



# Než řeknete ano AI: rizika, která by mělo vedení univerzity znát

---

Digitální transformace univerzit 2026

**Seyfor**

**Just sey it!**

Datum

**08.09.2026**

Prezentující

**Lucie Jahnová**

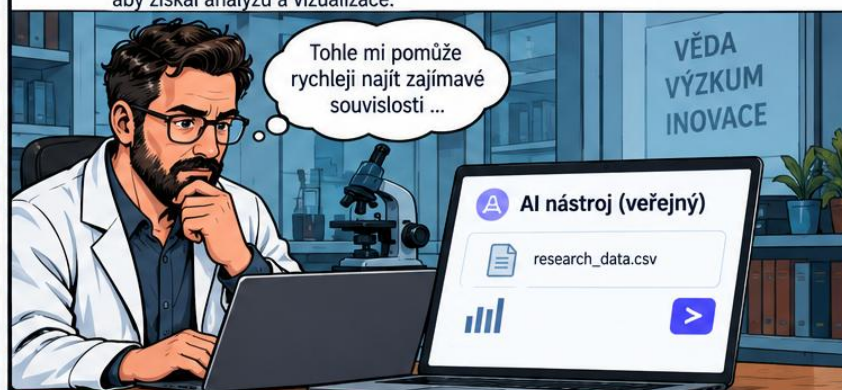
# AI na univerzitě nezakazovat. Ale vědět, čemu vlastně říkáme ANO

**1 STUDENT: DIPLOMOVKA DO CHATGPT**  
Student UPCE vloží svou diplomovou práci do veřejného AI nástroje.



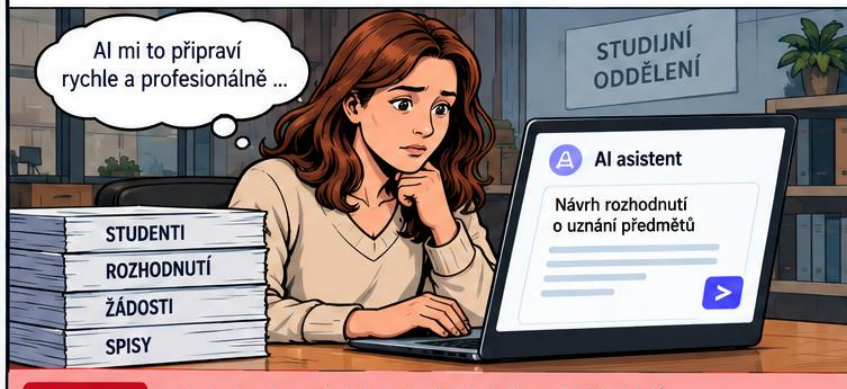
**RIZIKO:** Únik citlivého obsahu, ztráta originality, porušení pravidel pro závěrečné práce.

**2 VÝZKUMNÍK: DATA DO VEŘEJNÉHO AI NÁSTROJE**  
Výzkumník nahrává výzkumný dataset do veřejného AI nástroje, aby získal analýzu a vizualizaci.



**RIZIKO:** Únik výzkumných dat, porušení smluv (granty, NDA), ztráta konkurenční výhody, nejasné nakládání s daty.

**3 REFERENT: AI PŘIPRAVÍ ROZHODNUTÍ**  
Referent nechá AI připravit návrh rozhodnutí pro studenta.



**RIZIKO:** Chybná nebo neúplná rozhodnutí, únik osobních údajů (studenti), právní odpovědnost, nejasnost, kdo rozhodl.

**4 REKTORÁT: ENTERPRISE COPILOT**  
Vedení univerzity pořídí enterprise Copilot pro zaměstnance.



**RIZIKO:** Příliš široký přístup k datům, nechtečné sdílení informací, objevitelnost citlivých dokumentů, compliance a GDPR.



Všechno je AI. Ale  
riziko rozhodně není  
stejně.

## Trocha čísel...

Zdroj: [The Impact of AI on Work in Higher Education](#) | EDUCAUSE

94 %

použilo AI pro pracovní účely během posledních šesti měsíců

73 %

je těch, kteří AI používají, ji používá denně nebo týdně

86 %

chce AI používat i v budoucnu

54 %

ví o pravidlech či guidelines své instituce pro pracovní použití AI

56 %

použilo pro práci AI nástroje, které jim jejich instituce neposkytla

# ZÁKAZ AI



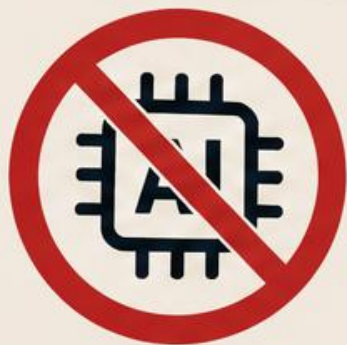
# AI SE NEPOUŽÍVÁ



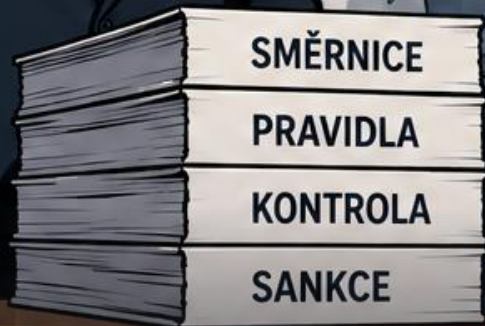
UNIVERZITA  
PARDUBICE



- ✗ ŽÁDNÝ ChatGPT
- ✗ ŽÁDNÉ AI NÁSTROJE
- ✗ ŽÁDNÉ AI AGENTY
- ✗ ŽÁDNÁ GENERATIVNÍ AI



POUŽÍVÁNÍ AI  
**ZAKÁZÁNO**



SMĚRNICE  
PRAVIDLA  
KONTROLA  
SANKCE

STEJNÍ LIDÉ. JEN JINÉ KANÁLY.



UNIVERZITA  
PARDUBICE

Gemini

Claude

Copilot

ChatGPT

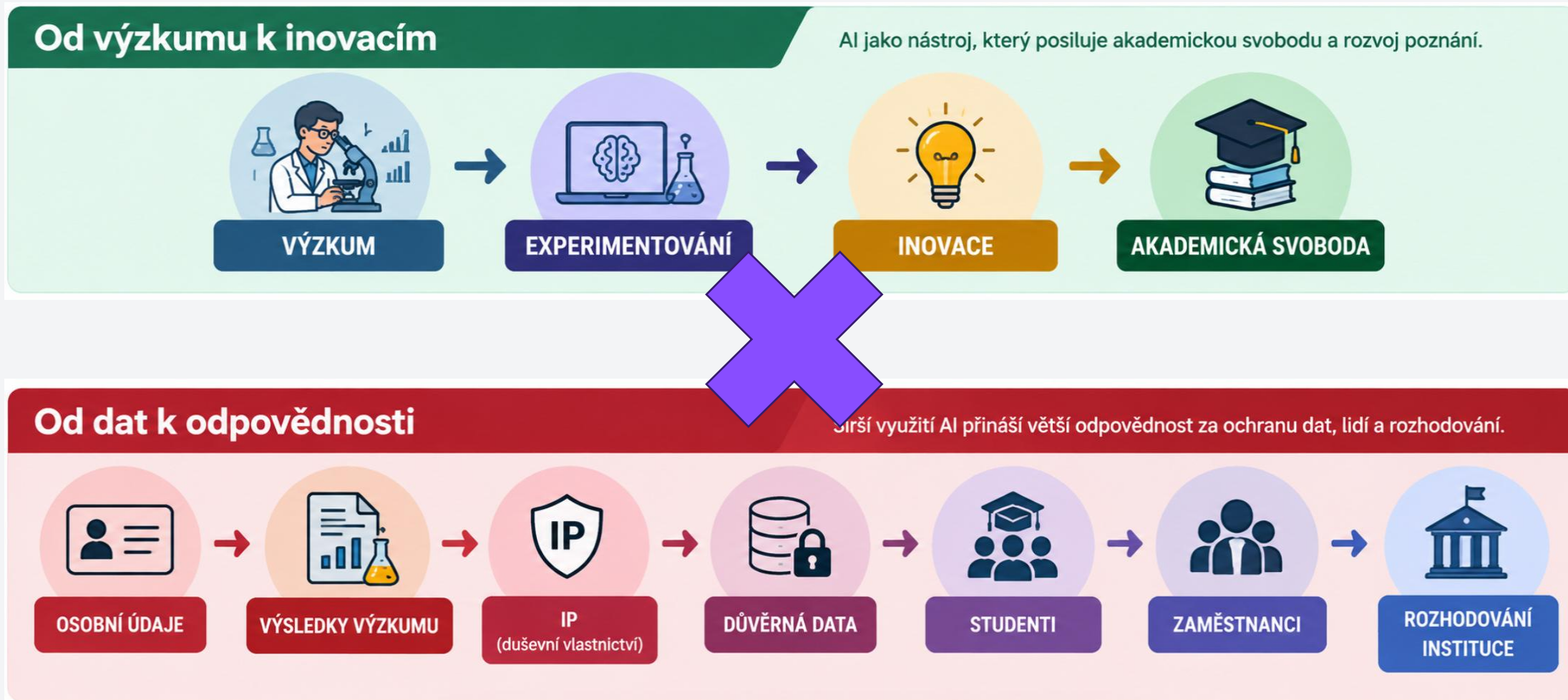
UPCE



Seminární práce  
- Nápady  
- Shrnutí  
- Rešerše

RYCHLEJI  
SNADNĚJI  
LEPŠÍ VÝSLEDKY  
STEJNÁ UNIVERZITA

# / Co musí univerzita zvažovat?



# / Co musí univerzita zvažovat?

## Od výzkumu k inovacím

AI jako nástroj, který posiluje akademickou svobodu a rozvoj poznání.



## Od dat k odpovědnosti

Širší využití AI přináší větší odpovědnost za ochranu dat, lidí a rozhodování.



„Kde může univerzita AI nechat volně žít a kde už musí začít řídit její použití?“



Výzkum



Vzdělávání



HR



Administrativa



Rozhodování



Bezpečnost

# / Jak k tomu přistoupit?

## ■ 3 ZÓNY POUŽITÍ:



### Experimentuj

Nízké riziko – otevřené možnosti

Vhodné pro kreativní práci, učení se s AI a objevování nových možností.



**Veřejná data**  
(všeobecně dostupné informace)



**Brainstorming**  
(nápady, koncepty, kreativní úvahy)



**Rešerše**  
(vyhledávání a sumarizace veřejných zdrojů)



**Pomoc s textem**  
(úpravy, překlady, stylistika)



**Prototypování**  
(nápady, návrhy, testování konceptů)



**Obecně povoleno**  
Při dodržení základních zásad bezpečného a odpovědného používání.



### Používej za podmínek

Střední riziko – s rozvahou a podle pravidel

Možné využívat, pokud dodržíš pravidla univerzity (např. ochrana dat, licence, citace, schválené nástroje).



**Interní data**  
(neveřejné, ale nevysoce citlivé)



**Výzkumné podklady**  
(analýza, sumarizace, příprava)



**Práce studentů**  
(např. podpora při učení, příprava materiálů)



**Administrativa**  
(dokumenty, procesy, komunikace)



**Pouze za stanovených podmínek**  
Dodržuj interní pravidla, zvaž typ dat a používej schválené nástroje.



### Řízené použití

Vysoké riziko – vyžaduje kontrolu a schválení

Pouze v rámci schválených nástrojů a procesů, často s dohledem odpovědných osob (IT, právní oddělení, vedení).



**Citlivá data**  
(osobní údaje, důvěrné informace)



**Významné rozhodování o lidech**  
(přijímání, hodnocení, stipendia, kariérní postup)



**Hodnocení**  
(známkování, posuzování práce studentů)



**Autonomní akce AI**  
(agenti, automatizované procesy)



**Napojení na univerzitní systémy**  
(IS, LMS, e-mail, finanční a personální systémy)



**Pouze řízeně a se schválením**  
Vyžaduje posouzení rizik, jasný účel, kontrolu a odpovědnost.

# / Na co si dát pozor?

„Licence ≠ governance“

To, že univerzita koupila bezpečnější enterprise platformu, ještě neznamená, že vyřešila **způsob jejího používání**.

- A šest otázek, které musí VŠ znát před použitím:

Jaká data do AI  
posíláme?

Kde jsou zpracována a  
kdo k nim může mít  
přístup?

Používají se pro další  
trénování?

Jak dlouho se  
uchovávají?

Kdo má k AI přístup a  
pod jakou identitou?

Co smí výstup AI  
ovlivnit nebo způsobit?

# Jak se nám v budoucnu a možná už teď musí změnit perspektiva?

1

## Klasická GenAI

Odpovídá na pokyny člověka.



2

## Agentní AI

Nejen odpovídá, ale samostatně plánuje a vykonává kroky k cíli.



3

## Multiagent

Více agentů spolupracuje, komunikuje a koordinuje práci na složitých úlohách.



# / Praktický příklad



```
<> Python

agent = PARDubot()

students = agent.find_students(
    risk="study_failure",
    threshold="HIGH"
)

for student in students:
    context = agent.get_context(
        student,
        sources=["SIS", "LMS", "attendance", "results"]
    )

    recommendation = agent.ai.recommend_next_steps(context)

    agent.contact(
        student,
        message=recommendation
    )

    agent.log_action()
```



„Najdi studenty ohrožené studijním neúspěchem, kontaktuj je a navrhnij jim další postup.“

K jakým datům agent přistupuje?

Smí něco změnit v IS?

Co smí odvodit?

Kdo schválil jeho rozhodovací logiku?

Smí studentovi napsat?

A kdo odpovídá, když se splete?

# Proč je to zásadní?

## Před AI

**Uživatel musí sám najít, připravit a odeslat data.**

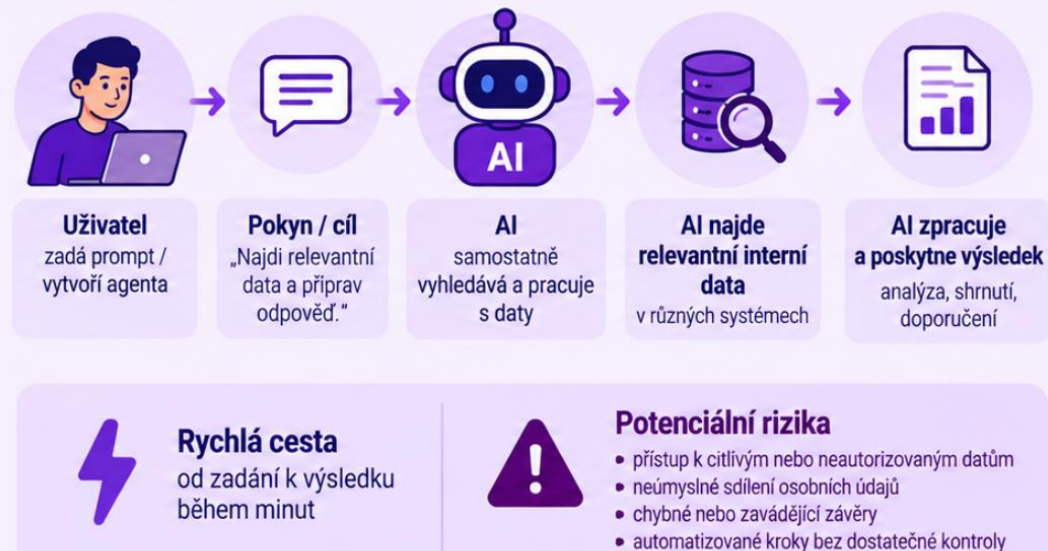
Proces je časově náročný, vyžaduje úsilí a výsledek je nejistý.



## S AI

**Uživatel zadá pokyn nebo vytvoří agenta a AI sama najde, zpracuje a využije relevantní interní data.**

Proces je rychlý a pohodlný, ale přináší nová rizika – AI může získat i data, která by uživatel měl nebo neměl mít.



# Univerzitní SharePoint

2 miliony dokumentů z různých oblastí



Diplomové práce



Grantové žádosti



Výzkumná data



Zápisy



Personální dokumenty



Smlouvy



Tabulky



Studentská data



Např.:  
Plán personálních  
změn 2025.docx

Oprávnění:  
Everyone



## Před Copilotem

Uživatel měl k dokumentu technicky přístup, ale musel ho sám najít.



**Uživatel**  
má přístup k SharePointu



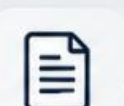
**1. Musel vědět, že existuje**

(tušit správné  
klíčové slovo /  
adresář)



**2. Najít ho**

(procházet složky,  
vyhledávat)



**3. Otevřít**

(a mít zájem  
ho číst)



**4. Pochopit jeho význam**

(že obsahuje  
citlivé informace)



### Praktická pravděpodobnost expozice

Relativně malá – vyžaduje aktivní hledání, čas a kontext.

## S Copilotem

Uživatel položí přirozený dotaz a Copilot vyhledá, shrne a předloží relevantní informace napříč celým úložištěm.



**Uživatel**  
zadá dotaz v Copilotu

„Shrň mi všechny  
informace o připravovaných  
personálních změnách  
na fakultě.“



**Copilot**

vyhledá v celém SharePointu  
(včetně dokumentů s oprávněním  
„Everyone“)

### Shrnutí klíčových informací:

- Plánované personální změny ...
- Harmonogram ...
- Dotčené fakulty ...
- Související dokumenty ...



### Výrazně vyšší pravděpodobnost expozice

Copilot může rychle najít a syntetizovat informace z dokumentů, o kterých uživatel ani nevěděl, že existují.



**Stejná data. Jiná cesta. Jiná míra rizika.**

Jde o to, co je technicky dostupné,  
ale i o to, jak snadno se k tomu lze s AI dostat.

# / Takže co je vlastně to riziko?

- Univerzita může mít:

**A:** jméno studenta  
**B:** výsledky studia  
**C:** aktivitu v LMS  
**D:** komunikaci  
**E:** informace o stipendiu  
**F:** knihovní data  
**G:** předchozí žádosti  
**H:** další administrativní údaje

- AI ale umí:

**A + B + C + D + E + F + G → inference**

A najednou se můžete ptát:

„Kteří studenti mají pravděpodobně finanční problémy a současně jim hrozí ukončení studia?“

**Nevznikla nová data. Vznikla nová informace.**

- AI nezvyšuje pouze množství dat.
- AI zvyšuje jejich dosažitelnost, propojitelnost a využitelnost.
- AI Agenti přidávají datům schopnost vyvolat akci.

# Úrovně AI schopností a související rizika

Čím vyšší autonomie a propojení, tím vyšší potenciální riziko.

| Úroveň | AI schopnost                                                                        |               | Co se mění                    | Riziko                                                                                |              |
|--------|-------------------------------------------------------------------------------------|---------------|-------------------------------|---------------------------------------------------------------------------------------|--------------|
| 0      |    | Bez AI        | Data existují                 |    | Nízké        |
| 1      |    | Chatbot       | Uživatel data vloží do AI     |    | Nízké        |
| 2      |    | RAG / Copilot | AI data sama vyhledává        |    | Střední      |
| 3      |   | Connected AI  | AI koreluje více zdrojů       |    | Vyšší        |
| 4      |  | Agent         | AI může jednat                |  | Vysoké       |
| 5      |  | Multi-agent   | AI systémy jednají mezi sebou |  | Velmi vysoké |

# Čím větší autonomie AI, tím silnější governance.

Vyšší schopnosti přinášejí větší užitek – ale také větší rizika. Proto roste potřeba řízení, kontroly a odpovědnosti.

**L0**

**AI radí**

Člověk všechno provádí.



AI = poradce

**L1**

**AI připravuje**

Člověk schvaluje.

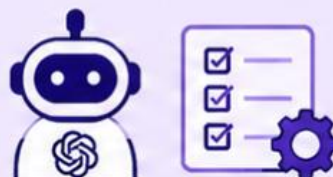


AI = příprava  
Člověk = rozhodnutí

**L2**

**AI provádí omezené akce**

V jasně definovaných mantinelech.



AI = provedení (v mantinelech)  
Člověk = nastavení pravidel

**L3**

**AI jedná samostatně**

Člověk vykonává dohled.

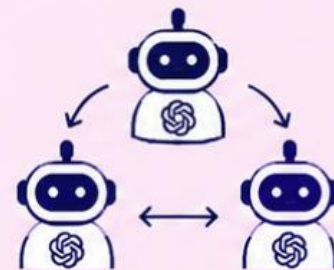


AI = samostatná akce  
Člověk = dohled

**L4**

**Více agentů spolupracuje**

A jedná autonomně.



AI = autonomní spolupráce  
Člověk = strategické řízení

Nižší potřeba governance

Vyšší potřeba governance



**Klíčové  
governance otázky**



**Kdo je  
vlastník?**

Kdo odpovídá za  
nasazení a provoz?



**Kdo schválil  
oprávnění?**

Jaké má AI přístupy  
a k jakým datům?



**Kdo kontroluje  
výstupy/akce?**

Jak ověřujeme kvalitu,  
spolehlivost a soulad  
s pravidly?



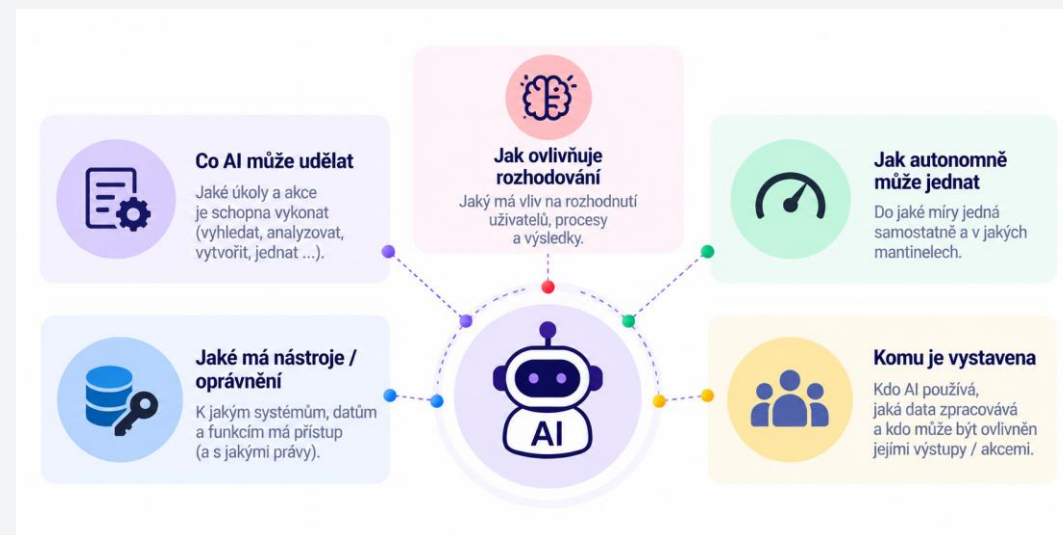
**Kdo nese  
odpovědnost?**

Kdo ručí za důsledky,  
pokud dojde k chybě?

# / Jak to řeší v zahraničí?

- Stanford u agentní AI už neklasifikuje riziko **jen podle citlivosti dat**.
- Univerzita přímo říká, že u osobních AI agentů je primárním rizikem **data exfiltration** – každý uživatel se může stát cestou, kterou citlivá data opustí hranice instituce přes prompty, přílohy nebo kontext.

## Co ještě hodnotí?



# / Co je tedy potřeba? 5 věcí, které byste měli udělat

1

## Inventář AI

Víme vůbec, kde AI používáme?



Přehled  
= základ pro řízení.

2

## Klasifikaci use-cases

Jaké případy použití jsou povolené, řízené a schvalované?



**Volné**  
Nízké riziko  
(možno používat)



**Řízené**  
Se stanovenými podmínkami



**Schvalované**  
Vyžadují formální schválení



Jasná pravidla  
= méně rizik.

3

## Pravidla pro data

Co do kterého AI prostředí smí?



Správná data  
= větší bezpečí.

4

## Odpovědnost

Každý významný AI use-case musí mít člověka / vlastníka.



Jasná odpovědnost  
= důvěra a kontrola.

5

## Pravidla autonomie

Co AI smí pouze navrhnout a co už smí sama vykonat?



**Pouze navrhnout**  
(člověk rozhoduje)

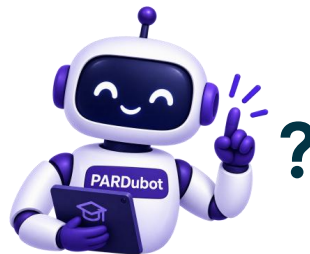
**Omezené akce**  
(v mantinelech)

**Samostatně vykonat**  
(dle pravidel)



Správně nastavené mantinely = bezpečná a užitečná AI.

# / A co by Vám na závěr poradil



## **Mějte přehled** (znáte svůj AI inventář?).

Víte, kde na univerzitě AI používáme, kdo ji používá a k čemu? Bez přehledu se těžko řídí rizika i příležitosti.



Co nevidíte,  
neřídíte.



## **Pracujte s daty** **zodpovědně.**

Ujasněte si, jaká data lze do AI prostředí vkládat a jak je chránit. Citlivá data vyžadují zvláštní pravidla.



Správná data  
= menší riziko.



## **Nastavte míru** **autonomie.**

Rozhodněte, co AI jen navrhuje, co provádí s omezeními a co může dělat samostatně. Čím větší autonomie, tím silnější governance.



Jasná pravidla  
= předvídatelné  
chování.



## **Určete odpovědnost.**

Každý významný AI use-case musí mít člověka – vlastníka, který nese odpovědnost za výsledky, dohled a soulad s pravidly univerzity.



AI má být  
pomocník,  
a anonymní autor.



## **Stanovte mantinely.**

Určete, kde už potřebujete řízení, schválení nebo dodatečné kontroly – podle rizik a dopadů na univerzitu.



Bezpečná  
inovace má  
jasné hranice.

# Děkuji za pozornost

AI není hrozba ani zázračné řešení.  
Je to nástroj – a s dobrými pravidly  
může být skvělým kolegou.

