

SCHOPNOSŤ VEĽKÝCH JAZYKOVÝCH MODELOV LOGICKY USUDZOVAŤ

Mgr. Marek Čiernik

Univerzita Komenského v Bratislave, Právnická fakulta
Katedra občianskeho práva
marek.ciernik@flaw.uniba.sk

Schopnosť veľkých jazykových modelov logicky usudzovať

Rozmach veľkých jazykových modelov (LLMs) priniesol do právnej vedy novú vlnu diskusií o tom, do akej miery sú tieto systémy schopné nahradiť ľudské premýšľanie, predovšetkým logické usudzovanie. V článku sa skúma, či LLMs dokážu konzistentne aplikovať základné princípy formálnej logiky pri interpretácii a subsumpcii právnych noriem. Metodologicky je výskum postavený na sérii experimentov, ktoré využívajú päť typov typických logických chýb v právnom usudzovaní. Testované boli viaceré modely (ChatGPT-5, Gemini 2.5 Pro, Copilot, Grok 4 a Perplexity), a to na dvoch úrovniach vstupov: laických a odborníckych promptoch. Výsledky ukazujú opakujúce sa logické nekonzistencie naprieč všetkými modelmi, pričom presvedčivejšie výstupy sa dosahujú len pri explicitne formulovaných odborných zadaniach. Zistenia podčiarkujú zásadný rozdiel medzi štatistickým generovaním textu a skutočným právnym myslením, ktoré si vyžaduje logickú dedukciu a transparentnosť. Záverom sa konštatuje, že LLMs síce môžu byť užitočnými nástrojmi na podporu právnikov, no zatiaľ nie sú schopné plnohodnotne nahradiť ich prácu.

The ability of large language models to reason logically

The rise of large language models (LLMs) has sparked a new wave of debate in legal scholarship about the extent to which these systems are capable of replacing human thinking, particularly logical reasoning. This article examines whether LLMs can consistently apply the basic principles of formal logic in the interpretation and subsumption of legal norms. Methodologically, the research is based on a series of experiments that use five types of typical logical errors in legal reasoning. Several models (ChatGPT-5, Gemini 2.5 Pro, Copilot, Grok

4, and Perplexity) were tested at two levels of input: lay and expert prompts. The results show recurring logical inconsistencies across all models, with more convincing outputs only achieved for explicitly formulated expert tasks. The findings underscore the fundamental difference between statistical text generation and actual legal reasoning, which requires logical deduction and transparency. In conclusion, while LLMs can be useful tools to support lawyers, they are not yet capable of fully replacing their work.

Kľúčové slová: umelá inteligencia, logika, usudzovanie

Key words: artificial intelligence, logic, reasoning

<https://doi.org/10.46282/afi.2025.2.04>

Úvod

Umelá inteligencia (AI) sa v poslednom desaťročí stala jedným z najfrekvencovanejších pojmov nielen v technologickej oblasti, ale aj v spoločenských a právnych vedách. Jej rýchly vývoj vyvoláva nielen očakávania spojené so zefektívnením ľudskej činnosti, ale aj obavy z možných negatívnych dôsledkov, vrátane ohrozenia tradičných profesií. Najmä veľké jazykové modely (LLMs), ako špecifická kategória generatívnej umelej inteligencie, dnes disponujú schopnosťami, ktoré dokážu navodzovať dojem autentického „premýšľania“. To prirodzene kladie otázku, či sú tieto systémy schopné logicky usudzovať spôsobom porovnateľným s človekom – a či teda v budúcnosti môžu nahradiť aj právnik.

Právne myslenie sa vo veľkej miere opiera o schopnosť vyvodzovať závery z právnych noriem prostredníctvom dedukcie, rozpoznávať logické vzťahy a udržiavať konzistentnosť argumentácie. Kým štatistické generovanie textu, na ktorom LLMs stoja, umožňuje plynulú jazykovú produkciu, otázne zostáva, či takýto proces možno považovať za logické usudzovanie v zmysle právnej metodológie. Diskusia sa ďalej komplikuje fenoménmi ako tzv. „black box efekt“, keďže vnútorné fungovanie modelov je pre človeka netransparentné, alebo „sycophancy“, teda tendencia modelov potvrdzovať predsudky používateľa.

Cieľom tohto článku je analyzovať hranice deduktívneho usudzovania veľkých jazykových modelov v právnom kontexte. Prostredníctvom experimentov, ktoré simulujú typické logické chyby v právnych

normách, testujeme, či a do akej miery sú LLMs schopné prekonať tieto nástrahy. Zároveň sa snažíme ukázať, aké dôsledky majú naše zistenia pre diskusiu o budúcnosti právnických profesií a o mieste umelej inteligencie v právnej praxi. Tento článok nadväzuje na predchádzajúci výskum, ktorý rozširuje a zlepšuje niektoré jeho nedostatky.¹

1 Umelá inteligencia a jej schopnosť premýšľať

Pred samotnou analýzou schopnosti umelej inteligencie premýšľať je potrebné vymedziť základné pojmy. Pojem umelá inteligencia (AI) patrí k najširším a najdiskutovanejším v súčasnej informatike i spoločenských vedách. Slovníkové definície ho spravidla opisujú ako využívanie alebo skúmanie počítačových systémov a strojov, ktoré disponujú niektorými vlastnosťami pripisovanými ľudskému mozgu – napríklad schopnosťou interpretovať a generovať jazyk zdanlivo ľudským spôsobom, rozpoznávať obrazy, riešiť problémy či učiť sa z dát, ktoré im sú poskytované.

Na konkrétnejšej úrovni možno hovoriť o generatívnej umelej inteligencii. Tento pojem odkazuje na systémy schopné vytvárať originálny obsah – text, obraz, zvuk, video alebo softvérový kód – ako odpoveď na podnet používateľa. OECD a ďalšie autoritatívne zdroje definujú generatívnu AI v zásade totožne.² Ešte špecifickejšou kategóriou sú veľké jazykové modely (Large Language Models, LLMs), ktoré predstavujú typ generatívnej AI trénovaný na masívnych korpusoch textu s cieľom produkovať nové jazykové výstupy. Pre zjednodušenie budeme pojmy AI a LLM používať ako synonymá, hoci ide o odlišné kategórie.

Otázka, či AI dokáže „premýšľať“, je predmetom intenzívnych diskusií nielen v informatike, ale aj vo filozofii a právnej teórii. V právnom diskurze sa pod „premýšľaním“ spravidla rozumie schopnosť logicky usudzovať, vyvodzovať závery z východiskových premís a aplikovať právnu normu na konkrétny skutkový stav. V kontexte umelej inteligencie sa na označenie tejto schopnosti používa pojem „reasoning“, teda schopnosť systému prepájať informácie a dospieť k riešeniu.

¹ ČIERNIK, Marek. Logické usudzovanie a umelá inteligencia v práve, In: *Bratislavské právnické fórum 2025*. Bratislava : Právnická fakulta Univerzity Komenského v Bratislave, 2025.

² Dostupné definície online na: <https://www.ibm.com/think/topics/generative-ai>, <https://www.oecd.org/en/topics/sub-issues/generative-ai.html>, <https://news.mit.edu/2023/explained-generative-ai-11109> [cit. 2025-09-09].

Je potrebné odlíšiť LLMs od novšieho konceptu tzv. „Large Reasoning Models“ (LRMs). Zatiaľ čo LLMs sú koncipované predovšetkým ako nástroje predikcie ďalšieho slova na základe štatistických distribúcií v texte, LRMs sa usilujú o explicitné a krokové uvažovanie, ktoré má pripomínať matematický dôkaz či formálnu logickú úvahu. V právnom kontexte by práve takáto schopnosť mala zásadný význam, keďže právne rozhodovanie nevyžaduje iba znalosť jazykových vzorcov, ale predovšetkým schopnosť konzistentnej subsumpcie a udržania vnútornej logickej koherencie.

LLMs bývajú často kritizované pre svoj štatistický charakter. Vstupný prompt je spracovaný tak, že model na základe pravdepodobnostných distribúcií odhaduje najpravdepodobnejšie pokračovanie textu. Tento proces síce pôsobí sofistikovane, no pripomína skôr pokročilé automatické dokončovanie než deduktívne usudzovanie v logikom zmysle. Marcus a ďalší autori preto upozorňujú, že tieto modely síce operujú s pravdepodobnosťou, no nechápu význam slov; ich výkon je výsledkom manipulácie s dátami, nie skutočného porozumenia.³ Protiargumentom je, že aj bez „chápania“ môže byť výsledok funkčne ekvivalentný ľudskému mysleniu – AI tak môže dosahovať podobné závery odlišným mechanizmom.

Osobitnú výzvu predstavuje tzv. „black box efekt“. Neurónové siete, ktoré tvoria jadro LLMs, obsahujú miliardy parametrov a ich vnútorná dynamika je pre človeka netransparentná. Aj keď poznáme architektúru modelu, nedokážeme presne rekonštruovať, prečo konkrétny vstup vyústil do konkrétneho výstupu.⁴ Pre právnu oblasť ide o zásadný problém: spravodlivé rozhodnutie musí byť odôvodniteľné a spätne kontrolovateľné, čo je pri netransparentnom procese AI problematické.

V zjednodušenej analógii možno povedať, že neurónové siete fungujú podobne ako biologický mozog iba na úrovni štruktúry. Umelé „neuróny“ sú matematické funkcie spracúvajúce vstupy a prenášajúce ich ďalej. Model sa učí z veľkých dátových súborov a rozpoznáva vzorce, no tieto vzťahy sú pre človeka nečitateľné. Preto sa často zdôrazňuje, že AI iba napodobňuje myslenie, zatiaľ čo skutočné

³ MARCUS, Gary. *The Next Decade in AI: Four Steps Towards Robust Artificial Intelligence*. Online. Eprint: *arXiv:2002.06177*. 2020. Dostupné z: <https://doi.org/https://doi.org/10.48550/arXiv.2002.06177>. [cit. 2025-09-09].

⁴ BATHAEE, Yavar. *The Artificial Intelligence Black Box and the Failure of Intent and Causation*. *Harvard Journal of Law & Technology*. 2018. Dostupné z <https://jolt.law.harvard.edu/assets/articlePDFs/v31/The-Artificial-Intelligence-Black-Box-and-the-Failure-of-Intent-and-Causation-Yavar-Bathae.pdf>. [cit. 2025-09-09].

premýšľanie – chápané ako vedomé, logicky konzistentné a transparentné usudzovanie – zatiaľ absentuje.

2 Opis metód

Navrhujeme experiment, ktorý nám umožní sledovať schopnosť AI logicky usudzovať. V experimente využijeme spolu až 5 rôznych chýb, ktorých sa môže dopustiť či už nesprávne uvažujúci adresát normy, advokát alebo nepozorný zákonodarca. Tieto chyby sú:

1. Chyba nesprávneho zavedenia ekvivalencie pri nutnej podmienke
2. Nesprávne použitie spojky „alebo“ pri negácii disjunkcie
3. Chyba v právnej norme
4. Popretie antecedentu a nesprávne použitie *a contrario*
5. Popretie antecedentu a nesprávne použitie *a fortiori*⁵

Vytvoríme vždy dva typy promptov pre umelú inteligenciu. Prvý typ nazveme „laický prompt“ a druhý „odbornícky prompt“. Laický prompt bude jednoduchý. Využíva hovorový jazykový štýl, opisuje nejakú situáciu, žiada LLM o radu. Nehovorí však nič o vykonaní logickej analýzy. Tá ostáva v prompte skrytá a na ramenách LLM. Tento typ sa snaží napodobniť prístup bežného používateľa AI, ktorý sa zatiaľ nezačínal zaoberať problematikou vytvárania čo najvhodnejších promptov.

Odbornícky prompt naopak jasne deklaruje zámer, že budeme vykonávať logickú analýzu. Ďalej žiada o prepis do formálneho jazyka logiky, overenie správnosti úsudku, pravdivosti premís, pravdivosti záveru a vzťah vyplývajúci medzi nimi. Takýto prompt sa v kontraste k tomu laickému snaží napodobniť prístup právnik, ktorý ovláda základy logiky.

Oba tieto prompty budú spojené rovnakou zvolenou právnou normou, na ktorej možno demonštrovať niektorú z vyššie uvedených chýb. Celkovo pracujeme teda s piatimi právnymi normami, desiatimi rôznymi promptmi a testovať budeme schopnosť logicky usudzovať až piatich rôznych veľkých jazykových modelov od rôznych spoločností. Vybranými modelmi sú ChatGPT 5, Gemini 2.5 Pro, Copilot (Think Deeper), Grok 4 (Expert), Perplexity. Pri každom modeli sme sa snažili

⁵ Body 1-3 ostávajú spracované ako v: ČIERNIK, Marek. Logické usudzovanie a umelá inteligencia v práve, In: *Bratislavské právnické fórum 2025*. Bratislava : Právnická fakulta Univerzity Komenského v Bratislave, 2025. Výskum je rozšírený a vysvetlenie týchto bodov slúži len pre ozrejmienie.

o dodatočné zvolenie takého modelu, ktorý implementuje hlboké zamyslenie sa (deep thinking), premýšľanie (reasoning) a neskľzne iba do generovania čo najrýchlejšej možnej odpovede. V tomto smere je dôležité upozorniť na limit výskumu, ktorý mohol ovplyvniť výsledky - tzv. „routing“.⁶ Ide o spôsob presmerovania niektorých zadaní pre LLM v snahe odľahčiť nápor na servery spoločnosti, ktorá danú umelú inteligenciu prevádzkuje. Môže sa teda stať, že jednoducho znejúci prompt bude presmerovaný na model, ktorý využíva menšiu kapacitu a nepracuje vyššie uvedenými postupmi. Modely, pri ktorých sa s týmto presmerovaním môžeme stretnúť sú ChatGPT 5 a Perplexity.

Keď komunikujeme s vybranou AI, tak využívame dočasné konverzácie, ktoré sa neukladajú do dlhodobej pamäte, aby sme odpovede čo najviac objektivizovali. Taktiež zadávame každý prompt pre vybranú umelú inteligenciu aspoň dva krát, kvôli sledovaniu konzistentnosti výsledkov. Týmto spôsobom sa snažíme eliminovať nedostatky, ktoré by mohla priniesť náhodnosť generovania odpovedí v LLM.⁷

V ďalšej časti postupne rozoberieme jednotlivé typy logických chýb a vysvetlíme, k čomu by podľa správnosti mala AI v závere dospieť. Čo môže z danej normy vyvodit’.

2.1 Chyba nesprávneho zavedenia ekvivalencie pri nutnej podmienke

Ako s prvou pracujeme s normou zo zákona č. 369/1990 Zb. o obecnom zriadení (vyhlásené znenie). Táto znie:

*„Rokovania obecného zastupiteľstva sú zásadne verejné. (...) Vyhlásiť rokovanie obecného zastupiteľstva za neverejné možno **iba vtedy, ak** sa prerokúvajú veci, ktoré majú byť predmetom utajovania v štátnom záujme.“*

Zaujímavá je pre nás druhá veta, ktorá obsahuje logickú spojku „...iba vtedy, ak...“, ktorá vyjadruje nutnú podmienku. Hoci podmienka sa nachádza v druhej vete a podmienené sa nachádza vo vete prvej, môžeme tvrdiť, že takáto stavba s nutnou podmienkou má formu

⁶ Deklarované samotnými vývojármi. Online. Dostupné na: <https://openai.com/index/introducing-gpt-5/>. [cit. 2025-09-09].

⁷ BOURAS, Andrew. Integrating Randomness in Large Language Models: A Linear Congruential Generator Approach for Generating Clinically Relevant Content. Online. Eprint: 2407.03582. 2024. Dostupné z: <https://doi.org/https://doi.org/10.48550/arXiv.2407.03582>. [cit. 2025-09-24].

implikácie. Často krát, najmä v práve, dochádza k mylnému zamieňaniu tejto spojky s ekvivalentom.⁸ Ak by sme chceli ekvivalenciu vyjadriť v prirodzenom jazyku, tak používame spojku „...vtedy a len vtedy, keď...“. Vidíme, že spojka v prezentovanej norme znie odlišne.

Samotný fakt, že „sa prerokávajú veci, ktoré majú byť predmetom utajovania...“ ešte nespôsobuje, že rokovanie obecného zastupiteľstva môže byť vyhlásené za neverejné. Z tohto dôvodu ide o podmienku nutnú. Nutná podmienka je taká podmienka, ktorej nesplnenie spôsobuje, že podmienené nenastane. Splnenie nutnej podmienky však ešte nemá za následok, že podmienené nastane. Ak by sme chceli ustanovenie preformulovať a zaviesť namiesto nutnej podmienky dostatočnú podmienku, splnenie ktorej by už logický dôsledok vyvolávalo, znela by nasledovne:

- **Ak** sa prerokávajú veci, ktoré majú byť predmetom utajovania v štátnom záujme, **tak** možno vyhlásiť rokovanie obecného zastupiteľstva za neverejné.

Správny spôsob použitia pravidla modus tollens pre pôvodnú normu v tomto prípade vyzerá nasledovne:

$$\begin{array}{c} A \rightarrow B \\ \neg B \\ \hline \neg A \end{array}$$

Pričom „A“ zastupuje časť „vyhlásiť rokovanie obecného zastupiteľstva za neverejné možno“

Implikátor vyjadruje spojku „...iba vtedy, ak...“

„B“ zastupuje časť „prerokávajú sa veci, ktoré majú byť predmetom utajovania“.

Na druhom riadku vidíme negáciu časti „B“, ktorá teda znie „**neprerokávajú sa veci, ktoré majú byť predmetom utajovania**“.

Na treťom riadku vidíme negáciu časti „A“, ktorá znie takto: „**vyhlásiť rokovanie obecného zastupiteľstva za neverejné nemožno**“

Teda slovné zjednodušenie môžeme toto pravidlo vyjadriť takto:

- 1) Vyhlásiť rokovanie obecného zastupiteľstva za neverejné možno iba vtedy, ak sa prerokávajú veci, ktoré majú byť predmetom utajovania.
- 2) Neprerokávajú sa veci, ktoré majú byť predmetom utajovania.

⁸ GAHÉR, František. *Logika pre každého*. 4. dopl. vyd. Bratislava: Iris, 2013.

- 3) Nemožno vyhlásiť rokovanie obecného zastupiteľstva za neveřejné.

Toto je správny spôsob usudzovania podľa pravidla modus tollens, pokiaľ norma obsahuje spojku, ktorá signalizuje nutnú podmienku. V prvom prípade teda budeme sledovať, či sa umelá inteligencia dopustí chyby a zmýli si implikátor s ekvivalentorom, prípadne, či bude považovať nutnú podmienku za podmienku dostatočnú.

2.2 Nesprávne použitie spojky „alebo“ pri negácii disjunkcie

Ďalšou chybou, s ktorou sa môžeme v písanom práve stretnúť je používanie spojky „alebo“ namiesto použitia spojky „ani“. Norma s ktorou pracujeme je ustanovená v § 137 ods. 2 zákona č. 40/1964 Zb. Občiansky zákonník a znie:

*„Ak **nie je** právnym predpisom ustanovené **alebo** účastníkmi (**nie je**) dohodnuté inak, (**tak**) sú podiely všetkých spoluvlastníkov rovnaké.“*

Právnu normu sme rekonštruovali doplnením slov v zátvorkách, pretože obsahuje výpusťky. Na tomto mieste môže uvažovať, čo je cieľom tejto normy. Môžeme predpokladať, že cieľom je nastaviť základný režim pre určenie veľkosti podielov spoluvlastníkov, a to v prípade, keď **nie je** právnym predpisom ustanovené **a zároveň nie je** účastníkmi dohodnuté inak. V takom prípade, keď absentuje zákonné ustanovenie a absentuje aj dohoda, tak sú podiely všetkých rovnaké.

Naskytli sa nám teda dva možné zápisy vo formálnom logickom jazyku:

- 1) $\neg A \vee \neg B$
- 2) $\neg A \wedge \neg B$

Prvý zápis reflektuje na znenie uvedenej normy, pretože spojka „alebo“ vyjadruje disjunkt. Sledovaním cieľa predmetnej normy však tvrdíme, že zamýšľaný a správny by bol v tomto prípade záporný konjunkt, ktorý je v prirodzenom jazyku vyjadrený spojkou „ani“.

V druhom prípade teda sledujeme, či LLM porozumie cieľu normy, identifikuje túto chybu a bude pracovať so správnym logickým operátorom.

2.3 Chyba v právnej norme

Niekedy dokážeme objaviť aj chybu, ktorej sa dopustil samotný zákonodarca. Napríklad v ustanovení § 39 ods. 6 zákona č. 618/2003 Z. z. o autorskom práve a právach súvisiacich s autorským právom (autorský zákon), ktoré znie:

*„Zmluvou o vytvorení diela ani odovzdaním veci, ktorej prostredníctvom je dielo vyjadrené, **nenadobudne** objednávateľ právo použiť toto dielo, ak nie je uvedené inak, **len ak súčasne** so zmluvou o vytvorení diela alebo po tomto momente **uzatvorí** s autorom licenčnú zmluvu.“*

Cieľom tejto normy je zjavne to, aby objednávateľ **nenadobudol** právo použiť dielo v prípade, že **neuzavrie** s autorom licenčnú zmluvu. V uvedenom prípade sa však zrejme zákonodarca dopustil chyby. Spojka „..., len ak...“ opäť signalizuje nutnú podmienku. Podľa vyššie uvedeného pravidla *modus tollens* z chybnej normy vyplýva nasledovné:

- 1) Súčasne so zmluvou o vytvorení diela objednávateľ **neuzatvorí** s autorom licenčnú zmluvu.
- 2) Takže objednávateľ **nadobudne** právo použiť dielo.

Upravená norma, aby vyvolávala žiadaný výsledok by teda musela znieť:

*„Zmluvou o vytvorení diela ani odovzdaním veci, ktorej prostredníctvom je dielo vyjadrené, **nenadobudne** objednávateľ právo použiť toto dielo, ak nie je uvedené inak, **pokiaľ** súčasne so zmluvou o vytvorení diela alebo po tomto momente **neuzatvorí** s autorom licenčnú zmluvu.“*

V tomto prípade sledujeme, či LLM je schopný identifikovať túto chybu a reflektovať na ňu.

2.4 Popretie antecedentu a nesprávne použitie *a contrario*

Pre demonštrovanie tejto chyby si vypožičiame Klugov príklad⁹ vychádzajúci z § 7 nemeckého Bürgerliches Gesetzbuch (BGB). Formulujeme úsudok, ktorý z predmetného paragrafu vychádza.

⁹ KLUG, Ulrich. *Juristische Logik*. 4., neubearbeitete Auflage. Berlin: Springer, 1982.

„Ak je niekto fyzickou osobou, tak môže mať viacero trvalých bydlísk. Nie som fyzická, ale právnická osoba. Nemôžem mať viacere trvalé bydliská/sídla.“

Takýto úsudok môžeme pripodobniť k použitiu argumentu *a contrario*. Gahér charakterizuje túto metódu výkladu nasledovne: „*Ak podľa normy niečo platí pre výslovného adresáta a konkrétny adresát nie je výslovným adresátom alebo pod neho nespadá, tak nemôžeme odvodiť, že právna regulácia platná pre výslovného adresáta je platná aj pre konkrétného.*“¹⁰ Deduktívne neplatná je však posilnená verzia argumentu *a contrario*, ktorá by znela: „*Ak podľa normy niečo platí pre výslovného adresáta normy a konkrétny adresát pod neho nespadá, tak z toho vyplýva, že právna regulácia, ktorá je platná pre výslovného adresáta, nie je platná pre konkrétného.*“¹¹ Deduktívna neplatnosť spočíva v tom, že popierame antecedent implikácie, z čoho nemôžeme zaručene vyvodiť pravdivostnú hodnotu konzekventu. Uvedené dokážeme vyčítať z tabuľky pravdivostných hodnôt pre implikáciu.

A	B	$A \rightarrow B$
1	1	1
1	0	0
0	1	1
0	0	1

Tabuľka 1: Implikácia

Antecedentom je časť „niekto je fyzickou osobou“ (A). Konzekventom je časť „môže mať viacero trvalých bydlísk“ (B). Vidíme, že pokiaľ je pravdivostná hodnota antecedentu 0, tak celá implikácia bude v oboch prípadoch pravdivosti konzekventu pravdivá. Tým pádom popretím antecedentu nedokážeme zaručiť prenos pravdivostnej hodnoty z premís nášho úsudku na jeho záver. Tento úsudok je logicky neplatný.

2.5 Popretie antecedentu a nesprávne použitie *a fortiori*

Základom poslednej chyby, s ktorou v našom experimente pracujeme, je opäť vyššie uvedené popretie antecedentu, avšak ho obohacujeme o zdanlivý argument *a fortiori*. Konkrétne *a minori ad maius*. Pracujeme s nasledujúcim úsudkom:

¹⁰ GAHÉR, František. Interpretácia v práve II. *Filozofia*. 2015, roč. 70, č. 10, s. 794.

¹¹ Tamtiež.

„Do parku je zakázaný vstup so psom. Ja chcem do parku vstúpiť s mačkou. Keďže mačka nie je pes a je menej nebezpečná ako pes, tak do parku vstúpiť s mačkou môžem.“

Klasickým príkladom pre vysvetlenie výkladovej metódy *a fortiori a minori ad maius* je návšteva parku so psom. Pokiaľ chránime bezpečnosť návštevníkov parku pred psom, tak o to väčší by sme ich mali chrániť pred tigrom, ktorý je nebezpečnejší.¹² V našom úsudku vychádzame z rovnakého základu, ale argument používame obrátene v náš prospech, že mačka je menej nebezpečná ako pes, a preto by vstup do parku mal byť povolený. V tomto prípade sledujeme schopnosť AI identifikovať účel normy, ktorým môže byť ochrana bezpečnosti návštevníkov pred útokom zvierat alebo aj ochrana návštevníkov pred chorobami, ktoré zvieratá prenášajú. V tomto smere výsledok, v ktorom sa LLM vôbec vysporiada s hľadaním účelu právnej normy (zákona), budeme považovať za dostatočný.

Ako uvádzame pracujeme opätovne aj s vyššie uvedeným popretím antecedentu, čiže LLM bude mať náročnejšiu úlohu – vysporiadať sa s dvoma rôznymi argumentmi, ktoré sa navzájom podporujú.

3 Diskusia k výsledkom

Pred tým, ako pristúpime k vyhodnoteniu výsledkov považujeme za nevyhnutné zmieniť ešte jeden limit LLMs, ktorý sa nazýva „sycophancy“. Veľké jazykové modely sa snažia utvrdzovať zadávateľa v jeho presvedčeníach, pretože používateľ je s takýmito odpoveďami spokojnejší a má tendenciu opätovne AI využívať. Teda v prípade, že otázka v prompte je formulovaná sugestívne, LLM bude hľadať spôsoby ako s používateľom súhlasiť.¹³ V laických promptoch sa môže „sycophancy“ v istej podobe prejavovať.

Pre každý model AI prezentujeme tabuľku, v ktorej len zjednodušene vyhodnotíme, či dospel model k správnej odpovedi, a to správnym postupom. Nepostačí nám ak v prípade chyby č. 4 len oznámi, že právnická osoba môže mať len jedno sídlo. A následne sa v ďalšom kroku vyjadríme k zaujímavostiam, ktoré sme pre daný model pozorovali.

¹² GAHÉR, František. Interpretácia v práve II. *Filozofia*. 2015, roč. 70, č. 10, s. 789-799.

¹³ SHARMA, Mrinank; TONG, Meg; KORBAC, Tomasz et. al. Towards Understanding Sycophancy in Language Models. In: *12th International Conference on Learning Representations, ICLR 2024*. Vienna, 2024. Dostupné z: <https://doi.org/10.48550/arXiv.2310.13548> [cit. 2025-09-10].

3.1 Odpovede ChatGPT 5

ChatGPT 5 sa hrdí využívaním pokrokového „reasoningu“, ale na druhej strane v ňom nevieme zvoliť takýto „premýšľajúci“ model, pretože sám podľa promptu určí model, ktorý nám odpovie (viď. „routing“). Správnosť odpovedí vyzerá takto:

Chyba	Laický	Odbornícky
č. 1	Nie	Áno
č. 2	Nie	Áno
č. 3	Nie	Áno
č. 4	Nie	Áno
č. 5	Nie	Áno

Tabuľka 2: ChatGPT 5

Čo sa týka zaujímavosti odpovedí, tak len v jednom prípade nás tento LLM upozornil, že budeme dlhšie čakať na odpoveď, pretože premýšľa, a to v prípade chyby v norme. Premýšľanie trvalo 2 minúty a 43 sekúnd, no dospel k výbornej odpovedi a zrejme aj najkvalitnejšej. V snahe navrhnúť úpravu zadanej normy však nebol až taký úspešný, pretože zmenil jej pôvodne zamýšľaný účel.

V prípade laických promptov bol neúspešný a uchýľoval sa skôr k vysvetľovaniu na príkladoch ako k vyvodeniu správneho záveru. Vo všetkých prípadoch sa dopustil presne tých chýb, ktoré sme vyššie opísali. Odpovede na odbornícke prompty boli veľmi dobré a môžeme povedať, že došlo k úspešnému napodobneniu premýšľania.

3.2 Odpovede Gemini 2.5 Pro

V prípade Gemini podobne ako s ChatGPT pracujeme s predplatnou verziou, takže očakávame lepšie a presvedčivejšie odpovede ako od zvyšných model AI. Správnosť odpovedí vyzerá nasledovne:

Chyba	Laický	Odbornícky
č. 1	Nie	Áno
č. 2	Nie	Áno
č. 3	Nie	Nie
č. 4	Nie	Áno
č. 5	Nie	Áno

Tabuľka 3: Gemini 2.5 Pro

Na prvý pohľad je zrejmé, že Gemini nebol tak úspešný ako Chat-GPT. V prípade chybnnej normy totiž túto chybu odignoroval. Áno, dospel k správnejmu výsledku, ale k tejto zrejmej logickej chybe sa nevyjadril. Táto skutočnosť môže potvrdzovať náš predpoklad, že model AI textu nerozumejú.

Ďalšiu chybu, ktorú sledujeme aj pri iných modeloch, predstavuje uvádzanie nesprávnych paragrafových ustanovení v zákonoch. Ustanovenia, ktoré LLMs uvádzajú sú častokrát nesprávne, alebo dokonca neexistujú.

Pri Gemini je možné badať zmätočnú prácu s pojmami, pretože často ich medzi sebou zamieňa. Napríklad niekedy nazýva nutnú podmienku nevyhnutnou podmienku. Tento jav však pozorujeme naprieč všetkými ďalšími modelmi AI. Domnievame sa, že pôvod tohto nedostatku bude v predvolenom jazyku, ktorý modely AI používajú. LLMs „uvažujú“ v angličtine a až následne dochádza k prekladu do slovenského jazyka.

3.3 Odpovede Grok 4

Grok 4 je prvým modelom, ktorý využívame v jeho bezplatnej verzii. Napriek tomu sa odpovede javia kvalitné. Využívame model, ktorý údajne premýšľa a preto sme na odpovede čakali dlhšie. Správnosť odpovedí:

Chyba	Laický	Odbornícky
č. 1	Nie	Áno
č. 2	Nie	Áno
č. 3	Nie	Nie
č. 4	Nie	Áno
č. 5	Nie	Áno

Tabuľka 4: Grok 4

Najkratšie „premýšľal“ Grok 19 sekúnd a najdlhšie 3 minúty a 39 sekúnd. Grok 4 odvážne načrel do predikátovej logiky. S tou však nepracoval správne. Z hľadiska výsledkov dopadol rovnako ako Gemini.

V prípade laických promptov bolo značne cítiť „sycophancy“, pretože vždy sa snažil zadávateľovi vyhovieť.

K odborníckym promptom možno uviesť, že odpovede boli veľmi komplikované a zbytočne dlhé. LLM uvádzal zbytočné informácie, ktoré znižovali zrozumiteľnosť odpovede a uvádzal ďalšie príklady podobných noriem/úvah, ktoré neboli analogicky správne.

3.4 Odpovede Perplexity

Perplexity vo svojej bezplatnej verzii neumožňuje užívateľovi vybrať model, s ktorým chce komunikovať. Predpokladáme teda, že sme v dôsledku „routingu“ komunikovali s jedným z menej výkonných modelov. Správnosť odpovedí:

Chyba	Laický	Odbornícky
č. 1	Nie	Nie
č. 2	Nie	Áno
č. 3	Nie	Nie
č. 4	Nie	Áno
č. 5	Nie	Nie

Tabuľka 5: Perplexity

Trend znižujúcej sa správnosti odpovedí pokračuje. Tento krát v prípade chyby č. 1 Perplexity presne túto nesprávnu ekvivalenciu zaviedla a spravila tým z nutnej podmienky podmienku dostatočnú.

V prípade chyby č. 5 s ľahkosťou prijala našu nesprávnu úvahu a formalisticky vyslovila záver, že zákaz sa vzťahuje iba na psy, preto s mačkou je vstup povolený.

Odpovede na laické prompty boli rovnako ako v ostatných prípadoch nesprávne, no tento krát sa javili aj menej presvedčivé. Napríklad namiesto vysvetlenia normy sa uspokojil len s jej preformulovaním.

3.5 Odpovede Copilot

A ako posledné odpovede analyzujeme odpovede modelu Copilot. Hoci opäť pracujeme s neplatenou verziou, tak Copilot umožňuje výber konkrétneho modelu. Zvolili sme teda model, ktorý „premýšľa“ na úkor rýchlosti generovania odpovedí. Správnosť odpovedí:

Chyba	Laický	Odbornícky
č. 1	Nie	Áno
č. 2	Áno	Áno
č. 3	Nie	Nie
č. 4	Nie	Áno
č. 5	Áno	Áno

Tabuľka 6: Copilot

Copilot môžeme považovať za veľmi prekvapivý model. Ako jednému sa mu správne podarilo odpovedať na laické prompty. Pre úplnosť uvádzame, že v prípade chyby č. 5 sme označili odpoveď za správnu, pretože rozbehol úvahu o účele daného pravidla a nepripustil formalistický prístup.

V prípade odborníckych promptov vo viacerých prípadoch využíva tabuľkovú metódu, aby nám vysvetlil možné kombinácie pravdivostných hodnôt aj s ich dôsledkami. Z hľadiska konzistentnosti sú však jednotlivé odpovede veľmi odlišné. Líšia sa v rozsahu, v množstve argumentov, ale aj v samotnom prístupe.

Ako sme povedali, Copilot nás prekvapil, ale nemôžeme ho považovať za spoľahlivého v logických analýzach.

Záver

Na základe predloženej analýzy možno konštatovať, že súčasné systémy umelej inteligencie, a osobitne veľké jazykové modely, nie sú ešte schopné spoľahlivo a konzistentne vykonávať deduktívne usudzovanie. Hoci ich výkon v oblasti spracovania prirodzeného jazyka pôsobí presvedčivo a neraz dokáže navodzovať dojem autentického logického myslenia, podrobnejšie skúmanie odhaľuje závažné nedostatky. Zaznamenali sme opakujúce sa logické nekonzistencie, ktoré sa prejavujú naprieč rôznymi modelmi a pri rozličných typoch úloh.

Tieto zistenia majú osobitný význam pre právnu vedu a prax. Právnické povolania stoja a padajú na presnosti argumentácie, konzistentnom uplatňovaní pravidiel a schopnosti aplikovať normy na konkrétne prípady. V tejto kľúčovej oblasti zatiaľ umelá inteligencia nedokáže poskytnúť záruku, ktorá by umožnila jej plnohodnotné nasadenie ako náhrady právnikovi, sudcovi či advokátovi.

To však neznamená, že takáto možnosť je vylúčená do budúcnosti. Technologický vývoj v oblasti umelej inteligencie napreduje mimoriadne dynamicky a je realistické predpokladať, že schopnosť logického uvažovania bude postupne zdokonaľovaná. Aj preto je nevyhnutné, aby právna veda už dnes reflektovala tieto trendy, analyzovala ich dôsledky a pripravovala teoretické i praktické rámce pre ich budúce využitie.

Zatiaľ môžeme povedať, že právnik ostáva nenahraditeľný – no nie je isté, dokedy tento stav potrvá.

Zoznam použitej literatúry

1. BATHAEE, Yavar. The Artificial Intelligence Black Box and the Failure of Intent and Causation. *Harvard Journal of Law & Technology*. 2018, roč. 31, č. 2, s. 890-938. ISSN 0897-3393. Dostupné z: <https://jolt.law.harvard.edu/assets/articlePDFs/v31/The-Artificial-Intelligence-Black-Box-and-the-Failure-of-Intent-and-Causation-Yavar-Bathae.pdf>. [cit. 2025-09-09].
2. BOURAS, Andrew. Integrating Randomness in Large Language Models: A Linear Congruential Generator Approach for Generating Clinically Relevant Content. Online. Eprint: 2407.03582. 2024, 13 s. Dostupné z: <https://doi.org/https://doi.org/10.48550/arXiv.2407.03582>. [cit. 2025-09-24].
3. ČIERNIK, Marek. *Logické usudzovanie a umelá inteligencia v práve*, In: Bratislavské právnické fórum 2025. Bratislava : Právnická fakulta Univerzity Komenského v Bratislave, 2025.
4. GAHÉR, František. *Logika pre každého*. 4. dopl. vyd. Bratislava: Iris, 2013. ISBN 978-80-89256-88-4.
5. GAHÉR, František. Interpretácia v práve II. *Filozofia*. 2015, roč. 70, č. 10, s. 789-799. ISSN 2585-7061.
6. JUNNAN Liu, HONGWEI Liu, LINCHEN Xiao et al. Are Your LLMs Capable of Stable Reasoning?. In: *Findings of the Association for Computational Linguistics: ACL 2025*, Vienna, Austria, s. 17594 – 17632. Dostupné z: [10.18653/v1/2025.findings-acl.905](https://arxiv.org/abs/2505.17594). [cit. 2025-09-10].
7. KLUG, Ulrich. *Juristische Logik*. 4., neubearbeitete Auflage. Berlin: Springer, 1982. ISBN 978-3-642-87157-3.
8. MARCUS, Gary. *The Next Decade in AI: Four Steps Towards Robust Artificial Intelligence*. Online. Eprint: [arXiv:2002.06177](https://arxiv.org/abs/2002.06177). 2020, 59 s. Dostupné z: <https://doi.org/https://doi.org/10.48550/arXiv.2002.06177>. [cit. 2025-09-09].
9. NAVEED, Humza; ULLAH KHAN, Assad a QIU, Shi. et. al. *A Comprehensive Overview of Large Language Models*. ACM Trans. Intell. Syst. Technol. 2025, vol. 16, no. 5, 47 s. ISSN 2157-6904. Dostupné z: <https://doi.org/10.1145/3744746> [cit. 2025-09-10].
10. SHARMA, Mrinank; TONG, Meg; KORBAC, Tomasz et. al. *Towards Understanding Sycophancy in Language Models*. In: 12th International Conference on Learning Representations, ICLR 2024. Vienna, 2024, 35 s. Dostupné z: [https://doi.org/10.48550/arXiv.2405.13548](https://arxiv.org/abs/2405.13548) [cit. 2025-09-10].
11. SHOJAEI, Parshin; MIRAZDEH, Iman; ALIZADEH, Keivan et. al. The Illusion of Thinking: Understanding the Strengths and Limitations of Reasoning Models via the Lens of Problem Complexity. In:

- NeurIPS*, 2025. 30 s. Dostupné z: <https://doi.org/10.48550/arXiv.2506.06941> [cit. 2025-09-10].
12. ŠTĚPÁN, Jan. *Logika a právo*. 2., dopl. vyd., Praha: C.H. Beck, 2004. 128 s. ISBN 80-7179-872-X.