

Univerzita Palackého v Olomouci
Filozofická fakulta
Katedra psychologie

POSTOJE STŘEDOŠKOLSKÝCH A
VYSOKOŠKOLSKÝCH STUDENTŮ
K UMĚLÉ INTELIGENCI

ATTITUDES OF HIGH SCHOOL AND UNIVERSITY STUDENTS
TOWARDS AI



Bakalářská diplomová práce

Autor: **Mgr. Bc. Eva Hašek Dóšová**
Vedoucí práce: **PhDr. Ondrej Gergely, Ph.D.**

Olomouc

2024

Velice si vážím a děkuji vedoucímu této práce PhDr. Ondreji Gergelymu, Ph.D., za jeho vřelý responzivní přístup, vynaložený čas a ochotu, poskytování cenných podnětů a podpory při psaní této bakalářské práce. Z celého srdce děkuji mé rodině, mým dětem Mie a Luně, mým rodičům, mým skvělým přátelům a zejména také mým výjimečným spolužákům Jamesovým dětem. Všem patří veliké díky za jejich lásku, péči a podporu, nejen během psaní této bakalářské práce, ale během celého studia i mimo něj.

Prohlášení

Místopřísežně prohlašuji, že jsem bakalářskou diplomovou práci na téma: „Postoj k AI u středoškolských a vysokoškolských studentů“ vypracovala samostatně pod odborným dohledem vedoucího diplomové práce a uvedla jsem všechny použité podklady a literaturu.

V Liberci dne 21. 3. 2024

Podpis

OBSAH

Číslo	Kapitola	Strana
	OBSAH	3
	ÚVOD.....	5
	TEORETICKÁ ČÁST.....	7
1	Postoje	8
1.1	Vymezení a vybrané definice postojů	8
1.1.1	Stručná historie chápání postojů.....	9
1.2	Vnitřní struktura postojů.....	10
1.2.1	Jednosložková struktura	10
1.2.2	Dvosložková struktura.....	10
1.2.3	Trosložková struktura.....	11
1.3	Funkce postojů.....	11
1.4	Měření postojů.....	12
1.4.1	Explicitní techniky měření	12
1.4.2	Implicitní techniky měření	13
1.4.3	Kvalitativní přístupy měření postojů	13
2	UMĚLÁ INTELIGENCE.....	14
2.1	Umělá inteligence.....	14
2.2	Historie AI	15
2.2.1	Stručný vývoj AI	16
2.2.2	Rozvoj po roce 2011.....	16
2.3	Velké jazykové modely a generativní AI	17
2.3.1	GPT a další jazykové modely.....	17
3	postoje k umělé inteligenci a technologiím.....	20
3.1.	Postoje k technologiím	20
3.1.1	Klíčové teorie a modely	20
3.2	Postoje k umělé inteligenci.....	25
3.3	Měření postojů k umělé inteligenci v psychologii	26
3.3.1	Vybrané modely měření postojů k AI	26
3.3.2	Aktuální studie s tematikou AI.....	29
3.3.3	Postoje vybraných škol k používání AI.....	31
3.3.4	Postoje vysokých škol a univerzit	32
	VÝZKUMNÁ ČÁST.....	34
4	výzkumný problém a cíle práce	35
4.1	Formulace výzkumných otázek	36

5	typ výzkumu a použité metody	37
5.1	Metoda získávání dat	37
5.2	Metoda zpracování a analýza dat	38
5.3	Tvorba dat.....	39
5.4	Výzkumný soubor	39
5.5	Etické hledisko a ochrana soukromí.....	41
6	Práce s daty a její výsledky	42
6.1	Individuální postoje k AI.....	42
6.2	Vnímání AI	45
6.3	Způsob používání a zkušenosti.....	51
6.4	Vzdělávání o AI.....	58
6.5	Obavy a hrozby z AI.....	61
6.6	Etika a regulace	63
6.7	AI jako entita	66
7	Diskuse	68
8	Závěr.....	76
LITERATURA.....		81
PŘÍLOHY		93

ÚVOD

V současné době umělá inteligence a strojové učení představují klíčové oblasti výzkumu a inovací s širokým dopadem na různé sféry lidské činnosti. Jedním z fascinujících fenoménů v této oblasti jsou jazykové modely, z nichž jeden z nejznámějších je ChatGPT. Tato bakalářská práce se zabývá zkoumáním postojů studentů středních a vysokých škol k umělé inteligenci, s důrazem na jazykové modely, včetně, ale neomezeně, na ChatGPT a jeho varianty. Zajímá nás, jakým způsobem vnímají tito studenti využití umělé inteligence v běžném životě, ve vzdělávání a ve společnosti obecně. Tato práce klade důraz na pochopení postojů, předsudků a očekávání studentů ohledně této inovativní technologie a přispívá k hlubšímu pochopení jejich role a významu v současném digitálním prostředí.

Není to tak dlouho, co byly jazykové modely poprvé představeny (ChatGPT byl uveden na trh v listopadu 2022), ani tak dlouho, co bylo lepší s nimi komunikovat v anglickém jazyce a s češtinou měly potíže, nebo informace, které poskytovaly působily neohrabaně. Možná vás hned nenapadlo, že první odstavec této kapitoly napsala samotná umělá inteligence (AI). Zadáni bylo prostě: „Ahoj Chate, píšu bakalářskou práci, napiš mi prosím první odstavec úvodu. Práce píše o tobě, zkoumá postoje studentů středních a vysokých škol k umělé inteligenci, konkrétně k jazykovým modelům – Chat GPT a dalším. Díky.“ Výsledkem je odstavec splňující všechny požadavky kladené na tento druh práce, dokonce velice vyvedený a trůfám si říci k nerozeznání od běžného akademického textu.

Tato bakalářské práce, se zabývá tímto současným fenoménem, kdy je skrze jazykové modely možné alternovat produkty lidského ducha, například sbírat a syntetizovat informace, nalézat souvislosti, protiklady, vytvářet nové modalitty a mnoho dalšího. Paradoxně to byla kreativita a schopnost vlastní lidskému druhu, která nás dovedla k objevu umělé inteligence, která tyto klíčové atributy začala do jisté míry nahrazovat. Stejně jako řada revolučních vynálezů v minulosti, i za vznikem AI stála snaha posunout lidskou invenci, usnadnit si život, a uspokojit poptávku po schopnostech, které odpovídají obsahu a povaze našeho světa. Stejně jako byla dříve potřeba zjednodušit a zautomatizovat manuální práci, tak dnes, ve světě s nekonečným obsahem informací, vznikla potřeba je třídit, vyznat se v nich, najít ty správné, pracovat s nimi a usnadnit si cestu skrze jejich nepřehledné množství. Příkladem je například psaní bakalářské práce v minulosti a dnes, nebo před pár lety. Dříve se studenti museli opřít o knihy, byly dokonce dostatečným a adekvátním

zdrojem. S nástupem internetu, snadnější dostupností zahraničních periodik, by návštěva univerzitní knihovny bez přístupu k online zdrojům sotva stačila.

Jazykové modely jako ChatGPT otvírají nové možnosti komunikace a interakce mezi člověkem a strojem. Obohacují virtuální prostředí, kde jsou využívány k řešení široké škály úkolů. Jejich schopnost přizpůsobit se na různé situace a uživatele a komunikovat jako člověk, s sebou nese řadu otázek týkajících se samotné lidské komunikace, důvěryhodnosti odpovědí, etiky, podoby pracovního trhu a celkově také budoucnosti fungování společnosti a interakce mezi lidmi. Cílem této práce je prozkoumat, jaký pohled mají na tato témata studenti středních a vysokých škol, jak je vnímají, jak je integrovali do svých životů, a co to pro ně znamená v každodenním školním, akademickém nebo i profesním životě.

TEORETICKÁ ČÁST

1 POSTOJE

Postoje hrají klíčovou roli v našem vnímání světa. Jsou nezbytné při našem rozhodování a mají přímý vliv na lidské chování. Skrze postoje hodnotíme svět okolo nás, hodnotíme objekty, jevy či myšlenky od negativních přes pozitivní. Naše sklony k takovým hodnocením, jsou podle jednoho z oblíbených pojetí v psychologii kombinací afektivních, kognitivních a behaviorálních prvků, které jsou ve vzájemné interakci (Hogg & Vaughan, 2014). Důležité je chápat postoje v jejich dynamické povaze, nejsou neměnné a mohou se měnit v závislosti na nových informacích nebo zkušenostech a učení. Vnímáme skrze ně okolní svět a jsou důležitou součástí sociálních interakcí. Ovlivňují naše vnímání světa, jak na něj reagujeme a jak se rozhodujeme.

Cílem této kapitoly je poskytnout základní strukturu a klíčové koncepty chápání postojů v psychologii. Úvod do problematiky je neopomenutelným prvkem, na který navážeme v dalších kapitolách, které se věnují postojům k technologiím a umělé inteligenci.

1.1 Vymezení a vybrané definice postojů

Význam slova postoj pochází z latinského výrazu „positura“, tedy poloha. Vzniklo odvozením od příbuzného slova „ponere“ s významem umístit nebo položit (Rejzek, 2012). Chápání výrazu v jeho dnešní podobě, zahrnuje nejen polohu fyzickou, ale zejména také tu duševní. Naše názory, přístupy k různým otázkám a životním situacím, hodnotám a následně také způsob jakým skutečně jednáme, to vše můžeme laicky zahrnout do tohoto termínu.

Podle Hartla a Hartlové (2010) je v psychologickém slovníku postoj definován jako: *„hodnotící vztah vyjádřený sklonem ustáleným způsobem reagovat na předměty, osoby, situace a na sebe sama; postoje jsou součástí osobnosti, předurčují poznání, chápání, myšlení a cítění; postoje se spolu s vědomostmi získávají v průběhu života, především vzděláváním a širšími sociálními vlivy, jako je veřejné mínění a sociální kontakt“* (Hartl & Hartlová, 2010, s. 431).

V českém prostředí se často objevuje také chápání postojů dle Řehana (2007), který uvádí, že je lze chápat jako *„relativně trvalé soustavy poznatků, pocitů a tendencí jednat určitým způsobem ve vztahu k některému předmětu našeho sociálního světa“* (Řehan, 2007, s. 7).

1.1.1 Stručná historie chápání postojů

V oblasti psychologie a sociologie se první akademický zájem o zkoumání postojů objevil ve druhé polovině 19. století, kdy Herbert Spencer tento termín v roce 1862 jako první použil ve vědeckém kontextu. Počátek 20. století v USA znamenal boom dalšího výzkumu postojů v psychologii a sociologii (Ajzen, 2020). Zasadili se o něj zejména velcí průkopníci jejich vědeckého zkoumání W. J. Thomas a F. Znaniecki (1918). Popsali jejich vznik jako reflexi sociálních hodnot, které vytvářejí vědomí jedince a jeho reakce a zásadně přispěli k pochopení vztahu mezi společností a jedincem (Nakonečný, 1997).

Ještě zásadnějším jménem spojeným s výzkumem postojů v první polovině 20. století je sociální psycholog Gordon Allport (1935). Význam postojů a jejich zkoumání považoval ve své době za nepostradatelný koncept americké sociální psychologie (Allport, 1968). Jeho definice je v různých obměnách platná dodnes. *„Postoj je mentální a nervový stav pohotovosti, organizovaný zkušeností a vyvíjející usměrňující či také dynamický vliv na reakce jedince, na všechny předměty a situace, se kterými je spojen“* (Allport, 1935, s. 798).

Postoje podle něj nejsou pasivním stavem, ale jsou výsledkem aktivního procesu vnímání a reakcí na svět kolem nás. Získáváme je skrze zkušenost a v interakcích jedince se sociálním i fyzickým prostředím hrají klíčovou roli (Allport, 1935). Ve druhé polovině 20. století, v roce 1975, představili Martin Fishbein a Icek Ajzen Teorii plánovaného chování (Theory of Planned Behavior – TPB). Ta naznačuje, že naše jednání je důsledkem našich subjektivních norem, přesvědčení a našich postojů. Postoj chápou jako naučenou tendenci reagovat na něco pozitivně nebo negativně, v závislosti na jejich hodnocení. Pokud věříme, že pro nás budou přínosem, bude náš postoj kladný a obráceně. Nemalou roli hraje také náš hodnotový žebříček a jak oceňujeme různé aspekty života.

Naše postoje jsou stále chápány jako výsledek životních zkušeností a sociálního učení ve společnosti. Nicméně význam postojů byl v souvislosti s reálným jednáním trochu přeceněn. Původní předpoklad, že chování lze na jejich základě předpovědět, se nakonec nepotvrdil (Fishbein & Ajzen, 1974; Fishbein & Ajzen, 1980). Podstatnějším faktorem při predikci lidského jednání je samotným záměr jednat určitým způsobem (Fishbein & Ajzen, 2010). Důležité je také zdůraznit, že postoje nechápeme jako neměnné aspekty našeho vztahování se ke světu, ale získáváme je během celého života na základě našich zkušeností, vlivem našeho okolí a skrze sociální učení (Fishbein & Ajzen, 1977).

1.2 Vnitřní struktura postojů

Definice postojů není zcela jednotná a na tento trend navazuje také povaha vnitřní struktury postojů. Tři odlišné přístupy jsou po desetiletí předmětem akademické debaty. Otázkou zůstává, který z těchto modelů nejlépe vystihuje jejich podstatu a umožňuje přesné měření a interpretaci. Tento dialog odráží pokračující snahu o proniknutí do hlubšího porozumění dynamice postojů. Výzkum v této oblasti pomáhá vytvářet teorie a metody, které lépe vysvětlují, jak se postoje formují, mění a ovlivňují lidské chování. Haugh a Vaughan (2014), stejně jako Řehan (2007) uvádějí tři základní rozdělení struktury postojů.

1.2.1 Jednosložková struktura

V jednosložkovém pojetí hrají hlavní roli emoce a pocity. V Thurstonově (1931) pojetí jsou postoje definovány „*afektem pro nebo proti psychologickému objektu*“ (Thurstone, 1931, s. 261). Tedy klíčovou otázkou je, zda se nám objekt líbí či nikoliv. Afektivní reakce jedince na podnět předpokládá, že postoje jsou jejich přímým výsledkem. Jednosložkový model podléhá časté kritice a je otázkou, zda poskytuje dostatečně komplexní rámec pro pochopení vytváření postojů a jejich vlivu na chování. Na Thurstona v pojetí jednodimenzionálního pojetí postoje navázal ve svých časných výzkumech Martin Fishbein (1963), který tvrdil, že postoje vycházejí z psychologického hodnocení objektu jako líbí či nelíbí¹.

1.2.2 Dvosložková struktura

K Thurstonově modelu přidal Gordon Allport další složku, kterou je mentální připravenost (Allport, 1935). Podle něj jde o predispozici, která má poměrně konzistentní vliv na naše rozhodování o tom, co je dobré a špatné. Postoje lze podle něj zkoumat introspekci vlastních mentálních procesů či vyvozovat na základě toho co říkáme, a jak se chováme (Allport, 1935 in Haugh a Vaughan, 2014). Dvosložkový modelem postojů se ve své práci zabýval i R. P. Bagozzi. V jeho pojetí není postoj formován pouze afektivní (emocionální) nebo kognitivní složkou, ale jejich kombinací. Obě složky se navzájem ovlivňují. Podle tohoto pojetí závisí postoj na situaci a aktuálních okolnostech, je ale důležité zmínit, že zároveň zamítá důležitost vnitřního postoje (Bagozzi, 1981).

¹ Martin Fishbein byl významnou osobností 20. století, který formuloval a silně ovlivnil náhled na postoje, jejich vytváření a důsledky. I jeho pohled na složky postojů prošel vývojem a později své chápání přehodnotil na multidimenzionální (Ajzen, 2012).

1.2.3 Třísložková struktura

Třísložkový model postojů předpokládá, že postoje se skládají ze tří základních složek: afektivní (emocionální), kognitivní (myšlenky a přesvědčení o objektu) a konativní (behaviorální). Vychází z předpokladu, že pro pochopení postojů je nutné zvážit všechny tři aspekty. Třísložkové struktury postojů se věnoval například Thomas M. Ostrom (1968), který zdůrazňoval význam jednotlivých složek, jejich interakce a vliv na chování. Martin Fishbein ve své práci společně s Icekem Ajzenem zpracovali řadu teorií vycházející z multidimenzionálního pojetí složek postojů. Mezi nejznámější patří teorie plánovaného chování a také teorie odůvodněného jednání, jsou oblíbené pro svůj komplexní pohled na to, jak postoje ovlivňují naše jednání a chování a berou v potaz více faktorů, než pouhou afektivní odpověď (Ajzen, 2012).

1.3 Funkce postojů

Postoje vznikají jako součást socializačního procesu (Fishbein a Ajzen, 1975; Oskamp & Schultz, 2014). Předpokladem jejich existence je užitečnost, mají své funkce. Jednou ze základních funkcí je přizpůsobení se prostředí (Festinger, 1964). Někteří autoři pojali tyto funkce detailněji, mezi nimi i Katz (1960), podle kterého jsou různé postoje s různými funkcemi. Ty nám pomáhají orientovat se ve společnosti a světě kolem nás. Rozdělil je na čtyři základní funkce: poznávací, instrumentální, hodnotově-expresivní a egobrannou. Každá z nich má dle Katze (1960) svá níže uvedená specifika, která jsou dle něj dále rozvedena:

- **Poznávací funkce** postoje nám pomáhá organizovat naše vnímání světa, kategorizovat informace a předvídat sociální realitu. Je pro nás výhodné roztrdit si naše zkušenosti a vytvořit si klasifikační systém, který nám umožní lépe se orientovat. Příkladem je negativní zkušenost, díky postoji, který si vytvoříme se můžeme takové zkušenosti příště například vyhnout.
- **Instrumentální funkce** postojů nám pomáhají dosáhnout cílů a odměn a minimalizovat nepříjemné zkušenosti. Příkladem může být pozitivní postoj k technologiím, který bude mít za následek větší chuť s nimi pracovat a překonávat překážky třeba v podobě složitého ovládání.
- **Hodnotově-expresivní funkce** postojů nám umožňují vyjadřovat naše nejhlubší hodnoty a jádro naší identity, skrze tento postoj dále komunikujeme směrem

k ostatním, kdo jsme a jaké máme hodnoty. Takové postoje pak formují i sociální interakci, skrze ně vytváříme příslušnost k sociálním skupinám.

- **Ego-obranná funkce** postojů, souvisí s ochranou naší sebeúcty a obrazu sebe sama. Funguje jako obranný mechanismus, který pomáhá například s racionalizací neúspěchů, nebo přijímáním nepříjemné pravdy o sobě samém. Pomáhá udržet pozitivní obraz sám o sobě.

Sharon Shavitt (1990) uvádí konkrétní funkce, které rozděluje na znalostní, utilitární, funkce sociální identity a udržování sebeúcty. Pokud je porovnáme s Katzem (1960), jejich význam je srovnatelný. Ještě předtím Fazio (1989) argumentuje, že vůbec největším přínosem je utilitární funkce postojů. Slouží k hodnocení objektu a není důležité, jestli má tento postoj pozitivní nebo negativní směr. Postoje tak slouží účelům jako třídění informací a mají vliv na naše rozhodnutí a jednání (Fazio, 1989).

1.4 Měření postojů

Měření postojů je složitý úkol a abychom je poznali, měli bychom ideálně zachytit všechny jejich složky – afektivní, kognitivní a behaviorální (Ajzen 2012). Zkoumání postojů je silně spjaté se způsoby jejich měření. Dva způsoby, explicitní a implicitní nabízejí rozdílné přístupy v kvantitativních metodách, nezapomínejme ale také na kvalitativní metody.

1.4.1 Explicitní techniky měření

Snaha o měření postojů se váže k Thurstonovi a jeho díle s velice explicitním názvem „*Attitudes Can Be Measured*“ (Thurstone, 1928). Předpokladem explicitního měření postojů bylo zjištění opravdových přesvědčení a názorů (Hogg & Vaughn, 2014).

Během první poloviny 20. století, s výjimkou sémantického diferenciálu, byly tyto nástroje hojně rozvíjeny a přinesly několik testovacích škál: Thurstonovu (1928), Likkertovu (1932), Guttmanova (1944) a Osgoodův sémantický diferenciál (1957). Tyto sebeposuzovací škály po respondentech vyžadují přímé, vědomé sebehodnocení a nejčastěji jsou využívány v kvantitativní analýze a poskytují dobře strukturovaná data, která jsou zpracována statistickou analýzou. Jak uvádí Hogg & Vaughn (2014), nebyla dříve sebraná data porovnávána se skutečnými jednáním, a i to vedlo vědecké zkoumání směrem k většímu využívání implicitních technik zjišťování postojů.

1.4.2 Implicitní techniky měření

Implicitní měření postojů přináší rozdílné techniky hodnocení subjektů. Na rozdíl od tradičních měření jako jsou výše zmíněné explicitní škály, které spoléhají na sebedposouzení, nabízejí nepřímé metody jako je evaluativní priming a Test implicitních asociací (IAT) nové přístupy. Tyto metody se zaměřují na reakční dobu a asociace, které se oproti explicitním metodám nedají tak snadno vědomě ovlivnit.

Evaluativní priming

Fazio et al. (1995) popsal evaluativní priming jako psychologickou metodu skrze kterou lze odhalit hlubší pohnutky, podvědomé hodnotící procesy, které nemusejí nebo neumějí respondenti vždy vyjadřovat explicitně. Principem této metody je rychlá podvědomá asociace mezi slovy nebo obrázky, takzvanými primy. Odpovědi na cílové stimuly (například slova nebo obrázky) jsou ovlivněny předchozími stimuly na základě hodnotícího (pozitivního nebo negativního) vztahu mezi nimi. Výzkumy ukázaly, že lidé reagují rychleji a přesněji, když prime a cílový stimul mají stejný hodnotový náboj (Kawakami et al., 2002).

Test implicitních asociací

Test implicitních asociací (IAT) který je v jistých ohledech podobný evaluativnímu primingu, vyvinul Greenwald a jeho kolegové v roce 2002. Je další oblíbenou metodou, která měří sílu asociací mezi koncepty a evaluací. Rychlost a přesnost s jakou respondenti párují slova či obrázky s pozitivním nebo negativními pojmy, může poukázat na podvědomý postoj (Greenwald et al., 2002). Greenwaldova analýza potvrdila oblíbenost testu implicitních asociací zejména při zkoumání sociálně citlivých témat (Greenwald et al., 2009).

1.4.3 Kvalitativní přístup měření postojů

Kvalitativní přístupy v měření postojů, jako jsou rozhovory a pozorování, nabízejí detailnější pohled na postoje a mohou lépe zachytit jejich jednotlivé vrstvy, lze lépe zkoumat interpretace a postoje vysvětlit a jejich velkou přidanou hodnotou je také zasazení do kontextu (Šafránková & Kocourková, 2011).

2 UMĚLÁ INTELIGENCE

Podstatou lidské rasy *Homo sapiens* – člověka rozumného je moudrost, intelligence. Po staletí jsme se snažili porozumět lidskému mozku a jeho schopnostem. Jak může něco tak malého vykonávat, tak zásadní úkoly, vnímat, rozumět, předpokládat a vykonávat řadu dalších úkolů, bylo a je předmětem zkoumání. Nyní přichází na řadu umělá intelligence (AI)², která je zajímavá zejména z pohledu jejího vytváření a vývoje entit, které dokážou samy reagovat ve zcela nových situacích. AI je univerzálním nástrojem, který se neustále rozvíjí a bude relevantní i v budoucnu.

2.1 Umělá intelligence

Chápání pojmu umělá intelligence se v průběhu času vyvíjelo. Jako první tento termín použil John McCarthy v roce 1955, který ji popsal jako vědu a techniku tvorby inteligentních strojů, tedy tak aby se stroje chovali chytře jako lidé. Logickou myšlenkou o povaze umělé intelligence, je její srovnání s tou lidskou, tady ale platí námitka, že lidé chybují. Popis, ke kterému se věda přiklání je kvalita umělé intelligence chápána skrze klíčový pojem *rationality*.

Russel & Norvig (2022) popisují dva přístupy. Jeden chápe inteligenci jako myšlenkové procesy a uvažování, tedy je interní, druhý je zaměřený na chování, tedy je vnější charakteristikou. Je také důležité rozlišit termíny *umělá intelligence* a *strojové učení*. Strojové učení je podoblastí spadající do AI, které se zabývá schopností vylepšovat svůj výkon na základě zkušeností (Domingos, 2015). Z rozdělení na dva základní směry, vycházejí čtyři základní přístupy k umělé inteligenci:

Myslí jako člověk (kognitivní přístup) – zaměřuje se na vývoj systémů AI, které nejen napodobují lidské myšlení a chování, vnímání, uvažování a řešení problémů, ale rozumějí tomu, jak lidský mozek funguje a abychom porozuměli principu myšlení programu skrze člověka, je nezbytné pochopit, jak myslí lidé (Mitchell, 2019). Právě tady se vědci obracejí na psychologii a neurovědy. K pochopení fungování lidského mozku se uplatňují

² Pro potřeby této práce používáme zkratku z anglického *artificial intelligence*, která se používá běžně i v českém jazyce.

tři strategie: introspekce, psychologické experimenty a metody zobrazování mozku, tedy jeho sledováním při aktivitě.

Jedná jako člověk (přístup Turingova testu) – podstatou slavného Turingovu testu je snaha o imitaci člověka tak, aby se hodnotitelé domnívali, že se jedná o lidskou bytost nikoliv AI. Pokud dokáže umělá inteligence v komunikaci s člověkem vytvořit iluzi, že je také člověk, je chápána jako inteligentní (Turing, 1950).

AI musí splňovat čtyři předpoklady, aby tohoto dosáhla: **zpracovávat přirozený jazyk** (Natural Language Processing – NLP) - musí umět **komunikovat** v lidském jazyce, **reprezentovat znalosti** – dokázat **uchovávat** informace, které již zná, a to ji umožňuje plnit komplexní úkoly, **automaticky uvažovat** – **odpovídat** na otázky a přicházet s novými závěry a využívat **strojové učení** – **učit** se z dat a zlepšovat svůj výkon na základě zkušenosti.

Myslí racionálně (přístup zákonů myšlení) - AI systémy jsou navrženy tak, aby na základě pravidel dokázaly simulovat ideálního racionálního myslitele, tedy dosáhnout pravdy na základě logiky a podaných informací. Často je toto využíváno u expertních systémů.

Jedná racionálně (přístup racionálního agenta) - vytvářejí se tzv. agenti, kteří mají podobu entity. Ta vnímá své okolí pomocí senzorů, na jejich základě by tito agenti a systémy měli jednat s cílem maximalizovat šance na úspěch.

S rozvojem technologií a jejím zasahováním do každodenního života, se rozšiřují i oblasti zájmu a využití, stejně jako samotná definice AI.

2.2 Historie AI

Historie umělé inteligence (AI) je fascinující a rozmanitá, sahající od teoretických počátků v matematice a logice až po moderní aplikace v robotice, analýze dat a strojovém učení a vytvoření velkých jazykových modelů jako je ChatGPT a jim podobné. Historie vzniku umělé inteligence je plná až neuvěřitelných a průlomových objevů, které daly vzniknout AI dneška. Z povahy této práce se zaměříme pouze na několik vybraných mezníků.

2.2.1 Stručný vývoj AI

Historie umělé inteligence sahá do 20. století. První prací, ač ne tou nejznámější, která se týká umělé inteligence byla teorie Warrena McCullocha a Waltera Pittse (1943), popisující neuronové sítě, které jsou doposud základem neuronových sítí v umělé inteligenci. Díky modelu neuronu, známého jako McCulloch-Pitts neuron, který mohl provádět jednoduché operace jako AND, OR a NOT (ano, nebo, ne) také přišli na možnost simulace jakékoliv výpočetní funkce, což vede k potenciálu vlastního učení a přizpůsobení (Nilsson, 2009).

Popularizace AI na jejím počátku je spojená zejména se jmény Alana Turinga nebo Johna McCarthyho či Marvinu Minského. Patří mezi průkopníky, kteří položili základy AI. Alan Turing a jeho Turingův test je nadále aktuálním nástrojem skrze který se hodnotí strojová, umělá inteligence (Turing, 1950). Definice Johna McCarthyho položila základy pro rozvíjení oboru AI, stejně jako Marvin Minsky zkoumající simulaci lidského myšlení.

V roce 1969 vydal Minsky společně s Papertem výzkum zaměřený na neurální sítě a strojové učení, který navazuje na práci McCullocha a Pittse (Minsky & Papert, 2017). Dalším zásadním průlomem bylo zjištění Allena Newella a Herberta Simona (1976), kteří formulovali hypotézu fyzických symbolů systému, která znamená, že jakýkoliv systém vykazujícího obecnou inteligenci, lidský či strojový, musí operovat s daty v podobě symbolů (Gugerty, 2006).

V následujících letech se pozornost obrátila ke konekcionismu, který se zaměřuje na simulaci neuronových sítí a proces učení, stejně jako je tomu u lidského mozku. Jeho cílem je vytvořit umělé sítě, bez přímého programování pro konkrétní úkoly. Je také klíčový pro rozvoj hlubokého učení.

Významným pokrokem byl nástup internetu, který umožnil vytvářet ohromné sady dat, fenomén zvaný jako velká data (big data). Tato velká data obsahují miliardy slov, obrázků, stejně jako hodin mluveného slova a videa, obsahu sociálních sítí a dalších představitelných podob dat.

2.2.2 Rozvoj po roce 2011

Ačkoliv hluboké učení hrálo svou roli již od 70. let 20. století (LeCun et al., 1995) bylo to až v roce 2011, kdy se objevily první rozpoznávání jazyka a vizuálních objektů (Krizhevsky et al., 2012). Hluboké učení v oblasti vizuálních úkolů, rozpoznávání jazyka,

strojových překladech, lékařských diagnózách, nebo herních schopnostech, překonalo ty lidské (LeCun et al., 2015). Tento významný pokrok zpopularizoval AI ve společnosti, vědě, mezi studenty, dostal se do zájmu společností i států.

Hluboké učení, posunulo programování k sebeučení, má velký význam pro rozvoj velkých jazykových modelů neboli large language models (LLM), mezi které řadíme i ChatGPT³ a další modely od jiných společností. Co bylo pro jejich rozvoj zcela zásadní je schopnost generovat přirozený jazyk a rozumět mu. Jak jsme zmínili výše, je to jeden ze zásadních předpokladů komunikace s člověkem (Turing, 1950).

Velké jazykové modely jako ChatGPT přinesli revoluci ve zpracování přirozeného jazyka. Jimi generovaný text je, až na výjimky, nerozeznatelný od lidského jazyka, dokáže rozumět nejen textu, ale i kontextu, stylům konverzace, umí informace hledat, syntetizovat, vyvozovat závěry a abstraktně uvažovat.

2.3 Velké jazykové modely a generativní AI

Generativní AI je technologie umělé inteligence, která umí vytvářet nová data. Modely generativní AI používají k učení vyspělé algoritmy, které generují nový obsah jako je text, obrázky, zvuky, videa či kódy. Příkladem generativní AI jsou velké jazykové modely ChatGPT, Gemini nebo třeba Dall-E či Midjourney.

2.3.1 Velké jazykové modely (LLM)

Velké jazykové modely jsou specifickým druhem generativní AI, které jsou trénovány na velkém množství textových dat. Ve svém principu fungují na předvídání dalšího slova v sekvenci, na základě toho předchozího (Radford et al., 2019). Díky tomu mohou generovat plynulý text. Dokáží dokonale zpracovat přirozený jazyk (NLP) a oblast jejich použití i složitost úkolů nadále roste (Bang et al., 2023; Bubeck et al., 2023).

2.3.2 GPT a další jazykové modely

První verze modelu GPT byla poprvé uvedena v roce 2018 a měla až 117 milionů parametrů (Bubeck et al., 2023; Radford et al., 2019). Zkratka GPT znamená „Generative Pre-trained Transformer“, což odkazuje na jeho fungování při zpracování požadavků

³ Vydala společnost OpenAI.

a následných odpovědí. Je schopný pochopit a reagovat na širokou škálu podnětů, poskytovat běžné i odborné odpovědi a s uživateli vede interaktivní konverzaci. Jeho výkon je mezi podobnými technologiemi dosud bezprecedentní (Bubeck et al., 2023).

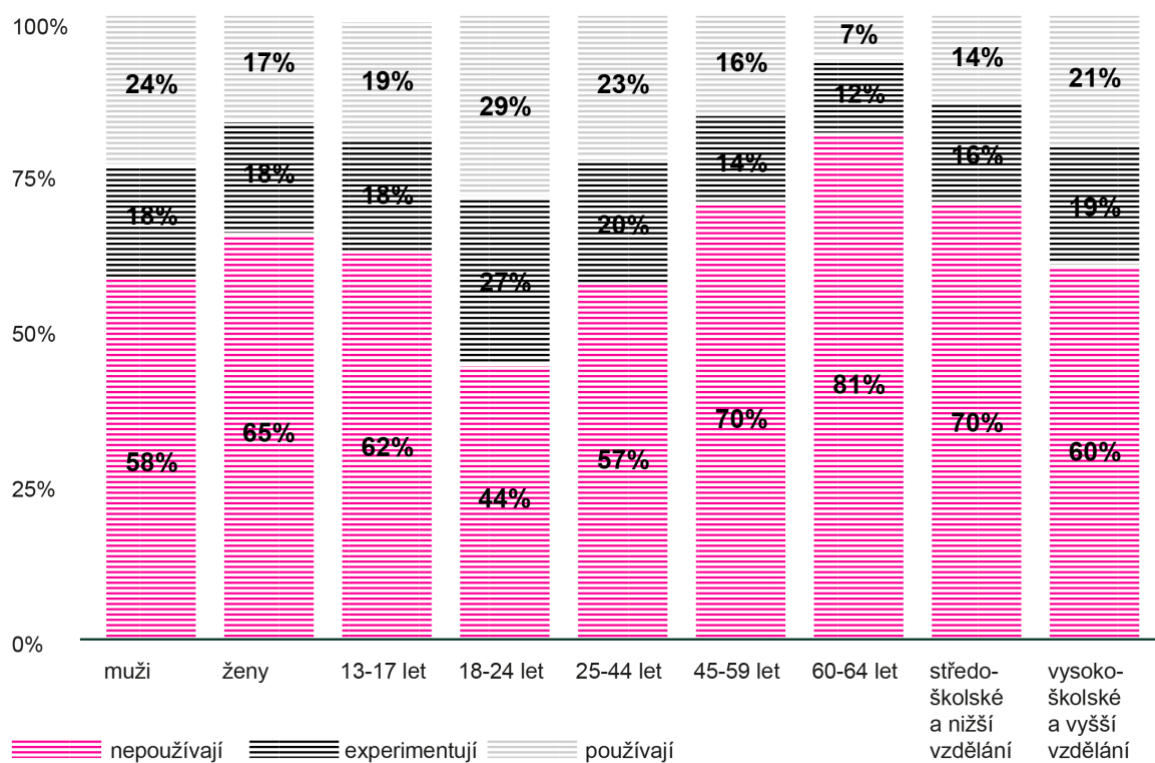
Jeho následovník GPT-2 byl o rok později v únoru roku 2019 s 1,5 miliardou parametrů vybaven ještě mnohem lépe, přesto tento model stále vykazoval problém s komplexním formováním názorů a chápáním kontextu (Radford et al., 2019). GPT-3 z června 2020 měl k dispozici do té doby nevidaných 175 miliard parametrů, jeho schopnosti převyšovaly předcházející modely zejména v kvalitě generovaného obsahu, nerozeznatelného od lidského textu (Dale, 2020).

Verze GPT-3.5 zlepšovala schopnosti tohoto modelu při zachování množství a posloužila k vývoji nástroje ChatGPT. ChatGPT-3.5, vyvinutý společností OpenAI, byl vydán 30. listopadu 2022 a je považovaný za přelomovou technologii. Oproti modelům GPT vylepšil ChatGPT-3.5 reakce na různé typy vstupů, což umožnilo podávat ještě lepší odpovědi (Zhang et al., 2023).

Model ChatGPT-4 byl oficiálně uveden společností OpenAI 14. března 2023. ChatGPT-4 je pokročilejší verze předchozích modelů, protože může zpracovávat jak textové, tak vizuální informace, což rozšiřuje jeho aplikace. Ačkoliv přesný počet parametrů nebyl veřejně oznámen, odhaduje se, že jich má biliony, což ho činí ještě sofistikovanějším. Model GPT-4 dokáže generovat text a reagovat na dotazy s vyšší přesností a je citlivější na nuance v komunikaci než jeho předchůdci. Kromě toho umožňuje procházení internetu a analýzu dat, což ChatGPT-3.5 nemůže. GPT-4 je také účinnější v detekování nebezpečných nebo zakázaných požadavků a ve vytváření přesných, fakticky založených odpovědí (ChatGPT-4, 2024).

ChatGPT získal během jednoho týdne od svého uvedení 30. listopadu 2022, milion uživatelů a do ledna 2023 jich dosáhl 100 milionů, v prosinci téhož roku měl přibližně 1,6 miliard návštěv (ChatGPT, 2024). Globální používání nástrojů AI, zejména ChatGPT podle pohlaví, věku a pouze dichotomického rozdělení vzdělání, ukazuje rozdělení v grafu níže Thormundsson (2023).

Graf č. 1: Globální používání AI v roce 2023, podle skupin



Zdroj: Thormundsson (2023; upraveno)

3 POSTOJE K UMĚLÉ INTELIGENCI A TECHNOLOGIÍM

Technologie jsou již dlouhou dobu součástí života každého z nás. Umělá inteligence byla donedávna přítomná spíše nenápadně, ale v posledních letech se její používání stalo stejně snadné a dostupné jako sám internet. Technologie a umělá inteligence ovlivňují naši komunikaci, práci, ale také jak se učíme a bavíme.

Jakým způsobem technologie a AI vnímáme, jak na nás působí, jaká jsou specifika našich postojů, která k nim zaujímáme a co na ně má vliv – to jsou vybrané otázky, které nám mohou otevřít cestu k pochopení společnosti, kde má AI své pevné místo. Jsou zásadní nejen pro lepší pochopení přístupu jednotlivců a společnosti k technologickému pokroku, současnému stavu, ale i ke způsobům a možnostem integrace nových technologií do našich životů. Přehled jednotlivých teorií a modelů zabývajících se přístupy k technologiím a umělé inteligenci a seznámení se s nimi, nabízí rámec pro pochopení postojů studentů v další kapitole této práce.

3.1 Postoje k technologiím

Všechny zásadní technologické pokroky v historii s sebou přinesly rozporuplné postoje (Gessl et al., 2019). Postoje k technologiím jsou specifickým odvětvím výzkumu postojů. V první kapitole jsme se zabývali postoji, historickým vznikem, jejich strukturou, funkcemi či měřením, v této kapitole se seznámíme s klíčovými teoriemi a modely přijímání nových technologií. V návaznosti na vývoj a zkoumání postojů v oblasti sociální psychologie jako teorii odůvodněného jednání (Fishbein & Ajzen, 1975) nebo teorii plánovaného chování (Ajzen, 1991) se rozvinuly další komplexnější modely a teorie vysvětlující postoje k technologiím.

3.1.1 Klíčové teorie a modely

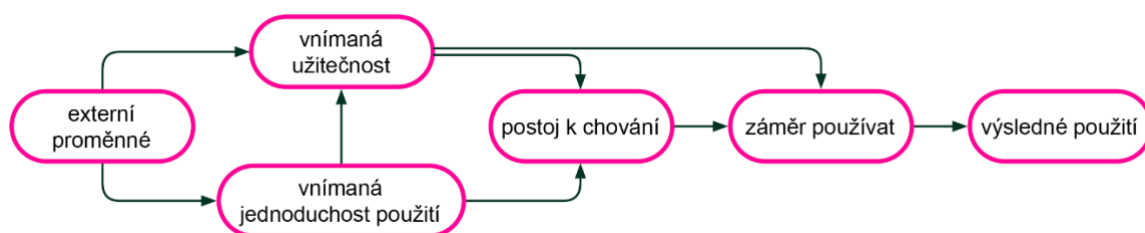
Klíčové teorie a modely se zabývají uživatelským přijetím technologií. Ta je definována jako ochota uživatele používat určitou technologii (Dillon & Morris, 1996). Podle DeLonea a McLeana (1992) lze skrze míru jejího přijetí a užívání, hodnotit její úspěch.

Postoje k technologiím se v psychologii zkoumají za použití různých metodologických přístupů, ať už kvantitativních či kvalitativních.

TAM

Oblíbeným teoretickým rámcem pro vysvětlení způsobů, jakým uživatelé přijímají nové technologie je tzv. Technology Acceptance Model (TAM) vyvinutý na konci 90. let minulého století (Davis, 1989). Model zkoumá jak užitečnost a jednoduchost technologií ovlivňují postoje vůči technologiím, ale i záměr je skutečně používat na nakonec i reálné použití (Davis, 1989). Zároveň má zkušenost s používáním zpětný vliv na postoje. Cílem modelu TAM je predikce a vysvětlení chování v kontextu používání informačních technologií. Tato teorie navazuje na Fishbeina a Ajzena (1975) a jejich teorii odůvodněného jednání⁴, kterou jsme již zmiňovali v první kapitole. Původní model se dočkal ještě dalších modifikací a vylepšení TAM2 a TAM3.

Obrázek č. 1: Schéma TAM



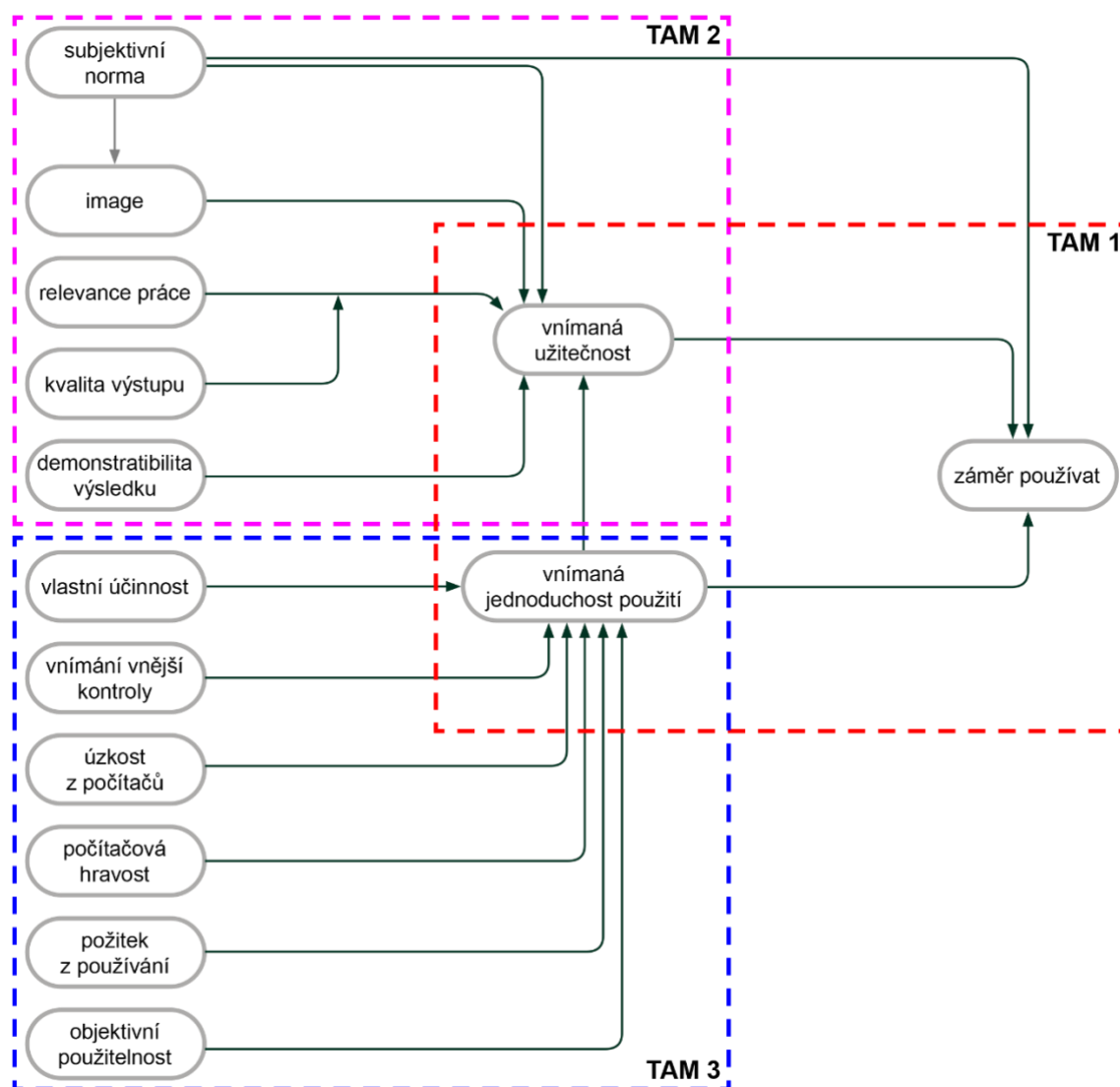
Zdroj: Davis (1989; upraveno)

TAM 2 a TAM 3

Venkatesh a Davis (2000) navrhli rozšíření modelu TAM na TAM2. Identifikovali obecné determinanty vnímané užitečnosti – subjektivní normu, image, relevanci pro práci, kvalitu výstupu, demonstratibilitu výsledků a vnímanou snadnost použití. K nim přidali dva moderátory, zkušenost a dobrovolnost (Venkatesh & Davis, 2000). První dva aspekty spadají do kategorie sociálního vlivu a zbývající determinanty jsou charakteristiky systému. V roce 2008 tento model TAM2 Venkatesh & Bala (2008) dále rozvinuly. V TAM3 je klíčovým měřítkem úspěchu systému jeho dlouhodobé používání (Venkatesh & Bala, 2008).

⁴ V originálním znění Theory of Reasoned Action.

Obrázek č. 2: Zjednodušené schéma TAM, TAM2 a TAM3

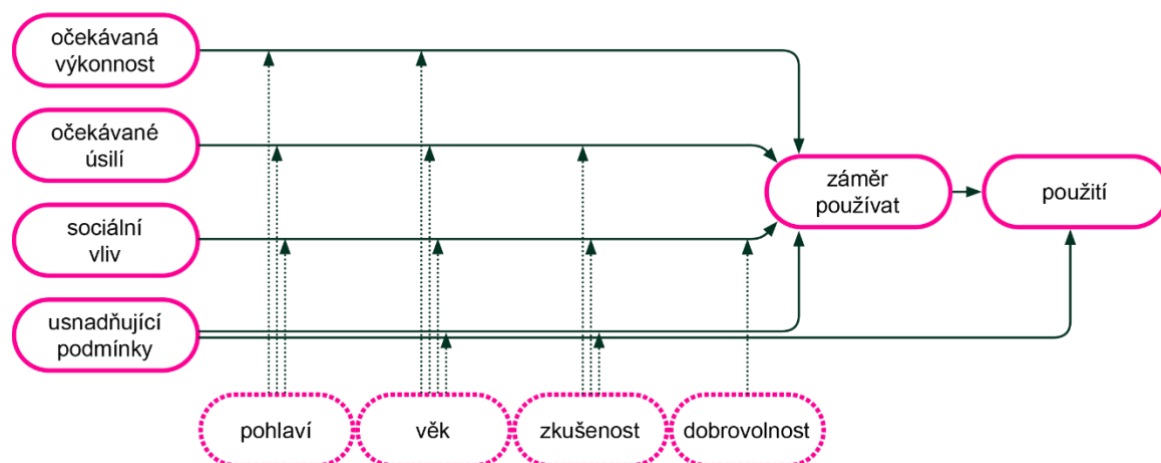


Zdroj: Davis (1989; upraveno); Venkatesh & Davis (2000; upraveno); Venkatesh & Bala (2008; upraveno)

UTAUT

Jednotná teorie přijetí technologie UTAUT (Unified Theory of Acceptance and Use of Technology) je shrnutím různorodých teorií přijímání inovací a technologií⁵, které Venkatesh et al. (2003) spojil do jednoho rámce, jehož podstatou je používání technologie a záměr našeho chování. Tento model navrhuje, že čtyři základní proměnné: očekávaný výkon, očekávaná snaha, sociální vliv a podmínky usnadňující přijetí, přímo určují záměr chování, a nakonec i jeho realizaci. Tyto faktory zároveň ovlivňuje pohlaví, věk, zkušenosti a dobrovolnost použití. Jeho možné ho využít pro různé kontexty a prostředí, ve kterých lze nalézt zásadní prvky, které uživatele vedou k jejich používání, jak v jejich záměru, tak ve skutečnosti (Williams et al., 2015)

Obrázek č. 3: Schéma UTAUT



Zdroj: Venkatesh et al. (2003; upraveno)

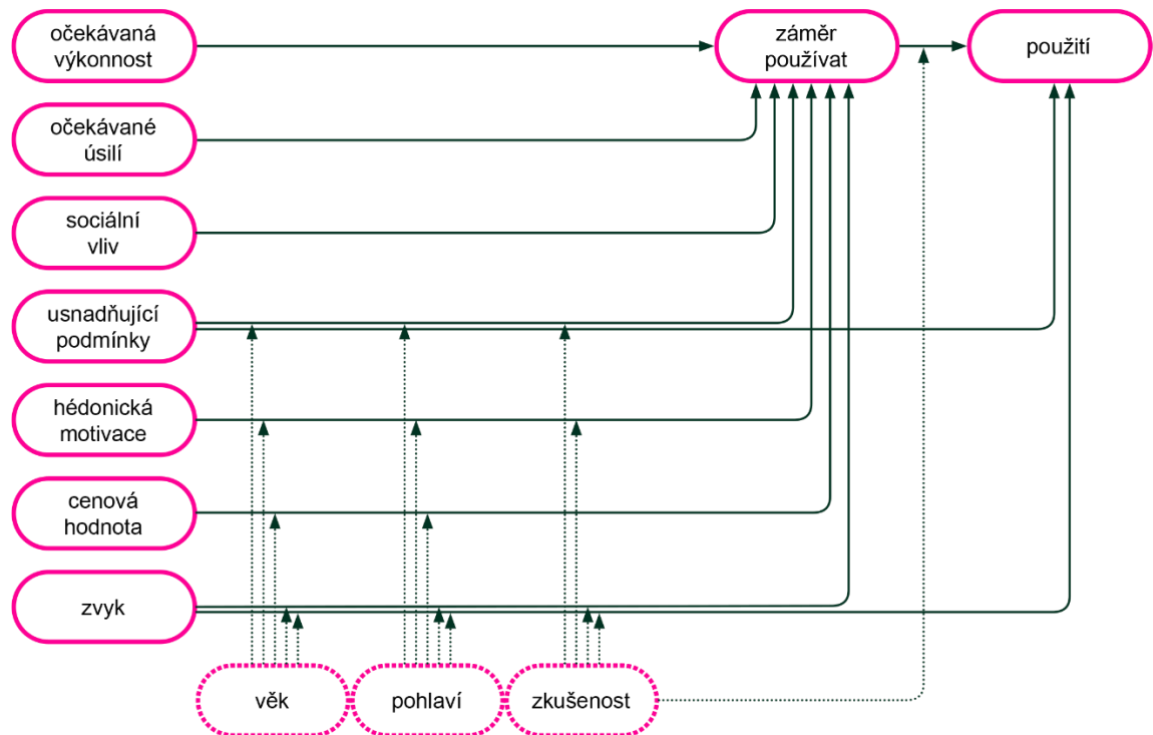
UTAUT2

Autoři rozšířili model UTAUT o hédonickou stránku motivace, cenovou hodnotu a zvyk. Všechny proměnné jsou přímými determinanty záměru používat, přičemž zvyk je zároveň přímým determinantem použití. Hédonická motivace odpovídá té části motivace,

⁵ Venkatesh et al. (2003) vycházel z těchto teorií: Teorie zdůvodněného jednání (TRA) - Theory of Reasoned Action; Model přijetí technologie (TAM) - Technology Acceptance Model; Motivační model - Motivational Model; Teorie plánovaného chování (TPB) - Theory of Planned Behaviour; Kombinovaný model TBP/TAM - Combined TBP/TAM; Model využívání počítačů (Model of PC Utilization); Teorie difuze inovací (IDT) - Innovation Diffusion Theory a Sociálně kognitivní teorie (SCT) - Social Cognitive Theory (Williams et al., 2015).

která zahrnuje chuť technologii používat, zábavu či potěšení, které uživatelé při používání pocítují (Venkatesh et al., 2012).

Obrázek č. 4: Schéma UTAUT 2

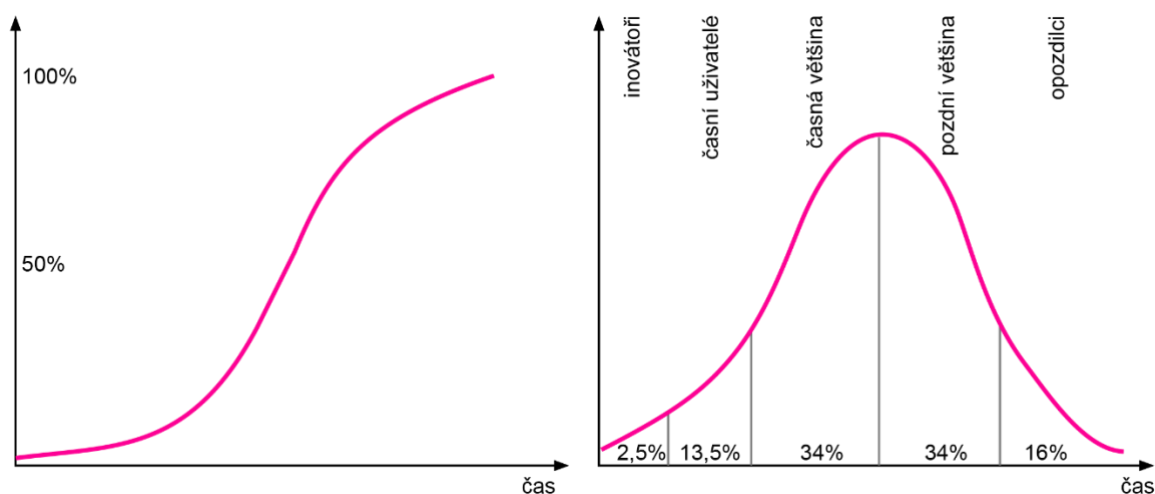


Zdroj: Venkatesh et al., (2012; upraveno)

DIFUZE INOVACÍ

Ač se jedná spíše o sociologický koncept, přesto má i v této práci své místo. Rogersova teorie difuze inovací z roku 1962 vysvětluje šíření technologií, nových nápadů a produktů, mezi lidmi a organizacemi (Rogers, 2003). Neptá se pouze jak, ale také proč a jakou rychlostí se nové myšlenky šíří v závislosti na čtyřech prvcích: inovaci, způsobů komunikace, času a sociálním systému. Pro jeho teorii je zároveň zásadní rozdělení na pět kategorií uživatelů, kterými jsou inovátoři, časní osvojitelé, časná většina, pozdní většina a opozdilci (Rogers, 2003).

Obrázek č. 5: Křivka adopce a Pět kategorií uživatelů



Zdroj: Rogers, (2003; upraveno)

3.2 Postoje k umělé inteligenci

Postoje k umělé inteligenci a postoje k technologiím by se snadno daly zaměnit, a předpokladem je, že metody měření postojů k technologiím se aplikují i na AI. Ta je ale v něčem přece jen specifická a zejména v oblasti obav má, jako poměrně čerstvá novinka aktuálně větší závažnost než existence počítačů, internetu nebo mobilních telefonů. Během pár let pronikla do mnoho oblastí běžného života bez našeho souhlasu či možnosti výběru. Objevil se termín artificial intelligence anxiety – obava z umělé inteligence (AIA) - fenomén, který s nástupem umělé inteligence musel logicky následovat. Wang & Wang (2019) v návaznosti na předchozí studie určili čtyři oblasti obav: obavy z nutnosti naučit se pracovat s AI; obavy z nahrazení míst pracovního místa AI; obavy z přílišné moci AI nebo

jejího zneužití a v neposlední řadě obavy o možnou budoucí podobu AI (např. humanoidní roboti).

3.3 Měření postojů k umělé inteligenci v psychologii

Jak umělou inteligenci vnímáme a co si pod ní představujeme, případně jaké o ní máme informace a kdo nám je předává, zásadním způsobem ovlivňuje naše postoje. Stejně jako u přijímání technologií a jejich úspěchu, hodnotíme úspěch AI na základě jejího přijetí a užívání (Dillon & Morris, 1996). V nejobecnější rovině lze předpokládat, že pokud má jedinec pozitivní postoj k technologiím a novinkám, jakou je i umělá inteligence, bude je používat častěji a snadněji je zapojí do svého běžného života. Negativní postoje budou pravděpodobně vést k odporu nebo váhání při jejich přijímání.

Podobně jako výzkumy postojů a přijímání technologií se i rozpoznávání a měření postojů k AI provádí pomocí různých metod a nástrojů, dotazníků, škál, rozhovorů, které hodnotí různé aspekty postojů. Díky nim bychom měli být schopni lépe identifikovat klíčové faktory, které ovlivňují vztah k AI, tedy jejímu přijímání a používání. Zároveň mohou nabídnout cenné informace, jak lze tyto postoje zlepšit, jak formulovat stanoviska vůči AI či jak se vůči ní vymezit.

3.3.1 Vybrané modely měření postojů k AI

Výzkumné metody měření postojů k umělé inteligenci, zahrnují několik druhů škál, které se liší úhlem pohledu a konkrétní oblastí zkoumání. Často nemají obecnou uplatnitelnost a jsou navázány na ad hoc výzkumy, proto se blíže věnujeme jen GAAIS, NAAIS, ATAI a ATTARI-12.

Tabulka č. 1: Škály postojů k umělé inteligenci

zkratka	název	autoři
GAIIS	Škála obecných postojů k umělé inteligenci	Schepman & Rodway, 2023
NAAIS	Škála negativních postojů k umělé inteligenci	Persson et al., 2021
ATAI	Škála postojů k umělé inteligenci	Wang & Wang, 2019
AIAS	Škála úzkosti z umělé inteligence	Kieslich et al., 2021
	Dotazník obav z autonomních technologií	Stein et al., 2019
ATTARI-12	Škála ATTARI-12	Stein et al., 2024

GAAIS

Schempan a Rodway (2023) vytvořili dotazník posuzující vnímání umělé inteligence. Pozitivní pohled chápe AI jako nástroj pomoci, negativní postoj vystihuje obavy z AI. Zajímavým zjištěním bylo, že pozitivní postoje se vážou k činnostem AI jako je věda, právo, nicméně v oblastech, kde je vyžadován lidský úsudek jako je zdravotní péče nebo terapie, jsou postoje negativní.

NAAIS

NAIIS, zkratka pro Negative Attitudes towards Artificial Intelligence Scale (Persson et al, 2021) je upravená verze dotazníku NARS, který byl původně navržen pro průzkum postojů k sociálním robotům. Tento nástroj byl přizpůsoben pro širší studium postojů k AI systémům obecně.

ATAI

Dalším významným nástrojem je Attitudes towards Artificial Intelligence Scale (ATAI), (Sindermann et al., 2021), která poskytuje kvantitativní měření pozitivních a negativních postojů vůči AI. Tato škála zahrnuje položky hodnotící jak přijímání AI, tak obavy z ní, což umožňuje komplexní pohled na postoj respondentů. Výzkumy, které využívají ATAI, například odhalily, že ženy mají průměrně nižší pozitivní postoje k AI než

muži, což naznačuje, že gender může být důležitým faktorem ve vnímání AI (Montag et al., 2024).

ATTARI-12

Výše zmíněné škály měření se potýkají s omezeními a nedostatky a podle Steina et al. (2024) je nutné vytvořit jednorozměrnou psychometrickou škálu, která bude nejen ekonomická, ale bude také zachycovat pozitivní a negativní aspekty postoje k umělé inteligenci. Ve snaze překonat tato omezení byla vytvořena škála ATTARI-12 (Stein et al., 2024). ATTARI-12 byla vytvořena s ohledem na možnost jejího interdisciplinárního uplatnění v podobě sady technologických možností, které jsou nezávislé na její podobě a kontextu použití. Důležitost zkoumání AI jako nehmotného konceptu potvrzují nedávné výzkumy (Stein et al., 2020), které předpokládají, že hodnocení a postoje k technologiím se změní ve chvíli, kdy budou mít vizuální podobu.

3.3.2 Aktuální studie s tematikou AI

Více než sto milionu uživatelů (UNESCO, 2023) jazykového modelu ChatGPT během prvních měsíců užívání je doposud nejrychlejší šíření technologie, jakého jsme byly svědky. U uživatelů získal velkou oblibu zejména díky svému jednoduchosti a snadné online dostupnosti (Raman, 2023 citovano v Dwivedi, 2023; Shoufan, 2023). Není překvapivé, že samotná podstata lidské komunikace je klíčem úspěchu ChatGPT, učení skrze lidský faktor umožnilo dát AI lidské atributy, a tedy ji otevřelo všem uživatelům, kteří ovládají jazyk (Bubeck et al., 2023). Je nezbytně také nutné si připomenout, že lidský prvek je pro práci s AI klíčový a nepostradatelný (Johnson & Verdicchio, 2017; Shoufan, 2023). Umělá inteligence má vliv na jedince, organizace i společnost, jaké to s sebou nese důsledky a výzvy zkoumá Dwivedi et al. (2023).

Samostatné studie, zabývající se přímo jazykovými modely jako je ChatGPT, pokrývají různá specifická témata. ChatGPT má velký potenciál při využití v mnoha oblastech například ve zdravotnictví (Biswas, 2023b; Eysenbach, 2023), vzdělávání (Cotton et al., 2024; King & chatGPT, 2023), výzkumu (Dis et al., 2023), financích (Dowling & Lucey, 2023; Saggu & Ante, 2023) v boji proti globálnímu oteplování (Rane et al., 2024; Biswas, 2023a) a dalších oblastech. Seznámíme se nicméně i s výzkumy, které se zabývají obecně umělou inteligencí.

Obavy

Využívání AI s sebou nese obavy o bezpečnost, spolehlivost, plagiátorství a další. Obavy a očekávání jsou základními tématy, které rezonují odbornou veřejností. Možnosti ChatGPT a širší využití na nejrůznější úkoly je živou oblastí, kde se střetává lidská invence běžného uživatele s neustále se posouvajícím vývojem a možnostmi použití. Obavy, které se často pojí s umělou inteligencí, jsou možný budoucí vývoj a inteligenční schopnosti, které přesáhnou ty lidské (Gabbiadini et al., 2023; Peters et al., 2023). Obavy plynou také z možného využívání AI k úkolům, které se týkají například výběru zaměstnanců, psychologického poradenství (Rodway, 2020) a dalších oblastí, kde hraje roli hodnocení a názor. Gherhes (2018) se zabývá i negativním vlivem AI na mezilidské vztahy, snížení pracovních míst, novou ekonomickou krizi, ale i hrozbami v podobě zneužití AI k výrobě inteligentních zbraní, zvýšení vojenských konfliktů či ovládnutí lidstva a jeho zničení (Gherhes, 2018).

Výčet obav spojených s AI přinesli také Wang & Wang (2022), řadí mezi ně ohrožení soukromí, zaujatost, zánik pracovních míst, existenciální hrozby a netransparentnost. Ukázalo se, že tyto obavy jsou spojené s celkově obezřetným postojem vůči AI. Další výzkumy souvisejí s obavami v různých specifitějších oborech jako je lékařská péče, již zmiňovaný proces náboru zaměstnanců nebo hodnocení finanční zodpovědnosti (Kieslich et al., 2021). Očekávání spojené s možnou změnou pracovního trhu se u výzkumu Braunera et al. (2023) překvapivě netýkají ztráty pracovních míst, ale výhod ekonomického růstu a rozvoje inovací. Skloňované obavy se týkají sběru a zneužití citlivých osobních i firemních dat (Gupta et al., 2023; Wu et al., 2023; Ray, 2023).

Etika

Obavy ohledně možného zneužívání rozvíjejí také potřebu vytvoření etických pravidel (Ray, 2023; Zhang et al., 2023; Tlili, 2023). Hadi Mogavi et al., (2024) zdůrazňuje důležitost vyváženého přístupu k možným přínosům a hrozbám spojenými s AI. Studie zdůrazňuje také důležitost větší diverzity a kulturního uvědomění v technologiích zpracování přirozeného jazyka (Cao et al., 2023). Morální otázky ke způsobu používání AI ve výuce i mimo ni (Dwivedi et al., 2023; Májovský et al., 2023) se týkají jak studentů, tak učitelů, lektorů. Na straně vyučujících je nedůvěra a podezírání z nepřiznaného použití jazykových AI modelů, na druhé straně jsou mylná obvinění studentů. V prostředí, kde použití AI nástrojů není jasně definované, ani není jasné, zda se smí používat, či jsou prostě

zakázané, vzniká morální dilema a tlak. Postoji studentů k využívání AI ke studijním účelům se zabývají dalších studie například Alzahrani, (2023).

Vzdělávání

Výzkumy s tematikou umělé inteligence s cílovou skupinou studentů, se zabývají otázkami užitečnosti AI ve vzdělávacím procesu. Na etické otázky spojené s oblastí vzdělání a vzděláváním, které nás zajímá, je nahlíženo z různých úhlů pohledu (Javaid et al, 2023;). Potřebou etického vedení při využívání jazykových modelů se zabývá (Cotton et al., 2024; Crawford et al., 2023). Anders (2023) přemýšlí o tom, jak AI jako ChatGPT mění pravidla ve vzdělávání, naznačuje, že AI by měla být vnímána jako učební nástroj, nikoli jako zdroj podvodu. Podle výzkumu Alzahrani (2023) studenti vidí přínos AI v každodenních aktivitách a věří, že je zjednodušuje, ale nemají zcela přehled o možnostech používání ve výuce a vzdělávání, nedostatečně rozumí etickým otázkám spojeným AI. Alzahrani (2023)⁶ vidí nutnost vzdělávání o etickém používání, jak na straně studentů, tak učitelů. Ne vždy jsou studentské názory na používání AI při psaní nejsou jednoznačně pozitivní (Cummings et al., 2024). Jak na příchod umělé inteligence reagují české školy u nás, se zabýval výzkum České školy a umělá inteligence (Kopecký et al., 2023), který byl zaměřený na měření postojů učitelů. Terzi (2020) poukazuje na potřebu vzdělávání v oblasti ochrany soukromí a citlivých dat (Terzi, 2020; Sebastian, 2023). Dwivedi et al. (2023) poukazuje na podceňování podpůrné role AI při práci a vzdělání.

Komunikace

Ojedinelý způsob komunikace je jednou z charakteristik umělé inteligence ve vztahu k lidem. To se může projevit ve vztahu k lidem pozitivně i negativně (Bowman et al., 2023; Gabbiadini et al., 2023; Mallick et al., 2023). Problematiku v komunikaci na straně AI i jejich uživatelů s sebou mohou nést otázky nepravdivých odpovědí AI (Dwivedi et al., 2023; Alawida et al., 2023; Cabitza et al., 2017). Mírou porozumění AI modelům se zabývá Brown et al. (2020). Schopnosti GPT-4 v řešení nových a obtížných úloh a to výrazně překonávání předchozích modelů ChatGPT zkoumá (Bubeck et al., 2023).

⁶ Možné pozitivní využití AI najde uplatnění například při řešení učebních obtíží, při podpoře učitelů, učení jazyků a dalších situací. Podstatou je možnost přizpůsobení učebních procesů potřebám jednotlivců, pomoci dětem se speciálními potřebami, snížit počet odchodů ze škol. Proto hovoří Alzahrani (2023) o nutnosti popularizace výhod AI a vytvoření legislativního a etického rámce pro použití ve vzdělávání.

3.3.3 Postoje vybraných škol k používání AI

Bezprecedentní rychlost, s jakou se generativní umělá inteligence začala šířit světem a našla odezvu mezi běžnými uživateli, příliš nenahrála okamžité reakci ve formě stanovisek či regulací. Od velkých správních jednotek jako je Evropská Unie, jednotlivých států, regionů, krajů, obcí přes jednotlivé instituce, nebylo v jejich možnostech plně zachytit široké pole možností využití tohoto nástroje. Efektivní formulace stanovisek k možným pravidlům jejího používání se objevují průběžně.

Jak je již zmíněno v kapitole dvě, nástup generativní umělé inteligence a velkých jazykových modelů rostl od jejího uvedení raketovou rychlostí a ChatGPT drží nadále prvenství jako vůbec nejrychleji rozšířená aplikace. Jak zmiňuje Příručka pro generativní AI ve vzdělávání a výzkumu UNESCO (2023)⁷ od listopadu 2022, kdy byl ChatGPT zveřejněn, do ledna 2023 dosáhl počet uživatelů sto milionů. Během následujícího roku přijala regulaci pouze jedna země a tou byla Čína (*Guidance for Generative AI in Education and Research* | UNESCO, 2023). Otázkou je, zda vůbec k regulaci přistupovat a jakou by měla mít formu. Příkladem mohou být autorská práva, ochrana produktů lidské činnosti. Regulace by se neměla týkat pouze způsobů používání, tedy etických pravidel, ale také práva a legislativních systému. Ujasnění hranic kreativní činnosti lidí versus umělé inteligence je palčivou otázkou, kterou je potřeba zodpovědět.

S ohledem na téma této bakalářské práce vyvstává zejména otázka využívání na školní, akademické půdě. Nástroje generativní AI se neustále zdokonalují a psaní textů, generování kódů, obrázků či videí, nutně vede vzdělávací instituce k přehodnocení dosavadních pravidel a přístupů ke vzdělání, zejména v oblasti hodnocení výstupů. Internet přinesl revoluci v podobě dostupnosti ohromného množství výzkumů a zdrojů. Umělá inteligence přináší možnost s daty efektivně pracovat, syntetizovat, v nasycenosti témat i lépe vyhledávat, porovnávat a orientovat se v nich. Vzdělávací instituce zasáhl tento pokrok nepřipravené a nyní čelí otázce jak, proč a co se učit. Dle mého názoru by nemělo být cílem jejich využívání omezit, ale nastavit pravidla jejich využívání, zařadit práci s nimi do osnov a naučit pedagogy, školáky, studenty s nimi efektivně a informovaně zacházet.

⁷ Vlastní překlad, dokument nebyl oficiálně přeložen do českého jazyka, v originále je to *Guidance for Generative AI in Education and Research* | UNESCO, 2023)

3.3.4 Postoje vysokých škol a univerzit

Používání umělé inteligence na univerzitách se odrazilo také v českém prostředí. K umělé inteligenci a jejímu využívání se vyjádřilo několik českých univerzit, řada z nich však oficiální vyjádření vůbec nevydala. Univerzita Karlova formulovala několik stručných bodů ve formě stanoviska, brněnská Masarykova Univerzita formuluje nejen svá stanoviska rozdělená pro studující a vyučující, ale také doporučení. Univerzita Palackého v Olomouci, konkrétně Pedagogická fakulta vydala v dubnu 2023 prohlášení k používání AI. V srpnu téhož roku zveřejnila webové stránky pro tematiku AI pro „budoucí i současní pedagogy“.

Masarykova univerzita definuje (MUNI, 2023a) podmínky pro bezpečné a etické využívání umělé inteligence v akademickém prostředí skrze své stanovisko a doplňující text navazující na toto stanovisko – Doporučení k využití nástrojů AI při plnění studijních povinností (MUNI, 2023b). Podporuje studenty ve zvědavosti, ale také ve svědomitosti a pragmatičnosti při práci s AI nástroji s cílem podpořit učení, při dodržování etických pravidel. Důraz je kladen zejména na transparentnost při využívání AI, ale také na opatrnost s ohledem na citlivost dat. Upozorňuje také na problémy, které s sebou nesou nespolehlivé nástroje pro odhalení původu textu. Postoj univerzity k AI je ale velice otevřený, zavazuje se k pokrokovému přístupu ve výuce a jejímu přizpůsobení novým technologiím a trendům. V nezávazném doporučení jsou například uvedeny konkrétní nástroje, které lze při tvorbě akademického obsahu, seminárních prací využívat a také jak tyto nástroje citovat a správně uvádět ve zdrojích. Jejich etické, tvůrčí, přínosné a bezpečné využívání je základním požadavkem univerzity (MUNI, 2023b).

V souvislosti s využíváním AI společností rezonovalo téma zrušení bakalářských prací na Vysoké škole ekonomické v Praze na Fakultě podnikohospodářské. Vedení fakulty se rozhodlo upustit od psaní bakalářské práce a nahradilo ji bakalářským projektem. Ten by měl mít různé podoby. Na webových stránkách fakulty uvádějí například možnost absolvovat odbornou stáž, zahraniční studijní pobyt, fakultní výzkumný projekt nebo také podnikatelský projekt. (Kruczková, 2024).

Technická univerzita v Liberci nezůstává pozadu. Na začátku září 2023 vydala směrnici, ve které naopak poukazuje na důležitost používání, a nalézání cest, jak spolupracovat s AI. Efektivní využívání AI v budoucím povolání studentů je jednou ze základních premis tohoto přístupu. Spolupráci studentů s AI podporují jak při akademické práci i psaní závěrečných prací. Podmínkou je zodpovědný přístup, jehož nedílnou součástí

je transparentnost ohledně používání a správné citování zdrojů. Vědomí rizik při práci s AI, jako může být klamavý obsah, je další z doporučení, které by mělo být při práci s nástroji reflektováno. Povzbuzení ze strany univerzity zahrnuje i práci pedagogických pracovníků. Aktivní zapojení studentů a nástrojů AI do výuky by mělo vést k oboustranné symbióze a nejlepším možným výsledkům (Kočárková, 2023).

VÝZKUMNÁ ČÁST

4 VÝZKUMNÝ PROBLÉM

V kontextu psychologie přináší rozvoj a integrace umělé inteligence (AI) do každodenního života řadu otázek týkajících se postojů, vnímání, očekávání, ale také dopadů na společnosti i akademické prostředí (Dwivedi et al., 2023; Cotton et al., 2024). Jak se v této nové realitě studenti orientují, jakým způsobem přistupují k otázkám využívání AI, v běžném i akademickém životě, jaká je jejich motivace a do jaké míry je neexistence jasných pravidel závazná či omezující v jejich rozhodnutích, nebylo dosud předmětem žádné studie v České republice zpracované kvalitativní metodou. Inovace v oblasti AI a neustálý rozvoj otevírají nová témata a oblasti pro výzkum.

Cíl práce, **mapování současných postojů středoškolských a vysokoškolských studentů k AI, se zaměřením na ChatGPT a jeho varianty, skrze přehled témat, která mezi studenty rezonují**, poskytne náhled a hlubší porozumění a propojení prožitků studentů s používáním nástrojů AI. Proměnlivá povaha možností AI i prostředí s sebou nesou i nové výzvy a úskalí (Gabbiadini et al., 2023; Gherhes 2018) a implikace v oblastech jako je například komunikace (Bowman et al., 2023; Mallick et al., 2023), vzdělání a vzdělávání (Alzahrani, 2023), etika a morálka (HadiMogavi et al., 2023) a dalších oblastí, které si zaslouží bližší prozkoumání zejména s ohledem na budoucí vývoj.

Většina výzkumů na téma postojů k AI jsou kvantitativně zaměřeny, s důrazem na jednotlivé aspekty problematiky používání AI (Biswas, 2023a; Rane et al., 2024; Rodway 2020; Peter et al. 2023). Kvalitativní metoda by měla mít ve výzkumu AI své pevné místo. Dokáže zachytit aktuální stav a náladu ve společnosti, názory a postoje. V rámci České republiky je počet výzkumů na téma postojů studentů k AI zanedbatelný, ač je veřejná debata na toto téma velice živá, nepodařilo se dohledat žádné relevantní studie z českého prostředí.

Díky vypovídající hodnotě této práce zasazené stále v období začátků masového používání nástrojů AI ji považujeme za přínosnou v debatě o implikacích používání nástrojů AI do budoucna. Přispět k debatě o směřování a regulaci jazykových modelů a jejich využívání v akademickém prostředí, skrze poznání témat, ale i míry odbornosti jakou studenti mají má zanedbatelnou hodnotu.

4.1 Formulace výzkumných otázek

Z našeho výzkumného problému vycházejí jednotlivé výzkumné otázky, které jsou obsahem následující analytické části této práce.

1. Jaké jsou individuální postoje středoškolských a vysokoškolských studentů k umělé inteligenci, vnímají ji spíše pozitivně nebo negativně?
2. Jakým způsobem vnímají respondenti velké jazykové modely?
- ~~3.~~ Jaká je povaha zkušeností studentů s AI?
4. Jaká jsou očekávání respondentů týkající se vzdělávání o AI?
5. Jaké jsou nejčastější obavy studentů a hrozby týkající se rozvoje a používání AI?
6. Jaké etické otázky respondenti identifikují v souvislosti s AI?

5 TYP VÝZKUMU A POUŽITÉ METODY

Kvalitativní výzkum v psychologii nabízí hlubší porozumění lidským postojům, chování, myšlenkám nebo emocím. Lze získat detailnější informace, odhalit témata, které bychom za pomoci kvantitativních metod mohli opomenout.

„Kvalitativní metody se užívají k odhalení a porozumění tomu, co je podstatou jevu, o nichž toho ještě moc nevíme. Mohou být také použity k získání nových a neotřelých názorů na jevy, o nichž už něco víme. V neposlední řadě mohou kvalitativní metody pomoci získat o jevu detailní informace, které se kvantitativními metodami obtížně podchycují“ (Strauss & Corbinová, 1999, s. 11).

Je mnoho faktorů, které ovlivňují postoje studentů k používání umělé inteligence. Výsledkem tohoto kvalitativního výzkumu by mělo být zpřehlednění zkoumané oblasti, přinesení témat, která zajímají uživatele AI, studenty, zmapování jejich současných postojů i výhled do budoucna. Jako metodu sběru dat jsme zvolili polostrukturované rozhovory, které jsou vhodné pro odhalení individuálních zkušeností i kontextů a umožňují komplexní vyjádření myšlenek, které mohou být pro výzkum klíčové.

5.1 Metoda získávání dat

Povaha výběru respondentů vycházela ze samotných cílů práce, před zahájením výzkumného projektu a analýzou dat. Zvolili jsme tedy respondenty způsobem, který odpovídá zamýšleným charakteristikám, kde je nejdůležitější stupeň aktuálního studia a typ vzdělávací instituce, ale také rozložení pohlaví. Byl zvolen účelový výběr.

Oblast výzkumu umělé inteligence nabízí mnoho neprozkoumaných témat. Pro potřeby této práce byly pro sběr a tvorbu dat zvoleny polostrukturované rozhovory. Polostrukturované interview nabízí velkou flexibilitu při realizaci, zároveň má i hranice, byť propustné. Kombinace předem připravených otázek s možností přizpůsobení rozhovoru v jeho průběhu a možností volně navazovat dalšími tématy, umožňuje diskusi prohlubovat a lépe individuální postoje prozkoumávat. Pořadí otázek, ale i jejich formulace, vzhledem k typu respondentů, lze podle potřeby měnit, vysvětlovat, vracet se k nim. Oproti pevně strukturovanému rozhovoru je tedy mnohem přirozenější a má větší potenciál odhalit nová i zamýšlená témata a oblasti. Polostrukturovaný rozhovor je zárukou pokrytí cílených témat,

lze tak snadno předejít opomenutí důležité otázky či problému, jak by se mohlo snadno stát v nestrukturovaném rozhovoru (Miovský, 2006).

Zvýšení validity zjištění bylo podpořeno triangulací získávání dat od různých respondentů, ale také časovou triangulací sběru dat v různých časech a kontextech.

5.2 Metoda zpracování a analýza dat

Výsledky sběru dat skrze polostrukturované rozhovory. Otevřené otázky umožnili respondentům bez omezení rozvíjet jejich myšlenky. Ty byly simultánně zaznamenávány na dvě nahrávací zařízení, výsledkem byly audio nahrávky, které jsme po ukončení výzkumu smazali. Následně byly rozhovory za pomoci transkripčních nástrojů Microsoft Office – Word přepsány do elektronické verze textu a podrobně editovány a kontrolovány opakovaným poslechem.

Odpovědi studentů byly editovány a analyzovány skrze studentskou placenou verzi programu Atlas.ti. Tento nástroj dokáže ve spolupráci s výzkumníkem analyzovat kvalitativní data jako jsou například rozhovory, přepisy textů či dotazníková šetření. Atlas umožňuje data z rozhovorů kódovat, označovat jeho části a přiřazovat jim kódy s cílem identifikovat opakující se témata. Množství kódů je neomezené, lze je přidávat či odebírat dle potřeby, jednotlivé části lze označit několika kódy. V závěrečné části byla data vyexportována do Excelu a dále analyzována. Témata se objevují skrze určení podobných kódů a jejich logického tematického spojování. Kvalitativní zpracování dat skrze tematickou atematickou analýzu umožnilo myšlenky respondentů kategorizovat, na základě opakujících se vzorců, tedy kódů a témat.

S ohledem na deskriptivní a analytické cíle této práce, kterým je seznámení se a mapování postojů na umělou inteligenci u středoškolských a vysokoškolských studentů (se zaměřením na jazykové modely jako ChatGPT) byla zvolena tematická analýza (Howitt, 2010). Postup tematické analýzy vychází z doporučení Braun a Clarke (2006). V rámci metody tematické analýzy jsme zvolili interaktivní induktivní postup, ve kterém jsou témat a oproti teoretickému neboli deduktivnímu postupu volena s větším propojením na sebraná data (Patton, 1990). V jeho průběhu probíhá souběžně několik procesů – tvorba vzorku, tvorba dat, jejich analýza, interpretace společně s reflexí celého výzkumu (Hendl, 2008). Cílem není aplikace dat na známé teorie, rozhodující roli v analýze hraje samotná podstata

našich dat. V průběhu práce byly data porovnávána a tříděna s cílem je popsat a dále kategorizovat.

5.3 Tvorba dat

Tvorba dat proběhla formou polostrukturovaných rozhovorů s deseti respondenty na téma postojů k umělé inteligenci, a to jak formou osobního rozhovoru bez přítomnosti další osoby nebo, skrze online nástroj Google Teams také při zachování soukromí. Nejdříve byly připraveny sady otázek, které byly rozděleny na šest témat souvisejících s AI: postoje; zkušenosti; informovanost; obavy a očekávání; vzdělání a povědomí a etické aspekty. Každé z témat obsahovalo dalších 4-5 otázek, celkem jich tedy bylo přibližně v rozmezí 24-30. Pilotní testování ukázalo, zda jsou otázky srozumitelné a relevantní. Na základě pilotáže byly upraveny a doplněny. Po pilotní fázi následoval sběr samotných dat, tedy jednotlivé rozhovory. Celkově proběhlo deset rozhovorů, z nichž byla polovina s respondenty ze střední školy, druhá polovina z vysoké školy. Délka rozhovorů se pohybovala v rozmezí 30-40 minut. Hodnocení dat probíhalo průběžně tak, aby byla data dostatečně saturována.

5.4 Výzkumný soubor

Nejprve jsme určili skupinu lidí, kteří jsou pro tento výzkum relevantní vzhledem k jejím cílům a zaměření, tedy středoškolské a vysokoškolské studenty. Podmínkou účasti na výzkumu nicméně nebyla přímá zkušenost s jazykovými modely, a i případná nezkušenost v tomto ohledu je v této práci chápána jako přínos do mozaiky zkoumání. Klíčový faktor pro výběr respondentů byla znalost jazykových modelů jako je ChatGPT a jiné.

Kategorie jsme nijak věkově neomezovali. Z řad středoškolských studentů se věk respondentů pohyboval v rozmezí 16–19 lety, u vysokoškolských studentů bylo věkové rozpětí 26–38 let odpovídající zastoupení prezenčních i kombinovaných studentů.

Respondenty jsme oslovovali většinou v návaznosti na osobní známost nebo skrze jejich blízké, prostřednictvím telefonu nebo jiného komunikačního online nástroje jako Messenger nebo aplikace WhatsApp. Samotné rozhovory probíhaly osobně, nebo skrze nástroj pro online konference Google Teams.

Výběr respondentů byl nicméně ovlivněn také jejich dostupností. Cílem práce bylo oslovit studenty různých oborů a zaměření. Jak z řad středoškoláků, tak vysokoškolských

studentů. Zda jsou vysokoškolští studenti na bakalářském nebo navazujícím magisterském programu nehrálo při jejich výběru roli. Mezi oslovenými jsou také studenti kombinovaného studia, kteří jsou již pracují v různých oborech. Náš výzkumný soubor tím kromě studentských zkušeností získal i přesah do firemních procesů s ohledem na využívání AI a také zkušeností jiného druhu. Z celkového počtu byli čtyři ženy a šest mužů. Výběr a počet respondentů proběhl s ohledem na cíle mé studie a zároveň také relevantnost a validitu výsledků studie a saturaci dat.

Kromě věku, pohlaví a vzdělávacího stupně, byly sledovány také obory studia a jeho typ (učiliště, střední odborná škola s maturitou, gymnázium, nástavba s maturitou, prezenční vysokoškolské studium, kombinované vysokoškolské studium). Rozložení respondentů v rámci oborů studia je následující – jazykové gymnázium, střední lesnická škola, nástavbové studium v oboru podnikání, průmyslová škola technického zaměření, střední škola řemesel a služeb (učiliště a střední škola); vysokoškolské zaměření oborů dalších respondentů je – informatika (softwarový vývoj), ekonomie, psychologie, humanitní studium vzdělanosti a občanské společnosti a obor matematiky a fyziky. Vybraní respondenti měli také zkušenosti na pracovním trhu, což přirozeně rozšířilo rámec zkoumání. Pro potřeby práce jsme jednotlivé studenty náhodně pojmenovali jinými jmény: Josefina, Marie jsou vysokoškolské studentky, Nad'a a Hana středoškolačky, Sebastian, Šimon, Pavel jsou vysokoškoláci, Marek, Karel a Adam jsou středoškolští studenti. Údaje, které by mohly vést k identifikaci respondentů nejsou předmětem výsledků sběru dat, byly úmyslně vynechány.

Tabulka č. 2: Deskriptivní charakteristiky respondentů

respondent	pohlaví	věk	studium	pracující
Josefina	žena	26	vysokoškolské	ano
Marie	žena	31	vysokoškolské	ano
Nad'a	žena	18	středoškolské	ano
Hana	žena	16	středoškolské	ne
Sebastian	muž	24	vysokoškolské	ano
Šimon	muž	38	vysokoškolské	ano
Pavel	muž	26	vysokoškolské	ano
Marek	muž	19	středoškolské	ne
Karel	muž	18	středoškolské	ne
Adam	muž	17	středoškolské	ne

5.5 Etické hledisko a ochrana soukromí

Ještě před zahájením rozhovorů byli respondenti seznámeni s cílem výzkumu, průběhem rozhovoru, jeho časovými nároky a také s pořizováním audiozáznamu z interview a *ujištěním*, že budou tato data pečlivě chráněna. Zároveň byla všem respondentům zaručena anonymita rozhovorů, jejich dobrovolná účast a možnost kdykoliv rozhovor bez udání důvodu ukončit. Účast na rozhovoru nebyla spojena s žádným nárokem na odměnu.

Všechna setkání se odehrávala v klidném a soukromém prostředí, a to jak při osobních setkáních, tak i při online rozhovorech či videohovorech. K rozhovorům nebyla přizvána žádná třetí osoba. V případě respondentů mladších 18 let byl informovaný souhlas získán verbálně, přičemž rozhovory s nimi probíhaly výhradně prostřednictvím online platformy. Audio nahrávky byly pečlivě archivovány a přístupné byly pouze členům výzkumného týmu, aby bylo zajištěno soukromí a anonymita účastníků. Po ukončení výzkumu budou smazány. Cílem dotazování bylo získat relevantní informace pro účely této studie.

6 PRÁCE S DATY A JEJÍ VÝSLEDKY

V této kapitole se budeme zabývat analýzou sebraných dat. Za pomoci tematické analýzy budeme postupně odpovídat na výzkumné otázky. Identifikovali jsme několik témat, která jsme následně blíže prozkoumali podle jednotlivých podoblastí. Celkové postoje k AI jsou rámcovým konceptem, ze kterého se v této výzkumné práci snažíme identifikovat jednotlivé vrstvy, jejich charakteristiky, odbočky, paradoxy či zdánlivě nepodstatná fakta.

6.1 Individuální postoje k AI

V této podkapitole budou představeny výsledky první výzkumné otázky – *Jaké jsou individuální postoje středoškolských a vysokoškolských studentů k umělé inteligenci, vnímají ji spíše pozitivně nebo negativně?*

Na základě rozhovorů s respondenty se vynořilo téměř dvacet hlavních témat, která odrážejí individuální postoje jednotlivých respondentů a jejich obecného **vnímání AI jako pozitivního či negativního** nástroje. Umístění na této škále vychází ze subjektivního rozhodování výzkumníků na základě výpovědí respondentů. Tato část práce není limitována názory respondentů na ChatGPT a další LLM⁸, vztahuje se na umělou inteligenci celkově. Přinesení každého tématu vypovídá o individuálním nastavení jednotlivých respondentů. Postoje jsou rozmístěny na pozitivně negativním spektru, proto jsme je rozdělili do tří kategorií pozitivní, neutrální a negativní. Respondenti identifikovali **pozitivně** laděná témata: zjednodušení a zrychlení práce ve školním i pracovním prostředí, úspora času při shromažďování a třídění informací, usnadnění rutiny, snadná adaptace a integrace do života, inovace a technologický průlom, součást historického momentu a fascinace objevem AI.

Mezi **negativní** témata řadíme: přílišné delegování kompetencí, rychlý a nekontrolovatelný rozvoj, omezení kompetencí AI, netransparentní investice do AI, podceňování etické regulace, produkce manipulativního dezinformačního obsahu a zneužití AI. Objevují se i **neutrální** témata: nezbytnost vzdělávání, důležitost porozumění, správné použití, legislativní rámec, ochranné známky.

⁸ Large language models (LLM) – velké jazykové modely, jak souhrnně nazýváme další nástroje fungující na stejném principu jako ChatGPT.

AI jako potenciální riziko a nebezpečí pro společnost i ohromný přínos dobře vystihují centrální téma u většiny respondentů – nalezení rovnováhy mezi riziky a pozitivními aspekty AI. Díky těmto odpovědím můžeme sledovat širší sociální a kulturní trendy v přijímání AI.

Tabulka č. 3: Přehled pozitivních, neutrálních a negativních aspektů postojů k AI

pozitivní	neutrální	negativní
• zjednodušení, zrychlení práce • rychlá adaptace • snadná adaptace • snadná integrace • repetitivní úkony • třídění, shromažďování informací • strukturování, tvorba textů • inovativní technologický průlom • unikátní historický moment	• vzdělávání o AI • porozumění správnému použití • legislativní rámec • etická pravidla	• delegování lidských kompetencí • rychlý a nekontrolovatelný rozvoj • netransparentní investice do AI • neexistence etické regulace • produkce dezinformačního obsahu • možné zneužití

Prvotní pozitivní postoj je spjat zejména s praktickým využitím jako je **zjednodušení** a **zrychlení práce** ve školním i pracovním prostředí. „*Já si myslím, že je to dobrý, že to urychlí práci. Může pomoci, že ti to něco vymyslí nebo když třeba na to nemáš čas a hodně toho zbývá, tak ti to pomůže*“ (Hana). Čistě pozitivní postoj se ukazuje u studentů, pro které je AI zejména **zdroj úspory času** při **shromažďování** a **třídění informací**, usnadňuje **rutinní** školní práci. „*Můj postoj je zatím pozitivní, protože pro mě to je velký ulehčovač práce. Já jsem se s tím dost naučila fungovat s takovým víceméně partákem, bez kterého si to už ani neumím představit. Je to pro mne hrozná úspora času, když mám zdrojovat z pěti zdrojů, to AI mi to vyplivne rovnou a já to nemusím hledat stránku po stránce. Takže je to usnadňovač každodenních školních povinností*“ (Naďa). Dalším významným faktorem čistě pozitivního hodnocení je **snadná adaptace** a **integrace nástrojů AI** do každodenního

života. „Celkově umělou inteligenci velmi vítám, jelikož mi hodně pomáhá jak v práci, tak ve škole. Takže za mě je to tady v tom ohledu určitě jako pozitivní. A vlastně jako teďka po pár měsících, už je to spíš jako dost měsíců, co ji hodně používám, vyloženě třeba ChatGPT, tak třeba teď už si jako moc nedovedu představit, že bych ho jako nepoužíval vůbec. Já už jsem si na to hodně zvyknul, že to tady s náma je a že to za mě může jako udělat spoustu věcí, které na to velice rád deleguji“ (Pavel). Pozitivně hodnocená je také **inovace a technologický průlom** a vůbec možnost být **součástí** tohoto vývoje. Ten s sebou ale nese obavy z přílišného **delegování** činností vlastních lidem. „Říkala jsem si nedávno, že je to zajímavý štěstí, že žijeme zrovna v době, kdy se taková velká věc rozvíjí, že to je náhoda, že zrovna jsme u toho a že to asi je dobrý. Ale je tam i nějaká obava takový jako s robotama, že má člověk strach, že by mohli někdy převzít moc nad lidstvem. Ne že bych z toho měla strach, že by mě to trápilo, ale že to je asi kombinace toho, že to nemáme úplně pod kontrolou. Zatím z mojí zkušenosti mám pocit, že to velmi zjednodušuje práci a co vidím kolem sebe mi připadá dobrý. Asi teď moc nedokážu nahlédnout nějaký negativa. Ale postoj bych řekla, že mám pozitivní“ (Josefina). Nedostatečná kontrola je ústředním tématem i konzervativnějších postojů k rozvoji AI, pojí se s mírnými obavami. Ty se týkají zejména **rychlého a nekontrolovaného rozvoje**. I přesto, že vidí přínosy, volají po regulaci a dodržování etických zásad. „Je to super nástroj k tomu, že si člověk může různě překládat, že ten program udělá spoustu věcí, umí jich spoustu vymyslet třeba v programování. Jako technologie je to super, ale nechal bych to jako fakt technologii, nástroj k něčemu. No jako nenechal bych to zajít nějak dál“ (Marek). Možnost, kdy AI dojde to bodu, kdy **nebude kontrolovatelná**, je spojena jak s jejím vlastním technologickým pokrokem, tak s externími faktory. **Kontrola** globálních nástrojů **úzkými skupinami**, společnostmi, lidmi, kteří do AI investují, rozvíjejí, a tedy jí svým způsobem kontrolují je možná tím nejpalčivějším problémem. **Etická regulace** je **podceňována** a podhodnocena. „Vnímám nějakou praktičnost tady těch nástrojů, ale mám pocit, což mne samotného překvapuje, že jsem spíš na konzervativnější straně. Jako, protože mě to osobně vlastně taky zasahuje do práce. Jsem konzervativnější a nemám rád věci, které se rozvíjejí moc rychle a nekontrolovaně a mám pocit, že tam je spousta hrozeb, který bychom měli ošetřit. A nezdá se mi, že se to děje, protože je to prostě takový závod mezi různými společnostmi, který to vyvíjejí a investují do toho peníze a myslím si, že se málo prostě řeší etika“ (Šimon). Absolutní nezbytnost **vzdělávání** a důležitost **porozumění** o jejím **správné použití**, by měly být efektivními nástroji, pro kontrolu, regulaci a proti zneužití. Nárůst **manipulativního obsahu** se záměrně nepravdivými informacemi jako jsou deep fake videa známých osobností, budí obavy.

Klíčovou roli hraje zavedení **legislativního rámce, ochranných známek** a samozřejmě vlastní **kritické myšlení**. „Myslím, že je to zároveň nebezpečný a zároveň to může být hrozně užitečný. Je nutný se vzdělávat o tom, jak ji správně použít. Mám hodně obavu i co se týče *deep fake*, protože to hodně řešíme i s lidma kolem mě, generovali jsme jak obrázky, tak texty a vím, že se na internetu objevují hodně videa třeba Elona Muska a věci, který nikdy neřekl a jsou to vlivný osobnosti. Ale to bylo třeba i u nás, na netu se objevila fotka někoho od Starostů s Putinem a vlastně lidi, který nemají kritický myšlení, tak vůbec nedokážou vyhodnotit, že to je fake. A vlastně si myslím, že bychom to měli i víc jako regulovat zákonem, aby musel být vodotisk třeba v těch obrázcích nebo v tom videu, aby bylo jasné, že to je jako fake. Protože si myslím, že jinak je to nebezpečný jako prase. Takže si myslím, že jako je nutný, asi jako u všeho, když se vynalezla atomová energie, tak to můžeš používat a může ti to být vlastně hrozně k užitku, ale musíš to dělat s rozumem, no“ (Marie). **Ohromení** nad samotným objevem AI a možnosti pozitivního rozvoje, se hned vzápětí střídají s hořkostí z nepozorovaných možností jejího **zneužití**. „Vidím tam spoustu rizik, které jsou s tím spojený. A tak teďka je to vlastně ta fascinace, že to je něco jako strašně zajímavýho, co nás může posunout, ale zároveň jako strašně velký riziko. Já mám pocit, že vždycky, jako když se objeví nějaký pozitivní jako boom, jakože ten chat GPT, tak hnedka s tím jako vyvstanou ty negativní pohledy, že mám pocit, že jako spíš to, co vždycky ovlivňuje, je ten pozitivní boom a jako jak vlastně se to pak dá zneužít“ (Sebastian).

AI jako potenciální riziko a nebezpečí pro společnost i ohromný přínos dobře vystihují centrální téma u většiny respondentů – nalezení rovnováhy mezi riziky a pozitivními aspekty AI. Díky těmto odpovědím můžeme sledovat širší sociální a kulturní trendy v přijímání AI.

6.2 Vnímání AI

Další podkapitola se věnuje širším souvislostem vnímání umělé inteligence, tentokrát již konkrétně velkým jazykovým modelům. Skrze jednotlivá témata hledáme odpovědi na druhou výzkumnou otázku – *Jakým způsobem vnímají respondenti velké jazykové modely?*

V této kapitole používáme zástupné výrazy, tedy pokud hovoříme o AI, nebo umělé inteligenci, máme na mysli velké jazykové modely jako je ChatGPT a další. Čtyři hlavní témat určují linku analýzy. Jsou to: **komunikace s AI, AI jako entita, důvěra v AI, porozumění fungování AI a kritické myšlení**. Na základě výzkumného šetření by se dalo

usuzovat, že způsob, jakým respondenti s LLM jednají, jak je oslovují, jaký jim přiřazují rod, ale i jakou k nim chovají důvěru, nebo se vůči nim emočně vymezují, jsou vzájemně provázané rozvětvené oblasti.

Komunikace s AI

Vztahuje se na nástroje umělé inteligence **nárok na slušné chování** a vůbec obecně k technologiím? Měli bychom je oslovovat a jednat s nimi s **úctou a respektem**? Dokud jsme nezačali s nástrojem jako je ChatGPT na přímo komunikovat a ptát se jej, interagovat, a tedy zahajovat aktivní komunikaci, kde dostáváme zpětnou vazbu, pravděpodobně tato otázka nebyla vůbec tématem. Ačkoliv studenti nahlíží na toto své chování **shovívavě** a uvědomují si, že je možná zbytečné, přece jen nalezneme také náznaky možných **obav, úcty, pokory** před tímto nástrojem. Rozdíl se ukázal mezi studenty středních a vysokých škol. U studentů středních škol jsme zřídka našli tento způsob jednání. Tam jsme se setkali spíše s jasně vymezenou pozicí uživatele a jazykového modelu. Nabízí se otázka, zda je tento postoj způsobený pouhým věkem, přirozeným vývojovým stupněm s jeho atributy (přibližný věkový rozdíl mezi respondenty je 10 let) a celkovým postojem v tomto životním období, nebo jinými faktory, kterými může být odlišná zkušenost s technologiemi, vystavení novým technologiím ve větší míře nebo jiné zkušenosti oproti starším studentům. Rozdíly u respondentů se objevily také chápání rodu jazykových modelů. Zatímco u středoškoláků je oslovování primárně rodu středního, u vysokoškolských studentů je to variabilní.

Jak se ukázalo v rozhovorech, není výjimkou, že se k AI chováme jako k entitě, která si oproti vyhledávacím nástrojům jako je třeba Google Search „zaslouží“ slušné chování. Běžné lidské způsoby komunikace jako oslovování, **prosba a poděkování** se ukazují jako validní forma komunikace i směrem k umělé inteligenci. Toto s sebou ale nese i jisté znaky **obhajoby chování** ze strany respondentů, a to v obou směrech – kladného i negativního jednání. Z rozhovorů je cítit, že ani sami respondenti si nejsou jistí, jaká je ta správná, individuálně a **společensky kladně hodnocená** odpověď. „*Určitě mu poděkuji i si uvědomuji, že je to takový zbytečný. Ale tím, jak to používám málo, tak mi to připadá, že mne to nezdržuje, že kdybych to používala každý den, tak podle mě bych se na to vykašlala, ale že to nemám tak zautomatizované, tak napíšu děkuji. Je to i tak trochu ze srandy, že on mi napíše není zač, jakože mě baví, že je jak člověk. To je podle mě tím, že to používám jednou za čas, takže si tam takhle s tím hraju, ale kdybych to používala denně, tak samozřejmě to je to asi jiné*“ (Josefina).

Slušné chování je také přece jenom o **vyjádření respektu**. I jazykové modely přirozeně evokují **autoritu**, dokonce jim je přisuzována vlastní osobnost. „*Jo jo oslovuji a říkám i prosím, mám tam ten respekt. Je to zvláštní. Beru to jako osobnost, autoritu, ačkoli vlastně jako vím, že to je umělá inteligence*“ (Sebastian). Další z vysokoškolských respondentů odpovídá, že ho sice nezdraví a neslovuje, ale prosí a někdy i děkuje. Je to způsobem vyjádření **ocenění za kvalitní práci**, kterou LLM odvádí. „*Píšu mu prosím, občas napíšu, i děkuju. Ale většinou ho neoslovuji ahoj Chatbote. Když ho o něco žádám, tak říkám jako please can you do something. Takže jsem k němu většinou zdvořilý, protože si opravdu vážím jeho práce*“ (Pavel). Jiný z respondentů popisuje, jak se jeho chování vůči AI měnilo na základě spokojenosti s odpověďmi. Způsob komunikace vůči LLM je nejen nástrojem odměny, ale také **trestem**. „*Ze začátku, než mě naštvale, tak se jsem na ní byl hodnej a normálně jsem ji oslovoval*“ (Šimon). Trochu s nadsázkou, inspirovaný vtípem⁹, zazněl od jedné respondentky také důvod, proč považuje za důležité respektovat pravidla slušného chování. „*Je to prostě budoucnost. A v budoucnu nás tady budou ovládat.*“ (Josefina).

U středoškolských studentů je vidět jiný trend. Hranice mezi AI jako nástrojem anebo samostatnou entitou, ale také nejsou zcela jasné. Je patrná **sebereflexe**, uvědomují si, že jejich jednání může působit **neslušně**. I zde lze zachytit emocionální vztah k AI a zvažování ideálního scénáře. „*Když jsem si psal s ChatGPT, tak jsem napsal, udělej mi to takhle, udělej tohle, fakt jako surově. Dá se říct, že jsem tomu zadával příkazy*“ (Adam). Větší odstup od AI potvrzují i další respondenti. „*Napíšu jenom strohé – napiš mi tento text XY a tečka, děkuji. No ani tomu možná neděkuji*“ (Nad'a). Komunikaci s AI berou mnohem **utilitárněji** s vědomím, že kroky navíc, po nich k dobrému fungování jazykové modely nevyžadují. „*Co potřebuju, tak tam prostě jako láduju. To je vše, co chce*“ (Marek).

Rozdíl v chování je patrný, s ohledem na typ nástroje umělé inteligence. AI jazykový bot, který je součástí aplikace Snapchat se chová interaktivně, vede aktivní komunikaci. „*Na Snapchatu je to tak, že vy se s ním musíte pobavit. Když napíšu příkaz, tak k tomu potřebuje vědět víc. Odepíše, že dobrý, vím, co po mně chceš, ale potřebuje ještě víc. Takže pak už se s ním bavíte jako s kámošem*“ (Karel).

Prívětivost a slušné chování vůči AI se projevuje i zpětně k uživatelům. Zároveň může být způsob komunikace také cestou, jak dostat co **nejlepší výsledky**. Tento jednoduchý krok by měl vykrystalizovat co nejlepší odpovědi LLM. „*Přemýšlím někdy*

⁹ Tento kreslený vtíp je součástí přílohy.

o tom, že si říkám, tak na něj budu milá a on na mne pak bude taky milý, říkám mu prosím, please. A taky mi kolegové poradili, že je dobré mu dávat výzvy. Když mi něco dá, řeknu děkuju, ale myslím, že to umíš ještě líp a výsledky jsou pak lepší” (Marie).

Způsob, jakým se vztahují respondenti k AI se liší také **zosobněním rodu**. Pro některé je to dozajista žena, pro jiné je to mužský rod jako program, pro někoho je to ono, neurčitý rod. Opět jsme našli relevantní rozdíl mezi studenty středních a vysokých škol. Zástupce středoškoláků má jasný postoj. „*Rozhodně to nevnímám jako osobu*“ (Adam). U všech středoškolských respondentů jsme zaznamenali vztahování se k AI jako k věci, je pro ně **neurčitým rodem**, nástrojem. U vysokoškoláků jsme se setkali s různým vztahováním i přímo s požadavkem na ženský rod. „*Sám jsem ji poprosil, ať o sobě mluví v ženském rodě, to pro mne bylo důležité*“ (Šimon).

Důvěra v ChatGPT a další jazykové modely

LLM jsou nástroje, které vzhledem ke své podstatě nejsou vždy spolehlivé v pravdivosti svých odpovědí. Zkušenosti uživatelů a našich respondentů do jednoho potvrzují, a sami se o tom přesvědčili při práci s jazykovým modelem, že si své odpovědi jazykové modely někdy „vymýšlí“, že nejsou založené na pravdivých informacích. Takové výstupy mohou vést uživatel ke ztrátě zájmu, byť dočasné. Víří také emocionální odpovědi a vzbuzují u respondentů nelibost. Logicky ale vedou k nezbytnosti ověřovat výstupy AI a informace, které jim jazykové modely poskytují. Tento model můžeme považovat za přínos v oblasti samostatného myšlení a může být jednou z odpovědí na obavy a riziko zneužívání AI ve vzdělávání. Studenty podpoří v kritickém myšlení, procvičí je v ověřování informací a pomůže jim informace utřídit vlastními schopnostmi. Tato vlastnost LLM může být paradoxně přínosem, protože dělat chyby je „lidské“ a to v respondentech vzbuzuje důvěru. Shovívavost k výmyslům AI, může souviset také s porozuměním způsobu jeho fungování. Těm, kteří mají povědomí o jeho principech je jasné, už z podstaty generativních jazykových modelů, proč dostávají takové odpovědi. Tím se dostáváme k otázce informovanosti a povědomí o fungování generativních jazykových modelů.

Neduhem jazykových modelů je buď **přílišná obecnost**, nebo riziko **smyšleného obsahu**. „*Pohybujeme se ve dvou extrémech, buď ten model mluví strašně obecně, což je pak korektní na cokoliv, nebo si vymyslí věci, které v těch zdrojích, co mu dáme, nejsou. A pak se ho zeptáš, kde k tomu přišel a on – no já jsem si to vymyslel. Je to sranda*“ (Marie). Nepravdivé informace mohou vést k **dočasné ztrátě důvěry** a její **nepoužívání**. „*Psala jsem*

*seminárku a potřebovala jsem jednu informaci. Doptávala jsem se na konkrétní informace a ty zdroje si to začalo vymýšlet. Takže to jsem měla pak chvilku období, kdy jsem si furt říkala, ne, ne, to si prostě vymýšlí ty zdroje, tomu člověk nemůže věřit.” (Josefina) Podobnou zkušenost, kde svou roli hrají také **emoce a zklamání**, popisuje také další respondent. „Hrozně halucinovala...hledal jsem práce pro moji bakalářku a sám jsem našel třeba osm studií. Říkal jsem si, že určitě existuje více vědeckých studií, tak jsem to zadal do ChatGPT a ona jich vypsala třeba dvanáct. Říkám si to není možné, vypadalo to nádherně, říkám, nemáš další a ona dalších dvanáct, tak to je bomba, tak pojď ještě další. Myslel, že už mám třicet studií a začal jsem je postupně hledat a prostě žádná neexistovala. Tím mě naštvala“ (Šimon).*

Takový přístup naznačuje, že vzdělání respondentů a uživatelů AI, hraje klíčovou roli v **efektivním a kritickém používání** jazykových modelů. Díky této zkušenosti k aplikaci studenti přistupují obezřetně a **informace ověřují**. „Když to člověk používá na nějaký odzdrojování, tak když se takhle něco dozvím, tak si to projedu ještě jedním zdrojem a kontroluji, jestli to opravdu existuje“ (Marek).

Nejde však jen o vymýšlení si zdrojů a informací, ChatGPT se snaží dle zkušeností se svou verzí 3.5, vždy odpovědět. Nikdy **nechce své uživatele zklamat a dělá chyby** a možná i proto působí **sympaticky**, a to je vlastně jeho pozitivem. Přílišná dokonalost je nedůvěryhodná¹⁰. „Hodně si vymýšlí. Můj tat'ka se snažil shánět informativní medailonky různých důležitých osobností a říkal, že u nějaké osoby, kterou to zná, tak to řekne přesně a u zbytku nenapíše, že neví, ale začne si vymýšlet.“ (Hana)

Porozumění fungování AI

Je otázkou, jestli si studenti a naši respondent dokážou díky **porozumění fungování** generativním jazykovým modelům dobře zdůvodnit, proč jim jazykové modely lžou nebo odpovídají za všech okolností, byť nepravdivě. Ukázala se shoda na nezbytnosti používání **kritického myšlení a selského rozumu** a nutnosti **porovnávat výsledky s realitou**. I poměrně **povrchové znalosti** o fungování generativních jazykových modelů jsou dostatečné pro pochopení výsledných nepravdivých odpovědí. **Různá hloubka znalostí** tedy není nijak zásadní. Jako základní informace by měly postačit pochopení, že podávání **fiktivních odpovědí není cílené**, ale je založené na **statistickém fungování** LLM. Mezi

¹⁰ Toto téma by určitě stálo za to dále prozkoumat, jestli je v zásadě způsob chybovosti generativních jazykových modelů úmyslný, aby působily důvěryhodněji, a tedy si získali větší oblibu.

respondenty jsou poměrně **velké rozdíly ve znalostech** a vycházejí jak z **oboru a stupně studia**, tak z **osobního zájmu**. Nejsou však přímou úměrou záměru je využívat. Pochybnosti se objevují vzhledem k široké paletě schopností jazykových modelů, již **je nemožné obsáhnout celé toto spektrum** znalostí a opravdu porozumět vnitřním procesům, proto přichází na řadu opět kritické myšlení. Rozpoznání kredibility odpovědí, které LLM dávají, je ale i při používání vlastního rozumu ošemetné. Pokud se nejedná o obor nebo oblast, o které má uživatel ponětí či přehled, tak lze výsledkům velmi snadno důvěřovat. Dalo by se říci, že zde platí přísloví důvěřuj, ale prověřuj.

Přínosem je uvědomění, že klamání není cílené, je pouze založené na statistickém fungování a sledu slov dle nejpravděpodobnější varianty. „*On to bere statisticky, takže říká věci na základě toho, jak to reálně děláme. Když to porovnáme s normálním algoritmem, tak by se tam musely dát jasná deterministická pravidla a on by se rozhodl, ale tady je to nahrazené tím, že to je spíš ve stylu – kde v tom normálním rozdělení se nacházíme a vezme to, co je nejčastější. Navíc aby neodpovídal to samé, tak má množinu věcí, ze kterých si náhodně vybere. Používá různé heuristiky*“ (Marie). Různá hloubka znalostí implikuje, že porozumění není pro efektivní používání zase až tak zásadní. „*Není to byl naprogramovanej falzifikát, ale je to nejpravděpodobnější další možnost, Takže mi to pomohlo pochopit, že mě nechce oklamat, ale že to tak funguje*“ (Josefina).

Nicméně rozdíly mezi jednotlivými studenty v chápání mechanismů fungování jsou poměrně veliké v závislosti na oboru jejich studia, zkušeností z práce a také potřebě. Jejich zájem o podrobnější pochopení AI je omezen specifickými obory nebo osobními zájmy. „*No tak něco málo vím od kámošek a z rodiny, když se o tom baví, ale moc ne, ale asi mi to ani nevadí*“ (Hana). Někteří uživatelé mají základní povědomí o AI, ale nezabývají se technologií hlouběji. Motivace dozvědět se o fungování není přímo úměrná jejich zájmu nástroj využívat. „*Nerozumím úplně těm vnitřnostem, jak to vzniká, jak ta data fungují–Mám málo informací, abych na to mohla mít nějaký fundovanější názor, ale není to můj obor, takže ani můj extra zájem, ale používám ho denně*“ (Naďa). Mnohem důležitější než samotné porozumění, je schopnost kritického myšlení. Společným tématem je nicméně skrze obory i věkové skupiny vědomí nedostatečného porozumění a vzhledu do fungování AI nástrojů, a to i vzhledem k širokým možnostem použití. „*Nějaké informace mám, ale prostě jako pořád ne tolik, protože umělá inteligence je hrozně širokej pojem. Minule jsme tvořili nějaké chatboty pro firmy tak asi vím, na jakým principu funguje NLP, to je nějaký rozpoznávání podobnosti nebo vím, že jak tam funguje doplňování, strojový učení a tak. Mám nějaký*

povědomí, ale i tak mě to přijde jako už hodně komplexní a složitý a vím, že toho hodně nevím“ (Sebastian).

6.3 Způsob používání a zkušenosti AI

V této kapitole se blíže seznámíme s důvody obliby ChatGPT či jiných jazykových modelů, také s intenzitou, s jakou jsou používány a adaptovány do každodenního života a vůbec s šíří jejich použití. V této části práce, kterou jsme rozdělily na dvě větší témata, **používání AI** a **zkušenosti s AI**, zodpovídáme třetí výzkumnou otázku – *Jaká je povaha zkušeností studentů s AI, z jakých oblastí, jak často a jaké nástroje AI používají?* Nejoblíbenější jazykové modely a způsoby, jak se osvědčili při práci, ve škole, při učení se novým věcem, jak usnadňují vyhledávání nových informací, přinášejí inspiraci nebo jsou zábavným společníkem budou předmětem následujícího text. V neposlední řadě analýza podkryje názory na možnost využití v terapeutické práci.

Zkušenost a četnost používání AI

Naši respondenti mají až na výjimku přímou zkušenost s využíváním nástrojů jazykových modelů. Jejich **četnost** a šíře souvisí pravděpodobně i s oborem jejich studia a požadavky, které jsou na ně kladeny. Zařazení do **každodenní rutiny** se objevuje jak u středoškoláků, tak vysokoškoláků. U respondentů jsme se setkali s **ovlivněním aktivity periodickými obdobími**, kdy na ně byly kladeny větší požadavky jako například zkouškové období, termíny odevzdávání seminárních prací, nebo účast na speciálních školeních týkajících se AI. Velká **intenzita** souvisela také s **novostí nástrojů jako ChatGPT** a dalších.

Z našich vybraných respondentů má devět z nich přímou zkušenost s používáním jazykového modelu, jedna respondentka pouze zprostředkovanou. Nulová zkušenost s LLM se se zdá možná jako překvapivé zjištění. Tato (ne)zkušenost se váže na absenci aktuální potřeby nástroj využít. *„Zatím žádnou přímou zkušenost nemám, ve škole po nás nechtěli nic psát, žádnou seminárku nebo referát, takže jsem neměla vlastně moc důvod to použít“ (Hana).* Vzhledem k výběru respondentů, kde jsou ze středních škol zastoupená různá zaměření a typy škol, by se dalo předpokládat, že prvotní motivace k využívání ChatGPT a podobných nástrojů vzniká na základě poptávky, školní či pracovní. Pro některé uživatele je již součástí běžného každodenního fungování. *„Přesně tady ty jazykové prostředky, tito pomocníci a umělá inteligence jsou mým každodenním pomocníkem. Já jsem se s tím dost*

naučila fungovat jako s takovým víceméně partákem, bez kterého si to už ani neumím představit“ (Naďa).

Využitelnost při školních povinnostech nebo v zaměstnání se odráží také v četnosti využívání u dalších respondentů. ChatGPT se dle těchto zkušeností profiluje jako nástroj pomoci při práci, pro její usnadnění. *„Používám ho, když si potřebuji něco ověřit nebo rychleji zjistit. Je to nárazové, třeba jednou za týden nebo za dva. Ted' jsme měli školení v práci, a to jsme mu celý den zkoušeli něco zadávat. Pak je zase období, kdy není nic do školy, nejsou žádné povinnosti, a to si na něj měsíc nevzpomenu“ (Josefina).*

Četnost používání se u jednotlivých respondentů liší v závislosti na jejich každodenních potřebách a také zaměření studia a práce. Není překvapivé, že počáteční zájem byl podpořen zvědavostí a seznámení se s tímto nástrojem, který byl přelomovým nástrojem. *„Když to bylo úplně nový, když vydali tu trojku, tak jsem to během prvních týdnů používal ze zvědavosti fakt intenzivně“ (Šimon).*

Druhy používaných jazykových modelů

Respondenti používají jako základní jazyková model ChatGPT. Další oblíbené LLM a jejich použití většinou souvisí s jejich vydáním či změnou (**Gemini**), požadavky v práci nebo zájem vyzkoušet kromě ChatGPT něco dalšího jako **Perplexity**, **Bing Chat**, **Deepl** nebo **My AI**, chatovací bot Snapchatu anebo AI na tvorbu obrázků **Midjourney**. Nejsnadněji dostupná verze ChatGPT-3.5 je neoblíbenější, je neplacená a uživatelé mnohem častěji dostávají odpovědi, dalo by se říci ve stylu lidové moudrosti. Vyšší placená verze ChatGPT-4 nabízí mnohem více možností, jako je instalace speciálních pluginů a přepínání na specificky zaměřené jazykové modely, a kromě textu také generování obrázků, ale ne vždy uživatele uspokojí jako nižší neplacená verze.

Všichni respondenti, kteří mají zkušenost s jazykovými modely, používají nejčastěji ChatGPT ve verzi 3.5. Jako výhodu vidí zejména v jeho prvenství ve zveřejnění. *„Chat GPT mě obrovský štěstí, že to jako první zveřejnili. I když jsou ted' i další vylepšený jazykový modely. No a ta poslední placená verze ChatGPT tam má ty přidané možnosti, takže už to není jenom ten jazykový model, tak to je taky výhoda“ (Sebastian).* Mezi další jim známé jazykové modely patří Perplexity, Bing Chat, Gemini (dříve Google Bard), Deepl nebo chatovací nástroj Snapchatu My AI. *„S vydáním Gemini jsem se o něj začal víc zajímat, vyzkoušel, koukal na videa, ale ta moje největší zkušenost je se ChatGPT a Snapchatem“ (Karel).* S jinými jazykovými modely mají často jen okrajové zkušenosti. *„Kromě ChatGPT*

jsem vyzkoušel Perplexity nebo používám Deepl, no a pak taky na obrázky Midjourney“ (Adam). Mezi respondenty se objevuje i zkušenost s placenou verzí Chat GPT-4, jejíž funkce mnohonásobně převyšují otevřenou verzi ChatGPT-3.5, nicméně uživatelská zkušenost vidí vylepšený model i z jiného úhlu. Oblíbenost verzí jazykových modelů nemusí jít nutně logikou, čím novější, tím lepší. Verze ChatGPT 3.5 je jednodušší ve svých funkcích. *„Překvapilo mne, že mi někdy ta nižší verze dává lepší odpovědi. Že mi třeba odpovídá víc takovým selským rozumem, kdežto ta čtyřka mi napíše, že to neumí nebo nemůže napsat. Ta hloupější verze je prostě taková odvážnější, a to je mi sympatičtější“* (Marek). Tlak na používání jiných jazykových modelů, než ChatGPT přichází také zvenčí v souvislosti s ochranou firemních dat. *„No v práci totiž teď nám přišla hláška o tom, že nemůžeme používat ChatGPT, že tam dávat citlivá data a že máme využívat Bing Chat. Ale s tím mám teda zkušenost minimální“* (Josefina).

Analýza odpovědí respondentů ukazuje velice komplexní odraz reálného používání AI. ChatGPT se během krátké doby od svého uvedení, což je o něco více než jeden rok, stal běžným nástrojem pokrývajícím široké spektrum aktivit. V této kapitole bychom je mohli rozdělit na několik dílčích témat: **efektivita a úspora času, přístup k informacím a vzdělávací nástroj, nástroj inspirace a zábava, programování**. Výstupy a použití nejsou nutně pouze v textové podobě.

Efektivita a úspora času při práci a učení

Jazykové modely nabízejí širokou škálu uplatnění a jedním z nich je **zefektivnění práce a učení**, kam spadá psaní emailů, pomoc se strukturou textu, formulacemi. Pokud jej používají na kreativní část práce, **vytvoření textu**, slouží tento jako **inspirace**, která je přepsána do vlastní podoby. Někdy je výsledný text jazykových modelů příliš květnatý, **příliš se drží zadání**. Proto je dobré se s ním **naučit pracovat** a vědět, jak se ho **správně ptát**.

Počínaje u každodenních jednoduchých komunikačních úkolů jako je psaní emailů, nebo pomoc s jejich strukturou a formulací. *„Párkrát jsem s ní něco konzultoval, a i když jsem chtěl něco někde napsat, třeba i text wireframu. Takže jsem to tam napsal svejma humpoláckýma slovama a řekl jsem ať to prostě přeformuluje tak, aby to bylo dobrý. A použil jsem to a vlastně to většinou stojí za to, jakože to je dobrý. Prostě když nechceš dlouhej čas trávit nad každým slovíčkem, tak jí řekneš, přeformuluj to ono tě to překvapuje“* (Šimon). Často jej respondenti, spíše než k přímému kopírování celého obsahu, využívají jako

konzultační nástroj, využívají zmiňované struktury textu. „*Když potřebuji vytvořit nějaký email a nechce se mi hodinu vymýšlet tu strukturu, tak prostě už si jednoduše poprosím o strukturu AI je to na světě a jenom si to přepíšu do svých slov*“ (Naďa). Při vytváření výsledného produktu, například emailu, jej přepisují do svých slov. Jedním z důvodů je i praktické použití. Jazyk ChatGPT či jiných modelů je přece jen někdy příliš věrný požadavkům a vyžadovanému stylu. Dalším z důvodů je využívání AI jako nástroje spolupráce, kde je člověk tím, kdo rozhodne o konečné podobě textu. Je efektivním nástrojem shrnutí obsahu textů, ač i tady jsou jistá omezení. Nespokojenost uživatelů je patrná v případě zpracování textu a porozumění některým jazykovým obrátům. Na povrch tak vyplouvá nutnost správného zacházení s ChatGPT, i při běžných úkolech můžete dostat lepších výsledků, pokud víte, jak se správně ptát. „*Používám ho, když je hodně dlouhý text. Já potřebuju jenom zjistit, o čem to tak cca je. Dokonce jsem to používala i v rámci uživatelských rozhovorů, který jsem měla v práci, produktáka. Vzala jsem ty rozhovory a chtěla jsem po ChatGPT, aby mi řeklo, který témata byly jako nejčastější. Ale protože je to jazykový model, tak nezvládá tak úplně dobře synonyma. Ale člověk se s tím musí naučit zacházet, vlastně jako je to čím dál tím zajímavější*“ (Marie).

Přístup k informacím a vzdělávací nástroj

Další velkou podskupinou využití je kromě komunikace také samotné vzdělávání. Využití se týká zejména těchto oblastí, tvorby **textů slohových prací**, **inspirace** při psaní, **shrnutí obsahu** dokumentu, **vyhledávání zdrojů** ale i **vysvětlování učiva** jako tutor.

V zájmu je ChatGPT jako primární nástroj, které pomáhá s tvorbou textů a slohových prací. „*Pomáhá mi třeba se školou, nebo když mě naopak něco samotnýho zajímá. Ona jako šílená věc je i to, že ví a umí v podstatě úplně o všem. Máme psát slohovku, tak jsem si řekl, porad' mi, jak bych mohl pokračovat. Nejdřív jsem mu tam napsal ten svůj text a ono mi ho dokončilo a fakt to vypadalo i dobře*“ (Marek).

Výčet použití ChatGPT by nebyl úplný a opravdový, pokud bychom neuvedli také tvorbu plného znění textů, například seminárních prací. Byť ještě procházejí překladem do jiného jazyka, tedy také vyžadují lidskou invenci. „*Když třeba teď mám úkol na dějepis, tak to použiju. Prostě nenechám, ať mi to celý napíše. Jenom si to přeložím do češtiny a tím, jaké si to přeložím do češtiny, do vlastních slov, tak tím pádem ani jako není dohledatelný, že napsal ChatGPT*“ (Naďa).

Zkušenosti středoškolských studentů doplňuje i Sebastian, pracující vysokoškolský student se, zkušenostmi se zpracováním odborných textů. „*Když potřebuju něco pro odborný text, tak najdu nějaký článek a nechám ho to třeba přepsat trošku vlastními slovy, ale aby to dávalo smysl*“ (Sebastian). Stejně oblíbený je také na vyhledávání zdrojů, prozkoumávání aktuálního výzkumu na určená témata. „*Používal jsem ji odborně, když jsem si chtěl získat zdroje ke své bakalářské práci, chtěl jsem si ji nějakým způsobem otestovat, jestliže my najde víc zdrojů, než jsem si našel sám*“ (Šimon). A možná překvapivě, slouží také jako tutor k samotnému učení a vysvětlování probírané látky. Možnost přímé interakce a doptávání se na detaily učiva, kterým respondenti nerozumí, se ukazují jako velice efektivní. „*Na co je to dobrý, je ta škola. Třeba ve fyzice, kde pokulhávám, tak tam mi pomalu dokáže vysvětlit látku líp jak učitel, takže jako to je, to je taky super*“ (Karel).

Inspirace a zábava

Svoje pevné místo získávají ChatGPT a další LLM jako nástroje vyhledávání informací, které nemusí být nutně školního či pracovního charakteru. Často jsou využívány jako **alternativa k vyhledávačům**. Jejich výhodou je nejen formulace odpovědí, ale i **dohledávání informací**, které nelze vygooglovat. Příkladem jsou **specifické dotazy** a využívání jako **zdroj zábavy**.

Víckrát byl zmiňovaný jako nástroj první volby ve chvílích, kdy by se donedávna respondenti obrátili na Google nebo Wikipedii. „*Můžeš to použít třeba místo Googlu, když tě něco zajímá.*“ (Hana) A tam kde jsou vyhledávače nedostatečné, má dnes své pevné místo LLM. „*Potom taky určitě další zásadní věc, my se bavíme o hrách a kolikrát ten program o těch hrách ví víc jak my, takže se dá říct, že ho používám často a na různé činnosti*“ (Karel). Povaha dotazů může mít kromě zjišťování informací, které souvisí s naším výkonem, také formu speciálních dotazů pro inspiraci, které jsou které by možná dříve také byly úkolem vyhledávačů. „*Já se ptám někdy na takový divný dotazy, který mi nejdou zformulovat do Googlu a dost často mi to zodpoví, což je super. Někdy to jsou takové otázky, který, víš, že by ti Google řekl blbě. Takže tam mám takový svoje občas otázky, třeba aby vymyslel vtip. Protože já jsem já jsem potřebovala specifický vtip, protože přítel hraje tenis a já jsem jednou řekla, že se tenis hraje s pálkama a jeho to štvě hrozně, že to nejsou páčky. A já jsem potřebovala specifický vtip o tom, že se tedy hraje s pálkama, tak jsem tam zadala takhle AI prostě vyhodilo mi to 10 vtipů o tenise s pálkama. Nebo jsem hledala jeden obraz a popisovala ho a chtěla jsem, aby mi to řeklo autora, a to by se mi taky blbě dávalo do Googlu*“ (Marie).

Práce s AI je rozhodně i formou zábavy. I na začátku jeho uvedení se v médiích hojně objevovaly příklady komunikace s AI, způsoby, jak ji ošálit, co všechno umí, jak dokáže být vtipná, co říká a jak reaguje na určité podněty. „*Kromě toho, že mi pomáhá s chápáním Excelu je to samozřejmě i jako nástroj prokrastinace. Myslím si, že když jsem to poprvé zaregistrovala, tak jsem si s tím hodně hrála, jakože prostě na nudu*“ (Josefina).

Nástroj kódování a programování

Využívání pro **generování kódů** při programování si zaslouží vlastní podtéma. Jedná se totiž o **revoluční nástroj**, který dokáže do nedávna nemyslitelné. Tato jeho schopnost **ovlivňuje** totiž ve velké míře také **bezpečnost**, a to se pak ve výsledku může týkat každého z nás, zejména v digitální podobě světa, ve kterém žijeme, kde jsou naše osobní a citlivá data většinou online dostupná. Některé platformy LLM nabízejí i možnost **vytváření vlastních jazykových modelů** v rámci těch existujících, většinou ve zpoplatněných verzích. Respondenti, kteří se pohybují v oblasti informatiky a podobných oborů se shodují, že je to ohromná pomoc a výhoda, která, kromě ulehčení práce, otevírá další možnosti kódování, byť i tady je nutná znalost, vzdělání a kritické myšlení. Respondentka Marie popisuje vlastní zkušenost a údiv. „*Ted’ mě hrozně fascinovalo, když jsem byla v úterý na semináři o bezpečnosti sítí, o kyberbezpečnosti a pánové to tam používali na to, že dokonce jako z binárky tahali původní kód, což je pro mě jako někoho, kdo se zabývá bezpečností mazec, že to jde. Protože s tímhle jsem strávila šest měsíců na škole, abych se to naučila dělat*“ (Marie). ChatGPT-4 nabízí také možnosti vytvoření vlastní jazykových modelů, které se dají trénovat. „*My jsme dávali ChatGPT, udělali jsme svůj model a pak jsme ho učili některé situace a dali jsme ho do aplikace a pak jsme ho používali na support, abychom ušetřili čas*“ (Marie). Výhoda při generování scriptů, kódů a je nepopíratelná. ChatGPT zná různé kódovací jazyky. „*V práci jsme to používali na kódování. Konkrétně ono je to super, jak generuje skripty jako programovací skripty, takže když já jsem třeba potřeboval nějaký script, tak jsem říkal, potřebuju v Pythonu napsat to a to a ono to vygenerovalo. A třeba to jako mělo chyby, ale dalo se tam dostat*“ (Sebastian).

AI v terapii

Možnost využití v terapii je často diskutovaným tématem. Kromě samotného zamyšlení nad způsobem využití AI nástrojů, se ale objevuje také **důraz na mezilidské vztahy**. Lidský kontakt a **autenticita** jsou pro terapii důležitými složkami. Respondenti je považují za klíčové pro vytvoření funkčního a efektivního terapeutického vztahu. Může AI

plnit funkci terapeuta a pokud ano, měla by být přiznaná nebo nikoliv, i nad tím se zamýšlí. Zároveň by AI mohla být zejména v **akutních situacích** funkčním řešením nedostatku lidské síly. Objevuje se problematika antropomorfismu. Při dlouhodobé terapii je ale možné vidět **technologie jako bariéru**. Možné negativní vlivy technologií na náš život přináší zamyšlení nad přesunem komunikace a našich životů do virtuálního světa.

AI by mohla být dostatečnou podporou v určitých situacích a oblastech, zejména v akutních situacích, jako první pomoc v krizi nebo při nedostatku terapeutů. „*Při nějakých problémech to může pomoci. Ještě když je péče nedostupná, dlouhá čekací doba, nebo máš nějaký akutní problém a tohle můžeš mít úplně na dosah třeba v mobilu*“ (Josefina).

Druhým úhlem pohledu, který doplňuje možnou roli AI v terapeutickém světě či v krizových situacích, je důraz na podstatu mezilidských vztahů. Ty hrají klíčovou roli při vytváření efektivního terapeutického prostředí. „*Já věřím že ten léčivý vztah je mezi dvěma bytostmi. Zároveň vím, že se lidé na seznamkách zamilovali do umělé inteligence. Takže to asi jde*“ (Pavel). Ač je to spíše zbožné přání, protože skutečnost a zkušenosti dokazují opak. „*Myslím, že to může pomoci na tu nejnutnější krizovku třeba při nějaký panický atace nebo v nějakých jasně ohraničených psychických problémech, kde se dá poradit třeba v KBT stylu. To si myslím, že se dá a že vlastně těch terapeutů prostě je pořád nedostatek a myslím si, že někomu může pomoci, když to bude dobře udělané, vlastně si promluvit jenom s tím Chatem*“ (Sebastian).

Využití AI v dlouhodobé terapii, i ve smyslu vytvoření blízkého terapeutického vztahu mezi klientem a terapeutem, je ve své podstatě v rozporu s používáním umělé inteligence, ovšem záleží, jakou by měla podobu. Vracíme se tak k antropomorfismu. Technologie je vnímána jako překážka, ovšem pokud by vypadala jako člověk a uměla tak jednat, tak by to potenciálně mohlo fungovat dobře. „*Nemyslím si, že by to mohlo nahradit úplně terapii, je tady ta technologická bariéra. Vlastně může přispívat i k nějakým duševní nepohodě. a jak víme i z výzkumu a tak dále, že jo, tak nejvíc, co funguje v terapii prostě je vztah mezi dvěma lidmi, přijetí a tak dále. Tudiž myslím si, že takovou tu hlavní terapii to jako nahradit nemůže a myslím si, že je to naopak jako jedna z mála oblastí, že se tam umělá inteligence úplně nedostane, zatím podle toho, co víme teď o terapii. Na druhou stranu, když si představím, že by tam vlastně byl někdo, kdo by vypadal jako člověk a mluvil jako člověk z obrazovky a reagoval na tebe, tak by to asi šlo*“ (Šimon).

Jedna z respondentek, studentka vysoké školy, popsala svoje zkušenosti z pracovního prostředí, kde je v roli produktového manažera aplikace pro psychickou pohodu. „*Naučili jsme ho, jak má reagovat na některé věci, že jsme mu řekli například, pokud člověk chce spáchat sebevraždu tak reaguj tak a tak. Naučili jsme ho nějaké situace, aby fungoval jako první kontakt s uživatelem*“ (Marie). Důležitým momentem ale byla možnost propojit se se skutečným terapeutem, lidskou bytostí. Možnost výběru a zhodnocení aktuální situace a její akutnosti vypovídá o postoji jednotlivce vůči technologiím a způsobům, jak může pomoci. „*Důležité bylo upozornění, že pokud stojí o kontakt s živým člověkem, musejí třeba den čekat, než bude mít volnou kapacitu. Někdy stačila uživatelům pouze rada, co si počít, když mají úzkosti*“ (Marie).

6.4 Vzdělávání o AI

Tato kapitola odpovídá na čtvrtou výzkumnou otázku – *Jaká jsou očekávání respondentů týkající se vzdělávání o AI ve školách a mimo ně?* Nutnost vzdělávání v oblasti používání AI a uvědomění si jeho možností, omezení, pochopení principů fungování AI, by mělo být součástí vzdělávacích programů. Jeho nedílnou součástí je také kritické myšlení, které umožňuje zdravý úsudek a kritické hodnocení výstupů AI. Díky bližšímu prozkoumání těchto témat, se objevují další otázky a témata, která otevírají bohatou diskusi o směřování a formě vzdělávání o AI. Zajímalo nás, jak respondenti vidí potřebu vzdělání o umělé inteligenci, jak by mělo probíhat, ale i na jaké úrovni školství by již mělo být běžnou součástí osnov. Zároveň jsme chtěli zmapovat jejich vlastní situaci, kde si stojí jejich škola a zdali vůbec zachytili snahu o nějakou aktivitu. V této kapitole nelimitujeme naše chápání na jazykové modely, ale chápeme AI i v širším pojetí.

Mezery ve vzdělanosti

Nová technologie s sebou nese nové poznatky a vznikají **mezery ve vzdělání**. Pružná reakce školních systémů je jednou z odpovědí, ale neměli bychom zapomenout ani na starší generace. Do jejich životů umělá inteligence rovněž zasahuje a **možnosti získávání informací** jsou poměrně **omezené**. I sami studenti cítí **potřebu vzdělávání** zejména v oblasti původu dat, která jsou nahrána v LLM. Jedním z důsledků absence vzdělání v oblasti AI je **přebírání rizikových informací**. Školní systémy mají velké nedostatky a mezery. Často se jeví jako **nepružné**, jejich změna a odraz aktuálních trendů trvá často roky. Stejně jako by mělo být součástí vzdělávání o používání internetu, mělo by být součástí i práce s AI, které by řešilo **neaktuálnost školních osnov**. Vzdělání by mělo být

přízpůsobené věkové kategorii a potřebám studentů dle jednotlivých oborů a jejich využitelnost v praxi.

Samotní studenti cítí, že by se měli více vzdělávat. I přesto, že někteří si myslí, že ví dokonce více než průměrná populace. Možná právě proto dokážou nahlédnout do dalších oblastí a je jim jasné, že i jejich znalosti stále pokulhávají. Nenaplněným zájem je původ dat. *„Trochu si říkám, že bych se sama mohla vzdělávat, ačkoli si připadám, že když se porovná s průměrnou populací, tak toho vím víc. Ale říkám si, že bych se mohla víc vzdělávat o tom kde ty data berou, nebo jestli už je to někde nějak jako právně ukotvený a jako regulovaný“* (Marie). Důsledky nevzdělanosti s sebou nesou riziko přijímání nebezpečných nebo zavádějících informací. Rozhodně by informace o AI neměly být limitovány typem škol, ale měla by být zahrnutá učiliště, a to zejména v kontextu kritického myšlení. ***„Mělo by se o AI určitě vzdělávat i na školách, který nejsou gymnázia a střední školy, ale i ty jiné formy, jako jsou třeba učiliště, prostě i v souvislosti s kritickým myšlením. Protože věřím a nechci je odsoudit, ale třeba děti, který nemají na nic víc než učňák, takže lehce všemu naskoče, protože to vůbec vlastně neřeší, že by to třeba nemusela být pravda a pak vlastně se z nich stávají naivní jedinci, který jsou obětí umělé inteligence a třeba dezinformací“*** (Naďa).

Starší generace je celkově v nevýhodné pozici v přístupu a v porozumění novým technologiím. Bariérou pro motivaci občanů, ale může být cena takových školení. Vzdělávání o AI by mělo být dostupné. *„Určitě si myslím, že by mělo být školení dostupné. Nedávno jsem viděla, že Česká hospodářská komora vypsala školení v Brně o tom, jak zacházet s AI a to školení stálo asi 4500 Kč a myslím si, že by to mělo být automaticky. Prostě, že by třeba mohl být nějaký veřejný seminář. Že to jakoby povědomí je a není, a když vlastně se nabízí, tak jako draze zpoplatněný, to mi nepřijde úplně dobře nastavený“* (Josefina). Na základních a středních školách je téměř až zoufalá situace s ohledem na aktuálnost učiva. Nabalující se problémy, psychické obtíže dětí, jsou možným důsledkem. *„Celkově chybí ve školách předmět, který řeší aktuality, že ani ne jako dějiny, ale to, co člověk reálně potřebuje. Od daní, finanční gramotnost, umělou inteligenci, práci na sociálních sítích a pak je vidět, jak všechny problémy narůstají. I třeba poslední výzkumy s tím, jak děti trpí depresivitou. Myslím, že to je přesně jeden z důvodů, že vůbec nemáme jako práci s novými technologiemi zaopatřenou v rámci studia“* (Sebastian). Kritické myšlení je nedílnou součástí a jednou z nejdůležitějších složek vzdělávání o AI. Vzdělávání by nemělo opomínat etický rozměr. *„Myslím, že by se k tomu každý stupeň vzdělání postavil, už třeba od druhého vstupně třeba nějakým způsobem. Vysvětlovat klidně i nějak“*

zjednodušeně ten princip, jak to jako funguje a k čemu se to dá využít, a hlavně a primárně prostě opakovat těm dětem etický zásady. To si je strašně důležitý“ (Šimon). Většinový názor, na kterém se respondenti shodují, je zavedení do školních osnov již na základní škole a určitě i na střední. „Určitě nějaký semináře na střední škole, roli hraje už i povědomí děcek o AI. Bylo by fajn mít třeba jednou týdně hodinu s nějakým specialistou, který bude vysvětlovat na co to je, jak s tím zacházet a pracovat do budoucna, jak to co nejlíp využít, i třeba v povolání, protože už teď na tom lidi dost vydělávají“ (Karel).

AI má široký záběr, a i ve vzdělání by měl být zohledněn studijní oborem. „Na oborových školách jsou jednotlivé obory a ty to můžou využívat různě. Mohlo by být fajn to oborově zaměřit a využít v praxi“ (Marek).

Jak by se měly vzdělávací instituce postavit k používání AI nástrojů

Morální dilema používání nástrojů umělé inteligence ve výuce a vypracovávání úkolů má snadné řešení. Jasným signálem je **nikoliv zakazovat** jejich používání, ale **začlenit** je adekvátním způsobem jako pracovní a „legální“ nástroje ve vzdělání. Je nutné si uvědomit, že efektivní práce s těmito nástroji je **hodnotou sama o sobě**. Vývoj se zastavit nedá, ani technologický pokrok. Můžeme si položit otázku, zda není tato schopnost dokonce důležitější než primární požadavek týkající se obsahu úkolu. Pozadu by nemělo zůstat ani **vzdělání pedagogických pracovníků**, kteří by měli mít adekvátní kvalifikaci.

Respondenti beze studu odrývají své názory na nesmyslný přístup některých institucí. Některé totiž dosud ani nereflektovali existenci internetu a nové povahy médií, nestačí držet krok. „Myslím si, že by bylo dobrý, aby lidi byli vzdělávaný o tom, co umělá inteligence dokáže a kriticky vyhodnocovali ty věci, které vidí, ale bohužel, jako to se mělo stát i s médii a víme, jak to dopadlo. Takže to moc nadějně nevidím, ale mělo by se to dít“ (Marie). Nalezení způsobů spolupráce s AI, při zachování stejného standardu výuky, by mělo být jedním ze základních požadavků. „Hodně se mluví o podvádění při psaní seminárek. Tak spíš aby byly existující cesty, jak to správně využít a nesnižovat jako nároky na tu práci těch studentů“ (Josefina). Nezbytnou součástí adekvátního vzdělávání na školách, je tak vzdělanost samotných učitelů, která může být další faktorem a zdržením v implementaci takových programů. „Přijde mi správný směr o tom vzdělávat, ne zakazovat. Jenom je to asi složitý, celkově nějaká transformace našeho „pružného“ školství. Museli by o tom pořádně vědět samotní učitelé a to je taky proces. Snad se to bude časem dít“ (Pavel).

Jak se aktuálně staví studenti k používání AI nástrojů

Nejasná situace na školách a často nevyřešené možnosti či zákazy používání LLM při školních a akademických povinnostech nechávají sice studenty v lehkém napětí, ale i přes výslovný zákaz by je používali. „*Kdyby nám to úplně zakázali, tak bych se asi stejně ptal na nějaký věci, kterým nerozumím. Pořád si myslím, že je lepší si ty věci udělat sám, aby proces vyučování neztrácel smysl. Ale jako jo, když je něco, u čeho nevím, mu napíšu ať mi to vypočte, ale snažím se zeptat, i jak to udělal a proč to dělá takhle, aby mi převyprávěl ten proces a já se to naučil*“ (Karel).

Jako nástroj inspirace si LLM už nechtějí nechat vzít. „*Ne jako že bych si myslela, že to napíšu nějak výrazně líp, ale spíš že mne to baví. Ale no ale jako bych se nebála toho zákazu, to klidně bych to na inspiraci používala*“ (Josefina).

6.5 Obavy a hrozby z AI

Obavy z umělé inteligence se pojí s pokrokem, a stejně jako u jiných technologií, také s rozporem mezi tím, co přináší dobrého a jaká s sebou nese rizika. „*Úspěch ve vytvoření AI by mohl být největší událostí v historii naší civilizace. Na druhou stranu by mohla být také poslední, pokud se nenaučíme, jak předcházet rizikům*“ (Stephen King). Citát Stephena Hawkinga může být více než trefný. V této kapitole se zaměřujeme na zodpovězení páté výzkumné otázky – *Jaké jsou nejčastější obavy a hrozby studentů týkající se rozvoje a používání AI v současnosti a v budoucnu?*

Nástroje AI mají potenciál zlepšit naši společnost k lepšímu, zároveň s sebou nesou hrozby, zejména v případě nedostatečné kontroly, regulace a zodpovědnosti při jejím používání. Obavy a hrozby pro potřeby této práce spojíme do jedné kapitoly. Obavy o **ztrátu vlastního rozumu, nedostatek kritického myšlení** ve společnosti, které může vést ke **snadné manipulaci, zneužití dat**, AI jako **nezastavitelný pokrok**, který nemáme pod kontrolou, společně s hrozbou vytvoření **vlastního vědomí a deteriorace mezilidských vztahů**, přivádí respondenty na různé scénáře budoucnosti i s ohledem na jejich individuální směřování. Možné zneužití LLM jako nástroje diskriminace nebo spojení AI se zbraněmi, už se týká problémů s dalekosáhlým dosahem, globálního měřítka. Problémem při vyhodnocování situací na základě fungování jazykových modelů je právě statisticky nejpravděpodobnější možnost při rozhodování. Ukazuje se, že při simulaci potenciálního vývoje, například v mezinárodní oblasti, je řešením jazykových modelů **agrese a válečný konflikt**.

Prvotní obavy z jejího používání plynou z omezení schopnosti lidí samostatně přemýšlet. „*Když ji někdo bude používat moc, tak pak už nebudou moct třeba samostatně přemýšlet*“ (Hana). Běžnou obavou je také zneužití dat. Pro používání je nutné se přihlásit pod emailem a tvoří se historie našich dotazů. „*Mám obavu ze zneužití dat, ale zároveň mne to neodrazuje od používání a mít to prostě v počítači pod svým jménem*“ (Josefina). Obavy se vztahují k také budoucnosti a její podobě, v případě, že nás AI v rozumových schopnostech. „*Teda trochu se bojím, že to je věc, která se valí. Nikdo si nestihá uvědomovat následky. Jestli přestaneme kompletně používat náš vlastní rozum a dáme kompletní oddanost AI přístroji, tak nevím, jestli budu chtít do tohoto světa vůbec přivést jako děti. Trochu se bojím, že prostě nevím, co za tím stojí, a aby to nebylo tak chytrý, že už by si to umělo udělat vlastní vědomí*“ (Naďa). Narušení autentických vztahů, deteriorace kvality mezilidských vztahů nebo jejich vznik směrem k AI, se také řadí mezi obavy studentů. „*Vidím jako hrozbu, že bysme si mohli myslet, že to může nahradit vztahy. Takže jako potenciálně jako hrozba je, že by někdo vytvořil vazbu, dávat by tam nějakou energii, která by se mu pak vlastně nevracela, z toho by mohla přijít nějaká frustrace*“ (Josefina). Nechybí ani obavy z diskriminačního chování. V minulosti se při používání nástrojů umělé inteligence ve výběrovém procesu zaměstnanců objevila diskriminace žen. „*Mohlo by to ubližovat nějakým skupinám, mohla by se rozšířit na základě jazykových modelů diskriminace. Když třeba Amazon nějak chtěl pomoci umělé inteligence vybírat zaměstnance, tak zjistili, že je to strašně diskriminační nástroj. Kdy měla nějaký data a z těch dat to vyhodnotilo například, že není výhodný přijmout ženu*“ (Marie). Velké obavy, hrozby s globální dopadem se týkají spojení umělé inteligence s dalšími technologiemi jako jsou zbraně. Válečné konflikty současnosti zbrojení a jejich neustálému vývoji nahrávají. „*Ta hlavní obava je spojení autonomních zbraní a umělé inteligence¹¹, protože myslím že to je tak dva, tři roky zpátky se to vlastně řešilo, že autonomní zbraně by neměli podléhat umělé inteligenci, aby se prostě nestalo, že se AI zblázní a odpálí třeba jadernou bombu*“ (Sebastian). S tím souvisí i obavy s ovlivňováním umělé inteligence, špatný gatekeeping a role ve strategickém řízení. „*Vnímám to jako hrozbu, když prostě se umělá inteligence dostane do nějakého strategického řízení. A budou ji mít v rukou vlastně špatný lidi*“ (Marek). Opatrně se bude muset přistupovat k rozhodování na základě minulých zkušeností

¹¹ „*Ono to bylo nějak v souvislosti kauzou, kdy Zuckerberg spustil dvě umělé inteligence proti sobě a oni si pak začali jako vyvíjet svůj vlastní jazyk a už to vlastně jako začalo být docela jako děsivý, takže to museli zastavit a na základě i tohoto myslím, bylo to zastavení, nebo jednalo se o zastavení autonomních zbraní.*“ (Sebastian)

lidstva. „*Tak ona si přečetla spousta historických knížek a má data ze Studené války a jiných konfliktů a přišlo jí to jako funkční řešení, takže to převzala*“ (Šimon).

6.6 Etika a regulace

Jaké etické otázky respondenti identifikují v souvislosti s AI? – je poslední výzkumnou otázkou, na kterou hledáme odpověď. Etické otázky spojené s používáním umělé inteligence jsou možná ty nejpálčivější¹². **Autorská práva a copyright, používání v akademickém či pracovním prostředí, potřeba větší regulace a transparentnosti, jak AI ovlivňuje společenskou i individuální zodpovědnost**, je jen výběrem témat, kterými se respondenti v souvislosti s etickou odpovědí na ně zabývají.

Etická a bezpečnostních opatření a regulace AI, ať už pozitivní nebo negativní, jsou velkým tématem, kterým je nutné se bezpodmínečně zabývat, aby výhody které s sebou AI přináší, převládaly nad hrozbami. Na otázky etiky se dá nahlížet z mnoha úhlů. Nejen jako na **používání nástrojů AI** v oblastech zmiňovaných v kapitole 7.4., jako je zefektivnění studia či práce a etického chování v jejich rámci, ale celkového zacházení při tvorbě obsahu vůči dalším subjektům (autorská práva). Patří sem i otázky o **povaze etických pravidel**, jejich vynucování a zacházení s nimi. Nelze opomenout ani problematiku **AI v nesprávných rukách**, tedy jejího **zneužívání** jako je šíření dezinformací, politický nástroj, nástroj manipulace, ale i obavy zneužití ve válečných konfliktech s opravdu závažnými důsledky.

Je otázkou, zda je používání AI etické samo o sobě, zejména s ohledem na duševní vlastnictví, vizuální a literární díla „*Mně k tomu napadá, jestli nejsou zneužitý autorský práva. Zaznamenala jsem to s téma obrázkama. AI nástroje něco generují a bylo by mi líto, kdyby umělci, ať už jako vlastně nějaký literární nebo ty výtvarný, nebo jakýkoliv fotografický trátili. Snažila bych se je nějak podpořit v tom, aby jim to nebylo kradený. Prostě no, ale to už se asi možná děje, to vidím jako etický problém*“ (Josefina).

V souvislosti s autorskou tvorbou a AI je problémem nejen samotné používání autorských děl a nalezení hranice mezi inspirací a krádeží, ale zejména také etický problém obohacování se na úkor umělců. „*Problém umělé inteligenci je, že spoustu lidí na ní hrozně*

¹² Poznámka k zamyšlení: Jednou z oblastí jsou omyly AI, který se při strojovém učení snadno mohou přihodit. Příkladem je zkušenost z přepisu audio rozhovorů do textové podoby. V jednom z rozhovorů zaznělo vcelku srozumitelně jaderná bomba, nicméně tato dvě slova byla přepsána jako „internet.“

jako benefituje, čerpá z toho spoustu peněz, ale reálně je umělá inteligence jako taková krádež, protože nahráli spoustu jako známých děl, který jsou veřejně dostupný a z těch děl se vlastně jako tvoří hodnoty“ (Sebastian). Jako problém se jeví i otázka regulace a nedostatek informací, se kterými se potýkají i respondenti, kteří mají umělé inteligenci dobré povědomí. Umělá inteligence může působit jako černá skříňka, do které nevidíme. Otázkou je pak i kdo nese odpovědnost a kdo rozhoduje o její podobě. „Ty jazykové modely používají prostě všichni, celá zeměkoule víceméně. Berou to z těch stejných zdrojů. No ale prostě mám pocit, že některý ty etické otázky, které se týkají i obsahu, diskriminaci nebo rasismu například, že o výstupech nakonec rozhoduje jenom nějaká jako malá parta lidí, prostě jako Microsoft nebo někde, která určuje vlastně co je dobrý a co je špatný. To mi přijde jako vlastně zvláštní a vlastně neetický“ (Šimon).

Problémem může být nedostatečné ošetření a regulace v oblastech, které AI předcházely. Úhel pohledu na jednotlivé problémy s sebou může nést řadu dilemat, jak tyto situace řešit a zda je ponechat v gesci příslušných oborů. „Kdybych to vztáhnul dětskou pornografií, jestli to vlastně, když už teďka máme nástroj na generování dětské pornografie, jestli vlastně je to dobře nebo špatně, protože například v oblasti sexuologie to může být přínosem, jako kde se mohou devianti realizovat v rámci zákona. Ale vlastně to je jako nezákonný, ale tímhle to najednou umožňuje“ (Sebastian). Jednou z palčivých oblastí jsou i humanitní obory na univerzitách, které v otázkách regulace a nastavování etických pravidel, za technickými obory zaostávají. A to i přesto, jak jsou pro společnost důležité. „My vyloženě používáme umělou inteligenci a teď i rušej bakalářské práce na některých fakultách. A je tam i tendence to používat. Je tam jako práce s firmami, aby vlastně ten člověk šel i do té firmy a zkusil si to. Jsou dva různé světy a myslím si, že to je pak přesně ten problém, kterej je i ve společnosti, že spoustu jako i odborníků z těch jako technicejch oborů nadává na ty humanitní obory, protože humanitní obory prostě vůbec nedržej krok. Tohle je jeden z těch důvodů, kterej si myslím, že pak hraje roli. Přitom je pak vidět, že ty humanitní obory jsou jako strašně moc důležitý pro tu společnost, ale oni nějak jako vůbec to nezvládaj“ (Sebastian).

Etických otázek se týká i obecná pravidla, kterými bychom se měli řídit, při vytváření regulací, nicméně jejich legislativní ukotvení bude velice těžké i vzhledem k současnému stavu společnosti, kdy může být liberální nastavení problémem. „Myslím, že nakonec jako možná to univerzální etický pravidlo, že jako nedělej, co nechceš, aby ubližovalo ostatním nebo co bys nechtěl, aby dělali ostatní tobě. A zároveň kde končí hranice jednoho a začíná

hranice druhýho. Tak nějak zhruba takhle. Jako legislativně si myslím, že to bude brutálně těžký. Takový, dotýká se mě to, pokud ne, tak to asi není problematický, ale myslím si, že s tím má celkově společnost jako s tou liberálností problém a že teďka na tomhle se to přesně projeví“ (Sebastian). Omílané používání AI ve škole s sebou nese etické otázky, které by se ale podle respondentů daly snadno vyřešit lepší komunikací a integrací AI do samotného vzdělávání.

Na to je často nahlížené jako na zaostalé a nereflektující současný a vlastně ani minulý pokrok. „Co se týče podvádění ve škole, přemejšlim bych to udělala z pohledu učitele, nebo spíš bych se snažila to začlenit a naučit se s tím pracovat. Že stejně jako se může podvádět u testu, tak se může podvádět u seminárky a nepřijde mi to, že by to byla nějaká jakoby vlastně velká změna. Že všechno je to takový dostupnější, ale taky to za tebe mohl předtím to někdo napsa“ (Naďa). Problémem je celkově nejasnost v tom, jak by se měly nástroje AI využívat. Neexistence celospolečenské shody, jak jsem již zmínili výše, může být zásadní překážkou při vytváření pravidel. „Že ono no jasně, asi jo, protože mám pocit, že ono to není nějaký úzus právě tím, jak je to takový ještě nový a každej to používá jinak a na něco jinýho a vlastně se moc jako nebo aspoň takhle v mé bublině mám pocit, že se moc jako neví tak tak vlastně asi není vytvořená shoda na tom, jak to používat. Takže tohle určitě by bylo fajn, kdyby jako nějak nějaká shoda byla, nebo kdyby třeba bysme měli nějakou informaci o tom, jak to můžeme používat“ (Karel). Zároveň je to zodpovědnost každého z nás, když nástroje AI využíváme, tak se také dozvědět, jaké mají implikace, což by mohlo usnadnit jejich regulaci do budoucna. „Lidi budou dávat fotky s Putinem na síť a pak nevíš, jestli to je fake, nebo ne. A mě to děsí jako obzvláště s téma jako vážnýma konfliktama, který probíhají“ (Marie). Problém spojeným s etikou je také neustále se valící množství témat, oblastí, které by se měly regulovat. „Nejhorší je, že se zastavíme nad jednou věcí, ale už vlastně přibývají další věci. Myslím, že těch oblastí je tolik, že není možný je pojmut, že by si někdo sednul a napsal 60–100 bodů, který se s AI pojí. Je to úplně revoluční nástroj, přináší hrozně moc nových příležitostí a hrozeb. Takže, takže to bude muset vznikat asi v jednotlivých odborech. A myslím, že nikdo asi moc nemá univerzální pravidlo, jak vlastně si dávat pozor a pracovat s tím. Bude teď znít děsivě, ale mám pocit, že to bude mít nějakou dohru, že to určitě nejde donekonečna. Že se nám to někde vymstí, v oblasti autonomie AI“ (Sebastian).

AI jako entita

Otázka budoucí možné podoby umělé inteligence, a tedy i budoucí možné interakce lidí s umělou inteligencí je spíše na úrovni filozofické. Co může způsobit antropomorfizace a co by znamenalo umělou inteligenci prezentovat v podobě kdy je **samostatnou entitou**, s sebou nese řadu **morálních a právních** otázek. V této kapitole se respondenti vyjadřují jak celkově k umělé inteligenci, tak k velkým jazykovým modelům.

Antropomorfizace je problematická, kromě právních otázek bychom se museli pustit do filozofické sféry o **podstatě lidské duše**, o povaze **chápání vlastního uvědomění**. Lidská podoba pro umělou inteligenci není favorizovanou možností. „*Mám trochu strach z toho, že by se možná nemělo úplně té umělé inteligenci dávat ta za první lidská podoba v těch robotech a za druhé vlastně jako dělat z toho entity. Že se pak musí řešit práva a vlastně pak už jako je ta tenká hranice, co je vlastně jako uvědomění sebe sama*” (Sebastian). Objevuje se velká nevole, až **strach**, kde hraje roli známý fenomén **uncanny valley**¹³. To potvrzuje studentka v maturitním ročníku ohledně lidské podoby AI. „*To je blbý, to je hrozný, to je to, čeho já se nejvíc bojím, já se strašně bojím robotů. Mám z nich prostě strach a neumím si představit, že jde člověk na terapii s robotem. Já bych ani nedokázala být v jedné místnosti, já tomu strašně nevěřím*“ (Nad'a).

S ohledem na praktické využití si další z respondentů dokáže představit, že by i antropomorfní robotická verze umělé inteligence mohla **přinášet pozitivní výsledky** například v oblasti domácí péče. Kde by zastupovala lidskou **asistentskou a pečovatelskou službu**. Zároveň by ale mělo být vždy jasně vymezené, že se **nejedná o lidskou bytost** a že komunikují, interagují s umělou inteligencí. „*U lidí, kteří na tom nejsou zdravotně nejlíp a kvůli tomu třeba jsou sami, si dokážu představit, že robot je tam 24/7 a vypadá jako člověk, že to pro ně může být příjemnější. Nicméně si myslím, že by se mělo držet to, že nám budou říkat, že to není reálný člověk, protože kolikrát jsou furt případy, že člověk si může říct, jo, vlastně vždyť to je úplně jako můj kámoš a ve finále pak zjistí, že je to robot. Asi takhle, je to padesát na padesát, když lidi jakoby jsou na tom zdravotně špatně a potřebují pomoci, tak*

¹³ Fenomén uncanny valley odkazuje na pocity neklidu nebo nejistoty, které lidé zažívají v případech, že se roboti nebo počítačově generované postavy jeví téměř jako lidské, ale mají nějaké malé nedokonalosti, které prozrazují jejich neživý původ. Představte si to jako křivku přijatelnosti: když se podoba s lidmi zvyšuje, naše pohodlí s touto podobou roste, až do bodu, kdy je podobnost téměř dokonalá – a pak najednou klesá. To je moment, kdy se objeví pocit uncanny valley – kdy něco vypadá téměř lidsky, ale něco na tom není stále úplně v pořádku, což vyvolává nepříjemný pocit.

si myslím, že by se klidně dalo to udělat, že to je jakoby člověk. Ale i přesto by to pořád měli vědět si myslím“ (Karel).

Na tomto tématu se setkávají jak středoškolští, tak vysokoškolští respondenti a je otázkou, jak se bude vyvíjet do budoucna.

7 DISKUSE

Výsledky výzkumu postojů k umělé inteligenci s užším zaměřením na velké jazykové modely jako je například ChatGPT identifikovaly řadu témat, která jsou poměrně otevřená a rozvětvená. V této kapitole je pro lepší porozumění zasadíme do širšího rámce zkoumání postojů k technologiím a umělé inteligenci. Nedílnou součástí je zhodnocení potenciálu této práce, jejího přínosu a dalšího možného využití v návaznosti na další možné výzkumy. Neopomeneme zhodnotit její limity a nedostatky.

Chápání úspěšnosti technologie na základě míry jejího přijetí a užívání (DeLone a McLean, 1992) řadí velké jazykové modely mezi **nejúspěšnější technologie** vůbec (Hu & Hu, 2023). **Individuální postoje** k umělé inteligenci, potažmo jazykovým modelům, v našem výzkum ukazují, že jejich přijímání nemusí být vždy spojeno pouze s kladnými postoji. Pozitivní i negativní aspekty postojů k LLM se často objevují současně. Jak potvrzuje i Gessl et al. (2019), technologický pokrok často evokuje smíšené pocity. V našem výzkumu jsme jich zachytili téměř dvě desítky a ač jsme je rozdělili na spektru pozitivních, neutrálních a negativních, nesouvisí tyto faktory s rozhodnutím využívat AI. To znamená, že pokud jsme našli negativní faktory, nejsou důvodem pro odmítnutí ChatGPT a dalších jazykových modelů. Můžeme tak zamítnout rané pojetí chápání postojů Martina Fishbeina (1975) implikující, že naše chování je ovlivněno zejména kladným nebo záporným postojem k určitému předmětu.

Naše výzkumná zjištění lze z části interpretovat za pomoci modelu technologického přijetí TAM (Davis, 1989) kde jsou postoje přímo ovlivněny užitečností a snadností použití daného nástroje. Model navíc zohledňuje i záměr a následnou reálnou aplikaci. Identifikovali jsme řadu faktorů, které na tyto dva atributy TAM odkazují. **Snadnost použití** implikuje rychlou adaptaci nástrojů LLM, jejich **integraci do běžného života**, každodenní využívání při **repetitivních úkonech** a prolíná se s **užitečností** rychlého a efektivního **třídění informací, strukturování a tvorby textů**. Oba znaky modelu TAM (užitečnost i jednoduchost) souvisí s pozitivním pohledem na **inovativní technologický průlom** a zásadní historický moment.

Dle Davise (1989) mají zkušenosti s používáním technologií také **reverzní** vlivy na postoje k ní. To znamená, že samotný zážitek s používáním, zkušenost a jeho dojem z něj,

dále formují budoucí pozitivní nebo negativní postoje. Jednoduše řečeno technologie a uživatel se vzájemně ovlivňují, působí mezi nimi zpětná vazba. To se samozřejmě odráží jak v pozitivně laděných zkušenostech, tak ve vzniku obav z přílišného **delegování lidských kompetencí** nebo rychlého a **nekontrolovatelného rozvoje** bez etické regulace. Aby nástroje AI nadále splňovaly požadavky užitečnosti a snadného použití, ukázalo se, že bude nutné zavést i preventivní opatření týkající se způsobu uplatnění těchto nástrojů s důrazem na **vzdělávání**, porozumění eticky správnému zacházení a nutnost vytvoření **legislativního rámce** a etických pravidel.

Přehled těchto témat až téměř dokonale splňuje podmínky modelu snadného přijímání technologií UTAUT (Venkatesh et al., 2003). Výzkum ukázal, že zamýšlené chování vede ke skutečnému používání jazykových modelů. Uživatelé jsou spokojeni s **očekávaným výkonem** AI i s **očekávanou snahou** na jejich straně. Usnadňují a rychleji plní své pracovní i školní povinnosti. Významnou roli hraje **vliv společnosti**, technologický skok je vnímán jako výjimečný moment v historii a **usnadňující podmínky**, snadná dostupnost a jednoduchost ovládání jazykových modelů (Raman, 2023; Sedaghat, 2023; Shoufan, 2023).

Postoje k AI se odráží ve **vnímání LLM** a způsobech jakými s nimi komunikujeme. Výzkum odhalil nejen emoční zabarvení a projevený respekt nástrojům AI, ale i rod jaký je jim přisuzován (ženský, mužský, střední). Na základě přímé zkušenosti s AI náš výzkum odhalil korelaci s důvěrou v souladu s chápáním postojů Fishbeina & Ajzena (1977) a Brendla & Higginse (1996) kde hraje roli proces sbírání zkušeností, učení.

Motivací uživatelů projevujících respektující chování vůči AI, které jsme v našem výzkumu zaznamenali, je i očekávání stejně respektující zpětné vazby od AI nástrojů. Můžeme tedy potvrdit dříve zjištěný přínos slušného chování v interakci s AI podle Bowman et al. (2023). Opět se také potvrzuje, že LLM dobře pasují do konceptů úspěšného přijímání technologií podle UTAUT a UTAUT2 Venkateshe et al. (2003; 2012). V rozporu s tvrzeními Gabbiadiniho et al. (2023) nelze odsouhlasit, že by podobnost lidské komunikaci s AI, přinášela negativní emoce. Respektující a slušná interakce mezi lidmi a AI je naopak kladná a v našem výzkumu přináší lepší, byť subjektivně hodnocené, výsledky. Mallick et al. (2023) potvrzují ve snaze vykresat z komunikace mezi AI a člověkem to nejlepší, pozitivní dopady stejně laděné komunikace. Milá umělá inteligence má v našem výzkumu jednoduše pozitivní vliv na lidské vnímání i chování a jak říká Bubeck et al. (2023), využívá

přirozené lidské potřeby komunikovat, a je dalším z faktorů, které ji otevřely tak širokému publiku.

Nejednoznačný etický rámec a hledání utilitárního způsobu komunikace s AI, nabízí otázku, zda je takové chování žádoucí, přijatelné, zcela nesmyslné nebo situačně a oborově podmíněné. Rozdíl mezi respondenty na vysokých a středních školách je v této otázce způsobu komunikace výrazný. Středoškolští žáci preferují strohé, věcné, někdy až **direktivní** jednání. Je důvodem vývojový stupeň adolescentů nebo je generace vysokoškoláků ovlivněna více popkulturou, o to více u nás, v zemi, kde Karel Čapek ve hře R.U.R. přivedl roboty k životu, s matnou vzpomínkou na legendu o pražském Golemovi, zůstává prozatím pouze předpokladem pro další zkoumání.

Jazykové modely typu ChatGPT přinášejí nové možnosti, ale také nové výzvy. AI dokáže neuvěřitelně rychle generovat odpovědi, ale zkušenosti uživatelů ukazují, ve shodě s výzkumným šetřením Alawida et al. (2023), že nejsou vždy pravdivé a mohou být zavádějící. Zatímco AI nemá konkurenci v rychlosti zpracování velkého objemu dat, nepostradatelnost lidského zásahu a kritického myšlení jsme v tomto bodě jsme shledali za nevyhnutelnou. Podle Enriquez et al. (2024) je to bezesporu lidský faktor, který zaručuje přesnost a kvalitu výstupů. Naše šetření ukazuje, že výstupy mohou kromě faktické nesprávnosti zaujímat například **diskriminační postoje**. Touto problematikou se zabývají Cao et al. (2023) upozorňující na přílišnou identifikaci jazykového modelu ChatGPT s americkou kulturou. Jak potvrzuje i náš výzkum, faktické chyby, nepřirozený jazyk, chyba v generovaných kódech, ale nerozeznání kulturní předpojatosti a možné diskriminace jsou oblasti, kde je nutné uplatnit **kritické myšlení**.

Shodně s předpoklady modelu přijímání technologií UTUAT 2 (Venkatesh et al., 2012), kde je kladen velký důraz na záměr používat technologie a kde oproti předešlým teoretickým schémátům hraje roli také **hedonická stránka motivace**, se v našem výzkumu ukázalo, že tyto nesprávné odpovědi mohou vest k frustraci a negativním postojům vůči AI a následně dočasné ztrátě zájmu a důvěry v AI s následným nepoužíváním.

Tato skutečnost má další **pozitivní implikace pro využívání ve vzdělání**. Paradoxně je to v našem výzkumu i rozvoj samostatného myšlení a kritického přístupu k informacím, motivovaný právě zavádějícími odpověďmi. Ač mohou být výstupy zprvu vnímány jako nedostatek, v průběhu času mohou zvyšovat důvěru/oblíbenost, tím, že připomínají lidskou omylnost, jak také uvádí Bubeck et al. (2023).

V našem výzkumu hraje roli při akceptování špatných odpovědí také **pochopení fungování** jazykových modelů. Minulé studie poukazují na povahu jazykových modelů, které jsou „černou skříňkou“ algoritmů, která je nepřehledná i pro vývojáře a nenabízí tak možnost zcela porozumět logice výstupů jazykových modelů (Cabitza et al., 2017), to potvrzuje i naše zkoumání. Ukázalo se, že hloubka znalostí není zásadní. Míra znalostí, které jsou pro uživatele dostačující odpovídá generování odpovědí podle statistického modelu a vytváření vět s logikou dalšího nejpravděpodobnějšího slova o které mluví Brown et al. (2020). Náš výzkum ukazuje, že znalost technologie se liší na základě vzdělání, oboru či osobního zájmu, ale tyto rozdíly ve znalostech nesouvisí se záměrem LLM používat. Podle modelu UTAUT2 mají faktory jako sociální vliv a podmínky usnadnění na rozhodování jedince větší vliv než jeho znalosti (Venkatesh et al., 2012).

Pokud výsledky našeho šetření opět aplikujeme na modely přijímání technologií jako je TAM s rozšířením na TAM2, shodně poukazujeme na externí faktory mající vliv na **zkušenosti** uživatelů s nástroji LLM. V našem výzkumu se ukázalo, že zkušenost a **četnost používání** nástrojů AI je ovlivněna stupněm a oborem studia a pracovním zaměřením a tím jaké vyvíjejí požadavky na výkon. Ty jsou zvýrazněny a zintenzivněny v obdobích vyšších nároků, které souvisí s relevantností práce v modelu TAM2 i výzkum Enríqueze (2023)

Součástí je i reflexe používaných typů LLM. Zatímco ChatGPT 3.5 je nejoblíbenějším, objevují se i další jako je Bing Chat, DeepL, Perplexity, ChatGPT4 nebo program na tvorbu obrázků Midjourney. Jejich preference v používání souvisí s jednoduchostí a jejich potřebou, jak shodně uvádí model TAM – vnímaná užitečnost a jednoduchost je pro použití klíčová. ChatGPT však své prvenství obhájí zejména v souvislosti s časným uvedením a verze ChatGPT3.5 souvisí i s jeho volným používáním, oproti pokročilejší verzi. Roli hraje bezesporu i jeho cena za poskytnutou hodnotu, jako jeden z faktorů, který identifikuje model UTAUT2 (Venkatesh et al., 2012). Není však bez vlivů dalších faktorů, které odpovídají spíše než modelu TAM, jeho vylepšení verzi TAM2 (Venkatesh & Bala, 2008). V našem výzkumu je zásadní také aspekt dobrovolnosti, používání nesouvisí s nutností, ale osobním či studijně-profesním přínosem. Objevují se požadavky na používání specifického nástroje, jiného než ChatGPT, které souvisí s obavou o zneužití citlivých firemních dat, shodná zjištění potvrzuje i Ray (2023). Jak jsme zmiňovali výše, ve snaze používat nástroje AI kontinuálně je potřeba také adekvátně reagovat na potřeby uživatelů. Kvalitní výstupy a demonstratibilita výsledku v modelu TAM2 analyzuje také faktory související s variabilitou úkolů, pro které jsou LLM používány. Od inspirace

pro kreativní nápady, po nástroj vzdělávání, kódování či jako zdroj zábavy, je tato variabilita úkolů spojená nejen se subjektivními potřebami, ale i faktorem zábavy. Kombinace těchto faktorů se odráží i v modelu TAM3.

Na rozdíl od Rodwayova (2020) výzkumu, zdůrazňuje náš výzkum potřebu lidského faktoru v terapeutické praxi, AI je v našem výzkumu chápána jako technologická bariéra. Jeho využívání by mělo zůstat v mezích asistenčního nástroje pro „první pomoc“, které snadno řeší problematiku nedostupnosti terapeutů. Zde teorie TAM2, ani TAM3 (zkontrolovat) nelze uplatnit. Vzhledem k vývoji modelů uživatelského přijetí technologických nástrojů bude nutné provést další revizi, která se bude specificky zaměřovat na oblasti ve kterých je nezbytné autentické rozhodování, mezi které patří takové, které mají přímý vliv na ovlivňování zásadních parametrů lidského života.

Rychlý nástup technologií umělé inteligence a jeho přijetí studenty a žáky, kteří je velice rychle začali používat, vznesl otázky, jakým způsobem by se s těmito nástroji mělo zacházet, aby se následkem jejich používání nenarušovala akademická integrita (Dwivedi et al., 2023). Skloňování témat jako plagiátorství a jiné neetické strategie naplňování školních povinností, vyvolaly v akademickém a školním prostředí urgentní poptávku nejen po zavedení pravidel jejich používání (Cotton et al., 2024 ; Anders, 2023), aby se zabránilo nečestnému jednání. Shodujeme se s výzkumy Dwivedi et al. (2023) a Kamalov et al. (2023), kde stejně jako v našem výzkumu má na straně žáků a studentů v oblasti vzdělávání prioritu změna způsobu vzdělávání, přizpůsobení školního kurikula, osnov, materiálů i vzdělání učitelů a dalších akademických pracovníků tak, aby odpovídaly rychle se rozvíjejícím technologiím a jejich dostupnosti. Ty i podle výsledků našeho výzkumu již našly pevné místo mezi technologiemi dneška a jejich přijetí a další přijímání podle porovnání s dostupnými teoriemi (modely TAM a UTAUT s jejich dalšími verzemi) není ohroženo, naopak stále upevňují své místo.

Domníváme se, že aktualita a velká pozornost výzkumů upřená/zaměřená k tématům vzdělání, zejména k oblasti plagiátorství, brzy pomine, nalezne řešení a budou se objevovat další, mnohem závažnější témata. Je až s podivem, že tolik reflektované téma plagiátorství není v našem výzkumu v hledáčku stěžejních téma u studentů. Stejná skutečnost se projevila i ve výzkumu Shoufana (2023), samotné vzdělávání stojí na žebříčku hodnot skutečných uživatelů výše, jak se shodneme. Nutnost vytvoření etických pravidel ve vzdělávání jak na straně studentů, tak učitelů zdůrazňuje Alzahrani (2023).

Jako by používání AI už bylo etablovanou skutečností a teď je řada na institucích vzniklou situaci reflektovat, protože na rovině studentů se to již stalo a velice rychle a pružně. V tomto ohledu bychom opět mohli otočit pozornost k modelu UTAUT2 (Venkatesh et al., 2012) a jednomu z důležitých faktorů, kterým je zvyk. Ten je jediným z aspektů modelu, který přímo souvisí s používáním technologie, nikoliv pouze záměrem. Je ovlivněn věkem, pohlavím a zkušeností. Jak potvrdilo statistické měření studenti ve věku 17-24 let jsou hned za věkovou kohortou 25-44 let, nejčastějšími uživateli, mají tedy bohaté zkušenosti a věkově jsou druhou nejčastěji ovlivněnou skupinou (Thormundsson, 2023).

Na základě našeho výzkumu tvrdíme, že efektivní používání AI nástrojů má hodnotu samo o sobě a lze předpokládat, že v budoucnu bude mít větší hodnotu než samotný obsah vypracovaného úkolu, což jsou shodná zjištění s výzkumem Javaid et al. (2023).

Pokud máme být upřímní, touto logikou jde vzdělání již dnes, seminární práce, eseje, bakalářské a diplomové práce jsou zkušební formou, kde jde ve velké míře i o formu a praxi v jejich vytvoření, aby je bylo možné následně aplikovat do praxe v pracovním prostředí. Jde tedy spíše o překonání obav v této technologii, které také mají svou podstatnou roli a význam v přijímání nových technologií.

V návaznosti na problematiku ve vzdělávání se obavy z umělé inteligence často točí okolo ztráty vlastních rozumových schopností a nedostatku kritického myšlení, vedoucí k manipulaci, problematická je také oblast **zneužití a sběr osobních dat** (Gupta et al., 2023; Wu et al., 2023). V našem výzkumu nicméně obavy a hrozby nejsou určující proměnnou, která by měla vliv na jejich používání, výhody použití převažují. Dwivedi et al. (2023) naproti tomuto tvrzení podceňují přínos AI a zdůrazňují obavy ze ztráty soukromí a šíření dezinformací. Na takové postoje lze namítnout, že ani v minulosti jsme zavedením online vyhledávačů neztratili schopnosti přemýšlet, ani nezmizely knihy. Naopak zastáváme názor, že internetová generace čte mnohem více, než ty předešlé, jenom se změnila povaha nástrojů. Stejně tak ani AI nedokáže nahradit lidskou inteligenci, jak potvrdily i Shoufanova (2023) zjištění.

Naše zjištění s tématem obav, jsou v rozporu ze nálezy Wang & Wang (2019), kteří mezi obavy řadí ohrožení soukromí, zánik pracovních míst, existenciální hrozby a netransparentnost. Tyto se v našem výzkumu neprojeví. Naopak, téma ztráty pracovních míst se ukázalo rovněž jako ve výzkumu Braunera et al. (2023) v pozitivním světle, v předpokladu vytvoření nových pracovních míst a ekonomického růstu.

Pokud bychom porovnali témata našeho výzkumu se zjištěním Sindermana et al. (2021) záporné hodnocení AI souvisí se ztrátou pracovních míst, to jsme nepotvrdili. Další je obava o lidstvo (Sinderman et al., 2021), ta se v našem výzkumu odrazila v případě předání kompetencí nástrojům AI zejména v oblastech bezpečnostních otázek a řešení konfliktů,

Posledním je strach ze samostatnosti AI, který, ač je v rozporu se zjištěními¹⁴ v jiné části práce, se ukázaly i jako existující součást obavy. Obavy vyvstávají z potenciálního využití umělé inteligence pro úkoly jako je psychologické poradenství (Schepman & Rodway, 2020) nebo výběry pracovníků, v našem výzkumu můžeme tato tvrzení uplatnit obecně v souvislosti se specifickými oblastmi, kdy jde o lidský život jako například záchranná linka, ale i náborů zaměstnanců, jak potvrzuje i Kieslich et al., (2021). S podobou AI souvisí také emoce, které v respondentech vyvolává. Vůbec představa AI jako samostatné entity, nástroje, který se vlastně projevuje zdánlivě vlastním rozumem a vůlí, je fascinující a zároveň doprovázená bázni a nejasností dalšího vývoje nástrojů AI, jak shodně tvrdí i (Johnson & Verdicchio, 2017). Obavy z humanoidní podoby robotů, kterou identifikoval Wang & Wang (2019) jsou také jednou z obav v našem výzkumu.

Etické otázky **práva a copyrightu**, používání v akademickém či pracovním prostředí, **regulace a transparentnost**, **ovlivňování společenské i individuální zodpovědnosti**, jsou výběrem témat naší analýzy. Nárůst obrovského množství etických hodnot, neexistence jasných pravidel a otázka jejich nastavování se odráží ve společnosti i různých oborech. Objevuje se množství nových etických otázek, které by měly být adresovány. Mají přednost celospolečenská pravidla, nebo jednotlivé obory? Problémy spojené s etikou obecně ve větší míře a hloubce reflektovali vysokoškolští studenti. Ukázalo se, že mají opravdu dobrý vhled do problematiky a dokážou přinášet nová témata do diskuse. U některých středoškolských studentů se tato otázka nesetkala ani s přílišným zájmem, nevěděli, co si mají o pojmu etika představit, což má vypovídající hodnotu samo o sobě a je nutné tato témata reflektovat také ve vzdělání. Což navazuje na kapitolu nezbytnosti vzdělávání o umělé inteligenci.

Morální dilema používání nástrojů umělé inteligence ve výuce a vypracovávání úkolů má snadné řešení. Jasným signálem je nikoliv zakazovat jejich používání, ale začlenit

¹⁴ V rámci našeho výzkumu je zásadní existence lidského rozumu, jak jsme již zmínili výše u studie Shoufana (2023).

je adekvátním způsobem jako pracovní a „legální“ nástroje ve vzdělání, jak potvrzuje i Anders (2023).

Etická pravidla jsou navázaná zejména na obavy, které jsou spojené s umělou inteligencí. Je zajímavé, že vypovídající hodnota o komplexnosti postojů jednotlivých studentů v této jedné otázce není dostatečná, ani jednoznačná. Samotný **termín umělá intelligence** odkazující na lidské charakteristiky může vést k jejich antropomorfizaci, ve vědomí se může utvářet obraz AI jako **samostatné entity**, které přisuzujeme lidské charakteristiky a která v nás vyvolává emoce (Bubeck at al., 2023).

Jedním z limitů této práce je novost výzkumu vycházející z poměrně krátké doby používání nástrojů AI širokou veřejností. Výzkumné problémy dostupných studií jsou tak přirozeně omezeny jak časem, tak oblastmi, které zkoumají. V době vydání této práce (začátek dubna 2024) uběhlo pouhých šestnáct měsíců od představení ChatGPT široké veřejnosti (v listopadu 2022). To se odráží i v dílčích oblastech této práce, které prozatím nejsou ve vědeckém diskurzu zpracovány. Jedním z cílů této práce je identifikace těchto témat k dalšímu možnému výzkumu.

Ačkoliv jsme v rámci tematické analýzy prošli několikanásobným kódováním a hledáním témat, dozajista by této studii prospělo jejich hodnocení jiným výzkumníkem. To však v tomto výzkumu nebylo možné provést. Povaha výzkumu neumožnila podrobné hledání korelací mezi jednotlivými kódy a tématy. Tato oblast otevírá možnost pro další výzkum.

8 ZÁVĚR

Cílem této práce bylo zmapovat současné postoje středoškolských a vysokoškolských studentů k umělé inteligenci a nalézt a přiblížit konkrétní témata, která mezi nimi rezonují. Postoje jsou pro výzkum umělé inteligence zcela zásadní. Zejména ve jejich fázi vývoje, kdy se stále jedná o technologii novou, která proniká do sfér osobního i komerčního života a každý den ovlivňuje miliony lidí po celém světě a stejně tak v České republice.

Mapování postojů studentů středních a vysokých škol přineslo řadu témat a poukázalo na jejich urgenci, prioritu nebo naopak okrajový zájem. Výsledky výzkumu tak přinesly odpovědi na zamýšlené otázky, témata. Obohatily současnou živou debatu o povaze a směřování AI.

- Identifikovali jsme šest hlavních témat s dalšími dílčími podtématy, která jsou poměrně otevřená a rozvětvená:
 - Individuální postoje k AI (pozitivní; negativní; neutrální)
 - Vnímání AI (komunikace s AI; důvěra v ChatGPT a jiné LLM; porozumění fungování AI)
 - Způsob používání a zkušenosti (zkušenost a četnost používání AI; druhy používaných jazykových modelů; způsob používání – efektivita a úspora času, přístup k informacím a vzdělávací nástroj, inspirace a zábava, nástroj kódování a programování; AI v terapii)
 - Vzdělávání o AI (mezery ve vzdělanosti; postoj institucí ke vzdělávání; postoj studentů k používání nástrojů AI)
 - Obavy a hrozby z AI (ztráta vlastního rozumu; nedostatek kritického myšlení; manipulace a zneužití dat; nezastavitelný pokrok; deteriorace mezilidských vztahů; agrese a válečný konflikt)
 - Etika a regulace (autorská práva; regulace a transparentnost; individuální zodpovědnost; zneužívání AI; AI jako entita)
- Na základě jejich povahy témat jsme identifikovali dvě oblasti:
 - Témata spojená s vlastní zkušeností a osobním prožitkem s AI
 - Témata spojená s AI a společností a jejím směřováním

Obecně se v našem výzkumu ukazuje, že starší výzkumy, které se datují před rokem 2022 často přicházejí s tvrzeními, která nejsou v souladu se zjištěními našeho výzkumu. Jak jsme již zmiňovali, proměnlivá povaha a rychlost změn, kterou s sebou nese nový svět s přítomností umělé inteligence v běžném životě, běžných lidí, přináší otázku jak dlouhou budou výsledky těchto výzkumů platné i v budoucnu.

Naplnili jsme nicméně základní cíl této práce, kterým bylo zachytit zmapovat postoje studentů a zachytit momentum v historickém vývoji umělé inteligence, zejména s důrazem na velké jazykové modely, jako je ChatGPT. Současné postoje studentů středních a vysokých škol a jejich mapování, je důležité nejen pro pochopení současných trendů a povahy jejich znalostí, ale právě i s ohledem na budoucí směřování vzájemné spolupráce lidí a umělé inteligence. Je to právě tato generace, na které leží v budoucnu zodpovědnost za výsledky a další směřování lidstva. Jen budoucnost ukáže, jakým směrem se budeme ubírat.

9 SOUHRN

Bakalářská práce se zabývá postoji středoškolských a vysokoškolských studentů k umělé inteligenci s důrazem na jazykové modely jako je ChatGPT, včetně, ale neomezeně, na ChatGPT a podobné jazykové modely.

První kapitola teoretické části se blíže zabývá postoji. Skrze postoje hodnotíme svět okolo nás, hodnotíme objekty, jevy či myšlenky od negativních přes pozitivní. Naše sklony k takovým hodnocením, jsou podle pojetí v psychologii kombinací afektivních, kognitivních a behaviorálních prvků, které jsou ve vzájemné interakci (Hogg & Vaughan, 2014). Zkoumání postojů se objevilo již v 19. století (Spencer, 1867). Počátek 20. století znamenalo boom ve výzkum postojů (Thomas & Znaniecki, 1918; Allport, 1935). V současnosti patří mezi nejčastěji skloňované výzkumy Fishbeina a Ajzena (Fishbein & Ajzen, 1974; Fishbein & Ajzen, 1980; Fishbein & Ajzen, 2010). Definice postojů není zcela jednotná a na tento trend navazuje také povaha jejich vnitřní struktury. Tři odlišné přístupy jsou po desetiletí předmětem akademické debaty – jednosložková struktura (Thurstone, 1931), dvojsložková (Bagozzi, 1981) a trojsložková (Ostrom, 1968; Ajzen 2012). Postoje vznikají jako součást socializačního procesu (Fishbein a Ajzen, 1975; Oskamp & Schultz, 2014) a mají své specifické funkce (Katz, 1960; Shavitt, 1990; Fazio 1989). Mezi nejčastější měření postojů kvantitativních metodách patří explicitní a implicitní měření, větší hloubku porozumění, ale nabízí kvalitativní metody zkoumání, zejména rozhovor (Šafránková & Kocourková, 2011).

Druhá kapitola teoretické části se věnuje umělé inteligenci (AI), která je univerzálním nástrojem se schopností reagovat na zcela nové situace. Jako první tento termín použil John McCarthy v roce 1955, který ji popsal jako vědu a techniku tvorby inteligentních strojů. Russel & Norvig (2022) popisují dva přístupy. Jeden chápe inteligenci jako myšlenkové procesy a uvažování, druhý je zaměřený na chování (Russel & Norvig, 2022). Popularizace AI na jejím počátku je spojena zejména se jmény Alana Turinga nebo Johna McCarthyho či Marvinu Minského. Přibližně od poloviny 20. století se AI neustále vyvíjela po linii významných milníků (Turing, 1950; Minsky & Papert, 2017; Newella & Simona, 1976). Ačkoliv hluboké učení, které má velký význam pro rozvoj velkých jazykových modelů (LLM), hrálo svou roli již od 70. let 20. století (LeCun et al., 1995) bylo to až v roce

2011, kdy se objevily první rozpoznávání jazyka a vizuálních objektů (Krizhevsky et al., 2012). Revoluci přinesly velké jazykové modely, které dokáží zpracovat přirozený jazyk (NLP) (Bang et al., 2023; Bubeck et al., 2023). První verze modelu GPT, zkratka GPT znamená „Generative Pre-trained Transformer“, byla poprvé uvedena v roce 2018 a měla 117 milionů parametrů (Bubeck et al., 2023; Radford et al., 2019) a je předchůdcem ChatGPT-3.5, dnes volně dostupné verzi velkého jazykového modelu, která má 175 miliard parametrů (Dále, 2020). ChatGPT zaznamenal raketový nástup a během prvních dvou měsíců získal 100 milion aktivních uživatelů.

Další část práce vymezuje postoje k technologiím (Dillon & Morris, 1996; DeLone & McLeana, 1992; Gessl et al., 2019)) a seznamuje s teoriemi a modely uživatelského přijímání technologií. Mezi ty nejpodstatnější patří TAM a jejich vylepšená verze TAM2 a TAM3 (Davis, 1989; Venkatesh & Davis, 2000; Venkatesh & Bala, 2008). Podstatná je také teorie přijetí technologie UTAUT a UTAUT2 (Venkatesh et al., 2003; Venkatesh et al., 2012) i Rogersova teorie difuze inovací (Rogers, 2003). Součástí této kapitoly je také přehled měření postojů k umělé inteligenci, často nemají rozšířené použití. Je to například škála obecných postojů k umělé inteligenci GAAIS (Schepman & Rodway, 2023) nebo škála negativních postojů k umělé inteligenci NAAIS (Persson et al., 2021). Součástí kapitoly je představení současných výzkumů v oblasti ChatGPT a AI. V poslední části se seznámíme s postojem vybraných škol a univerzit k AI v České republice.

Integrace umělé inteligence (AI) do každodenního života řadu otázek týkajících se postojů (Dwivedi et al., 2023; Cotton et al., 2024). Cílem výzkumné části práce je mapování současných postojů středoškolských a vysokoškolských studentů k AI, se zaměřením na ChatGPT a jeho varianty, skrze přehled témat, která mezi studenty rezonují. Od cílů práce se odvíjely dílčí výzkumné otázky s cílem blíže prozkoumat témata jako individuální postoje, vnímáním zkušenosti, očekávání, obavy a etické otázky týkající se AI.

S ohledem na cíle práce byl zvolený kvalitativní výzkum. Data byla získána skrze polostrukturované rozhovory s předem stanovenými otázkami. Triangulace byla zaručena získáváním dat od různých respondentů. Data byla zpracována doslovnou transkripcí. K analýze dat jsme s ohledem na cíl výzkumu a výzkumných otázek zvolili tematickou analýzu (Braun a Clarke, 2006).

Výzkumný soubor tvořilo deset respondentů, kteří byli pro tento výzkum relevantní vzhledem k jejím cílům a zaměření, tedy středoškolské a vysokoškolské studenty.

Podmínkou účasti na výzkumu nicméně nebyla přímá zkušenost s jazykovými modely. Cílem práce bylo oslovit studenty různých oborů a zaměření. Z celkového počtu byli čtyři ženy a šest mužů. Výběr a počet respondentů proběhl s ohledem na cíle mé studie a zároveň také relevantnost a validitu výsledků studie a saturaci dat.

Výsledky jsme rozdělili do šesti kapitol, která jsou také hlavními tématy. Individuální postoje k AI (pozitivní; negativní; neutrální), Vnímání AI (komunikace s AI; důvěra v ChatGPT a jiné LLM; porozumění fungování AI), Způsob používání a zkušenosti (zkušenost a četnost používání AI; druhy používaných jazykových modelů; způsob používání – efektivita a úspora času, přístup k informacím a vzdělávací nástroj, inspirace a zábava, nástroj kódování a programování; AI v terapii), Vzdělávání o AI (mezery ve vzdělanosti; postoj institucí ke vzdělávání; postoj studentů k používání nástrojů AI), Obavy a hrozby z AI (ztráta vlastního rozumu; nedostatek kritického myšlení; manipulace a zneužití dat; nezastavitelný pokrok; deteriorace mezilidských vztahů; agrese a válečný konflikt), Etika a regulace (autorská práva; regulace a transparentnost; individuální zodpovědnost; zneužívání AI; AI jako entita).

Na základě jejich povahy témat jsme identifikovali dvě oblasti: témata spojená s vlastní zkušeností a osobním prožitkem s AI a témata spojená s AI a společností a jejím směřováním.

Naplnili jsme nicméně základní cíl této práce, kterým bylo zachytit zmapovat postoje studentů a zachytit momentum v historickém vývoji umělé inteligence, zejména s důrazem na velké jazykové modely, jako je ChatGPT. Současné postoje studentů středních a vysokých škol a jejich mapování, je důležité nejen pro pochopení současných trendů a povahy jejich znalostí, ale právě i s ohledem na budoucí směřování vzájemné spolupráce lidí a umělé inteligence. Je to právě tato generace, na které leží v budoucnu zodpovědnost za výsledky a další směřování lidstva. Jen budoucnost ukáže, jakým směrem se budeme ubírat.

LITERATURA

- 1) Allport, G. W. (1935). Attitudes. In C. Murchinson (Ed.), *Handbook of social psychology* (pp. 798-844).
- 2) Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)
- 3) Ajzen, I. (2012). Martin Fishbein’s Legacy: The Reasoned Action Approach. *The Annals of the American Academy of Political and Social Science*, 640, 11–27. <https://www.jstor.org/stable/23218420>
- 4) Alawida, M., Mejri, S., Mehmood, A., Chikhaoui, B., & Isaac Abiodun, O. (2023). A Comprehensive Study of ChatGPT: Advancements, Limitations, and Ethical Considerations in Natural Language Processing and Cybersecurity. *Information*, 14(8), Article 8. <https://doi.org/10.3390/info14080462>
- 5) Anders, B. A. (2023). Is using ChatGPT cheating, plagiarism, both, neither, or forward thinking? *Patterns*, 4(3), 100694. <https://doi.org/10.1016/j.patter.2023.100694>
- 6) Bagozzi, R. P. (1981). Attitudes, intentions, and behavior: A test of some key hypotheses. *Journal of Personality and Social Psychology*, 41(4), 607–627. <https://doi.org/10.1037/0022-3514.41.4.607>
- 7) Bang, Y., Cahyawijaya, S., Lee, N., Dai, W., Su, D., Wilie, B., Lovenia, H., Ji, Z., Yu, T., Chung, W., Do, Q. V., Xu, Y., & Fung, P. (2023). A Multitask, Multilingual, Multimodal Evaluation of ChatGPT on Reasoning, Hallucination, and Interactivity (arXiv:2302.04023). arXiv. <https://doi.org/10.48550/arXiv.2302.04023>
- 8) Bowman, R., Cooney, O., Newbold, J. W., Thieme, A., Clark, L., Doherty, G., & Cowan, B. (2023). Exploring how politeness impacts the user experience of chatbots

- for mental health support. *International Journal of Human-Computer Studies (IJHCS)*. <https://www.microsoft.com/en-us/research/publication/exploring-how-politeness-impacts-the-user-experience-of-chatbots-for-mental-health-support/>
- 9) Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3, 77–101. <https://doi.org/10.1191/1478088706qp063oa>
 - 10) Brendl, C. M., & Higgins, E. T. (1996). Principles of Judging Valence: What Makes Events Positive or Negative? *Advances in Experimental Social Psychology*, 28(C), 95-160.
 - 11) Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D., Wu, J., Winter, C., ... Amodei, D. (2020a). Language Models are Few-Shot Learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901. <https://ysu1989.github.io/courses/au20/cse5539/GPT-3.pdf>
 - 12) Bubeck, S., Chandrasekaran, V., Eldan, R., Gehrke, J., Horvitz, E., Kamar, E., Lee, P., Lee, Y. T., Li, Y., Lundberg, S., Nori, H., Palangi, H., Ribeiro, M. T., & Zhang, Y. (2023). *Sparks of Artificial General Intelligence: Early experiments with GPT-4* (arXiv:2303.12712). arXiv. <https://doi.org/10.48550/arXiv.2303.12712>
 - 13) Cabitza, F., Rasoini, R., & Gensini, G. F. (2017). Unintended Consequences of Machine Learning in Medicine. *JAMA*, 318(6), 517–518. <https://doi.org/10.1001/jama.2017.7797>
 - 14) Cao, Y., Zhou, L., Lee, S., Cabello, L., Chen, M., & Hershcovich, D. (2023). Assessing Cross-Cultural Alignment between ChatGPT and Human Societies: An Empirical Study. https://www.researchgate.net/publication/369655572_Assessing_Cross-Cultural_Alignment_between_ChatGPT_and_Human_Societies_An_Empirical_Study/reference

- 15) Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
- 16) Crawford, J., Cowling, M., & Allen, K.-A. (2023). Leadership is needed for ethical ChatGPT: Character, assessment, and learning using artificial intelligence (AI). *Journal of University Teaching and Learning Practice*, 20. <https://doi.org/10.53761/1.20.3.02>
- 17) Cummings, R. E., Monroe, S. M., & Watkins, M. (2024). Generative AI in first-year writing: An early analysis of affordances, limitations, and a framework for the future. *Computers and Composition*, 71, 102827. <https://doi.org/10.1016/j.compcom.2024.102827>
- 18) Dale, R. (2021). GPT-3: What's it good for? *Natural Language Engineering*, 27(1), 113–118. doi:10.1017/S1351324920000601
- 19) Davis, F. D. (1989). Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- 20) Dillon, A., & Morris, M. G. (1996). USER ACCEPTANCE OF INFORMATION TECHNOLOGY: *THEORIES AND MODELS*.
- 21) Dis, E., Bollen, J., Zuidema, W., Rooij, R., & Bockting, C. (2023). ChatGPT: Five priorities for research. *Nature*, 614, 224–226. <https://doi.org/10.1038/d41586-023-00288-7>
- 22) Domingos, P. (2015). *The master algorithm: How the quest for the ultimate learning machine will remake our world*. Basic Books.

- 23) Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). Opinion Paper: “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- 24) Enriquez, B. G. A., Ballesteros, M. A. A., Jordan, O. H., Roca, C. L., & Tirado, K. S. (2024). *Analysis College Students' attitude towards the use of ChatGPT in their academic activities: Effect of intent to use, verify information and responsible use*. <https://doi.org/10.21203/rs.3.rs-3563928/v1>
- 25) Fazio, R. H. (1989). On the power and functionality of attitudes: The role of attitude accessibility. In *Attitude structure and function* (s. 153–179). Lawrence Erlbaum Associates, Inc.
- 26) Fazio, R., Jackson, J., Dunton, B., & Williams, C. (1995). Variability in Automatic Activation as an Unobtrusive Measure of Racial Attitudes: A Bona Fide Pipeline? *Journal of Personality and Social Psychology - PSP*, 69, 1013–1027. <https://doi.org/10.1037//0022-3514.69.6.1013>
- 27) Gabbiadini, A., Dimitri, O., Baldissarri, C., & Manfredi, A. (2023). *Does ChatGPT Pose a Threat to Human Identity?* (SSRN Scholarly Paper 4377900). <https://doi.org/10.2139/ssrn.4377900>
- 28) Gessl, A. S., Schlögl, S., & Mevenkamp, N. (2019). On the perceptions and acceptance of artificially intelligent robotics and the psychology of the future elderly.

- Behaviour & Information Technology*, 38(11), 1068–1087.
<https://doi.org/10.1080/0144929X.2019.1566499>
- 29) Gherheș, V. (2018). Why Are We Afraid of Artificial Intelligence (Ai)? *European Review Of Applied Sociology*, 11(17), 6–15.
<https://sciendo.com/article/10.1515/eras-2018-0006>
- 30) Greenwald, A. G., Banaji, M. R., Rudman, L. A., Farnham, S. D., Nosek, B. A., & Mellott, D. S. (2002). A unified theory of implicit attitudes, stereotypes, self-esteem, and self-concept. *Psychological Review*, 109(1), 3–25. <https://doi.org/10.1037/0033-295X.109.1.3>
- 31) Greenwald, A. G., Poehlman, T. A., Uhlmann, E. L., & Banaji, M. R. (2009). Understanding and using the Implicit Association Test: III. Meta-analysis of predictive validity. *Journal of Personality and Social Psychology*, 97(1), 17–41. <https://doi.org/10.1037/a0015575>
- 32) Gupta, M., Akiri, C., Aryal, K., Parker, E., & Praharaj, L. (2023). *From ChatGPT to ThreatGPT: Impact of Generative AI in Cybersecurity and Privacy*.
- 33) Hadi Mogavi, R., Deng, C., Juho Kim, J., Zhou, P., D. Kwon, Y., Hosny Saleh Metwally, A., Tlili, A., Bassanelli, S., Bucchiarone, A., Gujar, S., Nacke, L. E., & Hui, P. (2024). ChatGPT in education: A blessing or a curse? A qualitative study exploring early adopters' utilization and perceptions. *Computers in Human Behavior: Artificial Humans*, 2(1), 100027.
<https://doi.org/10.1016/j.chbah.2023.100027>
- 34) Hartl, P., & Hartlová, H. (2010). *Velký psychologický slovník*. Portál.
- 35) Howitt, D. (2010). *Introduction to Qualitative Methods in Psychology*. Prentice Hall.

- 36) Hu, K., & Hu, K. (2023, únor 2). ChatGPT sets record for fastest-growing user base—Analyst note. *Reuters*. <https://www.reuters.com/technology/chatgpt-sets-record-fastest-growing-user-base-analyst-note-2023-02-01/>
- 37) Javaid, M., Haleem, A., Singh, R. P., Khan, S., & Khan, I. H. (2023). Unlocking the opportunities through ChatGPT Tool towards ameliorating the education system. *BenchCouncil Transactions on Benchmarks, Standards and Evaluations*, 3(2), 100115. <https://doi.org/10.1016/j.tbench.2023.100115>
- 38) Johnson, D., & Verdicchio, M. (2017). AI Anxiety. *Journal of the Association for Information Science and Technology*, 68. <https://doi.org/10.1002/asi.23867>
- 39) Kamalov, F., Santandreu Calonge, D., & Gurrib, I. (2023). New Era of Artificial Intelligence in Education: Towards a Sustainable Multifaceted Revolution. *Sustainability*, 15(16), Article 16. <https://doi.org/10.3390/su151612451>
- 40) Kawakami, K., Young, H., & Dovidio, J. F. (2002). Automatic stereotyping: Category, trait, and behavioral activations. *Personality and Social Psychology Bulletin*, 28(1), 3–15. <https://doi.org/10.1177/0146167202281001>
- 41) Katz, D. (1960). The Functional Approach to the Study of Attitudes. *Public Opinion Quarterly*, 24, 163-204. <https://doi.org/10.1086/266945>
- 42) Kieslich, K., Lünich, M., & Marcinkowski, F. (2021). The Threats of Artificial Intelligence Scale (TAI). *International Journal of Social Robotics*, 13(7), 1563–1577. <https://doi.org/10.1007/s12369-020-00734-w>
- 43) King, M. R. & chatGPT. (2023). A Conversation on Artificial Intelligence, Chatbots, and Plagiarism in Higher Education. *Cellular and Molecular Bioengineering*, 16(1), 1–2. <https://doi.org/10.1007/s12195-022-00754-8>
- 44) Kočárková, J. (2023). Umělá inteligence už může legálně pomáhat i při psaní studentských závěrečných prací | Technický týdeník. Získáno 26. březen 2024, z

https://www.technickytydenik.cz/rubriky/ict/umela-intelligence-uz-muze-legalne-pomahat-i-pri-psani-studentskych-zaverecnych-praci_59206.html

- 45) Kopecký, K., Szotkowski, R., Voráč, D., Krejčí, V., & Dobešová, P. (2023). *České školy a umělá inteligence – výzkum*.
- 46) Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). ImageNet Classification with Deep Convolutional Neural Networks. In F. Pereira, C. J. Burges, L. Bottou, & K. Q. Weinberger (Ed.), *Advances in Neural Information Processing Systems* (Roč. 25). Curran Associates, Inc.
https://proceedings.neurips.cc/paper_files/paper/2012/file/c399862d3b9d6b76c8436e924a68c45b-Paper.pdf
- 47) LeCun, Y., Jackel, L., Bottou, L., Brunot, A., Cortes, C., Denker, J., Drucker, H., Guyon, I., Muller, U., Sackinger, E., Simard, P., & Vapnik, V. (1995). Comparison of learning algorithms for handwritten digit recognition. In *International Conference on Artificial Neural Networks*.
- 48) LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444. <https://doi.org/10.1038/nature14539>
- 49) Májovský, M., Černý, M., Kasal, M., Komarc, M., & Netuka, D. (2023). Artificial Intelligence Can Generate Fraudulent but Authentic-Looking Scientific Medical Articles: Pandora's Box Has Been Opened. *Journal of Medical Internet Research*, 25, e46924. <https://doi.org/10.2196/46924>
- 50) Mallick, R., Flathmann, C., Lancaster, C., Hauptman, A., McNeese, N., & Freeman, G. (2023). The pursuit of happiness: The power and influence of AI teammate emotion in human-AI teamwork. *Behaviour & Information Technology*, 1–25. <https://doi.org/10.1080/0144929X.2023.2277909>

- 51) Marikyan, D., & Papagiannidis, S. (2023). Unified Theory of Acceptance and Use of Technology: A review. In S. Papagiannidis (Ed.), *TheoryHub Book*. Retrieved from <https://open.ncl.ac.uk> / ISBN: 9781739604400
- 52) Miovský, M. (2006). *Kvalitativní přístup a metody v psychologickém výzkumu*. Praha: Grada.
- 53) Minsky, M., & Papert, S. A. (2017). *Perceptrons: An Introduction to Computational Geometry*. The MIT Press.
<https://doi.org/10.7551/mitpress/11301.001.0001>
- 54) Mitchell, M. (2019). *Artificial intelligence: a guide for thinking humans*. Pelican.
- 55) Montag, C., Ali, R., & Davis, K. L. (2024). Affective neuroscience theory and attitudes towards artificial intelligence. *AI & SOCIETY*.
<https://doi.org/10.1007/s00146-023-01841-8>
- 56) Montag, C., Nakov, P., & Ali, R. (2023). *Considering the IMPACT Framework to Understand the Ai-well-being-complex From an Interdisciplinary Perspective* (SSRN Scholarly Paper 4584349). <https://doi.org/10.2139/ssrn.4584349>
- 57) Nakonečný, M. (1997). *Encyklopedie obecné psychologie 2., rozš. vyd., v Akademii vyd. 1. (1. vyd. v nakl. Vodnář pod náz. Lexikon psychologie)*. Academia.
- 58) Nilsson, N. J. (2009). *The Quest for Artificial Intelligence*. Cambridge: Cambridge University Press.
- 59) Ostrom, T. M. (1968). 1 - The Emergence of Attitude Theory: 1930 — 1950. In A. G. Greenwald, T. C. Brock, & T. M. Ostrom (Ed.), *Psychological Foundations of Attitudes* (s. 1–32). Academic Press. <https://doi.org/10.1016/B978-1-4832-3071-9.50007-6>
- 60) Oskamp, S., & Schultz, P. W. (2014). *Attitudes and Opinions* (3. vyd.). Psychology Press. <https://doi.org/10.4324/9781410611963>

- 61) Peters, M. A., Jackson, L., Papastephanou, M., Jandrić, P., Lazaroiu, G., Evers, C. W., Cope, B., Kalantzis, M., Araya, D., Tesar, M., Mika, C., Chen, L., Wang, C., Sturm, S., Rider, S., & Fuller, S. (2023). AI and the future of humanity: ChatGPT-4, philosophy and education – Critical responses. *Educational Philosophy and Theory*, 1–35. <https://doi.org/10.1080/00131857.2023.2213437>
- 62) Patton, M. Q. (1990). *Qualitative evaluation and research methods*, 2nd ed (s. 532). Sage Publications, Inc.
- 63) Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., Sutskever, I. & others (2019). Language models are unsupervised multitask learners. *OpenAI blog*, 1, 9.
- 64) Ray, P. P. (2023). ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. *Internet of Things and Cyber-Physical Systems*, 3, 121–154. <https://doi.org/10.1016/j.iotcps.2023.04.003>
- 65) Rejzek, J. (2015). *Český etymologický slovník* (Třetí vydání (druhé přepracované a rozšířené vydání). Leda.
- 66) Rogers, E. M. (2003). *Diffusion of innovations* (3rd ed). Free Press ; Collier Macmillan.
- 67) Russell, S. J., & Norvig, P. (2022). *Artificial intelligence: a modern approach* (Fourth edition). Pearson.
- 68) Řehan, V. (2007). *Sociální psychologie: studijní texty pro distanční studium*. Univerzita Palackého v Olomouci.
- 69) Saggi, A., & Ante, L. (2023). The influence of ChatGPT on artificial intelligence related crypto assets: Evidence from a synthetic control analysis. *Finance Research Letters*, 55, 103993. <https://doi.org/10.1016/j.frl.2023.103993>
- 70) Sebastian, G. (2023). Privacy and Data Protection in ChatGPT and Other AI Chatbots: Strategies for Securing User Information. *International Journal of Security*

- and Privacy in Pervasive Computing (IJSPPC)*, 15(1), 1–14.
<https://doi.org/10.4018/IJSPPC.325475>
- 71) Schepman, A., & Rodway, P. (2023). The General Attitudes towards Artificial Intelligence Scale (GAAIS): Confirmatory Validation and Associations with Personality, Corporate Distrust, and General Trust. *International Journal of Human-Computer Interaction*, 39(13), 2724–2741.
<https://doi.org/10.1080/10447318.2022.2085400>
- 72) Shavitt, S. (1990). The role of attitude objects in attitude functions. *Journal of Experimental Social Psychology*, 26(2), 124–148. [https://doi.org/10.1016/0022-1031\(90\)90072-T](https://doi.org/10.1016/0022-1031(90)90072-T)
- 73) Sindermann, C., Sha, P., Zhou, M., Wernicke, J., Schmitt, H. S., Li, M., Sariyska, R., Stavrou, M., Becker, B., & Montag, C. (2021). Assessing the Attitude Towards Artificial Intelligence: Introduction of a Short Measure in German, Chinese, and English Language. *KI - Künstliche Intelligenz*, 35(1), 109–118.
<https://doi.org/10.1007/s13218-020-00689-0>
- 74) Stein, J.-P., Appel, M., Jost, A., & Ohler, P. (2020). Matter over mind? How the acceptance of digital entities depends on their appearance, mental prowess, and the interaction between both. *International Journal of Human-Computer Studies*, 142, 102463. <https://doi.org/10.1016/j.ijhcs.2020.102463>
- 75) Stein, J.-P., Liebold, B., & Ohler, P. (2019). Stay back, clever thing! Linking situational control and human uniqueness concerns to the aversion against autonomous technology. *Computers in Human Behavior*, 95, 73–82.
<https://doi.org/10.1016/j.chb.2019.01.021>

- 76) Stein, J.-P., Messingschlager, T., Gnambs, T., Hutmacher, F., & Appel, M. (2024). Attitudes towards AI: Measurement and associations with personality. *Scientific Reports*, 14(1), 2909. <https://doi.org/10.1038/s41598-024-53335-2>
- 77) Strauss, A. L., & Corbin, J. (1999). *Základy kvalitativního výzkumu: postupy a techniky metody zakotvené teorie*. Sdružení Podané ruce.
- 78) Terzi, R. (2020). An adaptation of artificial intelligence anxiety scale into Turkish: Reliability and validity study. *International Online Journal of Education and Teaching (IOJET)*, 7(4). 1501-1515.
- 79) Thurstone, L. L. (1928). Attitudes can be measured. *American Journal of Sociology*, 33, 529–554. <https://doi.org/10.1086/214483>
- 80) Tlili, A., Shehata, B., Adarkwah, M., Bozkurt, A., Hickey, D., Huang, R., & Agyemang, B. (2023). What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Smart Learning Environments*, 15, 1–24. <https://doi.org/10.1186/s40561-023-00237-x>
- 81) Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59, 433–460. <https://doi.org/10.1093/mind/LIX.236.433>
- 82) Masarykova Univerzita, M. (2023). *Doporučení k využívání umělé inteligence ve výuce*. Zajišťování kvality na MUNI | MUNI. Získáno 19. březen 2024, z <https://kvalita.muni.cz/kvalita-vyuky/doporuceni-k-vyuzivani-umele-inteligence-ve-vyuce>
- 83) Masarykova Univerzita, M. (2023). *Stanovisko k využívání umělé inteligence ve výuce na Masarykově univerzitě*. Masarykova univerzita. Získáno 19. březen 2024, z <https://www.muni.cz/o-univerzite/uredni-deska/stanovisko-k-vyuzivani-ai>
- 84) Venkatesh, V., & Bala, H. (2008). Technology Acceptance Model 3 and a Research Agenda on Interventions. *Decision Sciences*, 39(2), 273–315. <https://doi.org/10.1111/j.1540-5915.2008.00192.x>

- 85) Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User Acceptance of Information Technology: Toward a Unified View. *MIS Quarterly*, 27(3), 425–478. <https://doi.org/10.2307/30036540>
- 86) Wang, Y.-Y., & Wang, Y.-S. (2019). Development and validation of an artificial intelligence anxiety scale: An initial application in predicting motivated learning behavior. *Interactive Learning Environments*, 30, 1–16. <https://doi.org/10.1080/10494820.2019.1674887>
- 87) Williams, M., Rana, N., & Dwivedi, Y. (2015). The unified theory of acceptance and use of technology (UTAUT): A literature review. *Journal of Enterprise Information Management*, 28, 443–488. <https://doi.org/10.1108/JEIM-09-2014-0088>
- 88) Thormundsson, B. (2023). Worldwide: Use of generative AI 2023. (2023). *Statista*. Získáno 29. března 2024, z <https://www.statista.com/statistics/1455933/generative-ai-use-worldwide-by-group/>
- 89) Wu, T., He, S., Liu, J., Sun, S., Liu, K., Han, Q.-L., & Tang, Y. (2023). A Brief Overview of ChatGPT: The History, Status Quo and Potential Future Development. *IEEE/CAA Journal of Automatica Sinica*, 10, 1122–1136. <https://doi.org/10.1109/JAS.2023.123618>
- 90) xibang. (2022, září 7). *Be polite to robots* [Reddit Post]. r/funny. www.reddit.com/r/funny/comments/x8bc2w/be_polite_to_robots/

PŘÍLOHY

Seznam příloh:

1. Abstrakt v českém jazyce
2. Abstrakt v anglickém jazyce
3. Seznam použitých obrázků
4. Seznam použitých grafů
5. Seznam použitých tabulek
6. Kreslená anekdota o budoucnosti AI a lidí
7. Ukázky počátečního kódování dat z programu Atlas.ti

ABSTRAKT DIPLOMOVÉ PRÁCE

Název práce: Postoje středoškolských a vysokoškolských studentů k AI

Autor práce: Mgr. Bc. Eva Hašek Dóšová

Vedoucí práce: PhDr. Ondrej Gergely, PhD.

Počet stran a znaků: 92 stran, 158 222 znaků

Počet příloh: 6

Počet titulů použité literatury: 90

Abstrakt (800–1200 zn.):

Tato práce se zabývá aktuálním tématem postoje středoškolských a vysokoškolských studentů k umělé inteligenci (AI) se zaměřením na generativní jazykové modely jako je ChatGPT. Tato bakalářská práce se zabývá zkoumáním postojů studentů středních a vysokých škol k umělé inteligenci, s důrazem na jazykové modely, včetně, ale neomezeně, na ChatGPT a jeho varianty. Zajímá nás, jakým způsobem vnímají tito studenti využití umělé inteligence v běžném životě, ve vzdělávání a ve společnosti obecně. Tato práce klade důraz na pochopení postojů, předsudků a očekávání studentů ohledně této inovativní technologie a přispívá k hlubšímu pochopení jejich role a významu v současném digitálním prostředí.

Klíčová slova: umělá inteligence, postoje, středoškolští a vysokoškolští studenti.

ABSTRACT OF THESIS

Title: Attitudes among secondary and university students towards AI

Author: Mgr. Bc. Eva Hašek Dóšová

Supervisor: PhDr. Ondrej Gergely, PhD.

Number of pages and characters: 92 pages, 158 222 characters

Number of appendices: 6

Number of references: 90

Abstract (800–1200 characters):

This bachelor's thesis investigates the attitudes of high school and university students towards artificial intelligence, with an emphasis on language models, including but not limited to ChatGPT and its variants. It explores how students perceive the use of artificial intelligence in everyday life, education, and society at large. The work focuses on understanding students' attitudes, prejudices, and expectations concerning this innovative technology and contributes to a deeper understanding of their role and significance in the current digital environment.

Keywords: artificial intelligence, attitudes, secondary school and university students.

Příloha č. 3: Seznam použitých grafů

Graf č. 1: Globální používání AI v roce 2023, podle skupin (Thormundsson (2023; upraveno)

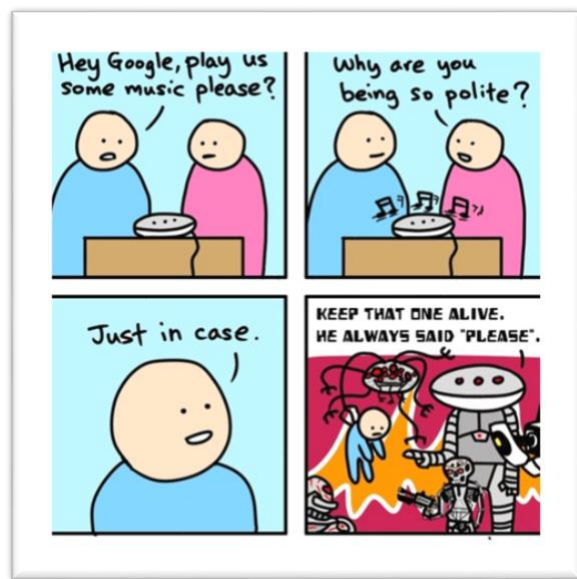
Příloha č. 4: Seznam použitých tabulek

Tabulka č. 1: Škály postojů k umělé inteligenci

Tabulka č. 2: Deskriptivní charakteristiky respondentů

Tabulka č. 3: Přehled pozitivních, neutrálních a negativních aspektů postojů k AI

Příloha č. 5: Kreslená anekdota o budoucnosti AI a lidí



Zdroj: xibang (2022)

Příloha č. 6: Ukázky počátečního kódování dat z programu Atlas.ti

Code Manager			
Search			
Name	Color	Groups ↓	Quotations
☐ informace z vlastní zkušenosti			6
☐ informovanost ve školách			10
☐ jazyk - AJ/Čj			7
☐ konzervativnější postoj			2
☐ kritické myšlení			13
☐ mám informace o fungování			3
☐ mám přímou zkušenost s AI			2
☐ marketing/copywriting			1
☐ nástroj inspirace			12
☐ negativní postoj k AI Details			8
☐ nemám dostatek informací o AI k rozhodování			5
☐ nemám přímou zkušenost s AI			3