

# Podvádění nepřišlo s AI. Měli bychom se soustředit na podporu studentů, kteří se díky AI rozvíjí, říká Koňas

4.12.2025 - | Vysoké učení technické v Brně

**Podvody a opisování jsou problém starý jako lidstvo (či minimálně) písmo samo. Nástroje umělé inteligence udělaly podvádění snadnější, respektive možnost odhalit je náročnější. Místo, abychom se soustředili na ty, kdo etickou hranici překračují, máme se podle AI ambasadora Petra Koňase zaměřit na studující, kteří umělou inteligenci využívají k tomu, aby se rozvíjeli. Kdy můžeme mluvit o zneužití AI a kdy o dobrém pomocníkovi? A v čem může studentům a zaměstnancům pomoci nový web AI@VUT?**

**Před půl rokem jste se stal AI ambasadorem VUT. Co se od té doby změnilo?**

Doufám, že hodně a že je to poznat. Vznikl seriál přednášek a školení, kterými jsme se snažili oslovit různé skupiny: od studentů, přes výzkumníky a pedagogy až po zaměstnance administrativy. A musím říct, že zejména u posledně jmenovaných to mělo velkou odezvu, protože cítí potenciál, že jim AI může výrazně zjednodušit práci. Už dokonce vznikl první nástřel pro poloautomatizované vypracování zpráv ze služebních cest, o kterém jsme se bavili posledně. Také tvoříme nástroje, které by nám měly pomoci při vytváření žádostí o granty, respektive obecně zvládnout projektovou agendu. A relativně čerstvá věc je portál AI@VUT ([link ai.vut.cz](https://ai.vut.cz)), kde se snažíme koncentrovat informace, aby každý, kdo si na VUT na AI vzpomene, věděl, kam zamířit.

**Vy jste si dal nelehký úkol, totiž více centralizovat živelný rozvoj AI na univerzitě, aby se například víckrát netvořily stejné nástroje. Daří se to?**

Jsme teď ve fázi, kdy si myslím, že to jinak než trochu živelně nejde. Pokud lidé chtějí něco dělat, tak jim nezbyvá, než si ve svém oboru najít nástroje a cestičky, které fungují. Na druhou stranu bych byl rád, kdyby existoval bod, kde bychom sbírali už ozkoušené a snižovali tak nároky, které souvisí se zmíněnou redundancí. Snažím se proto sloužit jako takový koncentrátor. A zdá se, že to začíná fungovat, lidé se mi ozyývají. Ale mít uzavřený seznam nástrojů a žádné jiné, to by mi přišlo zbytečně svazující. Když lidem necháme dostatečnou volnost a zároveň je upozorníme na rizika, která plynou z příliš velké benevolence nebo z úniku dat, jde myslím o rozumný kompromis a správnou cestu.

Otázka, kterou by si měl každý vždy položit je: „Jsem smířen s tím, že tato data může vidět kdokoliv?“. Protože jakmile opouští univerzitu, v zásadě se stávají veřejnými...

**Jaký je o oblast AI na VUT zájem?**

Velký a jsem strašně rád. Většina našich školení měla účast od šedesáti osob výš, na tom posledním jsme měli přes sto účastníků. To byla zrovna oblast AI agentů, která poměrně dost rezonuje s funkcionalitou, kterou AI nabízí v posledním půlroce. Spousta lidí v nich vidí potenciál, jak zautomatizovat procesy, které dělají.

**VUT je univerzita nejen technická, ale i umělecká, máme zde FaVU a Fakultu architektury. Co může AI nabídnout uměleckým oborům, které stojí primárně na lidské kreativitě?**

Všichni na univerzitě řeší problém byrokracie, to se týká i zmíněných fakult. Ale asi vím, kam směřujete... Na to je strašně těžká odpověď. Já si vždycky vzpomenu na Platonovou jeskyni, kdy člověk uzavřený celý život v jeskyni má za to, že je to celý svět. Jinými slovy uzavřenost někdy může vést k tomu, že nepoznáme, co je skutečně obsahem našeho světa a co je za ním. Já věřím, že pokud se umělecké směry dokáží oblasti AI otevřít, bude to pro ně přínos. A musím říct, že na zmíněných fakultách ta snaha je.

**Hodně se mluví i o možných rizicích AI. Nechci teď rozebírat rizika pro lidstvo, spíš se ptám čistě prakticky: na co si musíme všichni - od studentů po zaměstnance - dávat pozor, když AI používáme?**

## **AI nám vrací obraz nás samých**

Starší rozhovor s Petrem Koňasem "**AI nám vrací obraz nás samých**" z července 2025 si můžete přečíst [zde](#).

Rád bych zdůraznil jednu věc, která asi napadne každého, ale přesto na ni spousta lidí rezignovala: bezpečnost dat. Kardinální otázka, kterou by si měl každý vždy položit je: „Jsem smířen s tím, že tato data může vidět kdokoli?“. Protože jakmile opouští univerzitu, třeba nahráním do cloudu, v zásadě se stávají veřejnými; už jsme se mnohokrát setkali s tím, že poskytovatel služby sice něco tvrdil nebo uváděl ve smluvních podmínkách, ale přesto došlo k úniku, přesto s daty obchodoval a data tak byla z tohoto pohledu ztracena. Nechci poskytovatelům nasazovat psí hlavu, ale jim skutečně jde primárně o zisk, ne o naše soukromí. Soukromí používají jenom jako argument pro další zvyšování zisku. Měli bychom myslet na to, co je jejich cílem a že to není naše ochrana. Pokud k tomu budeme takto přistupovat, možná budeme trochu opatrnější.

Další věc je psychologická: čím dál častěji si budujeme na AI závislost. Tak jako u všeho by měl člověk hledat rozumný přístup, ne ji využívat 24/7. AI je sice zajímavý nástroj, ale není všespásný a může nás přehltnout.

A poslední je riziko řekněme geopolitické. Například u čínských modelů, které jsou často velmi kvalitní, si nemůžeme být jistí, jaké jsou informace, na kterých byly natrénovány. Nebo zda enginy, na kterých běží nemohou sloužit k vytahování či promazávání dat nebo k implementování škodlivého softwaru do vašeho počítače.

**Existují v této souvislosti nějaké běžně dostupné doporučené AI nástroje pro VUT a na druhé straně nějaký blacklist zakázaných nástrojů?**

Na CISu je uveden seznam dostupných nástrojů, který se průběžně aktualizuje (*pouze po přihlášení k VUT účtu, pozn. red.*). Patří tam určitě Copilot, ChatGPT a Gemini - tady stojí za zmínku, že Google pro studenty nedávno otevřel možnost využívat Gemini na rok zdarma a určitě doporučuji se s ním seznámit, je to dnes skutečně jednička mezi jazykovými modely.

Žádný blacklist vysloveně zakázaných modelů neexistuje. Nicméně Národní úřad pro kybernetickou a informační bezpečnost letos vydal varování před některými produkty čínské společnosti DeepSeek, které vláda následně zakázala používat ve státní správě. Ty bych označil jako silně nedoporučené a raději bych se jim vyhýbal.

## Studujeme pro sebe, AI je jen rádce

**Nedávno byl spuštěn zmíněný portál AI@VUT. Co na něm najdeme, ať už na VUT pracujeme či studujeme?**

Teď je tam naprostý základ informací, které s AI souvisí. To znamená, že máme vyhlášku, která určuje, jakým způsobem s AI pracovat, kde a jak ji můžeme využívat. Je tam rozcestník na AI nástroje, ať už komerční nebo open source. Najdete tam odkazy na spoustu tutoriálů, návodů a školení. Je to takový základní rozcestník.

Nyní pracujeme na dílčích návodech přímo pro konkrétní skupiny. Například pro studenty máme stránku, která řeší jejich nejčastější otázky a nabízí odkazy na vhodná školení. Odpovídáme na to, jak může AI studentům pomoci při tvorbě jejich práce, co mohou, co nemohou, jakou roli při zadávání práce hraje pedagog a kde už se při využívání AI dostanou do šedé zóny nebo za hranu. I toto plánujeme dále doplňovat, zatím třeba chybí forma citací pro jazykové úpravy, jako korektury a stylistické úpravy pomocí AI. Dávalo by mi smysl, mít je předem deklaratorně uvedeny v jednotné podobě, aby je studenti nemuseli sami formulovat.

**Vy jste dotknul oblasti důležité, a tou je etika. Studenti na webu najdou například checklist, kde si mohou sami zkontrolovat zapojení AI do tvorby své práce. Na co se mám sám sebe ptát, když při studiu využiji AI?**

Předně chci říct, že jsem strašně rád, že většina kantorů nebere studenty jako ty, kteří se a priori snaží podvádět. To je podle mne důležité, protože pokud studenti cítí podporu, pak hledáme cesty, jak AI rozumně využít. Zároveň je mylné myslet si, že podvádění a opisování začalo s AI. Pokud nějaký pedagog dává deset let stejná zadání, pak měli studenti už dřív spoustu příležitostí, jak vycházet ze starších prací a jak si tyto informace sdílet. Otázka etiky, podvádění či zneužívání tady byla vždycky a zpravidla ti samí studenti, kteří doteď hledali odpověď na Seminárky.cz, budou spíše zneužívat AI. My jako pedagogové bychom se měli soustředit na to, jak podpořit ty studenty, kteří se snaží využívat AI efektivním způsobem a k vlastnímu rozvoji. Já třeba díky AI můžu zadávat práce, které jsou mnohem komplikovanější, protože vím, že studenti si informace snadno dohledají.

Pokud AI použiju jako náhradu své vlastní práce, která se ode mne očekává, ocitám se za hranou. A řeším etické dilema, proč tady vlastně studuji, když ne kvůli sobě?

**Kdy tedy můžeme mluvit o zneužití AI a kdy o dobrém pomocníkovi?**

Na webu je jasně vysvětleno, že pokud AI použiju jako náhradu své vlastní práce, která se ode mne očekává, ocitám se za hranou. A řeším etické dilema, proč tady vlastně studuji, když ne kvůli sobě? Protože jako student se mám rozvinout v nějaké oblasti a AI mi má maximálně pomoci, například pokud nerozumím nějaké látce. V okamžiku, kdy mne nahradí a napíše za mne seminární práci, je to překročení červené linie.

Odhalovat podobné práce je obtížné, neexistuje spolehlivý nástroj, který rozezná výtvar AI od lidského. A s pokročilostí modelů se to bude jen zhoršovat. Pozoruji ale na učitelích, že jsou dost podezřívaví a jakmile mají dojem, že to není z hlavy studenta, začnou se ptát. Moje jednoduchá rada pro studenty je: pokud se chcete vyhnout tomu, aby se učitel příliš ptal, snažte se problematiku zvládnout sami, tvořte sami. AI může být rádce, oponent, učitel. Pak si ušetříte případné etické

dilema, ať už ze strany učitele nebo se svým svědomím.

**Není tedy AI i jistou výzvou a zrcadlem, zda věci, které děláme, dávají smysl? Protože pokud může práci vypracovat AI na jedno kliknutí, je otázka, zda má smysl ji v dnešní době studentům zadávat...**

Je to tak. Pedagogové by se měli soustředit na to, jaký prostor AI dávat, aby předmět posouvala. To znamená najít niku, která by například poukázala na novou oblast výzkumu nebo jiný posun. Pak nebudeme mít témata opakující se deset let, naopak mohou vznikat nové myšlenky. Snažím se kolegům ukazovat, že existují způsoby, jako například metapromptování, kdy se snažím zeptat AI, jak se mám zeptat, aby bylo dosaženo nějakého cíle. Takto si jako pedagog mohu pomoci při vytváření sad otázek nebo témat, která budou netradiční, budou vhodně navázaná na AI a budou nás nějakým způsobem posouvat.

To mi připomíná relativně nedávné články o AI, totiž že pokročilé modely začínají přejímat i negativní vlastnosti lidí a jedna z těchto vlastností je lhaní a fabulace. Už to není jenom halucinace, která vychází z pravděpodobnostního modelu. Najednou, i když implementujete zpětnou kontrolu, se jazykové modely jakoby snaží hledat cestu nejmenšího odporu, která pro ně bude nejméně „energeticky“ náročná. A je úžasné, že v zásadě kopírují nás lidi. Také se snažíme rozhodnout, jaké je optimum „energie“, kterou chceme investovat do jakékoliv činnosti, kterou dnes a denně děláme. Ale bohužel se ztrácí efekt, kdy bychom chtěli, aby AI byla naopak lepší než my. Ona se místo toho stává až moc dokonalou kopií nás samotných. A ačkoliv jsem obecně vzato technooptimista, v tomto moc optimistický nejsem. Lidstvo se snaží zlepšovat desítky tisíc let, a přesto jsme nenašli univerzální způsob, který by byl v čase neměnný a pro lidství typický, snad kromě naší schopnosti se adaptovat. Možná, že skutečně neexistuje žádná absolutní pravda, na druhou stranu bychom o ni pořád měli usilovat. A pokud se nám podaří AI přesvědčit, že snaha o absolutní pravdu je i pro ni přínosná, můžeme být optimističtí.

(ivu)

<https://www.zvut.cz/tema/-f38144/podvadeni-neprislo-s-ai-meli-bychom-se-soustredit-na-podporu-studentu-kteri-se-diky-ai-rozviji-rika-konas-d311114>