



# RADOVAN DÁVID

vzdělávací společnost

si dovoluje Vás pozvat na seminář

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

**Instituce akreditovaná u Ministerstva vnitra ČR  
Akreditace Ministerstva práce a sociálních věcí**

**Termín a místo konání:** 25. června 2026 od 9:00

Grandhotel International Praha, Koulova 15, Praha 6

- MHD Zelená (tramvaj 8 a 18)

- MHD Nádraží Podbaba (tramvaj 8 a 18, autobusy 107, 147, 160)

- MHD Čínská (autobusy 107, 147, 160)

- cca 4 minuty jízdy od sjezdu z městského okruhu, parkování na parkovišti před hotelem

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

sedmnáctiletá praxe asistenta soudce Ústavního soudu, osmnáctiletá praxe lektora právnických fakult v Brně a Bratislavě, patnáctiletá praxe v oblasti vzdělávání úředníků, 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 exekučnímu řádu)

Rozvoj umělé inteligence (AI) a datových technologií přináší nové možnosti i do oblasti sociálně-právní ochrany dětí. Moderní systémy dokážou analyzovat velké množství dat a upozornit na rizikové vzorce chování či prostředí, které mohou ohrozit dítě. Teoreticky tak mohou pomoci včas identifikovat případy zanedbávání, týrání, zneužívání či jiných forem ohrožení, dříve než se projeví v otevřené krizové situaci. Zároveň však vyvstává řada otázek: jaká data lze legitimně využívat? Jak zajistit ochranu soukromí dítěte a rodiny? A může se pracovník OSPOD spolehnout na algoritmus, aniž by potlačil vlastní odborný úsudek a empatii? Seminář se zaměří na to, jak se umělá inteligence již dnes uplatňuje v zahraničí, jaké jsou české legislativní limity a jaké etické a právní otázky vyvstávají. Účastníci se seznámí s konkrétními modely prediktivní analýzy, se způsoby, jak je možné (a jak není možné) je implementovat do praxe OSPOD, a s riziky, která přináší spoléhání se na algoritmičké rozhodování. Diskutovány budou i příklady dobré praxe a varovné kauzy, aby bylo možné zasadit AI do kontextu reálné sociální práce. Cílem je posílit schopnost pracovníků OSPOD orientovat se v nových

technologiích, využívat jejich přínos, ale současně kriticky hodnotit jejich limity a chránit nejlepší zájem dítěte.

**Po absolvování programu bude účastník schopen:**

- vysvětlit základní principy fungování AI v predikci rizik
- popsat možnosti využití dat k prevenci ohrožení dítěte
- identifikovat právní a etické limity využívání algoritmů
- rozlišit mezi podpůrným nástrojem a rozhodnutím v sociální práci
- posoudit rizika spojená se sběrem a zpracováním dat o dítěti a rodině
- využít dostupné poznatky k preventivní práci OSPOD
- orientovat se v judikatuře k ochraně soukromí a osobních údajů
- komunikovat s rodinou o tom, jaké technologie mohou být použity
- posoudit vhodnost zapojení AI v konkrétním případě
- spolupracovat s odborníky z IT a práva na nastavování hranic
- rozpoznat riziko diskriminace nebo stigmatizace algoritmy
- aplikovat princip nejlepšího zájmu dítěte při práci s AI
- vyhodnocovat výstupy technologií kritickým pohledem
- připravit metodická doporučení pro tým či obec

**Dotazy**, na něž by měl být kurz především zaměřen, lze posílat na kontaktní e-mail [radovandavid@radovandavid.cz](mailto:radovandavid@radovandavid.cz).

**Obsah kurzu:**

**1. Úvod do problematiky AI a predikce rizik**

- základní principy fungování umělé inteligence
- prediktivní modely v sociální práci
- přehled využití v zahraničí (UK, USA, Skandinávie)
- rozdíly mezi predikcí a diagnostikou
- výhody a přísliby pro OSPOD
- hrozby spojené s nesprávnou implementací
- pojem „algoritmická odpovědnost“
- role lidského faktoru vedle AI
- rozdíl mezi výzkumným a praktickým využitím
- kazuistiky ze zahraničních projektů

**2. Data a jejich využití**

- jaká data lze o dítěti a rodině sbírat
- právní limity dle GDPR a ZSPOD
- riziko neoprávněného zásahu do soukromí
- možnosti anonymizace a pseudonymizace
- validita a spolehlivost datových vstupů
- nebezpečí zkreslených nebo neúplných dat
- příklady typických datových zdrojů (školy, zdravotní zařízení, policie)
- riziko „profilování“ dítěte
- problematika citlivých údajů
- mezinárodní standardy v ochraně dat

**3. Přínosy AI pro OSPOD**

- včasná identifikace rizik

- podpora při rozhodování o intervencích
- odhalení vzorců chování v rodině
- možnost predikce eskalace násilí
- lepší cílení preventivních opatření
- úspora času pracovníků OSPOD
- vyšší transparentnost rozhodování
- příklady z praxe zahraničních služeb
- využití v krizových plánech
- role AI v podporovaných rozhodnutích

#### **4. Limity a rizika použití AI**

- falešně pozitivní a negativní výsledky
- riziko stigmatizace dítěte a rodiny
- přenos stereotypů do algoritmů
- etické dilema „predikce budoucího chování“
- ztráta důvěry klientů v OSPOD
- přílišná závislost na technologii
- snížení role lidské empatie a zkušenosti
- problém transparentnosti „black box“ algoritmů
- finanční a organizační limity implementace
- odpor odborné veřejnosti

#### **5. Právní rámec v ČR a EU**

- ZSPOD a hranice práce s daty
- GDPR a zákaz automatizovaného rozhodování
- Listina základních práv a svobod – právo na soukromí
- Úmluva o právech dítěte – nejlepší zájem dítěte
- připravovaný evropský AI Act
- stanoviska ÚOOÚ k prediktivním technologiím
- judikatura ESLP k ochraně soukromí
- role soudů při přezkumu zásahů OSPOD
- odpovědnost obce při zavádění AI nástrojů
- přehled českých právních limitů

#### **6. Role pracovníka OSPOD**

- kritické hodnocení výstupů AI
- zachování odborného úsudku nad technologiemi
- komunikace s rodinou o výsledcích predikcí
- zapojení dítěte do rozhodování o opatřeních
- prevence stigmatizace v přístupu k rodině
- odpovědnost za konečné rozhodnutí
- využití AI jako podpůrného, nikoli určujícího nástroje
- tvorba metodik a standardů pro tým
- vzdělávání pracovníků v oblasti AI
- etická reflexe využívání technologií

#### **7. Etické otázky a dilemata**

- souhlas dítěte a rodiny se zpracováním dat
- riziko „nálepkování“ dětí na základě predikce
- napětí mezi prevencí a právem na soukromí
- kdo nese odpovědnost za rozhodnutí – člověk nebo algoritmus?
- férovost a nediskriminace algoritmů

- ochrana nejlepšího zájmu dítěte v technologickém kontextu
- komunikace s veřejností o využívání AI
- etické kodexy sociální práce a AI
- princip opatrnosti při implementaci
- příklady etických pochybení

## **8. Kazuistiky a příklady z praxe**

- využití prediktivního modelu při odhalení domácího násilí
- selhání algoritmu vedoucí k neoprávněné intervenci
- případy falešně pozitivní predikce rizika
- zkušenosti z Velké Británie (projekt Troubled Families)
- zkušenosti z USA (Allegheny Family Screening Tool)
- příklady dobré praxe ze Skandinávie
- reflexe českého prostředí a jeho specifik
- co se můžeme naučit z cizích chyb
- přenositelnost zahraničních modelů do ČR
- inspirace pro metodickou práci OSPOD

## **9. Praktická doporučení pro praxi**

- jak nastavit hranice využití AI v OSPOD
- tvorba metodik a vnitřních pravidel
- spolupráce s IT specialisty a právníky
- pravidelný audit algoritmů
- vzdělávání pracovníků v oblasti AI a dat
- komunikace s veřejností a rodinami o využívání dat
- zavádění pilotních projektů s kontrolními mechanismy
- reflexe a evaluace zkušeností
- doporučení pro obecní a krajskou úroveň
- dlouhodobé trendy a vývoj

## **10. Diskuse a závěr**

- shrnutí možností a limitů AI
- doporučené postupy pro OSPOD
- reflexe etických a právních dilemat
- sdílení zkušeností účastníků
- diskuse nad konkrétními kazuistikami
- návrhy na metodické postupy
- doporučení pro praxi v obcích
- inspirace pro další vzdělávání