

Jaké překážky jsou kladeny vývoji umělé inteligence a jejímu využití při triáži?

Filip Hájek



Zdroje a kapacity zdravotnických pracovníků a zařízení nejsou neomezené a jedním z možných řešení se zdá být využití umělé inteligence. I přes všechny své výhody má AI spoustu nedostatků, které jsou technicky jen těžko řešitelné. Je to primárně otázka osobních zdravotních dat a samotná důvěra jak personálu, tak i pacientů. V tomto textu vám také předkládám problematiku implementace AI¹ do systému rozhodování při triáži. Tento aspekt se pro umělou inteligenci jeví s velkým potenciálem.

Klíčová slova: Urgentní medicína, triáž, umělá inteligence, zdravotní data

ÚVOD

Umělá inteligence a její potenciál v moderní medicíně je lékařům a odborníkům po celém světě čím dál zřejmější. Její rapidní vývoj nám však přináší několik etických a morálních problémů, jimž je nutné porozumět před samotnou implementací těchto systémů do klinické praxe. Cílem tohoto textu je představit některé významné etické aspekty začlenění a vývoje AI v lékařských oborech. Mezi tyto otázky patří: jakým způsobem nakládat se zdravotními daty pacientů? Je umělá inteligence věrohodným zdrojem a z čeho plyne případný nedostatek důvěry? Můžeme ji skutečně považovat za nestranného poradce při rozhodování? Podrobněji se podíváme na její využití při triáži v urgentní medicíně.

ETIKA VÝVOJE AI V MEDICÍNĚ

Klíčovým předpokladem pro vývoj funkční a plně kompetentní umělé inteligence, která bude schopna přesného a efektivního rozhodování, je dostatek dat. Při jejím užití v lékařství tomu není jinak, avšak v tomto případě je rozvoji kladena bariéra v podobě problému

osobních informací pacientů. Zdravotní data potřebná k vytvoření dostatečného základu pro vznik umělé inteligence často podléhají právním omezením, což značně limituje její vývoj. Jednou z možností je poté využití veřejných databází, kterých je velmi málo, či vytvoření vlastních, soukromých, databází. Ty však nedisponují velkým množstvím podnětů, a nelze z nich tedy vyvodit obecně platné závěry. Jako ideálním řešením se zdá být vytvoření mezinárodních směrnic pro depersonalizaci dat, což by vedlo k vytvoření nepřeborného množství zdrojů pro vývoj AI v medicíně (Piliuk & Tomforde, 2023).

PROBLÉMY VSTUPNÍCH DAT

Je umělá inteligence skutečně ve všech situacích morálním ideálem bez předsudků a předpojatosti? Na jednu stranu je neúnavná, a tím pádem budou její rozhodování konzistentní a přesná. Na druhou stranu tato přesnost vychází z dat, která ji poskytli lidé. Kvůli tomu mohou i její rozhodnutí být zaujatá či jinak chybná. Problém nastane ve chvíli, když si v průběhu učení umělá inteligence všimne stejného vzorce nesprávného chování či diagnózy v několika různých případech svého datového souboru. Ona toto chování za vadné samozřejmě neoznačí, nýbrž se naopak utvrdí v jeho správnosti (Mueller, a další, 2022).

DŮVĚRA V UMĚLOU INTELIGENCI

Klíčem k překonání morální překážky, která brání využití umělé inteligence při rozhodování o lidském zdraví a životech, je získání důvěry. Jedním z hlavních důvodů jejího nedostatku je skutečnost, že se nám AI často jeví jako taková *černá skříňka* (Stewart, a další, 2018). Její vysvětlení postupu při rozhodování často obsahuje jen nespočet číselných dat, která pro člověka nejsou dobře uchopitelná. Tento problém si vývojáři uvědomují se

¹AI = Artificial intelligence (umělá inteligence)

snahou o vytvoření takové umělé inteligence, s níž by interakce byla přirozenější a lépe pochopitelná.

TRIÁŽ

Triáž je proces, při kterém se zhodnocuje stav pacienta dle závažnosti a na základě toho přiděluje příslušná míra okamžité péče. Tento postup vyžaduje velké množství pozornosti a kapacity zdravotnického personálu, který je tím často i psychicky vyčerpáván. Při snaze rozlišit mezi podobně závažnými případy, které nejsou život ohrožující, poté často všechny spadají pod stejnou úroveň akutnosti, což vede k dlouhým a neoptimalizovaným čekacím dobám pacientů (Mueller, a další, 2022). Nezřídka dochází také k nedokonalému vyhodnocení stavu příchozího, tedy k tvoření případů, jimž byla přidělena příliš vysoká, či naopak nízká relevance.

VYUŽITÍ AI PŘI TRIÁŽI

V oblasti urgentní medicíny se potenciál umělé inteligence zdá být velmi vysoký právě při triáži. Její výhodou je neúnavnost a konzistence (vizte výše), která by uvolnila kapacity pracovníků. Taktéž by po dostatečném rozvoji byla schopna přesněji stanovit závažnost stavu pacienta a vytvořit tak daleko více úrovní relevance, než dovoluje lidský úsudek. Umožnila by tím přesnější stanovení čekací doby pacienta a primárně zamezila vzniku chybně posouzených případů. Nastolila by se rovnováha mezi pacienty, kterým se dostalo více prostředků, než potřebovali, a naopak (Raita, a další, 2019). Vystavající otázkou je, zda je morální nechat umělou inteligenci rozhodovat o záležitostech týkajících se lidských životů. Tento problém by mohl být částečně vyřešen zvýšením důvěry v AI (vizte výše) či vytvořením kontrolního systému zdravotních pracovníků, kteří by nad jejím rozhodováním měli poslední slovo.

DISKUZE

I přes všechny etické problémy, které jsou zmíněné výše, by primárním cílem měl být rozvoj umělé inteligence z hlediska funkčnosti a objektivní správnosti kroků, kterými by po implementaci do systému postupovala (Mueller, a další, 2022). Nefunkční stroj poháněný AI, který by sice byl vyvinut a choval se dle všech etických kodexů, ale podával nepřesné a falešné informace, by v praxi jen škodil.

ZÁVĚR

Účelem tohoto textu bylo postihnout některé etické problémy týkající se vývoje umělé inteligence a její implementace do procesu rozhodování při triáži.

Samotné tvoření AI pro medicínské účely je technicky i eticky komplikovaný problém, který vyžaduje značný počet opatření. Troufám si tvrdit, že ideálním řešením by v tomto případě bylo vytvoření jednotného souboru zdravotních dat, která byla depersonalizována, z nichž by se následně mohla umělá inteligence učit. Byla by tedy zapotřebí spolupráce několika mezinárodních lékařských institucí. Takto obsáhlý soubor dat by mohl umělou inteligenci zamezit učení nesprávných postupů a informací, kterých by v relativních číslech mezi daty tolik nebylo. V neposlední řadě bude pro aplikaci AI důležité, aby byla schopna srozumitelně vysvětlit a popsat svůj postup při rozhodování. Bez této schopnosti by bylo obtížné jí důvěřovat.

Myslím si, že umělá inteligence má obrovský potenciál a nemá smysl se jejímu rozvoji bránit. Její implementace do lékařských systémů uvolní kapacity pracovníků a optimalizuje chod celého zdravotnického zařízení. V momentálním stavu je však stále ve stádiu vývoje, při kterém by její využití bylo více než riskantní. Domnívám se, že klíčová bude stálá lidská kontrola a regulace.

Nahradí v budoucnu umělá inteligence lékařský úsudek úplně? Bylo by to dle morálních hodnot v pořádku? Tak by mohly znít otázky, kdyby se implementace AI v budoucnu podařila.

CITOVANÁ LITERATURA

- Mueller, B., Kinoshita, T., Peebles, A., Graber, M. A., & Lee, S. (2022). Artificial intelligence and machine learning in emergency medicine: a narrative review. *Acute Medicine & Surgery*, 9(1), stránka 740
- Piliuk, K., & Tomforde, S. (2023) Artificial intelligence in emergency medicine. A systematic literature review. *International Journal of Medical Informatics*, 180
- Raita, Y., Goto, T., Faridi, M. K., Brown, D. F. M., Carmago, C. A., Jr., & Hasegawa K. (2019) Emergency department triage prediction of clinical outcomes using machine learning models. *Critical Care*, 23
- Stewart, J., Sprivulis, P., & Dwivedi, G. (2018). Artificial intelligence and machine learning in emergency medicine. *Emergency Medicine Australasia*, 30(6), stránky 870–874