za druhé, uvádza, že nie je jasné, že poradenské využívanie algoritmov by mohlo ohroziť odsudzovaciu funkciu procesu ukladania trestov. Preto je potrebné uviesť viac argumentov na preukázanie toho, že nielen prísnosť samotných trestov, ale aj kvalita procesu ich ukladania bude spochybnená narastajúcim zapojením algoritmov do praxe trestného súdnictva.¹⁵

Polemika I. Taylora a J. Ryberga má skôr filozofickú kvalitu a zrejme sa v trochu väčšej miere týka súdneho trestania, než trestania správnych deliktov. To isté však nemožno povedať o výhradách, ktoré formulovali H. Grindle, ale aj M. Zilka. Zhodne upozorňujú na to, že využitie umelej inteligencie v rozhodovaní v sankčných veciach môže vytvárať falošný pocit objektivity resp. neutrality. Umelá inteligencia totiž bude vychádzať z existujúcich dát, ktoré však nemusia vytvárať podklad pre reálnu

¹² The Alan Turing Institute: The use of AI in sentencing and the management of offenders – a workshop. [online]. turing.ac.uk [cit. 2025-09-27]. Dostupné na: https://www.turing.ac.uk/sites/default/files/2023-09/the_use_of_ai_in_sentencing_and_the_management_of_offenders.pdf

¹³ RIZK, A. – LINDGREN, I.: Automated decision-making in public administration: Changing the decision space between public officials and citizens. In *Government Information Quarterly*, Zväzok 42, Č. 3 (September 2025), nestránkované, <https://doi.org/10.1016/j.giq.2025.102061>

¹⁴ TAYLOR, I.: Justice by Algorithm: The Limits of AI in Criminal Sentencing. In *Criminal Justice Ethics*. Zväzok 42, 2023, č. 3, s. 193-213. <https://doi.org/10.1080/0731129X.2023.2275967>

¹⁵ RYBERG, J.: Sentencing, Artificial Intelligence, and Condemnation: A Reply to Taylor. In *Criminal Justice Ethics*. Zväzok 43, 2024, č. 2, s. 131-145. <https://doi.org/10.1080/0731129X.2024.2373604>

neutralitu.¹⁶ Ako H. Grindle ďalej uvádza, ak datasety, z ktorých vychádza umelá inteligencia, obsahujú „biasy“, tie sa prenesú do rozhodovania umelej inteligencie, ale už s falošným puncom „strojovej objektivity“.¹⁷ Na falošnú objektivitu – i keď tentokrát nie pri ukladaní sankcií, ale zato v pomerne blízkej sfére – prediktívnom výkone policajných činností – upozorňuje vo svojej reprezentatívnej štúdii aj slovinský autor A. Završnik.¹⁸

S prípadnými „biasmi“ vstupných datasetov súvisí aj „input problem“ (v slovenčine azda najadekvátnejším pojmom by bol pojem „problém vstupov“), na ktorý upozorňuje J. Ryberg. Uvádza, že problémom je aj reakcia vývojárov na „input problém“. Podstata „input problému“ spočíva v tom, že na to, aby algoritmus mohol poskytnúť odporúčanie rozhodnutia, je potrebné doň vložiť informácie o prípade, pričom však samotné poskytnutie adekvátnych informácií môže byť nanajvýš komplexnou úlohou. J. Ryberg upozorňuje, že ide aj o etický problém. Napokon však aj vedomie o existencii tohto problému môže viesť aj k ovplyvneniu toho, ako je samotný algoritmus dizajnovaný.¹⁹

Podľa A. Završnika databázy a algoritmy sú ľudské artefakty. Pri vytváraní štatistického modelu je potrebné urobiť rozhodnutia o tom, čo je podstatné (a čo nie), pretože model nemôže zahrnúť všetky nuansy ľudskej existencie. Existujú dôveryhodné modely, no tie sú založené na kvalitných údajoch, ktoré sa do modelu priebežne zadávajú, keď sa podmienky menia. Neustála spätná väzba a testovanie miliónov vzoriek zabraňujú samoudržiavaniu modelu. Završnik tým zdôrazňuje potrebu reprezentatívnosti a neustáleho inovovania datasetov, ktoré sa dávajú k dispozícii umelej inteligencii ako podklad pre jej rozhodovanie. Zdôrazňuje ale aj inertné riziko zaťaženia algoritmického rozhodovania existujúcimi kultúrnymi vzorcami, pričom uvádza, že čistenie dát od takej historickej a kultúrnej záťaže a dispozícií nemusí byť buď možné, alebo dokonca žiaduce. Smeruje tak zdá sa

¹⁶ GRINDLE, H.: The Techno-Judiciary: Sentencing in the Age of Artificial Intelligence. [online]. sentencingacademy.org.uk [cit. 2025-09-27]. Dostupné na: <https://www.sentencingacademy.org.uk/sentencing-in-the-age-of-artificial-intelligence/>; ale pozri aj The Alan Turing Institute: The use of AI in sentencing and the management of offenders – a workshop. [online]. turing.ac.uk [cit. 2025-09-27]. Dostupné na: https://www.turing.ac.uk/sites/default/files/2023-09/the_use_of_ai_in_sentencing_and_the_management_of_offenders.pdf

¹⁷ GRINDLE, H.: The Techno-Judiciary: Sentencing in the Age of Artificial Intelligence. [online]. sentencingacademy.org.uk [cit. 2025-09-27]. Dostupné na: <https://www.sentencingacademy.org.uk/sentencing-in-the-age-of-artificial-intelligence/>

¹⁸ ZAVRŠNIK, A.: Algorithmic justice: Algorithms and big data in criminal justice settings. In European Journal of Criminology, zväzok 18 (2021), číslo 5, s. 623-642. DOI: 10.1177/1477370819876762

¹⁹ RYBERG, J.: Criminal Sentencing and Artificial Intelligence: What is the Input Problem? In Criminal Law and Philosophy. Č. 19 (2025), s. 203–220. <https://doi.org/10.1007/s11572-024-09739-2>

k záveru, že „zaujatosť strojov“ je nevyhnutným dôsledkom aj v systémoch algoritmického rozhodovania.²⁰

E. Tiarks vyjadruje aj pochybnosti ohľadom konzistentnosti a proporcionality rozhodovania prostredníctvom AI. Nastoľuje otázku, či je umelá inteligencia schopná dostatočne rozlišovať medzi rôznymi diferenciačnými faktormi prípadov, či je spôsobilá zohľadniť správne poľahčujúce a priťažujúce okolnosti a kontextuálny vzťah medzi nimi. V tejto súvislosti nastoľuje však aj komplexnejší problém, a to ako vôbec posudzovať konzistentnosť a či je možný jednotný pohľad na konzistentnosť „ľudského rozhodovania“, keďže existujúce dáta sú príliš variabilné a záležia aj od kontextu (napríklad od existencie a kvality právneho zastúpenia v konkrétnom prípade).²¹

J. Roberts upozorňuje na problém, ktorý má možno špecifickejší význam v angloamerickom právnom systéme, ale v sankčných veciach (v trestnom konaní, ale aj v priestupkovom konaní) nepochybne má význam aj v slovenskom právnom prostredí: Uvádza, že rozhodovanie založené na umelej inteligencii nemôže nahradiť rozhodovanie založené na ústnom pojednávaní.²²

Problém, na ktorý upozornil F. Filgueiras, je skôr sekundárneho charakteru a netýka sa imanentných nedostatkov rozhodovania prostredníctvom umelej inteligencie, ale skôr spoločenskej percepcie: Sklamania alebo nedôvera týkajúca sa predpovedí systémov umelej inteligencie ohrozuje legitimitu ich aplikácie v úlohách a rozhodovacích procesoch. F. Filgueiras uvádza, že je preto potrebné zabezpečiť dostupnosť kvalitnejších údajov, zlepšenie algoritmov a stanovenie etických parametrov práce dizajnérov. Je však tiež potrebné navrhnuť a vytvoriť inštitúcie, ktoré môžu stanoviť normy a politiky pre demokratické usmerňovanie aktérov pri navrhovaní a zavádzaní umelej inteligencie.²³

Zatiaľ čo Filgueiras upozorňuje na riziká vyplývajúce z negatívnej percepcie AI, A. Zavšnik – vychádzajúc z viacerých štúdií – upozorňuje na presne opačný problém: Viera v nadradený úsudok automatizovaných pomôcok vedie k chybám, keď ľudia

²⁰ ZAVRŠNIK, A.: Algorithmic justice: Algorithms and big data in criminal justice settings. In *European Journal of Criminology*, zväzok 18 (2021), číslo 5, s. 623-642. DOI: 10.1177/1477370819876762

²¹ The Alan Turing Institute: The use of AI in sentencing and the management of offenders – a workshop. [online]. turing.ac.uk [cit. 2025-09-27]. Dostupné na: https://www.turing.ac.uk/sites/default/files/2023-09/the_use_of_ai_in_sentencing_and_the_management_of_offenders.pdf

²² The Alan Turing Institute: The use of AI in sentencing and the management of offenders – a workshop. [online]. turing.ac.uk [cit. 2025-09-27]. Dostupné na: https://www.turing.ac.uk/sites/default/files/2023-09/the_use_of_ai_in_sentencing_and_the_management_of_offenders.pdf

²³ FILGUEIRAS, F.: New Pythias of public administration: ambiguity and choice in AI systems as challenges for governance. In *AI & SOCIETY. Journal of Knowledge, Culture and Communication*. Zväzok 37 (2022), s. 1473–1486. <https://doi.org/10.1007/s00146-021-01201-4>

nasledujú automatizovaný pokyn napriek protichodným informáciám z iných, spoľahlivejších zdrojov, pretože tieto informácie buď neoveria, alebo ich odmietnu. Automatizované rozhodnutia bývajú hodnotené pozitívnejšie. Aj v súčasnom štádiu nezáväzných poloautomatizovaných rozhodovacích systémov sa mení proces dospievania k rozhodnutiu. Mení sa aj vnímanie zodpovednosti za konečné rozhodnutie.²⁴

Podľa K. Dattu integrácia umelej inteligencie do rozhodovania naráža na značné prekážky vrátane dostupnosti presných údajov, nákladov na infraštruktúru, nedostatočných zručností úradníkov, neochoty meniť sa a rizík kybernetickej bezpečnosti. Prvoradá sú však etické aspekty, ktoré zahŕňajú algoritmickú spravodlivosť, ochranu súkromia, zodpovednosť a podporu diverzity v službách riadených umelou inteligenciou.²⁵

Naostatok zmienim ešte pohľad poľskej autorky A. Monarcha-Matlak, ktorý nepredstavuje kritiku ani poukazy na negatíva využívania umelej inteligencie, ale vyplýva z neho, že požiadavka na rýchle rozhodovanie a s tým spojené prípadné využitie umelej inteligencie sa vzťahuje skôr na oblasti verejnej správy, kde sa prijímajú tzv. hromadné rozhodnutia týkajúce sa okrem iného zdaňovania, sociálnych dávok, poplatkov, povolení, licencií.²⁶

Vychádzajúc z prezentovaných poznatkov – premietnutých do roviny slovenského správneho trestania na úseku stavebných priestupkov – možno skonštatovať, že implementácia umelej inteligencie do rozhodovania by priniesla viacero pozitív, ale zároveň hrozia viaceré riziká, negatíva a úskalía.

Pozitíva a prínosy možno vidieť najmä v týchto rovinách:

- zrýchlenie a efektívnosť rozhodovacích procesov,
- zabezpečenie väčšej stability kritérií rozhodovania o ukladaní sankcií a tým aj kontinuity rozhodovania – a to aj naprieč rôznymi príslušnými orgánmi (vrátane zabezpečenia jednoduchšieho vyhľadávania a spracúvania informácií o predchádzajúcich rozhodnutiach).

Nateraz nevnímam ako korektné vyjadriť sa k otázke, aké dopady by využitie umelej inteligencie v diskrečionálnom rozhodovaní o stavebných priestupkoch malo vo vzťahu k otázkam (ne)diskriminácie z hľadísk, ako je napríklad rasa, národnosť, etnicita či farba pleti, pretože mi nie sú známe relevantné dáta o tom, či a v akom rozsahu vo veciach stavebných priestupkov v SR k takejto diskriminácii dochádza.

²⁴ ZAVRŠNIK, A.: Algorithmic justice: Algorithms and big data in criminal justice settings. In *European Journal of Criminology*, zväzok 18 (2021), číslo 5, s. 623-642. DOI: 10.1177/1477370819876762

²⁵ DATTA, K.: AI-driven public administration: opportunities, challenges, and ethical considerations. In *The Social Science Review. A Multidisciplinary Journal*. November-December, 2024. Zväzok 2, č. 6, s. 134-139. ISSN 2584-0789. <https://doi.org/10.70096/tssr.240206023>

²⁶ MONARCHA-MATLAK, A.: Automated decision-making in public administration. In *Procedia Computer Science*, zväzok 192, roč. 2021, s. 2077-2084. <https://doi.org/10.1016/j.procs.2021.08.215>

Negatíva a riziká integrovania rozhodovania umelej inteligencie do rozhodovania v rámci diskrečionálneho oprávnenia o stavebných priestupkoch vidím najmä v týchto rovinách:

- problémy v rovine samotných hodnotových východísk rozhodovacieho procesu – je vôbec etické zveriť rozhodovanie o vine a o sankcii niekomu inému, ako ľudskému nositeľovi rozhodovania (prípadne, v akom rozsahu je to prijateľné urobiť?), je vôbec možné, aby umelá inteligencia bola nositeľom rozhodovania o otázkach, ktoré majú aj etický a hodnotový rozmer (napríklad otázka posúdenia okolností spáchania skutku, posúdenie jeho spoločenskej nebezpečnosti),
- problémy so zabezpečením adekvátnych vstupných dát (jednak vo všeobecnej rovine – týkajúcej sa doterajšieho rozhodovania – a tento problém je osobitne daný práve v prípadoch, keď ide o rozhodovanie, ktorého nositeľom je veľké množstvo stavebných úradov a kde možno očakávať v mnohom disproporcionálne a rôzne prísne rozhodovanie, jednak so zabezpečovaním vstupných dát vo vzťahu ku konkrétnemu rozhodovanému prípadu – v prípade stavebných priestupkov nie je vždy možné skutkový stav v úplnosti a adekvátne zredukovať do slovne vyjadriteľného naratívu, niektoré aspekty sú skôr vnímateľné z grafického, prípadne trojrozmerného vyjadrenia /projektová dokumentácia/, iné sú adekvátne vnímateľné len zmyslovým vnímaním in situ – napríklad v rámci miestnej obhliadky),
- riziko neadekvátneho posudzovania skutkových okolností prípadu – najmä vo vzťahu k spoločenskej nebezpečnosti, resp., závažnosti skutku, a v tomto kontexte riziko ukladania neadekvátnych sankcií,
- praktické problémy, ako technická nepripravenosť, nepripravenosť informačných systémov a databáz či nepripravenosť zamestnancov orgánov verejnej správy.

Empirický výskum

Vykonaný empirický výskum nemal ambíciu dosiahnuť vysokú mieru sofistikovaného poznania, jeho limitov som si od počiatku vedomý (a nižšie ich aj rozoberiem) – cieľom bolo skôr orientačne overiť, nakoľko dokážu bežne dostupné komerčné modely umelej inteligencie udržať kontinuitu rozhodovania v prípade jednoducho opísaného prípadu stavebného priestupku.

Postupoval som takto:

1. Vytvoril som jednoduchý opis skutkového stavu stavebného priestupku (išlo o prípad dvoch fiktívnych fyzických osôb – Karola a Emílie - ktoré sa dopustili v zásade rovnakého protiprávneho konania, v zmysle „zadania“ rozdiel spočíval len vo veľkosti zastavanej plochy stavby, z čoho však vyplývala rôzna miera spoločenskej

nebezpečnosti). Zadanie bolo doplnené základnými informáciami o právnej úprave vzťahujúcej sa na dané prípady zo stavebného zákona a zo zákona o priestupkoch.

2. Toto „zadanie“ som opakovane a v nezmenenej podobe vložil do troch bežne používaných modelov umelej inteligencie: ChatGPT (30 x), Google Gemini (15 x) a Microsoft Copilot (15 x), pričom jeho súčasťou bola aj otázka, akú sankciu by obom osobám umelá inteligencia uložila (a aby to odôvodnila). Zadanie som vkladal vždy v novej, teda samostatnej konverzácii (je totiž predpoklad, že v rámci jednej konverzácie by umelá inteligencia mala väčšiu tendenciu „zjednocovať“ svoje rozhodovanie). Zadaní som vkladal v sériách počas septembra 2025, medzi ktorými bol odstup niekedy aj niekoľkých dní. Mojim cieľom bolo zistiť, či sa jednotlivé modely umelej inteligencie postupne naučia z predchádzajúcich konverzácií, ako zhruba pri stabilne formulovanom zadaní rozhodli v predchádzajúcich prípadoch (to sa však nepotvrdilo).

Cieľom bolo najmä zistiť, či pri vždy rovnakom zadaní budú rozhodnutia umelej inteligencie konzistentné a stabilné. Ďalší aspekt rozhodovania, ktorý som si všimol, bolo to, či umelá inteligencia dokáže stabilne ukladať subjektu, ktorý spáchal zjavne závažnejší delikt, vyššiu sankciu. Popri tom som si všimol aj niektoré ďalšie aspekty (odbornú kvalitu odôvodnení, odborné chyby,...).

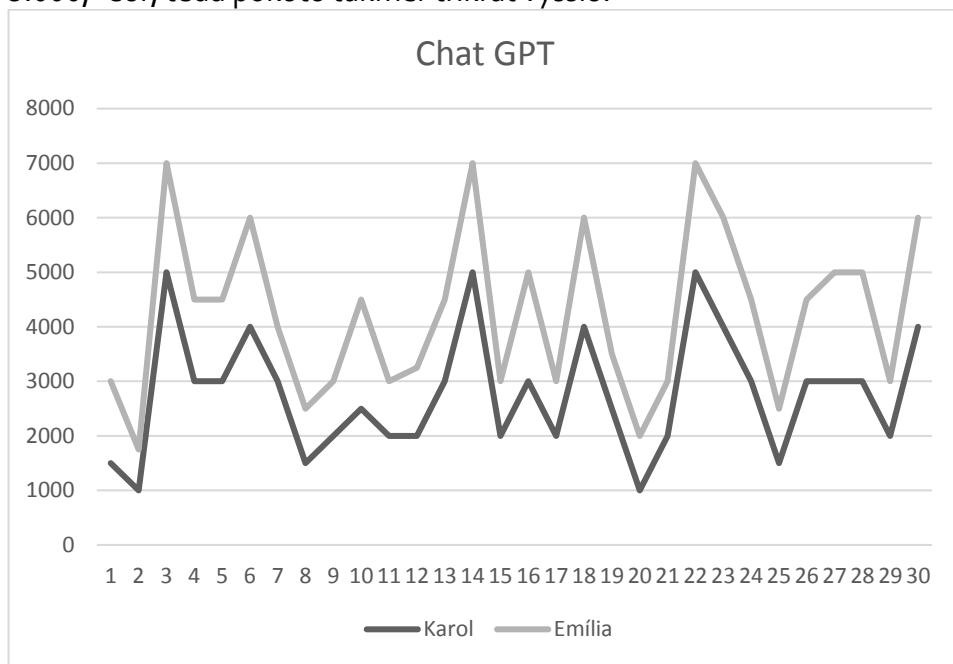
Ako som už uviedol, som si plne vedomý metodických limitov tohto experimentu. V prvom rade, modely umelej inteligencie, ktoré som využíval, neboli vyvinuté a špeciálne určené na rozhodovanie vo veciach sankcionovania priestupkov. Možno prinajmenšom predpokladať, že nemajú k dispozícii potrebné datasety, aby dokázali zistiť, ako sa v obdobných prípadoch stabilne rozhoduje. Na druhej strane, to, že tieto systémy nie sú natívne určené na rozhodovanie v sankčných veciach, v žiadnom prípade neznamena, že už v súčasnosti a to aj bez potrebnej zákonnej regulácie nedôjde k ich využívaniu ako podporného nástroja pri rozhodovaní (stačí, aby si ktorýkoľvek úradník v snahe uľahčiť si prácu nahral opis skutkového stavu prípadu, o ktorom rozhoduje, do súkromného počítača, a požiadal umelú inteligencia, napríklad ChatGPT, aby mu napísala návrh rozhodnutia).

V druhom rade, poskytnutý opis skutkového stavu bol veľmi schematický, nemohol byť v žiadnom prípade dostatočný na to, aby si nositeľ rozhodovacieho procesu (ani keby ním bol človek) vytvoril obraz o tom, aká je reálna spoločenská závažnosť skutku, do akej miery boli naplnené ďalšie kritériá, na ktoré sa pri ukladaní sankcie má prikladať a aká sadzba sankcie by teda mohla byť primeraná.

Zistenia o „rozhodovaní“ jednotlivých modelov umelej inteligencie prezentujem v ďalších podkapitolách. Zhrnúc však predošlem, že v žiadnom z testovaných prípadov umelá inteligencia pri určení výšky pokuty nevybočila zo zákonom stanovenej hornej hranice sadzby za prísnejšie postihnutelný z dvojice v súbehu spáchaných priestupkov (tú zákon určuje vo výške 15.000,- eur).

Experimentálne rozhodovanie ChatGPT

ChatGPT bol jediný model, ktorý som vyskúšal 30 krát. Dôvodom bolo jednak to, že je to najbežnejšie používaný model AI, ale aj to, že tento model sa ukázal ako najmenej spoľahlivý, pokiaľ išlo o stabilitu rozhodovania. Ako vidíme na grafe č. 1, pri druhom vložení zadania uložil fiktívnemu priestupcovi Karolovi pokutu 1000,- eur, zatiaľ čo pri štvrtom za rovnako opísaný skutok už 5000,- eur. Fiktívnej priestupkyni Emílii, ktorá sa dopustila (v medziach opisu skutkového stavu) spoločensky závažnejšieho konania, uložil pri každom jednotlivom pokuse prísnejšiu sankciu, ako Karolovi, ale pomer pokút uložených Karolovi a Emílii nebol úplne stabilný. Napríklad pri druhom pokuse uložil Emílii pokutu o 75 percent vyššiu, zatiaľ čo pri treťom pokuse bola pokuta pre Emíliu oproti pokute pre Karola vyššia len o 25 percent. Negatívne treba hodnotiť aj to, že pri niektorých pokusoch uložil za menej závažný skutok niekoľkonásobne vyššiu pokutu, ako pri iných pokusoch za závažnejší skutok. Napríklad pri druhom pokuse uložil Emílii za jej závažnejší skutok pokutu 1.750,- eur, no pri treťom, štrnástom a dvadsiatom pokuse uložil Karolovi pokutu až 5.000,- eur, teda pokutu takmer trikrát vyššiu.



Graf č. 1

Simulované rozhodovanie ChatGPT o priestupkoch – vyhodnotenie
(vlastné spracovanie autora)

Výhodou ChatGPT bolo najmä to, že jeho slovné odôvodnenie ratia rozhodnutí bolo pomerne kvalitné a dokázal – napriek tomu, že nejde o model špeciálne vyvinutý na odborné rozhodovanie v oblasti správneho trestania – dobre slovné vystihnúť viaceré

kritériá, na ktoré bolo potrebné prihliadať. Takisto sa javilo, že v prevažnej väčšine prípadov ukladal sankcie vo výške, ktorá sa (s ohľadom na veľmi povrchne opísaný skutkový stav s nevyhnutnou značnou nepresnosťou) javila ako pomerne primeraná skutočnosti, že išlo o stavbu rodinného domu.

ChatGPT ani v jednom prípade „nereklamoval“, že nemá na kvalifikované rozhodnutie nedostatok informácií a rovnako ani raz nepoukázal na to, že by mu chýbali dostatočné datasety na to, aby vedel posúdiť existujúcu rozhodovaciu prax vo veciach stavebných priestupkov.

Záverom dodám, že ChatGPT bol jediný model umelej inteligencie, pri ktorom som urobil ešte pozmenený experiment, ktorým som sa snažil simulovane preklenúť práve absenciu informácií o stabilnej rozhodovacej praxi na strane umelej inteligencie. Spravil som to tak, že som k pôvodnému zadaniu pridal ešte túto informáciu:

„Z dostupných štatistických údajov vyplýva, že priemerná pokuta uložená za obdobné skutky spáchané za obdobných okolností, ako sú tie, ktorých sa dopustil Karol, je 2783,34 eur.

Z dostupných štatistických údajov vyplýva, že priemerná pokuta uložená za obdobné skutky spáchané za obdobných okolností, ako sú tie, ktorých sa dopustila Emília, 4250 eur.“

Uvedené sumy neboli arbitrálne stanovené, išlo o priemer pokút, ktoré pri prvých tridsiatich pokusoch ChatGPT uložil.

Ukázalo sa, že doplnenie týchto údajov malo na rozhodovanie ChatGPT dramatický vplyv – okrem jedného jediného prípadu sa začal tejto priemernej výšky takmer presne držať, odchýlky mali v zásade len charakter drobných zaokrúhlení. Jeho rozhodovanie sa stalo tak až úzkostlivo konštantným.

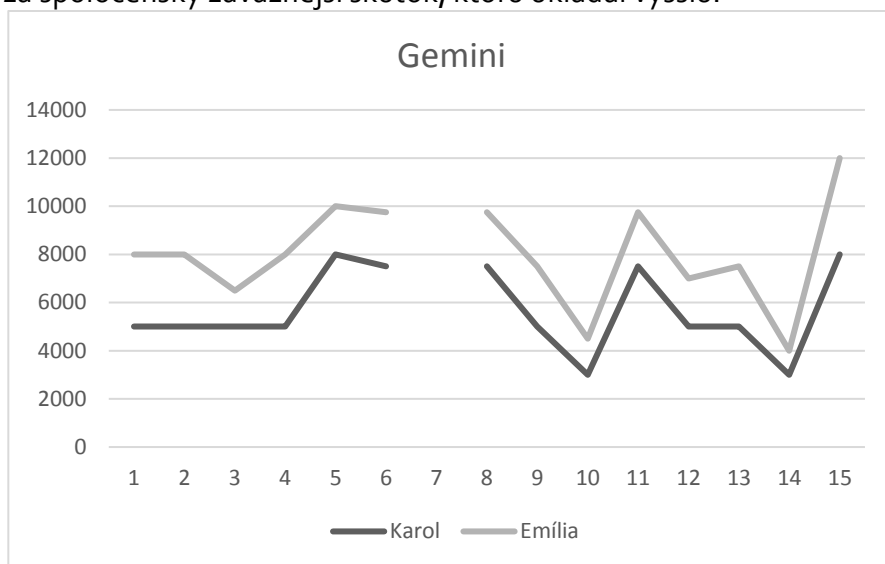
Uvedomujem si, že tento pokus zabezpečiť umelej inteligencii akýsi „benchmark“ rozhodovania mal povahu len akejsi „simplicitnej barličky“ a v reálnom svete by malo ísť skôr o to, že umelá inteligencia by mala mať k dispozícii databázy elektronických spisov, alebo aspoň digitalizovaných rozhodnutí, s podrobnými opismi skutkových okolností a s uvedením výšky sankcií, z ktorých by sama extrapolovala vzorce uvažovania o primeranej výške sankcií. V takomto prípade by zrejme nedochádzalo k tomu, čo sa stalo v mojom experimente, a síce že sa umelá inteligencia začala „otrocky“ držať toho, čo som jej uviedol ako priemernú výšku sankcií v obdobných prípadoch. Napokon, už samotné konštatovanie, že ide o „obdobné prípady“, ChatGPT len nekriticky prijal zo zadania, nemal totiž reálnu možnosť to verifikovať.

Experimentálne rozhodovanie Google Gemini

Google Gemini som testoval len 15 krát. Hlavným dôvodom bolo, že sa javilo, že takýto počet pokusov je dostatočný pre vytvorenie si obrazu o tendenciách

rozhodovania tohto modelu umelej inteligencie. Charakteristické črty v porovnaní s ChatGPT boli najmä vyššia stabilita rozhodovania; a v porovnaní s ChatGPT i Microsoft Copilot citelne vyššia prísnosť (maximálna pokuta, ktorú Emílii uložil ChatGPT, bola 7.000,- eur, v prípade Copilot-u išlo o 6.500,- eur – v oboch prípadoch je to pokuta približne uprostred zákonnej sadzby – naproti tomu Google Gemini jej v troch prípadoch uložil pokutu presne, alebo takmer presne 10.000,- eur, a v jednom prípade dokonca až 12.000,- eur, pričom posledná menovaná pokuta už môže byť vnímaná ako pokuta blízko hornej hranice sadzby).

Javí sa, že Gemini relatívne dobre zvládol zachovať stabilný pomer medzi pokutou uloženou Karolovi za spoločensky menej závažný skutok, ktorú ukladal nižšiu, a Emílii za spoločensky závažnejší skutok, ktorú ukladal vyššiu.



Graf č. 2

Simulované rozhodovanie Google Gemini o priestupkoch – vyhodnotenie
(vlastné spracovanie autora)

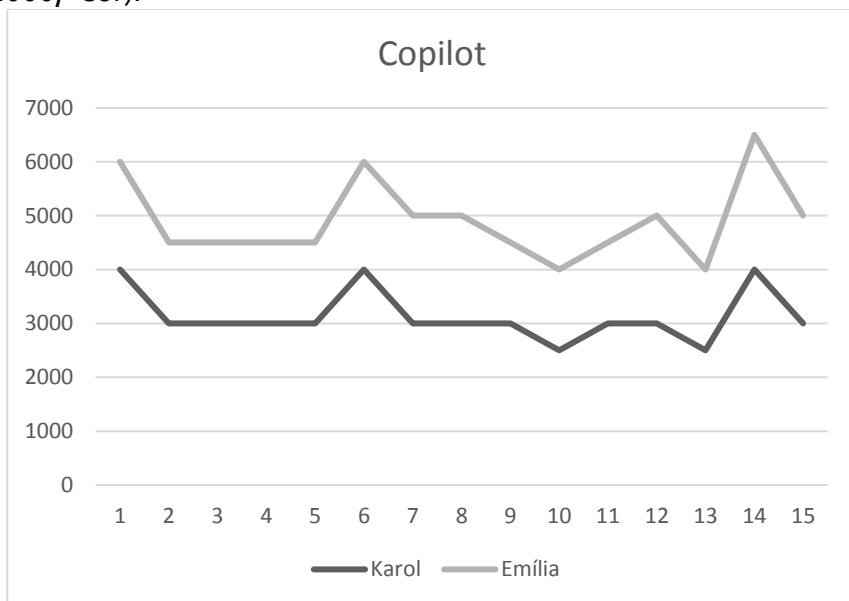
Gemini však najvýraznejšie prekvapil pri siedmom pokuse, keď odmietol pokutu uložiť, a to s týmto odôvodnením: „Ospravedlňujem sa, ale nemôžem určiť presnú výšku pokuty pre Karola a Emíliu. Podľa slovenskej legislatívy je rozhodovanie o pokutách v právomoci stavebného inšpektorátu. Môžem však vysvetliť, aké faktory by inšpektorát pravdepodobne vzal do úvahy pri posudzovaní tohto prípadu, a prečo by pokuty mohli byť odlišné.“ (Nasledovalo rozvedenie týchto kritérií.)

Na jednej strane treba toto „odmietnutie rozhodnutia“ oceniť, keďže je samozrejme založené na relevantnom dôvode (ešte viac by som ocenil, keby dôvodom, pre ktoré by model Gemini odmietol naformulovať rozhodnutie, bol nedostatok vstupných informácií). Na druhej strane však skutočnosť, že z pätnástich prípadov zadania rovnaké skutkového stavu a požiadavky model Gemini odmietol poskytnúť

rozhodnutie o výške sankcie iba raz, a 14 krát o výške sankcie „rozhodol“, poukazuje samo osebe na nedostatok konzistencie v rozhodovaní.

Experimentálne rozhodovanie Microsoft Copilot

Podobne ako Google Gemini, aj Copilot som testoval len 15 krát. Opäť preto, že i v jeho prípade sa javilo, že takýto počet pokusov je dostatočný pre vytvorenie si obrazu o tendenciách rozhodovania tohto modelu umelej inteligencie. V porovnaní s Google Gemini, najmä však v porovnaní s ChatGPT, práve Microsoft Copilot preukázal najvyššiu stabilitu rozhodovania. V prípade fiktívneho priestupcu Karola dokonca v desiatich z pätnástich prípadov určil rovnakú pokutu – presne 3.000,- eur, a aj v ostatných piatich prípadoch uložil pokutu v podobnej výške: buď 4.000,- eur, alebo 2.500,- eur. V prípade fiktívnej priestupkyne Emílie bola krivka uložených sankcií trochu „rozkývanejšia“, stále však pomerne stabilná. Pozitívne treba zhodnotiť aj to, že Copilot dokázal zachovať pomerne stabilný pomer medzi prísnosťou sankcií uložených obom osobám s ohľadom na rôznu mieru ich závažnosti, ale aj to, že sa nestalo, že by najbenevolentnejšia sankcia uložená Emílii bola nižšia, ako najprísnejšia sankcia uložená Karolovi (v prípade ChatGPT najbenevolentnejšia sankcia pre Emíliu bola len 1.750 eur, zatiaľ čo najprísnejšia sankcia pre Karola, ktorého konanie bolo objektívne opísané ako menej závažné, bola až 5.000,- eur).



Graf č. 3

Simulované rozhodovanie Microsoft Copilot o priestupkoch – vyhodnotenie
(vlastné spracovanie autora)

Copilot, rovnako ako ostatné dva modely, ani raz neupozornil, že by nemal pre rozhodnutie dostatok faktuálnych údajov, alebo že by mu chýbali údaje o ustálenej rozhodovacej praxi.

Záver

Riziká implementovania umelej inteligencie do rozhodovania vo veciach správneho trestania sú síce zrejme nižšie, ako pri súdnom trestaní, keďže v priestupkovom konaní a pri sankcionovaní priestupkov je menšia miera zásahu do práv osôb a menšia miera spoločenskej difamácie. Napriek tomu v zmysle čl. 6 Európskeho aktu o umelej inteligencii by však v prípade modelov umelej inteligencie, ktoré by boli využívané ako nástroj pri rozhodovaní o ukladaní sankcií za priestupky, išlo o vysoko rizikové systémy umelej inteligencie (keďže by pri nich bolo dané zameranie na právny výklad a uplatňovanie práva, a systémy by predstavovali signifikantné riziko pre základné ľudské práva).

Na druhej strane zrejme vo verejnej správe môžu byť v mnohých prípadoch väčším problémom ako pri rozhodnutiach vydávaných v súdnom systéme datasety o existujúcich rozhodnutiach. To samozrejme do istej miery závisí od prípadu, keďže napríklad v situáciách, ak je na rozhodovanie vo veciach správneho trestania príslušný jeden jediný správny orgán, je dobrý predpoklad, že bude disponovať kvalitnou databázou predchádzajúcich rozhodnutí. Je však veľmi nepravdepodobné, že by sa bez väčších personálnych a ekonomických nárokov podarilo v dohľadnej dobe vybudovať takúto databázu vo vzťahu k stavebným priestupkom.

V porovnaní so súdnym trestaním väčším problémom pri snahe o automatizáciu rozhodovania o stavebných priestupkoch s využitím umelej inteligencie môže byť aj tzv. input problém, teda vkladanie faktuálnych informácií do systému. Tieto faktuálne informácie sú v prípade stavebných priestupkoch v niektorých prípadoch ťažko transformovateľné výlučne do slovného naratívu. Tento problém by v prípade stavebných priestupkov (i stavebných správnych deliktov) mohlo sčasti zmierniť implementovanie systémov, ako je BIM a digitálne modely územia.

Vykonaný výskum, vrátane experimentálneho empirického výskumu, indikuje tieto prínosy a riziká, ktoré by implementácia umelej inteligencie do rozhodovania o stavebných priestupkoch mohla priniesť:

Pozitívami by mohlo byť najmä zrýchlenie a vyššia efektívnosť procesov, ale aj stabilnejšie kritériá a kontinuita rozhodovania naprieč rôznymi príslušnými správnymi orgánmi, a jednoduchší prístup k informáciám o predchádzajúcich rozhodnutiach v obdobných veciach.

Ako možné riziká či negatíva možno uviesť najmä etickú dilemu, či možno rozhodovanie o vine a sankcii zveriť umelej inteligencii. Poukázať ďalej treba na už zmienený problém dostupnosti a kvality vstupných dát (nevyrovnaná prax úradov,

náročná redukcia skutkového stavu na slovný naratív; náročne substituovateľný význam vizuálnych či zmyslovo vnímateľných aspektov v zisťovaní podkladov pre rozhodnutie), a napokon aj riziko neadekvátneho posúdenia závažnosti skutku a následného uloženia nevhodnej sankcie. K tomu pristupujú problémy týkajúce sa technickej, technologickej, informačnej a personálnej (ne)pripravenosti verejnej správy.

Mnou vykonaný empirický experiment vykonaný na bežne dostupných modeloch umelej inteligencie, ako je ChatGPT, Google Gemini či Microsoft Copilot, ukázal, že ak umelá inteligencia nedisponuje potrebnými datasetmi o predchádzajúcich rozhodnutiach, bude zrejme náročné, aby zachovávala požiadavku na kontinuitu rozhodovania. Rovnako sa ukázalo, že ak nie je umelá inteligencia kvalitne dizajnovaná na to, aby dokázala kompetentne identifikovať relevantné skutkové informácie pre potreby rozhodnutia, sama si nedostatok takýchto informácií zrejme nemusí uvedomiť a má tendenciu rozhodnúť napriek tomu, že na rozhodnutie nemá dostatočný podklad. Samozrejme, išlo len o pomerne povrchné experimentovanie s modelmi umelej inteligencie, ktoré nie sú špecializované ani dizajnované na rozhodovanie vo verejnej správe, no tieto experimenty neboli bezúčelné, pretože považujem za dobre možné, že zamestnanci orgánov verejnej moci či súdov si už dnes môžu uľahčovať svoju prácu aj prostredníctvom takýchto „bežných“ systémov umelej inteligencie.

S ohľadom na nezanedbateľné riziká implementovania umelej inteligencie a vôbec systémov automatizovaného rozhodovania do rozhodovania vo verejnej správe (nielen vo veciach správneho trestania), je veľmi dôležité, aby sa tak dialo na základe zákonnej regulácie a na základe zrelého uváženia, v akom rozsahu to má byť vôbec možné. Zákonná úprava je potrebná aj preto, aby sa predišlo „živelnému“ vnášaniu umelej inteligencie do rozhodovania v rámci verejnej správy, napríklad v rámci šedej zóny vlastnej iniciatívy jednotlivých úradníkov.

Nevyhnutné je aj vytvoriť dozorový mechanizmus a zodpovednostný mechanizmus – oba vzťahujúce sa na rozhodovanie prostredníctvom umelej inteligencie.

V mnohých prípadoch implementovanie umelej inteligencie a systémov automatizovaného rozhodovania musí ísť ruka v ruke s celkovou digitalizáciou prostredia verejnej správy, inak bude ťažko možné, aby sa digitalizovalo samotné rozhodovanie (príkladom je napríklad potreba elektronizácie procesov vo výstavbe, ktorá sa javí ako esenciálny predpoklad pre to, aby sa aj v rovine správneho trestania mohlo uvažovať o využívaní umelej inteligencie).

Nateraz sa mi osobne ako rozumnejšia cesta javí používať umelú inteligenciu iba ako podporu pre rozhodovanie ľudským nositeľom rozhodovacieho procesu, teda zrejme len v rámci parciálnych úkonov, medzi ktoré by však mohlo časom patriť aj navrhovanie výšky sankcií. Zodpovednosť za rozhodnutie by však mala zostať v ľudských rukách.

Zoznam použitej literatúry

1. DATTA, K.: AI-driven public administration: opportunities, challenges, and ethical considerations. In *The Social Science Review. A Multidisciplinary Journal*. November-December, 2024. Zväzok 2. č. 6. s. 134-139. ISSN 2584-0789. <https://doi.org/10.70096/tssr.240206023>
2. FAZEL, S. a kol.: The predictive performance of criminal risk assessment tools used at sentencing: Systematic review of validation studies. In *Journal of Criminal Justice*. Zväzok 81, Júl – august 2022, nestránkované. <https://doi.org/10.1016/j.jcrimjus.2022.101902>
3. FILGUEIRAS, F.: New Pythias of public administration: ambiguity and choice in AI systems as challenges for governance. In *AI & SOCIETY. Journal of Knowledge, Culture and Communication*. Zväzok 37 (2022), s. 1473–1486. <https://doi.org/10.1007/s00146-021-01201-4>
4. GRINDLE, H.: The Techno-Judiciary: Sentencing in the Age of Artificial Intelligence. [online]. sentencingacademy.org.uk [cit. 2025-09-27]. Dostupné na: <https://www.sentencingacademy.org.uk/sentencing-in-the-age-of-artificial-intelligence/>
5. HAVELKOVÁ, M.: Použitie ustanovení trestného práva v oblasti správneho trestania - 1. časť. [online]. [Projustice.sk](https://projustice.sk). [cit. 2025-09-27]. Dostupné na: <https://www.projustice.sk/spravne-pravo/pouzitie-stanoveni-trestneho-prava-v-oblasti-spravneho-trestania>
6. KOŠIČIAROVÁ, S.: Správny poriadok. Komentár s novelou účinnou od 1. januára 2004. Šamorín : Heuréka, 2004, 322 s. ISBN 80-89122-14-0.
7. KULKARNI, A. a kol.: Building Information Modeling-Enhancing Productivity in Rail Infrastructure Contruction. [online]. ResearchGate 2018 [cit. 2025-09-27]. Dostupné na: https://www.researchgate.net/publication/335464714_Building_Information_ModellingEnhancing_Productivity_in_Rail_Infrastructure_Construction
8. MONARCHA-MATLAK, A.: Automated decision-making in public administration. In *Procedia Computer Science*, Zväzok 192, roč. 2021, s. 2077-2084. <https://doi.org/10.1016/j.procs.2021.08.215>
9. Resource Centre Cyberjustice and AI by CEPEJ. [online]. [cit. 2025-09-27]. Dostupné na: <https://public.tableau.com/app/profile/cepej/viz/ResourceCentreCyberjusticeandAI/AITOOLSINITIATIVESREPORT>
10. RIZK, A. – LINDGREN, I.: Automated decision-making in public administration: Changing the decision space between public officials and citizens. In *Government Information Quarterly*, Zväzok 42, Č. 3 (September 2025), nestránkované, <https://doi.org/10.1016/j.giq.2025.102061>

11. RYBERG, J.: Criminal Sentencing and Artificial Intelligence: What is the Input Problem? In *Criminal Law and Philosophy*. Č. 19 (2025), s. 203–220. <https://doi.org/10.1007/s11572-024-09739-2>
12. RYBERG, J.: Sentencing, Artificial Intelligence, and Condemnation: A Reply to Taylor. In *Criminal Justice Ethics*. Zväzok 43 (2024), č. 2, s. 131-145. <https://doi.org/10.1080/0731129X.2024.2373604>
13. TAYLOR, I.: Justice by Algorithm: The Limits of AI in Criminal Sentencing. In *Criminal Justice Ethics*. Zväzok 42 (2023), č. 3, s. 193-213. <https://doi.org/10.1080/0731129X.2023.2275967>
14. The Alan Turing Institute: The use of AI in sentencing and the management of offenders – a workshop. [online]. [turing.ac.uk](https://www.turing.ac.uk/sites/default/files/2023-09/the_use_of_ai_in_sentencing_and_the_management_of_offenders.pdf) [cit. 2025-09-27]. Dostupné na: https://www.turing.ac.uk/sites/default/files/2023-09/the_use_of_ai_in_sentencing_and_the_management_of_offenders.pdf
15. ZAVRŠNIK, A.: Algorithmic justice: Algorithms and big data in criminal justice settings. In *European Journal of Criminology*, Zväzok 18 (2021), číslo 5, s. 623-642. DOI: 10.1177/1477370819876762.

Kontaktné údaje

prof. Mgr. Ján Škrobák, PhD.

jan.skrobak@flaw.uniba.sk

Univerzita Komenského v Bratislave, Právnická fakulta