



RADOVAN DÁVID

vzdělávací společnost

si dovoluje Vás pozvat na seminář

„Jak využít AI (umělou inteligenci) při vyhodnocování rizik dítěte“

Akreditace Ministerstva práce a sociálních věcí

Akreditace Ministerstva vnitra České republiky

Termín a místo konání: 13. listopadu 2025 od 9:00

online

Lektor: JUDr. Ing. Radovan Dávid, Ph.D.

lektor na právnických fakultách (Brno a Bratislava), šestnáctiletá praxe asistenta soudce Ústavního soudu, člen Světové organizace pro rodinné právo, čestný člen Profesní komory sociálně-právní ochrany dětí, člen redakční rady Juridical Analytical Journal, autor desítek článků, monografií (mimo jiné učebnice rodinného práva a civilního práva procesního) a komentářů (mimo jiné k občanskému zákoníku, občanskému soudnímu řádu, zákonu o zvláštních řízeních soudních a exekutivnímu řádu)

Umělá inteligence (AI) proniká do stále více oblastí veřejné správy, včetně sociálně-právní ochrany dětí. AI nástroje mohou pomoci při **vyhodnocování rizik ohrožení dítěte** – například při analýze velkého množství informací, identifikaci varovných signálů z datových zdrojů či podpoře rozhodování v krizových situacích.

Současně ale využívání AI přináší **zásadní právní a etické otázky**: ochrana osobních údajů a citlivých informací, transparentnost rozhodovacích procesů, odpovědnost za výsledek a riziko diskriminace či chybného vyhodnocení. Pracovníci OSPOD musí znát **limity použití AI**, umět posoudit spolehlivost výstupů a vždy rozhodovat na základě lidského úsudku, nikoli mechanického převzetí výsledků algoritmu.

Školení se zaměří na **praktické možnosti a právní mantinely** využití AI při vyhodnocování rizik dítěte, příklady z praxe, doporučené postupy a také na **zásady bezpečného a etického použití AI** v souladu s českým právem, evropskými regulacemi a mezinárodními dokumenty.

Na konci vzdělávacího programu bude účastník schopen:

- vysvětlit, co je AI a jaké má možnosti v oblasti SPOD
- rozpoznat situace, kdy je vhodné AI využít k vyhodnocení rizik dítěte
- znát právní rámec použití AI při zpracování osobních údajů
- aplikovat zásady nezbytnosti, proporcionality a ochrany soukromí
- interpretovat a ověřovat výstupy AI, aby byly použitelné v řízení
- identifikovat a minimalizovat riziko chyb či zkreslení výstupů
- vést dokumentaci o použití AI a odůvodnit její zapojení do procesu
- spolupracovat s IT odborníky a právníky na bezpečném nasazení AI
- vyhodnotit etické otázky a dopady na práva dítěte

Obsah kurzu:

1. Co je AI a jak funguje

- Definice umělé inteligence (strojové učení, neuronové sítě, LLM)
- Rozdíl mezi slabou a silnou AI
- Příklady využití AI v sociálních službách a SPOD
- Automatizovaná analýza dat vs. podpora rozhodování
- Možnosti integrace do informačních systémů OSPOD
- Příklady ze zahraničí
- Limity současných AI nástrojů
- Nutnost lidského dohledu
- Význam správného zadání (promptu) pro výsledek

2. Právní rámec použití AI při vyhodnocování rizik dítěte

- GDPR – zpracování osobních údajů dětí pomocí AI
- Zákon o SPOD – povinnosti při práci s informacemi o dítěti
- Povinnost chránit soukromí dítěte (OZ, Listina základních práv a svobod)
- Evropský akt o umělé inteligenci (AI Act) – rizikové kategorie AI
- Správní řád a zásada materiální pravdy
- Zásada zákazu diskriminace a rovného zacházení
- Mezinárodní závazky ČR (Úmluva o právech dítěte)
- Doporučení Rady Evropy k využití AI v ochraně dětí
- Odpovědnost za rozhodnutí při použití AI

3. Typy AI nástrojů použitelných pro OSPOD

- Analýza textových dokumentů a spisů
- Automatizované vyhledávání relevantních judikátů a právních norem
- Systémy pro detekci rizik v online prostředí (kyberšikana, grooming)
- Analytické dashboardy pro sledování indikátorů ohrožení
- Chatboty pro předběžné informování rodin a dětí
- Prediktivní modely pro riziko recidivy ohrožujících situací
- Systémy pro kategorizaci případů podle závažnosti
- Nástroje pro anonymizaci a pseudonymizaci dat

- Ověřovací nástroje pro analýzu důkazů

4. Přínosy a limity AI v práci OSPOD

- Rychlejší zpracování velkého objemu informací
- Možnost identifikace skrytých souvislostí
- Podpora při rozhodování ve složitých případech
- Zlepšení cílení preventivních opatření
- Riziko chybného vyhodnocení (false positive / false negative)
- Závislost na kvalitě vstupních dat
- Nedostatek transparentnosti algoritmů (black box)
- Riziko neoprávněného přístupu k citlivým datům
- Etické dilema mezi efektivitou a ochranou práv

5. Zásady bezpečného a etického použití AI

- Nezbytnost a přiměřenost použití AI
- Lidský dohled a konečné rozhodnutí
- Transparentnost vůči účastníkům řízení
- Právo dítěte a rodičů na vysvětlení použití AI
- Minimalizace zpracování osobních údajů
- Zajištění kybernetické bezpečnosti
- Prevence diskriminačních výstupů
- Etické kodexy pro práci s AI v sociálních službách
- Pravidelné vyhodnocování účinnosti nástrojů

6. Dokumentace a přezkoumatelnost

- Jak zaznamenat použití AI v konkrétním případě
- Uchování vstupních dat a nastavení parametrů AI
- Uvedení AI výstupu v záznamu nebo zprávě
- Oddělení doporučení AI od závěrů pracovníka
- Možnost nezávislé kontroly výstupu
- Dokumentace důvodů použití AI
- Uchovávání záznamů pro případ soudního přezkumu
- Přístup k dokumentaci pro účastníky řízení
- Archivace a likvidace dat

7. Příklady využití AI v praxi OSPOD

- Vyhodnocení rizika zanedbání péče na základě kombinace indikátorů
- Analýza předchozího vývoje případu pro volbu opatření
- Podpora při rozhodování o umístění dítěte do zařízení
- Automatická kontrola úplnosti spisové dokumentace
- Detekce vzorců v online chování dítěte (v součinnosti s policií)
- Sledování plnění uložených opatření v čase
- Včasné varování před opakovaným ohrožením

- Předvýběr relevantních podkladů pro soudní jednání
- Odhad dopadů různých scénářů zásahu

8. Rizika a preventivní opatření

- Závislost na technologiích a úbytek kritického myšlení
- Zneužití AI pro neoprávněné profilování
- Nedostatečné ověření spolehlivosti modelu
- Ztráta kontroly nad daty uloženými v externích systémech
- Chyby způsobené neaktuálními daty
- Omylné rozhodnutí při slepé důvěře ve výsledek
- Nejednotná metodika práce s AI v rámci OSPOD
- Reputační rizika při zpochybnění výstupu
- Preventivní školení a interní směrnice

9. Doporučení pro zavedení AI do praxe OSPOD

- Vypracování interních pravidel pro použití AI
- Pilotní projekty a testování na anonymizovaných datech
- Spolupráce s IT odborníky a právníky
- Pravidelné vyhodnocování účinnosti a rizikovosti nástrojů
- Zapojení etických komisí či poradních orgánů
- Školení pracovníků v oblasti digitální gramotnosti a AI
- Postupná implementace do jednotlivých agend
- Transparentní informování veřejnosti o využití AI
- Sdílení zkušeností mezi OSPOD a zahraničními partnery

10. Diskuse a závěr školení