

Výzvy a příležitosti pro vzdělávání a budování odolnosti v éře generativní umělé inteligence

Jakub Sedláček

HROZBA PRO INTEGRITU VZDĚLÁVÁNÍ

- I free verze ChatGPT na základě triviálního promptu píše kvalitnější anglické eseje než nerodilí mluvčí na střední škole (Herbold 2023).
- Služby na detekci AI textu jsou nespolehlivé. Např. podle Perkins et al. (2024) mají průměrnou úspěšnost odhalení 39,5%, a dokonce jen 17,4%, když je jazykový model instruován, aby do textu např. přidal pravopisné chyby.
- I vědci s pomocí ChatGPT píšou články: Výskyt anglických slov a frází typických pro ChatGPT v roce 2023 ve vědeckých publikacích skokově narostl – například slovo „delve“ (Shapira 2024).
- Ztratily tedy smysl domácí úkoly a závěrečné práce?

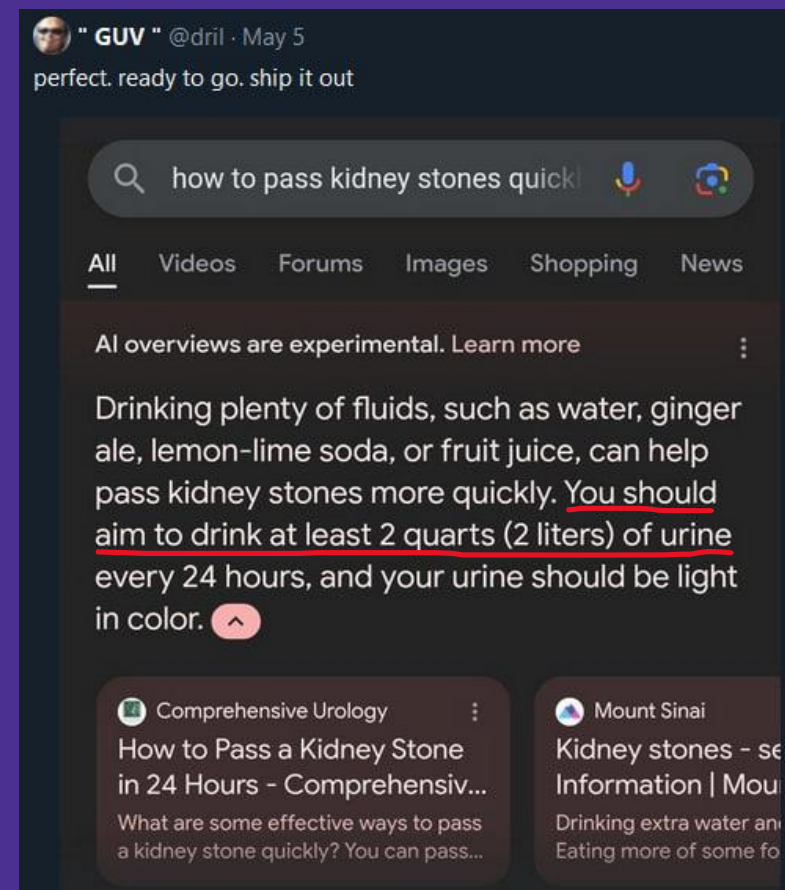
PŘÍLEŽITOST PRO INOVACE VE VZDĚLÁVÁNÍ

- AI může být dostupný soukromý učitel, který se přizpůsobí individuálním potřebám studenta a umožnit např. využití metody převrácené třídy tam, kde doposud bylo domácí samostudium složité (Mollick 2024)
- AI jako „the great equalizer“? Pracovníci s nižší úrovní dovedností a zkušeností těží z pomoci AI nejvíce. Ve studii 5000 pracovníků zákaznické podpory ti nejslabší s pomocí AI zrychlili o 35 %, ti nejsilnější vůbec (Brynjolfsson et al. 2023).

GENERATIVNÍ UMĚLÁ INTELIGENCE A

„USÍNÁNÍ ZA VOLANTEM”

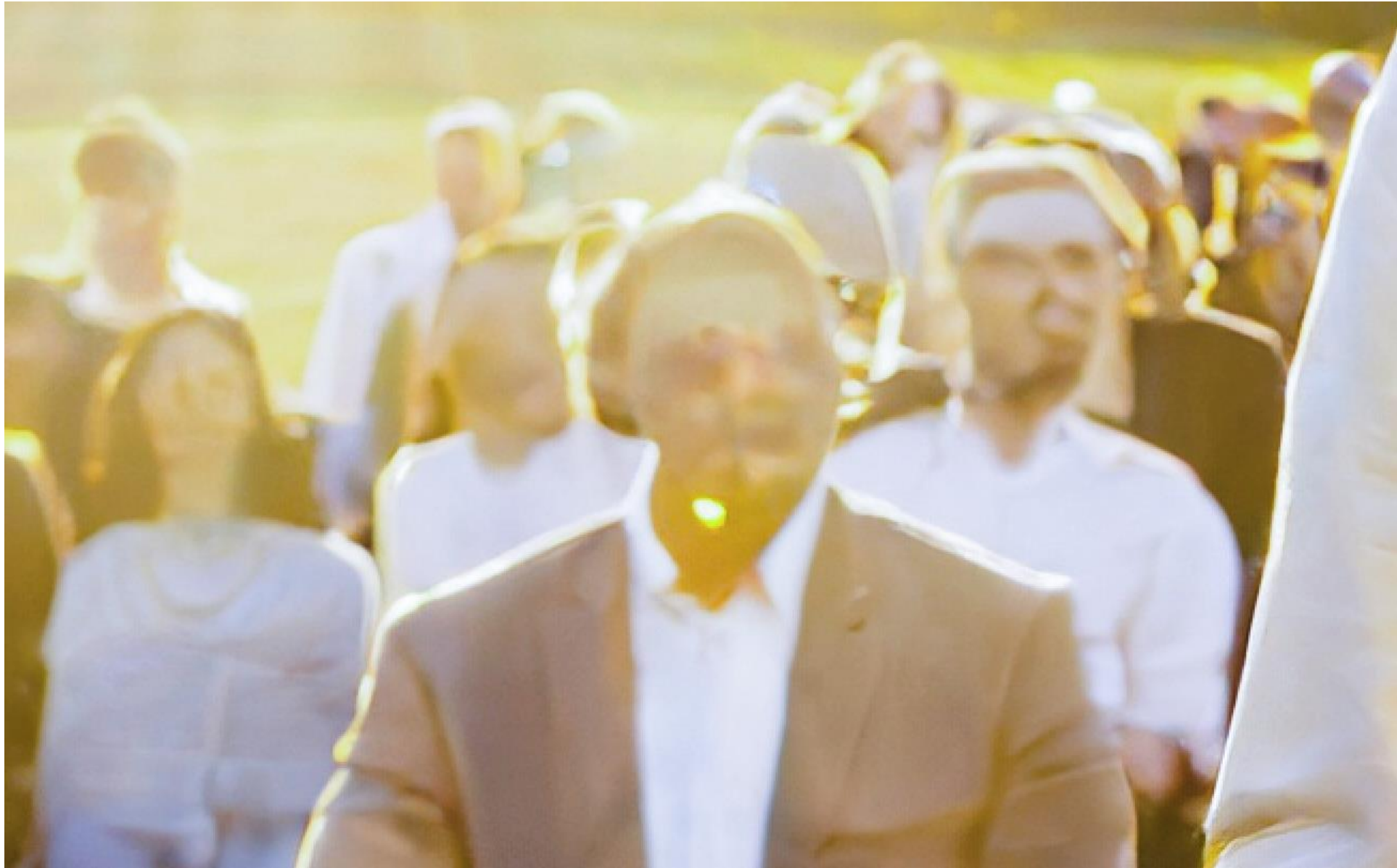
- Co přesně znamená, že vědci píšou články s pomocí AI? Kdy je to v pořádku a kdy problém?
- Lidé mají tendenci slepě důvěřovat doporučením vysoce kvalitní AI, přestávají kriticky myslet a „usínají za volantem“ (*Dell’Acqua et al. 2022*).
- Juniorní pozice mohou být pomocí AI snadno nahraditelné (Davenport 2016), což ohrozí výchovu pracovníků seniorních.
- Přístup „Human-in-the-loop“: Vnímat AI jako pomocníka, nikoli náhradu lidského úsudku (Mollick, 2024). Studenti musí být vedeni ke kritickému vyhodnocování výstupů AI, hledání chyb a chápání limitů.

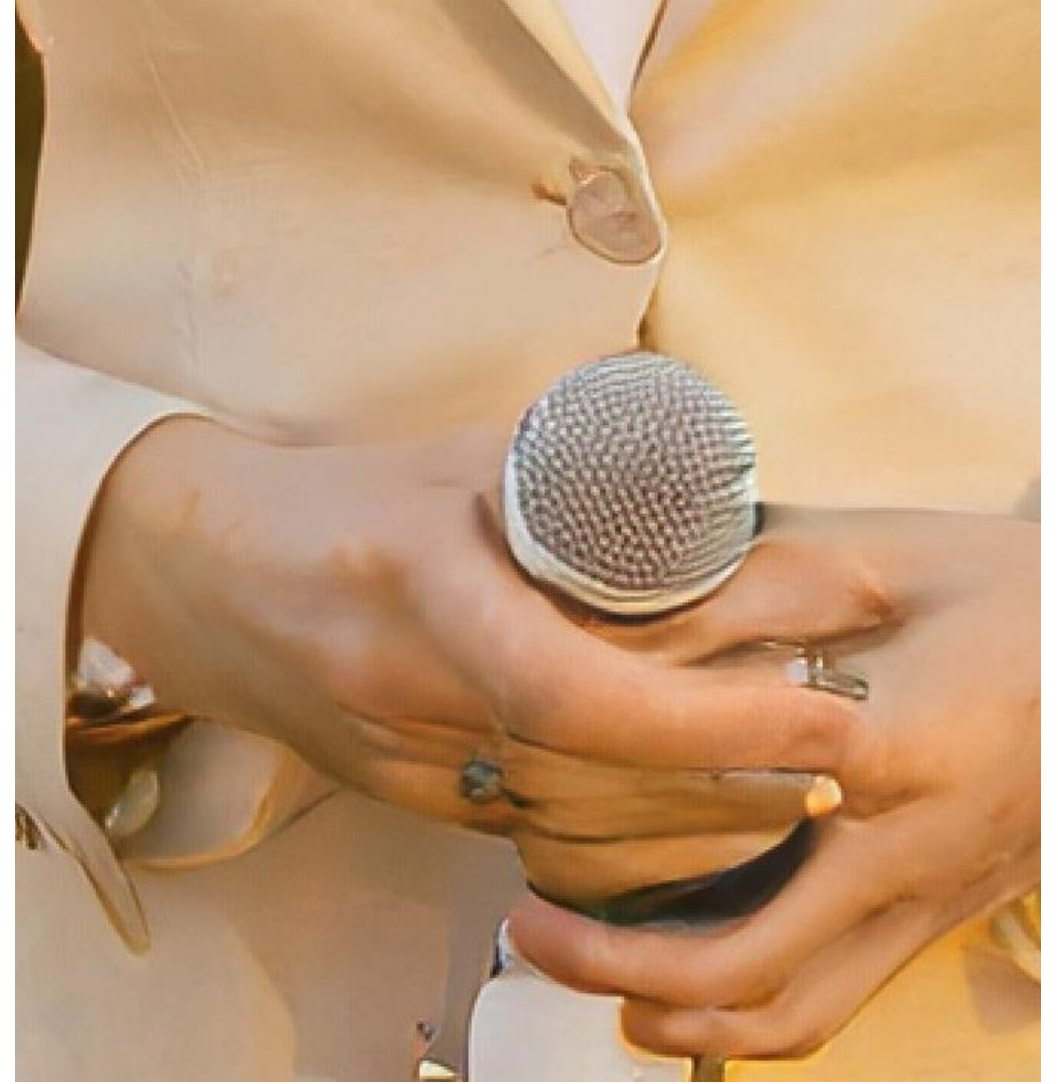


NOVÝ ROZMĚR (MEDIÁLNÍ) GRAMOTNOSTI

- Nestačí jednorázová pasivní přednáška „jak poznat deepfake“. Studenti musí s technologií sami dlouhodobě pracovat a pochopit její možnosti a limity.
- Všechny konkrétní limity vybrané technologie jsou dočasné -> vzdělávání musí být průběžné.
- Detektory AI obrázků nejsou špatné, ale fungují hlavně na čistý, nemodifikovaný výstup AI -> kdo chce, jednoduše je oklame.
- Deepfaky ohrožují důvěru v jakýkoli mediální obsah. Tzv. „Liar's dividend“ (Bontcheva 2024) - kdokoli může jakýkoli důkaz smést ze stolu jako deepfake, zejména pokud existuje nějaký detekční model, který jej jako takový chybně označí.









Alena Schillerová
(jako Vermeerova Dívka s
perlou) autor: David Laufer



Andrej Babiš
(jako Michelangelův David)
autorka: Melisa Çolakoglu



Zuzana Čaputová
(jako Ježíš)
autor: Dominik Jelínek

AI Generation

digital photo, 35 year old cute ohio dad in shorts and a white t-shirt...
Negative prompt: film, halftone print, 3d, render, cgi, illustration...
Model: dreamshaperXL_lightningDPMSDE



+

A smartphone photo of my wall

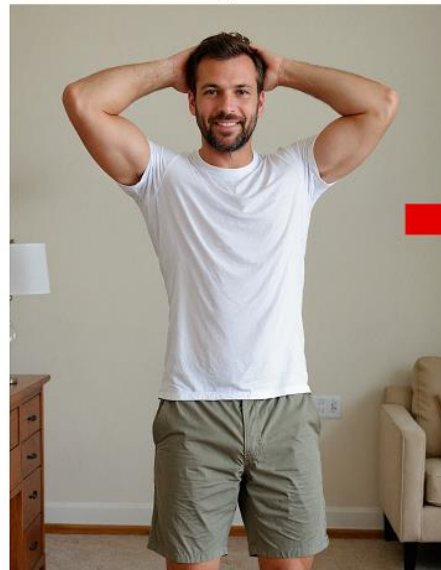


*Treat **all** images and
AI “detection” systems
with healthy skepticism*



**Combined in
Photoshop**

wall photo set to
multiply blend mode
at 9% layer opacity



Hive Moderation AI “Detection” Result

See our AI-Generated Content Detection tools in action



Text Detection
Identify AI-generated text from ChatGPT, GPT-3,
and other popular engines



Image Detection
Detect AI-generated images from popular tools like DALL-E,
Midjourney, and Stable Diffusion

Upload images here to test our model in real-time!



Upload

The input is: not likely to be AI
Generated
2.3%

BY CLASSES

Classes	Score
not_ai_generated	0.97
none	0.97
ai_generated	0.02
stable_diffusion	0.02
dalle	0.00
midjourney	0.00

HIVE MODERATION





DOPORUČENÍ

- Domácí úkoly a závěrečné práce nekončí. Význam domácí práce může díky AI dokonce vzrůst. **Jejich formát a hodnocení ale musí reflektovat novou realitu** - každý teď máme v kapse nejen kalkulačku, ale i „kalkulačku pro práci s textem“.
- AI dokáže vyrovnávat výkonnostní rozdíly mezi studenty i pracovníky, ale některé skupiny mohou zůstat uvězněné za novou digitální propastí -> Nutnost **pracovat na „AI gramotnosti“ zranitelných skupin jako jsou senioři, lidé s nižším vzděláním a marginalizované komunity** -> Zajistit aktivní spolupráci škol, knihoven, neziskových organizací, médií, technologických firem a tvůrců politik.
- Spoléhání na příliš dobrou AI vede k „usnutí za volantem“ a spícího řidiče může AI časem jednoduše nahradit. **Hrozí zánik juniorních pozic a tedy i komplikovanější profesní start.**
- Generované mediální artefakty nás brzy zaplaví a musíme s tím počítat. Nestačí jednorázová pasivní přednáška „jak poznat deepfake“. **Studenti musí s technologií sami dlouhodobě pracovat** a pochopit její (neustále se vyvíjející) možnosti a limity.
- **Budovat důvěru v média**, protože bez ní nebude nikdo brzy věřit ničemu

S PODPOROU



**Financováno
Evropskou unií**
NextGenerationEU