

Obr. 1. Algoritmus pro klinickou praxi: jak postupovat při konzultaci pacienta přicházejícího s výstupem z ChatGPT. Pětikrokový postup zdůrazňuje přijetí informace bez odsouzení, triage red flags, klinické ověření, edukaci a nastavení bezpečných hranic použití AI. Zlatá pravidla shrnují základní principy: pacient na prvním místě, klinický pohled a role lékaře a ochrana dat

Co říct pacientovi s výstupem z ChatGPT

Co říct pacientovi s výstupem z ChatGPT



Na straně zdravotníků nacházejí LLM uplatnění zejména v oblasti administrativní podpory (korespondence, shrnutí zpráv, strukturování propouštěcích zpráv), v edukaci (příprava výukových materiálů), jako podpora při řešeršní činnosti a jako asistent při formulaci diferenciální diagnostiky – s tím, že finální klinické rozhodnutí vždy zůstává na lékaři. Studie na lékařských licenčních zkouškách (v USA) opakovaně ukázaly, že moderní LLM dosahují výkonnosti srovnatelné s průměrným lékařem či ji překračují (1).

LLM, pokud jsou používány v dobře nastaveném pracovním rámci, mohou tedy pro zdravotníka představovat významný a efektivní posun v práci s informacemi. Pokud se lékaři a NLZP naučí pracovat s kontextem, budou moci zohlednit daleko více věcí, než to bylo doteď. A v ideálním případě, pokud vše bude správně nastavené, budou moci výsledkům i věřit. Velkou otázkou však i nadále zůstává, jak nastavit evaluace, aby tomu lékařům věřit mohli.

Rizika a limity: co LLM není a být nemůže

Přes výše popsané přínosy je nutné důrazně pojmenovat rizika, která jsou s nekritickým využíváním LLM spojena. Z klinického pohledu jsou zásadní především tato:

Absence certifikace LLM jako zdravotnických prostředků.

Veřejně dostupné LLM (ChatGPT, Gemini, Claude) nejsou certifikovány jako zdravotnické prostředky ve smyslu nařízení EU 2017/745 (Medical Device Regulation/MDR), amerického federálního úřadu pro kontrolu potravin a léčiv (Food and Drug Administration/FDA) ani nového Aktu EU o AI. Nemají validovanou klinickou bezpečnost, nejsou předmětem postmarketingového sledování a jejich provozovatel nenese odpovědnost spojenou se zdravotnickým prostředkem (5). To je zásadní rozdíl oproti certifikovaným AI systémům např. pro analýzu obrazových dat (např. Brainomix e-Stroke, Pixyl.Neuro), které prošly klinickou validací a jsou používány jako součást regulovaného diagnostického řetězce.

Halucinace a dezinformace. LLM generují přesvědčivě formulované odpovědi i tehdy, když nemají spolehlivá data – tzv. halucinace. U zdravotních dotazů může přesvědčivě znějící, ale fakticky nesprávná odpověď vést k odložení léčby, nevhodné samomedikaci nebo ke zbytečné zdravotní úzkosti. Jazykový model chce vždy najít odpověď na dotaz tazatele.

Úniky dat a ochrana soukromí. Vkládání osobních zdravotních dat, snímků nebo lékařské dokumentace do veřejně dostupných LLM systémů představuje závažné bezpečnostní riziko. Data opouštějí