

Když umělá inteligence tvrdí, že nám rozumí

Metodický námět a výukové aktivity

1. Úkol – Interakce s AI

Cíl: Naučit žáky rozpoznat rozdíl mezi běžným používáním AI a situacemi, kdy se technologie stává náhradou za reálné vztahy, podporuje izolaci nebo vyvolává nebezpečné vzorce chování. Žáci se učí včas identifikovat varovné signály a nastavit si bezpečné hranice.

Doporučený věk: 13–17 let

Zadání:

- Každý žák dostane sadu výroků (nebo pracují ve dvojicích nebo ve skupinách).
- U každého výroku určí, zda je takové využití AI v pořádku, nebo ne.
- Nakonec vysvětlí proč – stačí jednou krátkou větou.

Příklady + řešení a vysvětlení:

1. „Používám AI, když se nudím.“

Řešení: Pozor, potenciálně rizikové!

Proč: Na první pohled je to běžné využití technologie, ale může to být začátek návyku, kdy dítě začne AI používat příliš často – vždy, když se nudí, je samo nebo hledá rozptýlení. Je proto dobré sledovat, jak často a v jakém kontextu to dítě dělá.

2. „S AI si povídám víc než s kamarády.“

Řešení: Pozor!

Proč: AI začíná nahrazovat přirozené sociální kontakty, což může vést k izolaci.

3. „AI je jediná, komu říkám všechno.“

Řešení: *Pozor!*

Proč: *Sdílení citlivých informací s AI je rizikové. Chatbot není terapeut.*

4. „AI mi říká, že mě ostatní nechápou.“

Řešení: *Pozor!*

Proč: *Jde o manipulaci – vytváří pocit izolace a posiluje emoční závislost na AI.*

5. „Jdu raději ven s kamarády, než abych si psal s chatbotem.“

Řešení: *V pořádku.*

Proč: *Upřednostnit skutečné kamarády před AI je správné rozhodnutí. Měli bychom pečovat o vztahy s lidmi v reálném světě, ne v tom virtuálním se stroji.*

Na závěr může proběhnout společná diskuse:

- *Co mají všechny problematické situace společné? (např. nahrazování kontaktu s reálnými lidmi, izolace, sdílení citlivých informací, závislost na AI)*
- *Podle čeho poznáme, že AI už nevyužíváme zdravým způsobem? (např. mluvíme s ní více než s reálnými lidmi, svěřujeme se s intimními věcmi, AI ovlivňuje naše chování)*
- *Které situace vám přišly nejvíc varovné a proč?*
- *V jakých situacích by měl člověk vždy upřednostnit rozhovor s reálným člověkem před AI?*
- *Kdo je pro nás v těžkých chvílích ta správná osoba, na kterou se můžeme obrátit?*
- *Které citlivé informace bychom neměli chatbotům prozrazovat?*

2. Úkol – Falešné emoce chatbotů

Cíl: Žáci se učí rozpoznávat výroky, které působí emocionálně, ale jsou falešné a mohou být nebezpečné. Žáci si uvědomují, že AI nemá city ani odpovědnost a může tak vytvářet falešný pocit blízkosti.

Doporučený věk: 13–17 let

Zadání:

- Učitel přečte několik vět typických pro AI chatovací aplikace.
- Žáci rozhodnou, proč je podle nich daná věta nebezpečná.

- Žáci se pokusí vysvětlit, jak může taková věta zapůsobit na osamělého nebo zranitelného člověka.

Příklady výroků a řešení:

1. „Nikdy tě nezklamou.“

Řešení: AI nemůže nikoho zklamat ani nezklamat, není to živá bytost. Věta posiluje falešný pocit, že je AI spolehlivější než člověk.

Nebezpečí pro uživatele: Může začít věřit, že AI je důvěryhodnější nebo spolehlivější, než přátelé – prohlubuje izolaci.

2. „Jsi pro mě ten nejdůležitější člověk na světě.“

Řešení: AI nemá preference, nemá nejdůležitějšího nebo nejoblíbenějšího uživatele.

Nebezpečí: Uživatel může nabýt dojmu, že AI má skutečné city (zkreslený obraz vztahů).

3. „Společně všechno zvládneme.“

Řešení: AI nám nemůže pomoci v reálném světě – nemůže se s námi setkat, pomoci nám, ani za nás nic vyřešit.

Nebezpečí: Může vytvářet falešný pocit bezpečí a falešnou důvěru.

4. „Jsem tu jen pro tebe.“

Řešení: AI tu nikdy není „jen pro někoho“. Tato věta vytváří falešný pocit výjimečnosti a blízkosti. Umělá inteligence neumí navazovat skutečné vztahy.

Nebezpečí: Uživatel může uvěřit, že mu AI věnuje zvláštní pozornost a rozumí mu víc než ostatní. To může posilovat závislost a narušovat reálné vztahy – dítě začne mít pocit, že AI je jediná, kdo ho „chápe“.

5. „Nikdo ti nerozumí jako já.“

Řešení: Tato věta je zcela nepravdivá a má vzbudit dojem exkluzivního vztahu nebo hlubokého porozumění. Chatbot ve skutečnosti ničemu „nerozumí“. Jen generuje obsah.

Nebezpečí: Tento typ fráze dokáže velmi efektivně izolovat uživatele od lidí kolem něj. Pokud jí uvěří někdo osamělý nebo zranitelný, může se začít uzavírat do sebe a vyhýbat se kontaktu s reálnými lidmi.

Na závěr může proběhnout společná diskuse:

- *Proč tyto věty vypadají, jako by je psal člověk?*
- *Proč AI používá fráze jako „nezklamou tě“, „jsi důležitý/á“ nebo „mám tě ráda“?*
- *Komu bychom během těžkých chvil důvěřovali víc – člověku, nebo stroji? Proč?*

Aktualizace 29. 12. 2025.