

MICHAL HANÁK

KOMUNIKAČNÍ ZRALOST V INTERAKCI S GENERATIVNÍ AI

Empirická analýza uživatelských vzorců
v edukačním kontextu

Člověk.
Technologie.
Vzdělávání.
Smysluplná
budoucnost.

VĚDA
PRAXE
LIDÉ
POTENCIÁL



Michal Hanák

Uherské Hradiště 2026

KOMUNIKAČNÍ ZRALOST V INTERAKCI S GENERATIVNÍ AI

**Empirická analýza uživatelských vzorců v
edukačním kontextu**

Michal Hanák



Uherské Hradiště 2026

Obsah

Úvod	6
1 Základní přehled pro rychlou práci i orientaci	11
2 Udržování pořádku a struktury v komunikaci s AI	15
3 Jak se bavit s umělou inteligencí?	21
3.1 Uživatelská přívětivost	21
3.2 Přesnost a spolehlivost	25
3.3 Personalizace a adaptabilita	30
3.4 Transparentnost a důvěryhodnost	35
3.5 Eticko-sociální aspekty	40
3.6 Rychlost a výkonnost	46
3.7 Integrace a kompatibilita	53
3.8 Nákladovost a efektivita	60
3.9 Podpora a znalostní báze	66
4 Strategie pro bezproblémovou komunikaci s AI	74
4.1 Buďte konkrétní	75
4.2 Sdělte svůj záměr	77
4.3 Správný pravopis a gramatika	79
4.4 Řiďte formát výstupu	81
4.5 Pokládejte doplňující otázky	82
4.6 Výzva k ověření faktů	84
4.7 Experimentujte s různými frázemi	87
4.8 Učte se z odpovědí	88
5 Teoretický rámec MLEF-AI a koncept komunikační zralosti uživatele	91

5.1	Fragmentace teoretického aparátu výzkumu AI ve vzdělávání	91
5.2	Multi-Level Ecological Framework for AI Implementation in Education	92
5.3	AICP jako pedagogický rámec edukační platformy	97
5.4	Komunikační zralost uživatele jako nový teoretický konstrukt.....	99
5.5	Most k empirické studii: výzkumné otázky vyvozené z teoretického rámce	101
6	Empirická studie interakčních vzorců v komunikaci uživatelů s generativní AI	104
6.1	Status studie a metodologické rámování.....	104
6.2	Výsledky.....	111
6.3	Sedm pilotních případů — kvalitativní obohacení.....	119
6.4	Diskuse.....	121
6.5	Závěr a návaznost na Fázi 2.....	127
7	Vzorce pro perfektní prompt	129
7.1	Komponenty vzorců	132
7.2	Rozbor hlavních komponent pro vzorce	135
7.3	Nejdůležitější vzorce promptů	152
8	Správný zápis promptu	159
8.1	Klíčové prvky efektivního promptu	164
8.2	Aktivační slovesa.....	175
	Závěr.....	187
	Shrnutí klíčových přínosů.....	187
	Vztah mezi praktickou a teoreticko-empirickou rovinou	188

Limity monografie a otevřené otázky	189
Výhled: směry dalšího výzkumu.....	190
Závěrečné slovo	191
Věcný rejstřík.....	193
Jmenný rejstřík.....	199
Použitá literatura.....	202

Úvod

Komunikace představuje základní prostředek, jímž probíhá výuka a učení, **což je také ústředním tématem této vědecké monografie**. V pedagogické praxi slouží komunikace k předávání informací, rozvíjení porozumění a budování sociálního prostředí pro učení. Řada autorů zdůrazňuje, že komunikační dovednosti učitele jsou pro úspěšné vzdělávání zásadní (Kyriacou, 2009; Průcha, 2013; Vališová, 2011). **Kvalita komunikace mezi učitelem a žákem totiž ovlivňuje celý průběh i výsledek vzdělávacího procesu** (Provázková Stolinská, 2021). To znamená, že jasné vysvětlování, aktivní naslouchání a vhodná zpětná vazba mohou výrazně zlepšit porozumění učiva a studijní výsledky.

Za klíčový aspekt komunikace ve výuce lze považovat dialog – obousměrnou výměnu myšlenek. Dialogické vyučování staví na tom, že promyšlená diskuse a pokládání otázek rozvíjí myšlení žáků a prohlubuje jejich pochopení učiva (Alexander, 2008). V takovém pojetí nejsou žáci pouhými pasivními příjemci informací, ale aktivně se účastní rozhovoru, kladou dotazy a formulují své myšlenky. Výzkumy ukazují, že když jsou žáci vedeni k pokládání vlastních dotazů, zvyšuje se jejich aktivita a motivace k učení, což pozitivně ovlivňuje studijní výsledky (Shakurnia et al., 2018). Tvorba otázek studenty – například v rámci tzv. SIS dotazů (Sincerely Information Seeking questions) – odráží jejich vnitřní zvědavost a snahu porozumět problému do hloubky (Van der Meij, 1994; Jirout & Klahr, 2011). Bohužel, výzkumy také naznačují, že taková spontánní zvědavá otázka žáků se ve třídách objevuje spíše zřídka (Engel & Randall, 2009). Studium uživatelských (resp. žákovských) dotazů je proto důležité, aby pedagogové porozuměli, na co se žáci ptají, jaké mají mezery v porozumění a jak je efektivně vést k hlubšímu učení.

S rozvojem technologií se oblast komunikace rozšířila o interakci člověka s počítačem. Na rozdíl od mezilidské komunikace, která

zahrnuje bohaté kontextové a neverbální signály, komunikace s počítačem tradičně probíhala prostřednictvím formálních příkazů nebo strukturovaných dotazů. Studium uživatelských dotazů v kontextu lidsko-počítačové interakce (Human–Computer Interaction, HCI) si klade za cíl přiblížit tuto komunikaci co nejvíce přirozené konverzaci. Díky pokroku v oblasti zpracování přirozeného jazyka dnes mohou uživatelé komunikovat s počítači *přirozeným jazykem*, podobně jako by hovořili s člověkem. To má značný didaktický význam – technologie se stávají dostupnější i pro uživatele, kteří nejsou technicky zaměřeni, a mohou sloužit jako nástroje pro učení.

Je však zapotřebí porozumět specifikům takové komunikace. Uživatelé často přizpůsobují způsob kladení dotazů počítači; očekávají například přesné a rychlé odpovědi. Výzkumy poukazují na odlišné komunikační zvyklosti – například **studenti očekávají od počítačových tutorů přesnost spíše než společenskou zdvořilost**, kterou by naopak dodržovali lidský učitel z ohledu na sebedůvěru žáka (Person et al., 1995; Lepper & Woolverton, 2002). Pro pedagogy i designéry výukových software je důležité tato specifika znát. Studium uživatelských dotazů v HCI umožňuje navrhovat takové rozhraní a programy, které lépe rozumějí potřebám uživatele a reagují vhodným způsobem. Z pedagogického hlediska sem spadá i rozvoj tzv. **digitální gramotnosti** – tedy schopnosti efektivně komunikovat s digitálními nástroji. Učit žáky, jak klást srozumitelné dotazy vyhledávačům nebo vzdělávacím aplikacím, je novou výzvou didaktiky v 21. století (Redecker, 2017).

Dialogové modely umělé inteligence představují systémy (chatboti, virtuální asistenti, inteligentní tutorové), které zvládají vést konverzaci s uživatelem v přirozeném jazyce. Vzdělávání těží z těchto pokročilých nástrojů, neboť umožňují *individualizovanou komunikaci*: každý žák se může ptát vlastním tempem a dostávat okamžité odpovědi či rady. Je však zásadní zkoumat, **jakým způsobem uživatelé formulují své dotazy** a jak na ně dialogové modely reagují. Takový výzkum odhaluje silné a slabé stránky těchto systémů – například kde

dochází k nepochopení otázky nebo poskytnutí nesprávné odpovědi – a poskytuje podklady pro zlepšování jak algoritmů, tak metodické podpory uživatelů.

Příkladem přínosu studia dialogových interakcí je vývoj inteligentních učicích systémů typu *AutoTutor*. Tento systém simuluje rozhovor žáka s lidským tutorem a průběžně reaguje na dotazy a odpovědi studenta. Výsledky ukazují, že taková dialogová forma výuky může významně zvýšit efektivitu učení – AutoTutor v opakovaných studiích dosáhl zlepšení studijních výsledků o 0,3 až 0,8 sigma (tj. zhruba 30–80 % standardní odchylky) ve srovnání se samostatným studiem textu (Graesser et al., 2016). Dialogový model zde umožňuje přizpůsobit vysvětlení potřebám jednotlivce, poskytnout okamžitou zpětnou vazbu a napodobit sokratovský dialog, který stimuluje myšlení studenta. Podobně pedagogičtí agenti – animované postavy provázející výukou – využívají verbální i neverbální komunikaci k navázání vztahu se studentem a zvýšení jeho motivace (Johnson et al., 2000; Porayska-Pomsta et al., 2013).

Od konce roku 2022 zažíváme nástup velkých jazykových modelů (LLM) jako jsou ChatGPT, Claude a Gemini, které dokážou vést poměrně sofistikovanou konverzaci (Mollick, 2024; OpenAI, 2025) na libovolné téma. Jejich využití ve vzdělávání slibuje například nepřetržitou dostupnost konzultací – student se může zeptat na doplňující vysvětlení učiva kdykoli potřebuje. Studie zmiňují, že chatboti mohou odpovídat na otázky studentů, poskytovat návod k úkolům a podporu mimo vyučování (Atlas, 2023). Zároveň však odborníci upozorňují, že pro efektivní nasazení těchto AI nástrojů je třeba rozvíjet nové kompetence na straně učitelů i žáků (Kasneci et al., 2023). Uživatelé se musí naučit pokládat promyšlené dotazy (tzv. prompt engineering) a kriticky hodnotit získané odpovědi, zatímco učitelé potřebují porozumět možnostem a limitům těchto systémů, aby je dokázali smysluplně začlenit do výuky (Kasneci et al., 2023). Studium interakcí mezi studenty a dialogovými modely zde poskytuje

cenné poznatky – například jaké typy dotazů vedou k užitečným odpovědím a které naopak způsobují misinformace nebo zmatení.

Komunikace a její vliv na výsledek je tedy ústředním bodem jak tradiční pedagogiky, tak edukativních technologií. Z pedagogického hlediska nám výzkum komunikace – od třídního dialogu po konverzaci s umělou inteligencí – umožňuje lépe porozumět procesu učení a vytvářet takové výukové postupy a nástroje, které podporují aktivní zapojení a porozumění. Studium uživatelských dotazů a dialogových modelů je důležité, protože spojuje to nejlepší z obou světů: pochopení lidské komunikace a využití moderních technologií k jejímu rozvoji. Tento mezioborový přístup napomáhá zlepšovat výukovou praxi i vývoj inteligentních systémů tak, aby výsledkem bylo účinné a smysluplné učení.

Tato monografie postupuje od základů ke specializovaným tématům ve dvou propojených rovinách. **Praktická rovina** (kapitoly 1–4, 7–8) představuje soustavně budovaný aparát komunikačních dovedností pro práci s generativní AI: vymezuje klíčové pojmy, charakterizuje devět dimenzí kvalitní AI komunikace, představuje osm strategií efektivního dotazování a uzavírá ji rozбором vzorců promptů a aktivačních sloves. Tato část je primárně určena pedagogům, studentům a profesionálům, kteří hledají systematický průvodce praxí.

Teoreticko-empirická rovina (kapitoly 5–6) přesahuje rámec praktického průvodce a představuje vlastní vědecký přínos monografie. Kapitola 5 zavádí **rámec MLEF-AI** (*Multi-Level Ecological Framework for AI Implementation in Education*) z autorovy habilitační teze (Hanák, 2026) jako primární teoretickou kotvu výzkumu, doplňuje jej o **AI-centrickou pedagogiku (AICP)** z autorovy předchozí monografie (Hanák, 2025) jako kontextový rámec a operacionalizuje pojem **komunikační zralost uživatele** s návrhem nového měřicího nástroje *Communication Maturity Index* (CMI). Kapitola 6 prezentuje empirickou retrospektivní observační studii 147 uživatelů edukační

platformy Kurzologie.cz, která testuje výzkumné otázky vyvozené z teoretického rámce a generuje hypotézy pro plánovanou Fázi 2 výzkumu.

Empirická studie v Kapitole 6 je publikována jako **interní pracovní výstup** dokumentující průběžnou fázi výzkumu. Vychází z rekonstruovaných dat získaných v rámci provozu edukační platformy mezi březnem 2025 a lednem 2026 a přiznává otevřeně limity vyplývající z této rekonstruktivní povahy. Hlavní zjištění studie — silný a robustní vztah mezi dokončeností kurzu a deklarovanou změnou v interakci s generativní AI — slouží jako východisko pro řádný prospektivní výzkum, jehož metodologie je v Kapitole 5 (oddíl 5.4) a v Kapitole 6 (oddíl 6.5) předregistrována. Tato monografie tedy nestojí jako definitivní empirické tvrzení, ale jako **teoreticky ukotvený most mezi praktickými kompetencemi a budoucím výzkumem komunikační zralosti uživatelů AI ve vzdělávacím kontextu**.

1 Základní přehled pro rychlou práci i orientaci

Umělá inteligence (AI) se stala nedílnou součástí současného života, od zjednodušení běžných činností až po řešení komplexních vědeckých úloh (Russell, Norvig, 2021). Ve vzdělávacím kontextu představuje generativní AI nástroj, jehož účinnost je podmíněna kvalitou a přesností interakce mezi uživatelem a systémem. Aby bylo možné efektivně číst a využívat tuto publikaci, je nezbytné nejprve vymezit základní pojmy a principy AI.

Cílem této monografie je poskytnout uceleného průvodce, který umožní hlubší porozumění generativním modelům AI a jejich aplikacím ve vzdělávání. Snahou je eliminovat potřebu vyhledávat definice a výklady v externích, potenciálně nespolehlivých zdrojích. Text proto obsahuje jasné a odborně podložené vysvětlení klíčových pojmů, čímž vytváří pevný základ pro orientaci v problematice komunikace s AI, a to z hlediska didaktiky i praktické aplikace.

V oblasti vzdělávání má komunikace zásadní vliv na kvalitu dosažených výsledků. Podle Hattieho (2012) efektivní interakce mezi učitelem a žákem zásadně ovlivňuje míru porozumění a schopnost aplikovat získané poznatky. V kontextu AI hraje tuto roli uživatelův dotaz – tzv. **prompt**, který určuje směr a kvalitu výstupu. Prompt lze chápat jako formu instrukce, která v digitální komunikaci nahrazuje verbální výuku. Stejně jako v tradičním pedagogickém prostředí i zde platí, že jasné, strukturované a kontextově bohaté zadání vede k vyšší přesnosti a relevantnosti odpovědi (Liu et al., 2023).

Co je to prompt?

V oblasti generativních AI modelů (např. ChatGPT, Gemini, Copilot) označuje prompt textový vstup – pokyn nebo otázku – na jehož základě model generuje odpověď. Prompt může mít jednoduchou formu („Jaké je hlavní město Francie?“) nebo být složitější, zahrnující specifický kontext a omezení („Napište krátký příběh ve stylu sci-fi odehrávající se na Marsu“). Z pedagogického hlediska je prompt

analogií k učitelské otázce či úkolu, který má vést k určitému kognitivnímu výkonu žáka (Brookhart, 2010).

Pokud bychom tento pojem měli vysvětlit co nejjednodušeji, mohli bychom říci: „Prompt je otázka.“ Přesnost, s jakou ji položíme, určuje přesnost odpovědi – a tím i kvalitu výsledného učení.

V praxi to znamená umění i vědu, jak *správně položit AI otázku* – tedy jaké zvolit formulace, formát a detaily promptu – aby výsledná odpověď byla co nejpřesnější, relevantní a užitečná. Prompt engineering zahrnuje různé techniky a strategie **tvorby promptů a často vyžaduje kombinaci lingvistických schopností a kreativního myšlení**. Efektivně navržené prompty dokážou výrazně zlepšit výstupy AI modelů a proměnit je v užitečné nástroje – například dobře navržený prompt může z chatbota učinit mnohem spolehlivějšího a uživatelsky přívětivějšího asistenta. Tato dovednost nabývá na důležitosti zvláště v oblastech, kde AI asistuje při rozhodování, v tvůrčích činnostech, ve vzdělávání, výzkumu apod., neboť v těchto kontextech je kvalita generované odpovědi kriticky významná.

Co je to Prompt Engineering?

Prompt Engineering představuje proces cíleného navrhování a optimalizace promptů pro interakci s jazykovými modely. Tato disciplína spojuje technické i pedagogické principy s cílem zajistit co nejpřesnější a nejrelevantnější odpovědi. Její význam roste zejména v oblastech, kde AI podporuje rozhodování, kreativní procesy a vzdělávací činnost (White et al., 2023).

Mezi klíčové aspekty Prompt Engineeringu patří:

- **Jasná formulace dotazu:** Prompt by měl být formulován srozumitelně a jednoznačně, aby model přesně pochopil, na co se ptáme. Vyplatí se **explicitně vymezit, jakou odpověď očekáváme**, čímž předejdeme možným nedorozuměním. Například pokud chceme po AI *souhrn* textu, měli bychom to v

promptu jasně uvést, aby model neprodukoval příliš detailní analýzu (Liu et al., 2023). Stručně řečeno – **konkrétnost a jednoznačnost** zadání výrazně zvyšují šanci na kvalitní výsledek.

- **Kontext a specifičnost:** Čím více relevantních informací a kontextu modelu v promptu poskytneme, tím lépe. Uvedení kontextu (např. role, ve které má model vystupovat, cíle uživatele, odborná úroveň textu apod.) pomáhá AI odpovědět přesněji a vhodněji (Brown et al., 2020). Specifičnost zadání zahrnuje také požadavek na formát nebo rozsah výstupu. Pokud například potřebujeme odpověď ve formě seznamu či tabulky, měli bychom to explicitně zahrnout do promptu. Dobře definovaný kontext v promptu usměrní model, aby generoval odpověď relevantní dané situaci a požadavkům.
- **Využití vzorců a strategií:** Zkušený tvůrce promptů často využívá osvědčené **strategické postupy** pro získání lepších výsledků. Patří sem například poskytnutí *příkladů* v promptu (tzv. *few-shot prompting*), rozložení složitého dotazu na dílčí kroky či otázky, nebo určování role, v níž má AI odpovídat (např. „Představ si, že jsi lékař...“). Existují specifické techniky jako *chain-of-thought prompting* (řetězení myšlenek), kdy model postupuje krok za krokem a své uvažování si zaznamenává, což prokazatelně zlepšuje schopnost řešit komplexní úlohy. Využívání takovýchto vzorců může vést k **logičtějším, konzistentnějším a přesnějším** odpovědím AI.
- **Testování a iterace:** Návrh promptu je **iterativní proces**. V praxi to znamená, že prvotní formulace promptu často nevede hned k ideálnímu výsledku – je běžné prompt několikrát upravit na základě pozorované odpovědi modelu. Experimentování s různými zněními a strukturami promptu a následné doladování podle toho, jak AI reaguje, je zásadní pro dosažení optimálního výstupu (Reynolds, McDonell, 2021). Tato zpětnovazební smyčka (zadání → výstup → úprava zadání) umožňuje postupně vypilovat prompt tak, aby **maximalizoval přesnost a relevanci** odpovědi.

- **Kreativita a inovace:** Kromě metodických postupů hraje v prompt engineeringu roli také **kreativní myšlení**. Formulování opravdu užitečných dotazů někdy vyžaduje nápaditý přístup – například zkusit netradiční úhel pohledu, originální příměr nebo specifický styl vyjadřování, který může AI inspirovat k zajímavější odpovědi. Dobří *prompt inženýři* kombinují systematický přístup s kreativitou; nebojí se experimentovat s různými formáty a formulacemi, aby objevili, co funguje nejlépe. Koneckonců, generativní modely samy v sobě skrývají prvek kreativity – a vhodně konstruované prompty dokážou tuto jejich schopnost plně rozvinout.

Stejně jako v tradičním vzdělávání i zde platí, že způsob, jakým položíme otázku, zásadně ovlivňuje odpověď a tím i proces učení. Ve zkratce, *prompt engineering* je o tom, jak **mluvit s AI jazykem, kterému porozumí**, aby nám poskytla co nejlepší možnou odpověď.

V dalších kapitolách této publikace budeme detailně zkoumat, *jaké strategie při tvorbě promptů fungují*, ukážeme si praktické příklady efektivních promptů a upozorníme i na možná úskalí. Vybaveni znalostí základních pojmů a principů se nyní můžeme s jistotou pustit do objevování světa generativní AI a **efektivní komunikace s ní**.

Základním předpokladem úspěšné práce s AI je totiž porozumění tomu, jak správně formulovat své požadavky – a přesně tomu se budeme věnovat na následujících stránkách.

2 Udržování pořádku a struktury v komunikaci s AI

Historický vývoj komunikace s počítači: V počátcích vývoje počítačů byla komunikace s nimi velmi formální a odlišná od dnešního dialogového stylu. Programátoři museli dodržovat přísná pravidla syntaxe – jazyky jako FORTRAN a COBOL vyžadovaly naprostou přesnost v zápisu každého příkazu. Každý krok algoritmu musel být explicitně definován, což kladlo velké nároky na detailní porozumění fungování počítačového hardwaru a nízkourovňových operací. Například ve starším jazyce Fortran bylo nutné přesně dodržet syntaxi při indexování polí – zápis $4*1$ byl platný, zatímco $1*4$ vedl k chybě. Tato striktní syntaxe znamenala, že i drobná odchylka či překlep mohly způsobit selhání programu (Malik, 2000). V důsledku toho se kladl důraz na pořádek v kódu a programátorskou disciplínu: začínající programátoři se museli naučit psát kód bezchybně a logicky, protože počítač nerozuměl ničemu, co nebylo přesně podle pravidel.

S nástupem objektově orientovaného programování (OOP) v 80. a 90. letech se komunikace s počítačem posunula od ryze procedurálního psaní kódu k modelování světa pomocí objektů. Programovací jazyky jako C++ a Java umožnily programátorům definovat objekty a jejich vzájemné interakce, což přineslo větší přirozenost a srozumitelnost do tvorby programů. Syntaxe samozřejmě zůstávala důležitá, avšak těžiště se přesunulo na strukturu programu a na vztahy mezi jeho částmi. Cílem bylo, aby kód lépe odrážel reálný svět – jak poznamenal Niklaus Wirth, objektový přístup „věrně odráží strukturu systémů ve skutečném světě a je proto vhodný k modelování složitých systémů“ (Wirth, 2006). Jinými slovy, programy se stavěly tak, aby odpovídaly konceptům z praxe (např. objekty jako Zákazník, Objednávka v obchodní aplikaci), což usnadnilo pochopení kódu i jeho údržbu. Tento posun zdůraznil pořádek ve struktuře: důležité bylo navrhnout čistou architekturu a dodržovat osvědčené principy (např. zapouzdření, dědičnost). Formální syntaxe již nebyla tak neúprosná

jako u Fortranu či COBOLu, avšak bez dobré struktury mohl i flexibilnější kód být “nepořádný” a hůře udržitelný. Programátor tedy získal volnost, ale zároveň musel dbát na pořádek v návrhu – učil se přemýšlet v pojmech objektů a jejich vztahů, což představovalo novou úroveň abstrakce v komunikaci s počítačem.

V éře umělé inteligence a pokročilých jazykových modelů (LLM) — od OpenAI GPT-5 přes Anthropic Claude až po Google Gemini 3 (OpenAI, 2026; Google DeepMind, 2026; Anthropic, 2025) — dochází k dalšímu výraznému posunu ve způsobu, jak lidé s počítači komunikují. Tyto modely jsou schopny porozumět a reagovat na přirozený jazyk, což činí interakci s nimi mnohem intuitivnější a flexibilnější než kdy dříve. Uživatelé již nemusejí formulovat dotazy striktně formálně – místo psaní kódu kladou otázky či pokyny v běžné řeči, podobně jako by mluvili s člověkem. Modely navržené pro dialog (jako je ChatGPT) dokážou zpracovat kontext konverzace a záměr uživatele a na tomto základě generovat smysluplné a koherentní odpovědi. Jedním z významných pokroků u moderních AI je právě schopnost rozumět kontextu v textové konverzaci a zohlednit význam slov či frází při formulaci odpovědi (Ray, 2023). Díky tomu jsou interakce s AI přirozenější a více připomínají dialog než technické programování. Například moderní multimodální modely (například GPT-5 nebo Gemini 3) přijímají vstupy v podobě textu, obrazu, audia a videa a na základě svého tréninku na obrovském korpusu dat dokáže odpovídat velmi plynule a „lidsky“. Je schopen pochopit i neformálně položené otázky, dedukovat z kontextu a odpovídat s ohledem na uživatelské instrukce. To vše znamená, že bariéra komunikace mezi člověkem a počítačem se opět zmenšila – místo precizního kódu můžeme použít přirozený jazyk.

Přestože moderní jazykové modely (například GPT-5, Claude 4 nebo Gemini 3) nevyžadují striktní dodržování programovací syntaxe, je stále zásadní zachovávat určitý řád a strukturu v tom, jak s nimi komunikujeme. Volnost přirozeného jazyka totiž neznamená, že na formě dotazu nezáleží. Naopak – formulace vstupu silně ovlivňuje

kvalitu a relevanci generované odpovědi. Uživatelé proto musejí strategicky přemýšlet nad svými dotazy (prompty), mají-li od AI získat konzistentní a užitečné výstupy. Efektivní komunikace s těmito modely se podobá spíše dialogu s člověkem, avšak uživatel zde hraje roli tazatele i „instruktora“, který udává mantinely odpovědi. Má možnost experimentovat s různými způsoby formulace – a právě způsob, jakým otázku či úkol zadá, často rozhoduje o tom, jak kvalitní odpověď AI vrátí.

Klíčem k úspěchu je pochopit a naučit se, jak nejlépe formulovat dotazy, aby odpovědi co nejvíce vyhovovaly potřebám uživatele. Tato dovednost se dnes označuje jako *prompt engineering* – neboli umění sestavit vstup pro AI tak, aby přinesl požadovaný výstup. Například Cain (2023) označuje *prompt engineering* za „nově vznikající klíčovou dovednost“ v práci s generativní AI, a další autoři si všímají, že se rychle stává kritickou kompetencí pro studenty i profesionály (Raftery, 2023; Spasić & Janković, 2023). Již před masivním rozšířením těchto modelů poznamenali Hocky a White (2022), že „*prompt engineering, doslova učení se jasněji rozhraní s AI, je nyní předmětem výzkumu*“, což dokládá význam správné komunikace s AI. Úspěšné využití generativních modelů tedy **není jen otázkou technologie na straně AI, ale také dovedností na straně uživatele** – špatně položený nebo neurčitý dotaz může vést k nerelevantní či nesprávné odpovědi (princip *garbage in, garbage out* platí i zde). Naopak pečlivě strukturovaný a promyšlený prompt pomáhá AI lépe porozumět našemu záměru.

V praxi se osvědčuje udržovat při formulaci promptů určitý pořádek a systematickост. Jak bylo rozebráno v předchozích kapitolách, existují osvědčené strategie a vzorce pro psaní promptů, které mohou výrazně zlepšit kvalitu odpovědí AI. Tyto postupy nám pomáhají strukturovat naše dotazy tak, aby pro model byly co nejsrozumitelnější a jednoznačné. Patří sem například zásady jako *být konkrétní*, poskytnout potřebný kontext, jednoznačně formulovat požadovanou úlohu či výstup a rozdělit komplexní problém na menší

kroky (OpenAI, 2023). Příprava promptu “offline” – například v textovém editoru – umožňuje promyslet strukturu dotazu, použít členění (odrážky, číslované kroky, oddělení instrukcí a vstupních dat) a případně využít formátování pro zvýraznění klíčových částí. Někteří experti doporučují umístit hlavní instrukce na začátek promptu a oddělit je od kontextu například symboly #### pro přehlednost. Takový strukturovaný přístup usnadňuje jak revizi dotazu, tak i opakované experimenty s jeho úpravami.

Je užitečné si uchovávat záznamy různých verzí promptů a výsledků, které AI vrátí. Pomocí textového editoru můžeme vést jakýsi deník promptů, kde si poznamenáme, co jsme zkusili a s jakým efektem. Tento postup podporuje systematické zlepšování dotazů – postupnou úpravou formulací a porovnáváním odpovědí můžeme vypořádat, co funguje nejlépe. V přehledně vedených poznámkách lze také zvýraznit různé části promptu (např. barevně odlišit instrukce, kontext, příklady výstupu), což pomáhá udržet pořádek v myšlenkách při složitějších úkolech. Výsledný vyladěný prompt pak stačí zkopírovat do rozhraní AI (například ChatGPT, Claude nebo Gemini) a získat požadovanou odpověď. Tento iterativní a strukturovaný proces tvorby promptu dává uživateli větší kontrolu nad tím, co od AI získá, a zároveň zvyšuje opakovatelnou úspěšnost – když se k obdobnému problému vrátí později, může vycházet z osvědčených formulací.

Z pedagogického hlediska je zmíněný posun ve způsobu komunikace s počítači velmi významný. Znalost práce s AI prostřednictvím promptů se stává důležitou součástí digitální gramotnosti pro širší odbornou veřejnost. Zatímco dříve se v kurzech programování kladl důraz na syntaktickou správnost kódu nebo na návrh softwarové architektury, dnes se objevuje potřeba učit také efektivní interakci s AI. Aktuální výzkumy v oblasti vzdělávání potvrzují, že schopnost formulovat dobré prompty může zásadně ovlivnit kvalitu využití AI nástrojů a že tuto dovednost je možné a nutné rozvíjet (Lee & Palmer, 2025). Dokonce se začínají objevovat rámce a metodiky, jak prompt engineering vyučovat – od krátkých praktických kurzů pro studenty až

po začlenění do vysokoškolských osnov (Raftery, 2023; Liu, 2023). Dobře navržené prompty mají potenciál proměnit interakci s AI v efektivní nástroj učení i práce. Například přehledová studie Lee a Palmer (2025) zdůrazňuje, že rozvoj praktických dovedností v oblasti komunikace s AI (včetně smysluplného tvoření promptů) by měl být podporován pomocí promyšlených výukových rámců sladěných s cíli daného oboru. Walter (2024), v navazující systematické přehledové studii Walsh (2025), a Tuomi (2025) hovoří o AI gramotnosti jako o nezbytnosti moderního vzdělávání a řadí umění práce s prompty po bok kritického myšlení jako kompetence pro 21. století. To vše ukazuje, že udržování pořádku a struktury v komunikaci s AI není jen technickým detailem, ale stává se široce uznávanou součástí odborné kvalifikace napříč obory.

Pro lepší ilustraci vývoje komunikace s počítači a změny důrazu na strukturu lze uvést srovnání jednotlivých etap:

Éra vývoje	Způsob komunikace s počítačem	Klíčové požadavky a dovednosti
Rané programování (50.–70. léta)	Psaní formálního kódu v jazycích jako Fortran, COBOL	Přesné dodržení syntaxe; explicitní definování každého kroku algoritmu; detailní pochopení práce počítače na nízké úrovni.
Strukturované a OOP programování (80.–90. léta)	Definování strukturovaných bloků a objektů (C, Pascal; C++, Java)	Důraz na architekturu programu a modulárnost; modelování reálných entit v kódu (objekty); potřeba dobrých návrhových principů vedle syntaxe.

Interakce s AI modely (2020s)	Dialog v přirozeném jazyce (prompty pro modely GPT aj.)	Jasná formulace záměru a kontextu v dotazu; schopnost strukturovat prompt pro jednoznačné porozumění; iterativní vylepšování dotazu na základě výstupů AI.
--------------------------------------	---	--

Jak je patrné z výše uvedené tabulky, požadované dovednosti se postupně posunuly: od precizního psaní kódu přes schopnost abstraktně navrhovat struktury až po umění efektivně komunikovat s inteligentními systémy. Tato evoluce v komunikaci s počítači zdůrazňuje důležitost adaptace a kontinuálního učení novým způsobům interakce v rychle se vyvíjejícím světě technologií. Pořádek a struktura přitom zůstávají průběžným motivem: ať už šlo o pořádek v syntaxi kódu, v architektuře programu nebo dnes v pečlivě připraveném promptu pro AI, vždy platí, že uspořádaná a promyšlená komunikace vede k lepším výsledkům. Moderní AI nám sice odpustí překlepy a zvládne porozumět i neformálnímu jazyku, ale **kvalitu jejích odpovědí řídí kvalita našich dotazů**. Proto je na nás, abychom k interakci s AI přistupovali poučeně: využívali osvědčené postupy, udržovali strukturu a jasnost výchozích promptů a tím maximalizovali přínos těchto mocných nástrojů pro naši práci i učení.

3 Jak se bavit s umělou inteligencí?

Způsob komunikace s umělou inteligencí (AI) by měl být volen s ohledem na kontext dotazu a osobní preference uživatele. Moderní AI systémy (například konverzační modely typu ChatGPT, Claude nebo Gemini) jsou navrženy tak, aby porozuměly různým formám vyjadřování a dokázaly se jim přizpůsobit. Díky tréninku na rozsáhlých datových souborech dokážou tyto modely reagovat na formální i neformální oslovení, různé styly a tóny řeči (Brown et al., 2020). Tato flexibilita výrazně rozšiřuje možnosti jejich využití a činí interakci uživatelsky přívětivější – uživatel nemusí striktně dodržovat formální jazyková pravidla a může komunikovat přirozeně. Schopnost AI přijímat a efektivně zpracovávat různé komunikační styly z ní činí mimořádně adaptabilní nástroj, využitelný v mnoha různých situacích.

Komunikace je však jen jedním z aspektů interakce člověka s AI. Celková akceptace AI ze strany uživatelů závisí na širší paletě faktorů, které ovlivňují, jak je AI vnímána a používána (Rane et al., 2024). Při zvažování nasazení AI je třeba zohlednit rozdílné perspektivy dvou hlavních aktérů – na jedné straně společnosti (tvůrců a poskytovatelů AI nástrojů) a na druhé straně koncových uživatelů. Následující přehled shrnuje klíčové aspekty ovlivňující přijetí AI, vždy z obou uvedených pohledů, doplněné o příklady praktických aplikací.

3.1 Uživatelská přívětivost

Pohled společnosti (vývojářů a poskytovatelů):

1. **Intuitivní rozhraní:** Uživatelské rozhraní AI nástroje by mělo být navrženo maximálně intuitivně a srozumitelně. To umožní uživatelům začít systém používat efektivně i bez rozsáhlého školení či technických znalostí. Výzkumy v oblasti přijímání technologií dlouhodobě ukazují, že vnímaná snadnost použití patří mezi klíčové předpoklady akceptace nových systémů (Davis, 1989). Návrh rozhraní podle zavedených principů použitelnosti

(např. jasné ovládací prvky, konzistentní design) tak přímo podporuje ochotu uživatelů AI využívat. Například **Nielsenova heuristika použitelnosti** zdůrazňuje, že systém má uživatele vést přirozeně jejich úkolem a neměl by je zahltit nadbytečnou komplexitou (Nielsen, 1993).

2. **Přízpůsobení a personalizace:** AI by měla být schopna přizpůsobit se různým potřebám a preferencím uživatelů. To zahrnuje schopnost *učit se* z interakcí s uživatelem a upravovat své odpovědi či chování podle specifických požadavků. Personalizované uživatelské zážitky zvyšují spokojenost a angažovanost zákazníků, což potvrzují i empirické studie – například v oblasti smart healthcare byly **důvěra uživatelů a míra personalizace služby** identifikovány jako významné faktory ovlivňující veřejné přijetí AI nástrojů (Liu & Tao, 2022). Pro vývojáře AI to znamená navrhovat systémy, které dokážou reagovat na individuální vstupy a postupně vyladit svůj výkon podle preferencí uživatele.
3. **Přístupnost a inkluzivita:** AI produkty by měly být navrženy tak, aby byly přístupné co nejširší škále uživatelů, včetně osob se specifickými potřebami. To zahrnuje podporu technologií pro handicapované (např. hlasové ovládání pro zrakově postižené, titulkování pro sluchově postižené) a obecně **univerzální design** služeb. Inkluzivní přístup nejen rozšiřuje okruh potenciálních uživatelů, ale také zvyšuje společenskou hodnotu nabízeného AI řešení. Firmy, které dbají na přístupnost, mohou snáze získat důvěru veřejnosti a vyhnout se kritice, že jejich technologie zvýhodňuje jen určitou skupinu lidí.
4. **Jazyková a kulturní citlivost:** Na globálním trhu je zásadní, aby AI systémy dokázaly komunikovat v různých jazycích a zohledňovat kulturní odlišnosti. Pro společnosti to znamená investovat do lokalizace obsahu a trénování modelů na multijazyčných datech. AI nástroj, který rozumí lokálnímu kontextu a idiomatice jazyka

uživatele, poskytne relevantnější a přirozenější odezvy, což opět zvyšuje spokojenost. Například virtuální asistent vyvinutý s porozuměním místním reáliím (svátky, měnové formáty, společenské zvyklosti) bude působit mnohem přívětivěji než univerzální anglická verze téhož asistenta.

Pohled uživatele (konzumenta AI služby):

1. **Snadné a přímočaré používání:** Uživatelé očekávají, že interakce s AI bude co možná nejjednodušší. Technologie, která vyžaduje od běžného uživatele složité nastavování nebo technické znalosti, bude jen obtížně hledat široké uplatnění. Naopak, pokud je používání AI *přímocharé*, roste u uživatelů vnímaná užitečnost a ochota takovou technologii přijmout (Venkatesh et al., 2003). Intuitivní ovládání a srozumitelné instrukce proto výrazně zvyšují šanci na pozitivní uživatelskou zkušenost.
2. **Rychlá a přesná odezva:** Uživatelé od AI požadují, aby na jejich dotazy či příkazy reagovala pohotově a spolehlivě. Očekávají nejen rychlost, ale i věcnou správnost odpovědi – AI musí poskytovat relevantní informace či řešení jejich potřeb. Výhoda AI spočívá právě ve schopnosti zpracovat velké objemy dat v krátkém čase a nabídnout výstup prakticky okamžitě. Studie zaměřené na přijímání AI potvrzují, že **výkon a spolehlivost systému** významně ovlivňují postoj uživatelů: přesnost odpovědi a schopnost dosahovat konzistentně správných výsledků posilují u uživatele důvěru a spokojenost (Jiang et al., 2024). Naopak pomalé nebo nepřesné reakce mohou vést k frustraci a odmítnutí technologie.
3. **Důvěra a bezpečnost:** Pro uživatele je klíčové, aby mohli AI důvěřovat – nejen z hlediska správnosti výsledků, ale také ohledně zacházení s jejich daty a soukromím. Uživatelé musí mít pocit, že AI nástroj je bezpečný a nezneužije poskytnuté informace. Pokud má systém přístup k citlivým údajům (např. finanční data, osobní poznámky, zdravotní informace), očekává

se, že budou adekvátně chráněny. **Důvěra uživatelů v AI** je křehká – jediné selhání, únik dat či výrazná chyba mohou reputaci technologie výrazně poškodit. Proto platí, že čím transparentněji a odpovědněji AI nakládá s daty a čím konzistentněji se chová, tím více jsou jí uživatelé ochotni svěřit (Shin et al., 2022). Například šifrování komunikace, anonymizace osobních údajů a jasná pravidla ochrany dat jsou dnes minimem pro získání důvěry veřejnosti.

4. **Podpora a zdroje:** Uživatelé oceňují, když mají k dispozici snadno dostupné informace a podporu pro využívání AI nástroje. To zahrnuje přehledné návody, nápovědu zabudovanou přímo v rozhraní, často kladené dotazy (FAQ) či komunitní fóra, kde mohou hledat rady. Pokud nastanou potíže nebo nejasnosti, uživatel očekává rychlou pomoc – ať už formou automatizovaného průvodce, či možnosti obrátit se na živou podporu. Dostupnost podpůrných materiálů a služeb významně ovlivňuje ochotu uživatelů technologii dlouhodobě využívat. V rámci modelů přijímání technologií se tyto faktory označují jako **podpůrné podmínky** (angl. *facilitating conditions*) a prokazatelně usnadňují přechod uživatele na novou technologii (Venkatesh et al., 2003). Uživatel jednoduše potřebuje vědět, že na využití AI „není sám“ a v případě problémů najde oporu.

Příklady:

AI v zákaznické podpoře: Nasazení chatovacích chatbotů v zákaznických centrech je dnes rozšířeným příkladem uživatelsky přívětivé AI. Z pohledu firmy takový AI chatbot **snižuje náklady** a zvyšuje efektivitu – dokáže obsloužit mnoho dotazů současně a být k dispozici 24/7, čímž odlehčuje vytížení lidských operátorů. Odhaduje se, že využití AI asistentů může firmám snížit výdaje na zákaznickou podporu až o 30 % (Jobanputra, 2024). Z pohledu uživatele představuje chatbot rychlý a pohodlný způsob, jak získat potřebné informace či vyřešit problém bez dlouhého čekání na lince. Studie

ukazují, že pokud je chatbot navržen s lidskými prvky (např. smysl pro humor, osobní oslovení, drobné konverzační „odmlky“ napodobující lidskou komunikaci), uživatelé se cítí komfortněji a více se mu svěří s dotazy či problémy (Schanke et al., 2021). Takový **“polidštěný” chatbot** přináší firmě spokojenější zákazníky a uživatelům pocit přirozené konverzace.

Personalizovaná doporučení: Mnohé e-commerce platformy a streamingové služby využívají AI k personalizaci obsahu – například doporučování produktů či filmů na míru konkrétnímu uživateli. Z pohledu firmy tato personalizace **zvyšuje obchodní úspěšnost**: firmy, které excelují v personalizovaných doporučeních, generují podstatně vyšší tržby z těchto aktivit než průměr (až o 40 % více podle analýzy společnosti McKinsey; McKinsey, 2021). AI analyzuje minulou interakci uživatele (historii nákupů, hodnocení filmů apod.) a na základě toho nabízí položky, které by ho mohly zajímat – tím dochází ke zvýšení pravděpodobnosti nákupu nebo udržení zákazníka. Z pohledu uživatele personalizovaný systém poskytuje **příjemnější zážitek**: uživatel snadno objeví nový relevantní obsah, aniž by musel zdlouhavě procházet katalogem. To šetří čas a zvyšuje hodnotu, kterou ze služby získává. V obou pohledech je klíčové, že AI se dokáže adaptovat na individuální preference – čím déle uživatel službu využívá, tím lépe se systém „naučí“, co konkrétnímu jedinci vyhovuje, a tím kvalitnější doporučení (nebo jiné personalizované funkce) poskytuje.

3.2 Přesnost a spolehlivost

Pohled společnosti (tvůrce a provozovatele AI nástroje):

1. **Validace a testování:** Pro vývojářské firmy je zásadní, aby AI systém procházel důkladným testováním a ověřováním, než je nasazen do ostrého provozu. Je třeba zajistit, že model poskytuje správné a konzistentní výsledky v široké škále scénářů. To obnáší testování AI na **rozmanitých datech** i v krajních případech použití. Robustní validační proces pomáhá odhalit slabiny

systému (např. situace, kdy může selhat či podat chybný výstup) dříve, než by tyto chyby ovlivnily reálné uživatele. Vysoká úroveň přesnosti není důležitá jen z hlediska uživatelské spokojenosti, ale také kvůli reputaci společnosti – systém, který vrací chybné výsledky, může vážně poškodit důvěryhodnost značky.

2. **Tréninkové datasety:** Kvalita a diverzita trénovacích dat jsou klíčové pro vytvoření přesného a spolehlivého AI modelu. Pokud jsou data neúplná, nekvalitní či zkreslená, může AI přejímat **předsudky a systematické chyby** obsažené v datech. Například analýza word2vec modelů ukázala, že AI trénovaná na reálných textových datech může převzít genderové stereotypy (model přiřadil analogii „muž: programátor = žena: hospodyně“), což reflektovalo zkreslení v trénovacím korpusu (Bolukbasi et al., 2016). Pro firmy je proto nezbytné pečlivě volit a čistit datasety, hlídat jejich vyváženost a reprezentativnost. V praxi se uplatňují techniky jako augmentace dat, odstraňování anomálií či **debiasing** (odstranění zaujatostí), které mají zlepšit obecnou schopnost AI správně generalizovat a minimalizovat chybovost.
3. **Aktualizace a údržba:** Aby si AI systém udržel vysokou přesnost i v dlouhodobém horizontu, vyžaduje průběžné aktualizace a údržbu. Svět se mění – objevují se nová data, mění se trendy, pravidla i vnější podmínky – a AI na to musí reagovat. Pro poskytovatele AI nástroje to znamená **pravidelně přetrénovávat model** s novými relevantními daty, ladit jeho parametry a monitorovat jeho výkon v čase. Například modely pro detekci spamů nebo malwaru musí být neustále aktualizovány, aby zachytily nové typy útoků. Bez takové péče může přesnost AI časem degradovat, což by vedlo ke snížení spolehlivosti výsledků. Důsledná údržba je tedy klíčem k dlouhodobé spolehlivosti systému.
4. **Transparentnost algoritmů:** Pro tvůrce AI je často důležité rozumět tomu, *jak a proč* AI dochází ke svým závěrům – tedy

zajistit jistou míru transparentnosti. Z technického hlediska mnoho moderních AI modelů (např. hluboké neuronové sítě) funguje jako tzv. *černé skříňky*, u nichž je obtížné vysvětlit interní rozhodovací proces. Přesto by se společnosti měly snažit implementovat mechanismy, které umožní výsledky AI interpretovat a auditovat. **Vysvětlitelná AI (XAI)** se stala aktivní oblastí výzkumu právě proto, aby vývojáři i uživatelé mohli lépe porozumět chování modelu (Doshi-Velez & Kim, 2017). Provozovatelé AI nástroje mají rovněž právní a etickou povinnost dohlížet na to, že algoritmus funguje korektně a nediskriminačně – v některých odvětvích (finance, zdravotnictví) či regionech (EU) dokonce platí regulační požadavky na určitý stupeň vysvětlitelnosti a transparentnosti automatizovaných rozhodnutí (Goodman & Flaxman, 2017). Otevřenost ohledně fungování AI (např. zveřejnění použité metodiky, zdrojových dat, nebo poskytnutí *audit trail* jednotlivých rozhodnutí) pomáhá předcházet nečekaným selháním a zvyšuje důvěryhodnost systému.

Pohled uživatele (konzumenta AI služby):

1. **Očekávání správných výsledků:** Když uživatel využívá AI, přirozeně předpokládá, že získá správné a kvalitní výstupy. Ať už jde o odpověď chatbotu, doporučení produktu nebo diagnostický závěr, uživatelé vstupují do interakce s očekáváním, že AI „ví, co dělá“. **Percepce vysoké přesnosti** posiluje u uživatele ochotu systém používat – pokud opakovaně dostává užitečné a bezchybné výsledky, upevňuje se jeho pozitivní postoj. Studie uvádějí, že právě výhody AI v oblasti přesnosti a spolehlivosti patří mezi hlavní faktory, které utvářejí kladné přijetí AI služeb z pohledu zákazníků (Jiang et al., 2024). Naopak, pokud AI generuje časté omyly nebo nekonzistentní odpovědi, uživatelé začnou její výstupy zpochybňovat a mohou se od jejího používání odklonit.

2. **Důvěra ve výsledky:** Spolehlivost AI systému je úzce spjata s důvěrou, kterou k němu uživatelé chovají. Důvěra je zde křehká – uživatel ji buduje postupně na základě zkušeností s předchozími interakcemi. Pokud AI dlouhodobě poskytuje přesné a rozumné výsledky, důvěra uživatele roste. Naopak jeden zásadní omyl může tuto důvěru výrazně narušit. Výzkumy v oblasti *trust in automation* ukázaly, že lidé mají tendenci opustit automatizovaný systém hned po prvních výrazných selháních, zatímco lidským expertům dokáží chyby spíše odpustit (Lee & See, 2004). To klade na AI vysoké nároky – musí dosahovat velmi nízké chybovosti, má-li si udržet přízeň publika. Z hlediska uživatele je důvěra posilována i tím, když chápe silné a slabé stránky systému (viz níže) a má možnost získat vysvětlení k výsledkům. Například je-li AI schopna objasnit, proč doporučila určité řešení (byť třeba zjednodušeně), uživatel bude jejím radám důvěřovat více než v případě zcela netransparentního „černého boxu“.
3. **Pochopení omezení:** Součástí pozitivního vztahu uživatele k AI je také realistické povědomí o tom, co AI umí a co naopak není v jejích možnostech. Uživatelé by měli být informováni o možných omezeních systému, aby nepřeceňovali jeho schopnosti ani nebyli překvapeni jeho selháním v nestandardních situacích. Výzkum naznačuje, že často vzniká *propast očekávání* – uživatelé mají tendenci přisuzovat AI větší inteligenci a schopnosti, než jaké ve skutečnosti má, což vede k nesplnitelným očekáváním a následnému zklamání (Luger & Sellen, 2016). Proto je pro firmy důležité komunikovat jasně rozsah funkcionalit a účel AI nástroje. Z perspektivy uživatele je pak důležité chápat, že AI může mít určité *doménové omezení* (např. lékařský diagnostický model je výborný v rozpoznání pneumonie z rentgenů plic, ale „neumí“ už nic jiného) a že stále může docházet k chybám či nejistotě v odpovědích. Pokud jsou uživatelé s těmito možnostmi a limity předem seznámeni, mohou AI využívat efektivněji a bezpečněji –

vědět, kdy se na ni spolehnout a kdy je na místě zvýšená opatrnost či verifikace výsledku jiným způsobem.

Příklady:

AI ve zdravotnictví: Systémy umělé inteligence pro lékařskou diagnostiku (např. analýzu rentgenových snímků nebo snímků z MRI) musí dosahovat **velmi vysoké přesnosti**, aby lékaři mohli důvěřovat jejich doporučením. V diagnostice platí, že chybný závěr může mít vážné následky pro pacienta. Například pokud by AI nástroj přehlédl nádor na rentgenovém snímku, mohl by pacient přijít o možnost včasné léčby. Z tohoto důvodu jsou diagnostické AI modely trénovány a testovány na obrovských množstvích medicínských dat a jejich výkon je srovnáván s odborníky. Studie z poslední doby dosahují v některých úlohách srovnatelné úspěšnosti jako specialisté (např. v detekci určitých typů rakoviny na snímcích sítnice oka či mamografu), což naznačuje potenciál AI asistence ve zdravotnictví (Topol, 2019). Přesto odborníci zdůrazňují, že i vysoce přesná AI musí fungovat **pod dohledem lékaře**, který zná kontext pacienta a dokáže posoudit, zda výstup AI dává smysl. Kliničtí lékaři tak stále nesou konečnou zodpovědnost a potřebují mít možnost výsledky AI ověřit či se podívat, z jakých dat AI vychází – jinými slovy, spolehlivost AI je předpokladem, ale nikoli jedinou podmínkou jejího využití v praxi.

AI ve finančních službách: V oblasti financí se AI využívá například pro automatizované obchodování na burze nebo pro analýzy tržních dat a navrhování investic (tzv. *robo-advisory* služby). Zde je **spolehlivost a stabilita AI** kriticky důležitá – algoritmus, který by chybně vyhodnotil situaci na trhu, může během zlomků sekundy učinit nesprávné obchody a způsobit značné finanční ztráty. Známé jsou případy takzvaných *flash crash* situací na burze, kdy automatizované systémy reagovaly na nečekané podněty řetězovou reakcí prodeje, což vedlo k dramatickému propadu cen v krátkém čase. Pro finanční firmy je proto nezbytné implementovat v AI systémech kontrolní mechanismy (např. omezovače příliš velkých transakcí, pravidla pro

kill switch v případě podezřelého chování), které zabrání katastrofickým scénářům při potenciálním selhání modelu. Uživatelé – ať už investoři využívající robo-poradce, nebo běžní klienti bank spoléhající na AI v mobilním bankovníctví – musejí mít jistotu, že systém je náležitě otestován a za normálních okolností bude fungovat konzistentně. Když AI ve financích doporučí klientovi investici, očekává se, že tak činí na základě solidních datových analýz a nikoli náhodně; jinak by klient rychle ztratil důvěru a službu přestal používat. Příkladem může být nasazení AI v úvěrovém scoringu: pokud by AI model omylem vyhodnocoval bonitní klienty jako rizikové kvůli chybě v datech či biasu v algoritmu, vedlo by to k nespokojenosti klientů a pravděpodobně i k regulatorním postihům. Proto se finanční instituce velmi pečlivě věnují testování těchto AI systémů, často využívají jednodušší a **vysvětlitelné modely** (např. rozhodovací stromy, kde lze snadno doložit, proč byl úvěr zamítnut) namísto nepřehledných „černých skříněk“, a průběžně monitorují jejich výkon.

3.3 Personalizace a adaptabilita

Pohled společnosti (tvůrce a prodejce AI nástroje):

1. **Personalizované uživatelské zážitky:** Schopnost AI poskytovat uživatelům individualizovaný zážitek je pro firmy vyvíjející tyto technologie vysoce ceněná. Personalizace obsahu a funkcí na míru konkrétnímu uživateli zvyšuje **uživatelskou spokojenost a loajalitu**, což se typicky promítá i do obchodních ukazatelů (např. delší doba strávená v aplikaci, vyšší míra konverze či opakovaných nákupů). Společnosti proto investují do AI systémů, které umějí analyzovat uživatelská data (historie interakcí, preference) a na jejich základě dynamicky přizpůsobit rozhraní či nabídku. Příkladem jsou doporučovací algoritmy e-shopů: AI se z předchozího chování zákazníka naučí, o jaké typy produktů má zájem, a na úvodní stránce mu příště nabídne novinky přesně z těchto kategorií. Uživatel tak dostává pocit, že služba **“rozumí jeho potřebám”**, což je pro celkový pozitivní dojem klíčové. Pro

tvůrce AI nástroje ovšem taková personalizace klade nároky na **zpracování dat s ohledem na soukromí** – je nutné pečlivě vyvažovat využití osobních údajů pro personalizaci s respektováním zákonných regulací a etických zásad (Liu & Tao, 2022). Společnost, která personalizaci zvládne technicky i eticky, získává významnou konkurenční výhodu.

2. **Flexibilita a adaptace na měnící se potřeby:** Úspěšné AI nástroje by měly být navrženy tak, aby je bylo možné pružně upravovat podle měnících se požadavků uživatelů a prostředí. To vyžaduje modulární architekturu a robustní trénovací postupy. Pro vývojáře to znamená, že při návrhu systému počítají s **možností rozšiřování a aktualizací** – např. přidávání nových funkcí, podpory dalších jazyků nebo integrace nových typů vstupních dat. Adaptabilita se může týkat i samotného modelu: moderní AI může průběžně *online* trénovat (učit se z nových dat za provozu), čímž se přizpůsobuje aktuálním trendům. Například doporučovací platforma může automaticky zaznamenat změnu preferencí uživatele (začne vyhledávat jiný žánr zboží) a promptně tomu uzpůsobit nabízený obsah. Z pohledu firmy je taková flexibilita zásadní pro udržení relevance produktu – potřeby uživatelů i tržní podmínky se vyvíjejí a AI nástroj, který se nedokáže **rychle přizpůsobit**, by brzy zastaral.
3. **Personalizace na základě dat:** Společnosti těží z bohatství uživatelských dat, která jim moderní digitální služby poskytují, a AI je prostředkem, jak z nich získat hodnotu v podobě personalizace. Strojové učení umožňuje identifikovat vzorce v datech a **predikovat preferované volby** jednotlivých uživatelů. Například streamingové služby využívají AI k analyzování, které filmy či písně uživatel poslouchal, v jakou denní dobu, zda některé přeskakoval atd., a podle toho sestavují unikátní playlisty či doporučení šité na míru. To zvyšuje pravděpodobnost, že uživatel najde obsah, který se mu líbí, a bude službu více využívat. Z pohledu tvůrce AI je důležité zajistit, aby tato práce s

daty byla **transparentní a respektovala soukromí** – uživatel by měl mít možnost dozvědět se, jaká data o něm systém využívá, případně se rozhodnout personalizaci omezit. Firmy, které jsou v personalizaci úspěšné, dokážou zároveň udržet důvěru zákazníků ohledně správy jejich dat, díky čemuž jsou zákazníci ochotní svá data výměnou za lepší služby poskytnout (McKinsey, 2021).

4. **Různorodost a inkluzivita:** Při navrhování AI systémů je třeba myslet na to, aby vyhovovaly širokému spektru uživatelů s různými potřebami, zázemím a zkušenostmi. Personalizace nesmí znamenat diskriminaci – AI by měla nabízet přizpůsobení pro každého, nejen pro většinového uživatele. Z pohledu společnosti to znamená testovat a vyvíjet AI s důrazem na **inkluzivní datasety** i funkcionality. Například v jazykových modelech je žádoucí, aby rozpoznávaly vstupy od lidí s odlišným dialektem či přízvukem stejně dobře jako “standardní” projev; v edukačních aplikacích by AI měla dokázat přizpůsobit styl výuky jak velmi nadaným studentům, tak těm, kteří potřebují více podpory. Firmy tímto předcházejí riziku, že určitá skupina (například menšina) bude jejich AI nástrojem systematicky znevýhodňována nebo opomenuta, což by mělo jak etické, tak PR důsledky. Naopak AI navržená s ohledem na různorodost uživatelů může otevřít nové trhy a zvýšit celkovou uživatelskou základnu.

Pohled uživatele (konzumenta AI služby):

1. **Přizpůsobení individuálním potřebám:** Z pohledu uživatele je jedním z největších přínosů AI to, že se mu dokáže *přizpůsobit*. Uživatelé oceňují systémy, které „se učí“, jakým způsobem je preferují používat, a zjednodušují jim práci tím, že se adaptují na jejich návyky. Například pokud chytrý asistent rozpozná, že uživatel každé ráno žádá o zprávy a předpověď počasí, může mu tyto informace začít nabízet proaktivně sám od sebe. Nebo textový editor s AI doplňováním se může přizpůsobit stylu psaní

konkrétního člověka, nabízet personalizované fráze a tím urychlit tvorbu textu. To vše zvyšuje **efektivitu** – uživatelé se nemusí přizpůsobovat stroji, místo toho se stroj přizpůsobuje jim. Výsledkem je plynulejší a přirozenější interakce, která šetří čas a námahu.

2. **Interaktivní a responzivní systémy:** Uživatelé preferují AI, která reaguje **interaktivně** – tedy dokáže vést dialog, postupně upřesňovat požadavky a reagovat na zpětnou vazbu uživatele. Například konverzační AI, která umožní uživateli doplňujícími otázkami upřesnit, co přesně potřebuje, poskytne relevantnější odpověď než systém, který na jeden dotaz vrátí fixní výsledek bez možnosti další interakce. Adaptabilita zde znamená, že systém bere v potaz aktuální *kontext* komunikace a předchází kroky uživatele. Pokud se uživatel v navigační aplikaci odchýlí od navržené trasy, AI ihned zareaguje přepočítáním nové trasy; obdobně pokud se uživatel nečekaně zeptá hlasového asistenta na nesouvisející věc, AI se přizpůsobí změně tématu. Tyto vlastnosti vytvářejí dojem **“chytrosti”** a vstřícnosti systému, což je pro uživatele atraktivní.
3. **Zlepšování zkušenosti na základě učení:** Ideální AI z pohledu uživatele je ta, která se **s časem zlepšuje** – čím déle ji používá, tím lépe pro něj funguje. To uživatelé pociťují například u zmíněných doporučovacích systémů (čím více filmů ohodnotí, tím přesnější doporučení dostávají) nebo u virtuálních asistentů (časem se “naučí” uživatelův denní režim a nabízí informace proaktivně). Uživatelé toto postupné učení vnímají pozitivně, neboť vidí hmatatelný přínos trvalého používání služby. Samozřejmě, pokud by učení vedlo k nežádoucím efektům (např. AI se adaptuje špatně a začne vnucovat obsah, který uživatele nezajímá), může to být i na škodu – proto je důležité, aby personalizace probíhala korektně a měla možnost korekce. Celkově však platí, že uživatelé hodnotí adaptivní systémy kladně, zejména pokud mají nad personalizací jistou míru kontroly (např.

možnost upravit preference ručně nebo vypnout některé personalizované doporučování, pokud jim nevyhovuje).

Příklady:

Personalizované e-commerce platformy: Internetové obchody stále častěji nasazují AI, která analyzuje nákupní chování zákazníků a podle toho **personalizuje obsah** – od doporučených produktů přes pořadí zobrazovaných položek až po individuální slevové nabídky. Příkladem může být Amazon, jehož doporučovací algoritmus generuje podstatnou část tržeb tím, že zákazníkům navrhuje zboží související s tím, co si již prohlíželi nebo koupili. Z pohledu firmy tak AI pomáhá zvýšit objem nákupů na jednoho zákazníka (tzv. *upselling* a *křížový prodej*). Z pohledu zákazníka je výhodou, že na velkém e-shopu snadněji najde relevantní produkty a neztrácí se v nabídce – například pokud si často kupuje sportovní vybavení, e-shop mu na hlavní stránce nabídne nové kolekce sportovního zboží namísto obecné nabídky. Úspěšnost tohoto přístupu potvrzují data: firmy, které dokážou personalizaci využít naplno, získávají výrazně vyšší podíl tržeb právě díky personalizovaným doporučením (McKinsey, 2021). Pro uživatele je nákup personalizovaným způsobem příjemnější a méně časově náročný, což podporuje jeho loajalitu k platformě.

AI ve vzdělávání: Vzdělávací platformy a výukové programy stále více experimentují s AI pro **adaptivní učení** – to znamená, že vzdělávací obsah se přizpůsobuje úrovni, tempu a stylu učení jednotlivého studenta. Například systém dokáže detekovat, že student tápe v určité kapitole matematiky, a nabídne mu dodatečná vysvětlení či procvičovací úlohy zaměřené právě na tuto oblast. Naopak u látky, kterou student zvládá rychle, může AI zkrátit výklad nebo přeskočit snadné příklady. Taková adaptivita vede k efektivnějšímu učení, protože každý student dostává pozornost tam, kde ji potřebuje – ani se nenudí triviálním výkladem, ani není frustrován příliš těžkým obsahem. Studie skutečně ukazují zlepšení studijních výsledků při nasazení personalizovaných výukových technologií ve srovnání s

uniformním přístupem pro všechny (Ko et al., 2022). Pro školy a tvůrce edukačních software je to velký příslib: AI může pomoci škálovat *individualizovaný přístup*, který byl dříve možný jen při velmi nízkém poměru studentů na učitele. Z perspektivy studentů (a rodičů) je adaptivní výuka atraktivní v tom, že respektuje jedinečnost každého žáka – nadaní mohou postupovat rychleji, ti s obtížemi dostanou více podpory. Klíčové je ovšem zajistit, aby tyto AI systémy byly spravedlivé a *inkluzivní* (viz etické aspekty dále) – musí sloužit všem studentům bez rozdílu a nesmí například některé skupiny systematicky znevýhodňovat.

3.4 Transparentnost a důvěryhodnost

Pohled společnosti (tvůrce a provozovatele AI nástroje):

1. **Porozumění funkčnosti AI:** Firmy vyvíjející AI by měly svým zákazníkům a uživatelům poskytovat co nejjasnější vhled do toho, jak jejich systémy fungují. To neznamená prozradit detailní technické know-how nebo zdrojový kód, ale spíše nabídnout srozumitelné vysvětlení principů: jaké typy dat AI využívá, jaká je obecná logika rozhodování, jaké výstupy lze očekávat a s jakou mírou nejistoty. Tato transparentnost pomáhá budovat důvěru – uživatel (nebo klient z firemní sféry) má větší jistotu v produkt, o němž ví, jak funguje, než v “magickou černou skříňku”. Z pohledu tvůrce AI nástroje to může znamenat například publikování dokumentace vysvětlující algoritmy v běžném jazyce, poskytování nástrojů pro vizualizaci rozhodovacích procesů AI nebo možnost nahlédnout do vzorových datových sad. Kromě posílení důvěry toto také umožní externím expertům zkoumat a kontrolovat AI systém, čímž se zvyšuje šance odhalit případné chyby či bias v modelech dříve, než způsobí problém.
2. **Odpovědnost a dodržování pravidel:** Pro společnost je zásadní dodržovat platné zákony a nařízení týkající se AI a dat. V Evropské unii například obecné nařízení GDPR stanovuje přísná pravidla pro zpracování osobních údajů, včetně práva subjektů znát, jak

jsou jejich data využívána v automatizovaném rozhodování. To klade důraz na transparentní nakládání s daty a možnost vysvětlit rozhodnutí AI, pokud mají pro uživatele významný dopad (Goodman & Flaxman, 2017). Tvůrce AI nástroje musí tedy zajistit, že jeho systém splňuje požadavky na soukromí (minimalizace sběru dat, zabezpečení databází), umožňuje získat souhlas uživatelů se zpracováním dat a případně poskytuje mechanismus, jak automatizované rozhodnutí zrevidovat lidskou cestou. Kromě zákonných aspektů sem spadá i celková etická odpovědnost za dopady AI – firmy by měly proaktivně přemýšlet o možných rizicích zneužití nebo negativních důsledcích jejich technologie a podnikat kroky k jejich zmírnění.

3. **Prevence a řešení předsudků:** Jak již bylo zmíněno u datových aspektů, AI může nevědomky přejímat nebo dokonce zesilovat lidské předsudky (bias) obsažené v datech. Společnosti nesou odpovědnost za to, aby jejich produkty byly spravedlivé a nediskriminační. To vyžaduje transparentní přístup při vývoji – otevřeně řešit otázky fairness už v návrhu systému, testovat modely na různých segmentech populace a výsledky publikovat. V případě zjištění zaujatosti (např. že algoritmus hůře funguje pro určité etnikum či gender) by měla firma rychle zasáhnout – upravit trénovací data, modifikovat algoritmus či přidat kompenzační mechanismy. Transparentnost vývojového procesu (např. zveřejnění metrik spravedlnosti algoritmu) umožňuje budovat důvěru uživatelů i regulačních orgánů, že společnost bere etické dopady vážně. Některé firmy zavádějí i externí audity svých AI systémů a zveřejňují tzv. dokumentace modelů a dat (Mitchell et al., 2019), které popisují charakteristiky datové sady, účel modelu a známé limity – to vše přispívá k odpovědné a důvěryhodné AI.
4. **Budování důvěryhodnosti:** Celkově by společnosti měly aktivně komunikovat a demonstrovat, že jejich AI je důvěryhodná (*trustworthy AI*). Kromě výše uvedené transparentnosti a fairness

sem patří i vysoká bezpečnost a robustnost systému (ochrana před hackery, odolnost proti neoprávněné manipulaci výsledků), plnění slibů (nedělat reklamu AI, která nesnese praktickou zkoušku) a otevřený dialog s veřejností. Například pokud dojde k incidentu (AI nástroj pochybí nebo dojde k úniku dat), firma by měla situaci transparentně vyšetřit, informovat uživatele a přijmout nápravná opatření. Důvěryhodnost se buduje postupně a zahrnuje všechny zmíněné aspekty – od technických (přesnost, bezpečnost) přes právní (soulad s regulacemi) až po morální (etické zacházení s daty, nediskriminace). Firmy jako Google, Microsoft či IBM již publikovaly své etické směrnice pro AI, kde deklarují principy jako transparentnost, spravedlnost, soukromí a odpovědnost, čímž se snaží dát najevo závazek k důvěryhodné AI (Jobin et al., 2019).

Pohled uživatele (konzumenta AI služby):

1. **Důvěra ve výsledky a proces:** Z pohledu uživatele je transparentnost klíčová proto, aby mohl výsledkům AI věřit. Pokud uživatel nechápe, jakým způsobem AI k dosaženému závěru dospěla, může v něm klíčit nedůvěra – zejména pokud jde o citlivé záležitosti (zdravotní diagnóza, rozhodnutí, zda mu bude poskytnuta půjčka apod.). Uživatelé proto oceňují, když jim AI dokáže poskytnout zdůvodnění nebo kontext ke svým odpovědím. Například moderní AI-asistenti pro podporu rozhodování mohou uvádět i „důkazy“ – tj. odkaz na zdroj dat nebo faktor, který nejvíce ovlivnil jejich rozhodnutí. To uživateli umožňuje posoudit, zda je závěr důvěryhodný. Výzkumy ukazují, že poskytnutí takových vysvětlení zvyšuje ochotu lidí svěřit se rozhodnutí AI a sdílet s ní citlivé informace (Shin et al., 2022). Celkově platí, že pokud AI funguje průhledně a je zřejmé, že s uživatelem „hraje férově“, roste tím její akceptace.
2. **Pochopení možností a omezení:** Pro uživatele je důležité rozumět nejen výsledkům AI, ale i jejím schopnostem a limitům

(viz též výše pochopení omezení). Transparentnost systému pomáhá uživatelům vytvořit si mentální model toho, co AI umí, v jakých situacích ji použít a kdy raději ne. Například pokud je známo, že AI model funguje spolehlivě jen pro dospělou populaci, ale méně už pro děti, uživatel-lékař tomu přizpůsobí způsob, jak výsledky využívá. Nebo pokud chatbot transparentně sdělí, že jeho trénovací data dosahují pouze ke konkrétnímu datu (typicky v rozmezí let 2024–2025 u nejnovějších modelů), uživatel pochopí, proč mu neodpoví na dotaz týkající se aktuálního dění. Některé moderní modely tento limit překlenují přes integrované webové vyhledávání. Bez této otevřenosti hrozí, že uživatelé budou AI přeceňovat nebo používat nevhodně, což vede ke špatným zkušenostem. Transparentnost tak vlastně slouží i jako forma vzdělávání uživatelů – pomáhá jim nastavit správná očekávání a využívat AI racionálně.

3. **Ochrana soukromí a datová bezpečnost:** V neposlední řadě uživatele zajímá, jak AI nakládá s jejich osobními daty. Transparentnost v tomto kontextu znamená, že uživatel má přístup k informacím, jaké údaje o něm systém uchovává, k čemu je využívá a komu jsou případně dále zpřístupněny. Důvěryhodný AI systém by měl uživateli poskytnout kontrolu nad jeho daty (možnost záznamy smazat, upravit nastavení soukromí apod.). Rostoucí obavy veřejnosti z narušení soukromí AI technologiemi jsou dobře zdokumentovány – například analytické nástroje sledující online chování uživatelů či hlasoví asistenti nahrávající konverzace v domácnosti vyvolaly v minulosti značnou nedůvěru, pokud se ukázalo, že data mohla být zneužita. Uživatelé proto stále častěji požadují od poskytovatelů služeb jasné ujištění a důkazy o tom, že jejich data jsou v bezpečí. To může zahrnovat i nezávislé certifikace služeb, transparentní politiky ochrany dat a promptní informování o jakémkoli bezpečnostním incidentu. Jen za těchto podmínek budou

ochotni své údaje AI systémům svěřovat, což je podmínkou pro jejich kvalitní fungování (většina AI se bez uživatelských dat neobejde).

Příklady:

AI ve zdravotnictví: Představme si AI systém, který lékařům navrhuje diagnózy nebo léčebné postupy na základě analýzy pacientových údajů. Pro to, aby takovému systému lékaři i pacienti věřili, je nezbytná značná míra transparentnosti. Lékař potřebuje vědět, na základě, jakých dat či pravidel AI k danému doporučení došla – například zda u pacienta vyhodnotila vysoké riziko určité nemoci kvůli kombinaci konkrétních symptomů, výsledků testů a třeba i vzácné genetické informace. Pokud AI poskytne k svému závěru i toto zdůvodnění (např. formou vyjmenování klíčových rizikových faktorů), lékař mnohem ochotněji návrh přijme a zahrne do rozhodování (Topol, 2019). Z perspektivy pacienta zase vyvstává otázka soukromí: takový systém pracuje s vysoce citlivými údaji o zdraví, a pacient tedy právem očekává, že bude přesně informován, jak jsou jeho data uložena, kdo k nim má přístup a zda jsou v bezpečí před únikem. Důvěryhodný AI nástroj ve zdravotnictví by měl například umožnit audit přístupů k datům pacienta (aby bylo dohledatelné, kdo s nimi nakládal), měl by používat anonymizaci pro trénování modelů a umožnit pacientovi *opt-out* (nebýt zahrnut v některých typech analýz, pokud si to nepřeje). Pouze při splnění těchto podmínek – vysvětlené závěry a ochrana soukromí – lze očekávat širší přijetí AI v tak citlivé oblasti, jako je zdravotnictví.

AI ve finančním sektoru: **Robo-poradci** a další AI systémy v bankovníctví a investicích získávají na popularitě, ale jejich úspěch závisí na důvěře klientů. Představme si robo-poradce, který doporučuje klientovi investiční portfolio. Pokud klient obdrží jen strohé doporučení „Investujte 20 % do fondu X a 30 % do akcií Y...“ bez jakéhokoli vysvětlení, pravděpodobně bude váhat, zda návrhu vyhovět – zejména v období finanční nejistoty. Naproti tomu

poradenská AI, která ke každému návrhu připojí jednoduché zdůvodnění (např. „fond X byl zvolen kvůli diverzifikaci do technologického sektoru, kde vidíme potenciál růstu; akcie Y jsou zahrnuty pro svou stabilitu a dividendový výnos...“), působí mnohem **důvěryhodněji** a klient má pocit, že návrh je promyšlený (Komić et al., 2021). Dalším aspektem je **spravedlnost a zodpovědnost** algoritmů: v USA i Evropě se například klade důraz na to, aby úvěrové skoringové modely nebyly diskriminační (nesmí znevýhodňovat minoritní skupiny). Představme si, že by AI model pro schvalování půjček systematicky vyhodnocoval určitou menšinu jako rizikovější jen na základě korelace s poštovním směrovacím číslem bydliště – taková praxe by byla nejen neetická, ale i nezákonná. Odhalit a napravit podobné **biasy** jde jedině díky transparentnosti: banka musí být schopna vysvětlit, proč byl konkrétní úvěr zamítnut, a prokázat, že k tomu nevedly diskriminační faktory. Právní předpisy (např. v EU) dokonce dávají zákazníkům právo vyžádat si lidský přezkum automatizovaného rozhodnutí, což nutí instituce držet AI „na uzdě“ a nevyužívat netransparentní algoritmy tam, kde by mohly negativně ovlivnit život klienta (Goodman & Flaxman, 2017). Díky tomu mají uživatelé (klienti) větší jistotu, že se mohou spolehnout na férové zacházení – a pokud by pocítili opak, mohou se dožadovat nápravy.

3.5 Eticko-sociální aspekty

Pohled společnosti (tvůrce a prodejce AI nástroje):

1. **Etický design a vývoj:** Firmy vyvíjející AI by měly již od počátku zvažovat etické důsledky svých technologií. To zahrnuje zajištění spravedlnosti, nediskriminace a respektu k lidským právům v návrhu algoritmů (Jobin et al., 2019). Prakticky to znamená vybírat trénovací data tak, aby nepodporovala stereotypy nebo neznevýhodňovala určité skupiny, testovat modely z hlediska biasu (jak bylo zmíněno výše) a konzultovat vývoj s odborníky na etiku či zástupci dotčených komunit. Příkladem etického selhání, které firmu poškodilo, byl případ náborového AI nástroje

společnosti Amazon: model trénovaný na historických datech upřednostňoval životopisy mužů a znevýhodňoval ženy v technických pozicích, protože vycházel z předchozího stavu, kde dominovali muži (Dastin, 2018). Amazon po odhalení systém stáhl, ale příklad ukázal, jak důležité je na etiku myslet už při vývoji – prevence takových dopadů je mnohem lepší než následné hašení problémů. Zodpovědné firmy dnes zavádějí interní etické komise pro AI, definují závazné etické zásady pro vývojáře a některé dokonce uplatňují tzv. *ethics checklists* jako součást procesu nasazení nového modelu (Jobin et al., 2019). Takové iniciativy pomáhají předcházet situacím, kdy by AI mohla způsobit společenskou újmu.

2. **Dopady na pracovní trh:** AI má potenciál automatizovat řadu činností, které dříve vykonávali lidé, což přirozeně vyvolává obavy o budoucnost práce. Společnosti, které vyvíjejí a nasazují AI, by měly reflektovat, jak jejich technologie ovlivní zaměstnance a trh práce obecně. Od průmyslové revoluce víme, že nové technologie sice zvyšují produktivitu, ale také mohou krátkodobě vytlačovat lidskou práci – odhady v polovině rok 2010 naznačovaly, že až cca 47 % povolání zejména v administrativě, dopravě či výrobě by mohlo být do jisté míry automatizováno pomocí AI a robotiky (Frey & Osborne, 2017). Zodpovědný přístup firem spočívá v tom, že aktivně přispívají k řešení těchto dopadů: nabízejí rekvalifikační programy pro propuštěné zaměstnance, vytvářejí nové role zaměřené na dohled nad AI (např. tzv. *AI trainers* nebo *explainers*), případně spolupracují s vládami a vzdělávacím sektorem na přípravě pracovní síly na změny. Firmy by rovněž měly komunikovat realisticky – nepřehánět okamžité přínosy AI, ale ani nemlčet o tom, které rutinní pozice možná do budoucna zaniknou. Tím, že se podniky zapojí do debaty o budoucnosti práce, pomáhají zajistit sociálně udržitelnou implementaci AI technologií.

3. **Společenská zodpovědnost:** Širší sociální dopad AI je další oblast, kterou by tvůrci AI neměli opomíjet. Sem spadá například otázka, jak AI ovlivní komunitu, ve které je nasazena, či celou společnost – přinese převážně prospěch, nebo prohloubí nerovnosti? Společnosti by měly usilovat o to, aby jejich technologie byly prospěšné pro společnost jako celek (*beneficence*) a minimalizovaly potenciální škody (*non-maleficence*) – což jsou dva základní principy často zmiňované v etických kodexech AI (Floridi et al., 2018). Prakticky to může znamenat dělat si *socioekonomické analýzy* před nasazením AI (např. jak autonomní pokladny ovlivní komunitu pokladních v malém městě), podporovat nasazení AI i v oblastech veřejného zájmu (zdravotnictví, vzdělávání) a nejen tam, kde jde hlavně o zisk, a v neposlední řadě být otevřený dialogu s veřejností a regulátory. Firmy by měly přijmout zodpovědnost za negativní externality – např. platforma, na které AI doporučovací algoritmy nečekaně podporují dezinformace či společenskou polarizaci, by měla rychle zasáhnout a upravit systém, i když to třeba krátkodobě sníží míru prokliků. Takové **proaktivní a odpovědné chování** je dnes často požadováno nejen občanskou společností, ale i investory (kteří se zajímají o tzv. ESG kritéria).

4. **Transparentnost a zodpovědnost (accountability):** Kromě technické transparentnosti algoritmů (řešené výše) jde v eticko-sociálním kontextu o to, aby bylo jasné, kdo nese odpovědnost za jednání AI systému. Společnosti by měly deklarovat, že za svými AI stojí a v případě jejich selhání ponesou odpovědnost (právní i morální). To je důležité např. u autonomních vozidel – pokud způsobí nehodu samořídící auto, musí existovat rámec, kdo za to odpovídá (výrobce software? majitel vozidla?). Firmy by neměly alibisticky svalovat vše na “nepředvídatelné chování algoritmu”, ale spíše vytvářet systémy tak, aby umožnily dohledatelnost rozhodnutí a případně lidský zásah. Též by měly spolupracovat na tvorbě průmyslových standardů a regulací,

které nastaví pravidla pro bezpečný a etický vývoj AI. Například mnoho předních společností se zapojilo do iniciativ, které formulují obecné zásady pro **“Trustworthy AI”** (důvěryhodnou AI) – ty typicky zahrnují právě požadavky na spravedlnost, transparentnost, soukromí, ale i lidský dohled a accountability (Jobin et al., 2019). Všechny tyto snahy nakonec slouží k tomu, aby se veřejnost mohla cítit bezpečně s faktem, že AI proniká do stále více oblastí – a aby vnímala, že tvůrci AI berou vážně nejen komerční, ale i společenskou stránku svého podnikání.

Pohled uživatele (konzumenta AI služby):

1. **Důvěra a etické obavy:** Uživatelé jsou citliví na to, zda je AI využívána eticky správně a zda nepůsobí společenské škody. Pokud mají pocit, že určitá AI aplikace je **“nespravedlivá”** nebo ohrožuje hodnoty, na kterých jim záleží, budou se jí vyhýbat bez ohledu na její technickou vyspělost. Například veřejnost s nedůvěrou pohlíží na AI systémy pro rozpoznávání obličejů, pokud jsou nasazeny hromadně ve veřejném prostoru – lidé se obávají narušení soukromí a možného zneužití ke sledování obyvatel. Tyto etické obavy mohou brzdit přijetí jinak užitečných AI nástrojů. Z pohledu uživatele je proto důležité, aby viděl, že AI je používána ku prospěchu, a ne proti němu. Když například město zavádí AI pro prediktivní policing (předpovídání kriminality), místní komunita bude chtít mít jistotu, že to nepovede k policejnímu obtěžování nevinných lidí na základě zkreslených dat. Pokud takovou jistotu nemá, technologie odmítne a může dojít k veřejným protestům. Uživatelé tedy potřebují vnímat, že AI je eticky hlídaná a slouží jim, nikoli jen zájmům těch, kdo ji provozují.
2. **Ochrana soukromí a dat:** Již jsme zmínili význam soukromí – z pohledu uživatele je to jedna z největších oblastí obav u moderních technologií. Uživatelé se stále více zajímají o to, jaké informace o nich AI shromažďuje. Skandály spojené s úniky dat

či jejich zneužitím (např. kauza Cambridge Analytica využívající data z Facebooku) zvýšily povědomí o rizicích. Uživatelé proto mohou AI služby odmítnout, pokud mají pocit, že cena v podobě jejich soukromí je příliš vysoká. Například někteří lidé vypínají hlasové asistenty na svých telefonech, jelikož se obávají neustálého odposlechu. Z pohledu běžného občana je ideální stav, kdy AI sice využívá jeho data k personalizaci, ale činí tak anonymně a bezpečně – uživatel chce mít kontrolu a v žádném případě nechce, aby se jeho osobní údaje dostaly do nepovolaných rukou. Existence silných zákonů (jako je GDPR) dává uživatelům určitou oporu v prosazování svých práv, ale důležitá je i praktická zkušenost: pokud např. AI aplikace nabídne jednoduché soukromostní nastavení a srozumitelně vysvětlí, jak data používá, získá si tím přízeň uživatelů mnohem spíše než služba, která tyto věci skrývá v nepřehledných podmínkách.

3. **Sociální dopad a spravedlnost:** Uživatelé – coby členové společnosti – vnímají i širší dopady AI na sociální spravedlnost a kvalitu života. Například pracovníci mohou mít obavy, že jejich pracovní místa zaniknou kvůli automatizaci. Tato úzkost není pouze individuální, ale má celospolečenský rozměr: v regionech, kde je více rutinních pracovních míst, panuje větší nejistota z technologických změn. Uživatelé, kteří se cítí ohroženi (nebo jejich rodiny), pak mohou být vůči AI negativně naladěni a nechtějí ji podporovat. Dalším příkladem je vzdělávání – rodiče mohou mít strach, že nasazení AI ve škole (třeba pro hodnocení esejí či testů) bude nefér a poškodí jejich dítě, pokud systém nebude dokonale spravedlivý. Nebo mohou vnímat, že přílišná digitalizace ochudí děti o sociální interakce s učiteli. Tyto obavy mohou vést k tlaku veřejnosti na omezení využívání AI v dané oblasti. Z pohledu jednotlivce je důležité, aby měl pocit kontroly nad dopady AI na svůj život – chce mít možnost rozhodnout se, zda využije autonomní taxi nebo raději podpoří zaměstnaného řidiče; zda nechá AI vyučovat své děti, nebo trvá na lidském

přístupu. Pokud AI technologie budou lidem tyto volby upírat nebo je tlačit do nežádoucích změn, narazí na odpor. Naopak AI, která bude prezentována a nasazena tak, že pomáhá řešit skutečné společenské problémy (např. medicínská AI zlepšující péči, environmentální AI pomáhající v boji proti změně klimatu), získá mezi lidmi mnohem větší podporu.

Příklady:

AI v náboru zaměstnanců: Jak již bylo zmíněno, nasazení AI do oblasti HR (např. pro předvýběr uchazečů o zaměstnání) přináší značné etické otázky. Pokud by takový systém nebyl pečlivě navržen, může neúmyslně zvýhodňovat určité skupiny lidí. Například pokud trénovací data obsahují historickou převahu mužů v technických rolích, AI by mohla tyto vzorce považovat za normu a doporučovat k pohovoru častěji muže než ženy – aniž by to kdokoli explicitně zamýšlel. Z pohledu tvůrce je klíčové takovým situacím předejít (testovat modely na bias a opravovat je). Z pohledu uživatele (resp. uchazeče o práci) pak v opačném případě klesá důvěra v náborový proces firmy – pokud by vyšlo najevo, že algoritmus má předsudky, kandidáti se budou cítit podvedeni a firma může čelit diskriminačním žalobám. Příklad Amazonu je výstrahou: jejich interní AI náborový nástroj vykazoval bias proti ženám, což vedlo k negativní publicitě a jeho zrušení (Dastin, 2018). Tento případ zvýšil obecnou ostražitost veřejnosti vůči podobným systémům. Nyní kandidáti oprávněně požadují, aby měli možnost ucházet se o práci férově a aby případné automatizované filtry byly transparentní – někteří volají i po právu dozvědět se, proč je AI případně vyřadila z výběru. Dá se očekávat, že do budoucna budou přijaty standardy pro **audit algoritmů** v HR tak, aby kandidáti mohli AI důvěřovat (Jobin et al., 2019).

AI ve vzdělávání: Nasazení AI pro personalizovanou výuku a hodnocení studentů může přinést velký pokrok, ale nesmí opomenout princip rovného přístupu ke vzdělání. Představme si AI systém, který bude studentům automaticky zadávat různé obtížné

úkoly podle jejich výkonu. Pokud by byl špatně navržen, mohlo by se stát, že studenty z méně podnětného prostředí či s počátečními potížemi rychle zařadí do „pomalé“ větve, ze které už nebude úniku – čímž by jim de facto upřel šanci zlepšit se, kterou by jinak v tradiční třídě s dobrým učitelem třeba dostali. Společensky by tak AI mohla zvýšit propast mezi nejlepšími a slabšími žáky, namísto aby ji zmenšovala. Zodpovědný vývoj takových systémů musí proto zahrnovat mechanismy, které tomu zabrání (např. pravidelné „míchání karet“, možnost manuálního zásahu učitele, dodatečné podpurné lekce pro ty, které AI označí za slabé, aby měli šanci se dotáhnout). Z pohledu rodiče či studenta je klíčové mít jistotu, že AI není nespravedlivá a že všichni dostávají adekvátní příležitost se zlepšovat. Pokud by se rodiče dozvěděli, že výukový software jejich dětí má v sobě nějakou zabudovanou znevýhodňující systémovou chybu, okamžitě by volali po jeho odstranění ze škol. Dobrým příkladem inkluzivního přístupu je vývoj AI, která byla testována na různých školách v odlišných socio-ekonomických podmínkách, aby se ověřilo, že funguje univerzálně a nejen pro určitý typ studentů (Holstein et al., 2019). Tím se zajistí, že AI skutečně plní roli pomocníka učitele v **personalizaci** výuky, aniž by některé děti odsunula na vedlejší kolej.

3.6 Rychlost a výkonnost

Pohled společnosti (tvůrce a provozovatele AI nástroje):

1. **Optimalizace a škálovatelnost:** Pro technické týmy vyvíjející AI je výzvou dosáhnout vysokého výkonu systému a zároveň zajistit možnost jeho škálování na velké objemy dat či uživatelů. AI modely, zejména ty založené na hlubokém učení, mohou být velmi náročné na výpočetní výkon. Společnosti proto investují do optimalizace kódu a infrastruktury tak, aby zpracování probíhalo co nejrychleji. To může zahrnovat využití paralelních výpočtů na grafických procesorech (GPU) či specializovaných čipech (TPU), jakož i úspornější algoritmy (např. komprimované neurální sítě,

kvantizace modelů, efektivnější architektury). Důležitým aspektem je také škálovatelnost – systém by měl fungovat hladce, i když se dramaticky zvýší počet uživatelů nebo objem dat. Cloudové technologie a horizontální škálování (přidávání dalších serverů) umožňují firmám dynamicky reagovat na rostoucí zátěž. Příkladem je streamovací služba s AI doporučováním: musí zvládnout generovat personalizované doporučení pro miliony uživatelů současně, aniž by se zpozdlilo načítání obsahu. Firmy jako Netflix či YouTube investovaly značné prostředky do toho, aby jejich AI infrastruktura byla vysoce škálovatelná a robustní – to zahrnuje jak techniky ke zrychlení algoritmů, tak optimalizaci sítí pro datový přenos. Výsledkem těchto snah je, že moderní AI systémy dokážou v reálném čase analyzovat ohromná data (např. sociální sítě s miliardou příspěvků) a poskytovat okamžité outputy uživatelům.

2. **Výpočetní kapacita:** V souvislosti s výkonem hraje roli i zajištění dostatečné výpočetní kapacity a její efektivní využití. Trénování špičkových AI modelů (například velkých jazykových modelů nebo detailních obrazových analyzátorů) vyžaduje enormní množství výpočetních operací. Analýza trendů ukazuje, že od roku 2012 do 2018 rostlo množství výpočetních zdrojů použitých při tréninku největších modelů exponenciálně – zdvojnásobilo se zhruba každé 3–4 měsíce (Amodei & Hernandez, 2018). Pro firmy to znamená investovat do výkonného hardware (GPU farmy, specializované AI akcelerátory) a do efektivních algoritmů. Náklady na trénink modelu jsou dnes nezanedbatelné – vyšplhají se klidně do milionů dolarů u nejpokročilejších modelů s miliardami parametrů. Proto jsou firmy motivovány hledat co nejefektivnější cesty: využívat předtrénované modely a dotrénovávat je jen na své úlohy (*transfer learning*), dělit trénink mezi mnoho strojů (*distributed learning*) nebo aplikovat nové výzkumné poznatky v oblasti komprese modelů bez výrazné ztráty přesnosti. Současně poskytovatelé cloudu (Amazon AWS,

Google Cloud, Microsoft Azure) nabízejí pronájem výpočetní kapacity na vyžádání, což firmám umožňuje škálovat výkon podle potřeby bez masivních jednorázových investic do vlastního hardware. Správná kombinace dostatečné kapacity a chytrého využití zdrojů je zásadní pro to, aby AI systém běžel plynule a rychle i při složitých výpočtech.

3. **Minimalizace latence:** V mnoha aplikacích je požadována odezva v reálném čase, a tedy velmi nízká latence (zpoždění) odezvy systému. To je případ například autonomního řízení (kde AI musí ve zlomcích sekundy reagovat na podněty ze senzorů vozidla), hlasových asistentů (uživatel očekává odpověď téměř okamžitě poté, co domluví), či obchodních algoritmů reagujících na změny trhu. Vývojáři AI proto optimalizují nejen samotný model, ale i celý řetězec zpracování – od načtení vstupních dat, přes průchod modelem, až po přenos výsledku k uživateli. Kde je to možné, přesouvá se výpočet blíže ke zdroji dat (např. *edge computing*, kdy AI běží přímo v zařízení uživatele, jako je telefon nebo IoT senzor, aby odpadlo odesílání dat na server a čekání na odpověď). Také se používají rychlé datové spojení a protokoly, aby komunikace mezi komponentami systému byla co nejméně zatěžující. Inženýři analyzují tzv. *bottlenecky* v systému – hledají, která část procesu způsobuje největší zpoždění – a tu se snaží vylepšit. Například pokud je úzkým hrdlem samotná neuronová síť, může pomoci její zjednodušení (méně vrstev, menší vstupy); pokud je problémem přenos dat, řeší se komprese dat nebo lokální caching. HCI výzkum navíc ukazuje, že vnímání rychlosti systémem člověka má určitá mezní hodnoty: odezva do cca 0,1 sekundy působí jako okamžitá, do 1 sekundy je ještě považována za plynulou součást interakce, zatímco delší čekání (přes 3–5 sekund) už výrazně narušuje uživatelský tok myšlenek (Nielsen, 1993). Tyto poznatky motivují firmy k tomu, aby latenci držely co nejnižší – i nepatrné prodlevy mohou rozhodovat o tom, zda uživatel bude spokojen. V konkurenčním prostředí digitálních

služeb platí, že *rychlost* = *konkurenční výhoda*, proto je minimalizace latence pro firmy strategickým cílem.

4. **Efektivní algoritmy:** Výpočetní výkon hardware musí jít ruku v ruce s efektivitou samotných algoritmů. Pro tvůrce AI to znamená volit takové modely a metody, které poskytují potřebnou přesnost, ale s co nejmenší výpočetní komplexitou. Někdy lze dosáhnout zrychlení nasazením novějších architektur modelů – například v oblasti zpracování přirozeného jazyka došlo k velkému skoku zavedením *Transformer* architektury (Vaswani et al., 2017), která umožnila paralelizovat zpracování sekvencí a nahradit pomalejší rekurentní sítě. Dále se prosazují algoritmy pro tzv. *model pruning* (prořezávání modelů – odstranění neuronů či vah, které mají malý vliv na výstup, bez zhoršení kvality), *knowledge distillation* (převod znalostí z velkého modelu do menšího) a další techniky, které mohou dramaticky snížit počet operací nutných k dosažení podobného výsledku. Výzkumníci i inženýři tak neustále hledají rovnováhu mezi výkonem (např. co nejvyšší přesností) a náklady (časovými, energetickými). V praxi to může být i rozhodování, jaký typ modelu použít pro daný úkol: někdy stačí jednodušší algoritmus s menší datovou náročností (např. rozhodovací strom nebo lineární model) namísto hluboké sítě – zvláště pokud to přinese významnou úsporu času a zdrojů a uživatelský přínos vyšší přesnosti není zásadní. Firmy při vývoji produktů často experimentují s více variantami modelů a volí tu, která dává nejlepší **poměr výkon/výpočetní náročnost**.

Pohled uživatele (konzumenta AI služby):

1. **Rychlá odezva:** Z pohledu uživatele je rychlost reakce AI systému jedním z nejviditelnějších znaků kvality. Uživatelé jsou zvyklí na okamžitost moderních technologií – pokud zadávají dotaz vyhledávači nebo mluví na hlasového asistenta, očekávají odpověď téměř ihned. Když systém reaguje pomalu, uživatel se cítí frustrován a má tendenci interakci přerušit (či službu opustit). Už zmíněná pravidla v UX mluví jasně: pokud odezva překročí jednotky sekund, uživatelé získávají pocit, že je „počítač zdržuje“, a mohou ztratit o danou činnost zájem (Nielsen, 1993). Proto je pro pozitivní uživatelskou zkušenost nezbytné, aby AI držela tempo s uživatelem. U dialogových systémů to znamená reagovat v řádu milisekund až jedné sekundy, u webových služeb načítat výsledky do cca 2 sekund atd. Rychlost je dokonce tak důležitá, že mnozí uživatelé dají přednost systému, který odpoví rychle, i kdyby jeho odpověď byla o něco méně detailní, před systémem, který sice nabídne super-kvalitní výstup, ale nechá na něj čekat (to se řeší např. u vyhledávačů – obsáhlé, ale pomalé odpovědi vs. rychlé stručné výsledky). **Pružnost a pohotovost** AI je tedy základním předpokladem spokojenosti uživatele.
2. **Přesnost a efektivita:** Samotná rychlost však nestačí – pokud AI bleskově vrátí odpověď, ale ta bude chybná nebo nepoužitelná, uživatel spokojen nebude. Uživatelé vyžadují kombinaci rychlosti a přesnosti. V ideálním případě AI poskytne správné řešení problému a ještě dříve, než by ho uživatel stihl najít vlastními silami. Například u překladače očekává uživatel téměř okamžitý překlad věty, ale také chce, aby překlad byl kvalitní a dával smysl. Obdobně u navigace – aplikace by měla rychle spočítat trasu, ale ta trasa musí být rozumná (rychlá, bezpečná), jinak rychlá odpověď postrádá hodnotu. Uživatel tedy posuzuje výkon AI komplexně: hodnotí, jak rychle a jak dobře splnila jeho požadavek. Výzva pro designéry AI služeb je najít rovnováhu: někdy extrémní zrychlení může mírně snížit přesnost (např. když

se použije menší model), pak je otázkou, co je pro uživatele přijatelnější. Obecně pokud se přesnost blíží lidské úrovni, upřednostní uživatel rychlost – drobných chyb si možná ani nevšimne, pokud dostane výsledek pohotově. Jakmile by však kvalita výrazně utrpěla, trpělivější přístup s lepším výsledkem může být preferován. Každopádně cílem je dosáhnout, aby uživatel nemusel volit mezi rychlostí a kvalitou – špičkové AI systémy se snaží doručit oboje.

3. **Spolehlivost a konzistence:** Uživatelé také očekávají, že AI bude spolehlivá v každé situaci – tedy že si udrží svižnou reakci a funkcionalitu i při vyšší zátěži nebo v různých podmínkách. Není pro ně omluvou, že „server je přetížen“ nebo „dnes to funguje pomaleji, protože...“. Od robustní služby čekají konzistentní výkon. To znamená, že např. streamovací platforma musí fungovat stejně dobře večer, kdy je online nejvíce lidí, jako brzy ráno při nízkém provozu. Pokud by uživatelé zaznamenali, že v určité hodiny je AI aplikace nepoužitelně pomalá, naruší to jejich důvěru a začnou se poohlížet po alternativě. Dostupnost služby (uptime) a konzistence odezvy jsou tedy důležité metriky, které vnímají i laičtí uživatelé. Kromě toho spolehlivost zahrnuje i to, že AI nebude „zamrzat“ nebo padat – tedy že rychle nejen vypočte odpověď, ale vůbec ji vždycky vypočte a nezasekne se u určitých vstupů. Z pohledu uživatele je ideál takový, že AI funguje tak samozřejmě a hladce, že o její přítomnosti ani nepřemýšlí – plně splyne s uživatelskou zkušeností a nijak ji neruší. Toho se dá dosáhnout právě kombinací rychlosti, přesnosti a spolehlivosti.

Příklady:

AI v e-commerce (chatboti pro zákaznickou podporu): Zákaznické chatboty na webech obchodů nebo poskytovatelů služeb jsou dobrým příkladem, kde je rychlost a výkonnost AI přímo viditelná. Když má zákazník problém (např. chce reklamovat zboží, zjistit stav objednávky apod.), často preferuje využít online chat s podporou.

Pokud je na druhé straně AI chatbot, musí reagovat okamžitě – ideálně do pár sekund pozdravit a v řádu dalších sekund už poskytnout relevantní odpověď či požádat o upřesnění. Jakékoli delší čekání (např. „Moment, vyhledávám vaši objednávku...“) uživatel vnímá negativně. Moderní chatboty proto integrují přímo do prostředí e-shopu a jsou napojeny na databáze tak, aby mohly *v reálném čase* dotazovat potřebné informace (stav objednávky, skladová dostupnost apod.) a bleskově je vracet zákazníkovi. Pokud se chatbot zdrží nebo zpráva nedorazí, zákazník chat opustí a raději třeba zavolá na infolinku – čímž se mýjí účel automatizace. Z pohledu firmy je tedy zásadní, aby AI asistent stíhal odpovídat i při špičce (např. během velkých slevových akcí, kdy píše mnoho zákazníků naráz) – k tomu slouží škálování serverů a optimalizace dotazů do databází. Kromě rychlosti je důležité, že odpovědi musejí být věcně správné (podobně jako u lidského operátora). Jakmile AI vygeneruje mylnou informaci (třeba špatné datum doručení zásilky), zákazník ztrácí důvěru. Proto firmy někdy omezují nasazení AI jen na časté dotazy, kde je vysoká jistota správné odpovědi, a složitější požadavky přepojují na člověka – tím se snaží udržet jak rychlost, tak kvalitu služby. Když je chatbot dobře navržen, přináší firmě úsporu nákladů (může odbavit tisíce dotazů paralelně) a zákazníkům rychlou pomoc 24/7.

AI ve zdravotní diagnostice: V medicíně jsou dnes vyvíjeny AI systémy, které například analyzují radiologické snímky (RTG, CT, MRI) a pomáhají lékařům odhalit patologie. Rychlost těchto systémů je klíčová zejména v urgentních případech. Představme si nemocnici, kde AI v reálném čase sleduje snímky plic pacientů s podezřením na zápal plic nebo jiné akutní stavy. Pokud AI dokáže prakticky okamžitě po pořízení snímku označit podezřelé oblasti (např. výpadek obrazu svědčící o krvácení do mozku) a upozornit lékaře, může to urychlit léčbu a doslova zachránit život (Topol, 2019). Rychlost analýzy tu musí jít ruku v ruce s přesností – falešně pozitivní nálezy by zbytečně plašily lékaře, falešně negativní by byly ještě horší (neoznačit problémový snímek znamená riziko, že lékař něco přehlédne). Proto se v praxi AI

využívá tak, že slouží jako „druhý pár očí“ radiologa: upozorní ho v časové tísní na snímky, kterým by měl věnovat zvýšenou pozornost. Například při mamografickém screeningu velké množství snímků hodnotí AI předběžně a extrémně rychle vyřadí ty, kde s velkou jistotou není žádný nález (čímž šetří čas lékaře), zatímco u podezřelých snímků připojí vyznačení podezřelých ložisek (McKinney et al., 2020). Lékař pak finálně snímky posoudí – díky AI se však může soustředit tam, kde je to potřeba, a to bez prodlev. Celkově je tak diagnostický proces rychlejší, což pro pacienty znamená dřívější výsledky a v případě nálezu dřívější zahájení léčby. I tady ovšem platí, že kdyby AI generovala příliš mnoho omylů nebo nestíhala ve špičce (např. během pandemie analyzovat tisíce plicních snímků denně), lékaři by se na ni nemohli spolehnout. Proto se nasazení takových systémů pečlivě testuje v pilotních projektech a ladí se tak, aby přinášely skutečnou hodnotu: kombinaci rychlosti (analýza do několika vteřin) a kvality (citlivost a specifita na úrovni zkušeného odborníka).

3.7 Integrace a kompatibilita

Pohled společnosti (tvůrce a prodejce AI nástroje):

1. **Bezproblémová integrace do stávajících systémů:** Pro komerční úspěch AI produktu je často rozhodující, aby jej klient (např. firma) mohl snadno napojit na své současné IT systémy a procesy. Vývojáři AI nástrojů proto kladou důraz na otevřená rozhraní (API) a standardy, které usnadní integraci. Ideální AI systém je navržen modulárně a poskytuje jasně definované vstupy a výstupy, aby jej bylo možné „zapojit“ jako komponentu do většího celku. Pokud například firma nabízí AI modul pro analýzu zákaznických dat, měl by jít lehce propojit s běžnými databázemi, CRM systémy či datovými sklady, které klient už používá. V opačném případě, kdy by zavedení AI vyžadovalo kompletní přebudování infrastruktury klienta, by mnoho potenciálních zákazníků odradilo. Bezproblémová integrace snižuje

implementační náklady a časy, což zvyšuje ochotu organizací AI adoptovat. Studie z praxe ukazují, že jednou z hlavních překážek nasazení AI ve velkých podnicích bývá složitost integrace s legacy systémy (Singh et al., 2020) – proto se společnosti vyvíjející AI snaží tuto bariéru odbourávat tím, že jejich řešení podporují různé platformy a protokoly. Například AI nástroj pro zpracování přirozeného jazyka může nabízet pluginy do oblíbených komunikačních platforem (Slack, MS Teams), aby jej firmy mohly hned využít v prostředí, na které jsou jejich zaměstnanci zvyklí.

2. **Kompatibilita s různými platformami:** V dnešní heterogenní technologické krajině je výhodou, pokud AI řešení funguje napříč operačními systémy (Windows, macOS, Linux), zařízeními (desktohy, tablety, mobily) i prostředím (on-premise servery vs. cloud). Tvůrce AI nástroje se proto snaží nabídnout multiplatformní podporu. Například pokud vyvine mobilní AI aplikaci pro iOS, často brzy uvede i verzi pro Android, aby pokryl co největší trh. Na úrovni podnikových řešení je zase důležité, aby AI mohla běžet jak v cloudu poskytovatele, tak případně v privátním cloudu nebo lokálním serveru klienta (to bývá požadavek kvůli bezpečnosti dat). Poskytovatelé AI služeb proto volí flexibilní přístupy – nabízejí variantu formou cloudové API (kdy zákazník volá AI model vzdáleně) i variantu on-premise (kdy je model nainstalován u zákazníka). Dále je důležité, aby AI spolupracovala s již používaným software – například AI modul pro analýzu dat by měl umět číst data v běžných formátech (CSV, SQL databáze atd.) a exportovat výsledky do formátů, se kterými pak pracují jiné nástroje. Tím se eliminuje nutnost konverzí a ručních zásahů. Kompatibilita výrazně zvyšuje přívětivost AI produktu pro zákazníky, protože zapadá do jejich existujícího ekosystému místo toho, aby nutila zákazníka budovat vše kolem nového řešení.
3. **Modulární design:** Z pohledu vývojářské firmy je modulární architektura klíčová pro udržitelnost a *future-proofing* AI

systému. Modulární design znamená, že systém je složen z relativně nezávislých částí (modulů), které lze snadno vyměňovat, aktualizovat nebo rozšiřovat bez nutnosti překopávat celek. Například AI platforma může mít zvlášť modul pro zpracování vstupu (ETL pipeline pro data), zvlášť modul pro samotný prediktivní model a zvlášť modul pro vizualizaci výsledků. Když dojde k zastarávání jedné části (např. objeví se nová a lepší metodika modelování), firma ji může nahradit novou verzí modulu, zatímco zbytek systému zůstane nedotčen. To zjednodušuje jak údržbu, tak integraci s dalšími technologiemi. Pro klienty to znamená, že se mohou dočkat pravidelných vylepšení AI nástroje (protože dodavatel může pružně inovovat moduly) a také že případné úpravy pro specifické potřeby klienta jsou snazší (lze přidat custom modul nebo nahradit ten, který nevyhovuje). Modulárnost rovněž podporuje kompatibilitu – například firma může nabídnout několik variant vstupního modulu podle toho, jaký zdroj dat zákazník používá, a propojit to s jedním univerzálním jádrem AI. Celkově modulární a škálovatelná architektura umožňuje tvůrcům AI rychle reagovat na zpětnou vazbu trhu a udržet produkt kompatibilní s evolucí technologického prostředí.

4. **Standardizace a otevřená rozhraní:** Důležitým aspektem integrace je využívání zavedených standardů v odvětví. To zahrnuje datové formáty (např. JSON, XML), protokoly pro komunikaci (REST API, gRPC) či interoperabilitu modelů (např. formát ONNX pro přenos modelů mezi frameworky). Když AI nástroj podporuje standardizované vstupy/výstupy, významně to usnadňuje jeho propojení s jinými systémy. Například, pokud vývojář AI nabídne API dle REST konvencí s JSON výstupem, každý moderní programátor je schopen toto API zavolat z libovolné platformy bez nutnosti zvláštních klientů. Některé firmy jdou ještě dál a uvolňují část svých technologií jako open-source nebo poskytují detailní dokumentaci pro integrátory. To buduje kolem

produktu komunitu, která může vytvářet další konektory a pluginy. Příkladem je nástroj pro strojové učení TensorFlow od Googlu, který byl uvolněn jako open-source a díky tomu existuje množství knihoven umožňujících jeho použití v různých jazycích, i integrace do různých aplikačních serverů. Pro komerční AI produkty sice open-sourcing kompletního kódu často nepřipadá v úvahu, ale i tak mohou těžit z otevřených standardů. Například IBM Watson nabízí integrace do různých podnikových softwarů právě skrze oficiálně podporovaná rozhraní a SDK. **Standardizace** zkrátka snižuje bariéru vstupu pro zákazníky – mohou svou stávající infrastrukturu propojit s AI nástrojem bez obav, že narazí na proprietární komplikace.

Pohled uživatele (konzumenta AI služby):

1. **Snadné začlenění do pracovního procesu:** Pro uživatele, ať už individuálního či firemního, je výhodné, pokud může AI nástroj zapojit do svých běžných činností s minimálním úsilím. Například pro datového analytika je nejpříjemnější, když nový AI algoritmus může spustit v prostředí, které už zná (třeba jako plugin v jeho oblíbeném statistickém softwaru), místo aby se učil obsluhovat úplně novou aplikaci. Podobně v kancelářském prostředí – pokud je AI asistent integrován přímo do textového editoru či komunikačního nástroje, je daleko použitelnější než samostatná aplikace, do které by uživatel musel přepínat. Uživatelé tedy oceňují, když je AI kompatibilní s jejich workflow. Pro firmy zavádějící AI platí totéž: čím méně nutnosti měnit interní procesy a softwarovou výbavu, tím lépe. Například při zavádění AI do výrobní linky je ideální, když nový modul dokáže komunikovat se stávajícím systémem řízení výroby – operátoři pak dostávají AI doporučení přímo ve známém rozhraní a nemusí přecházet do jiného programu. Snadné začlenění se tedy promítá do vyšší ochoty zaměstnanců AI používat, protože to nevnímají jako *disruptivní* změnu, ale spíš jako vylepšení stávajících nástrojů.

2. **Křížová funkčnost (cross-functionality):** Uživatelé dnes často přeskakují mezi různými zařízeními a platformami – začnou úkol na notebooku, pokračují na mobilu atd. Proto od moderní služby očekávají, že jejich zkušenost bude plynulá napříč platformami. Z pohledu AI aplikace to znamená, že by měla mít konzistentní funkcionalitu a synchronizaci dat mezi verzí pro počítač, webovým rozhraním, mobilní aplikací apod. Pokud například uživatel využívá AI poznámkovou aplikaci s funkcí automatického shrnutí textu, chce mít k těmto shrnutím přístup jak na mobilu, tak na webu, a data by se mu měla synchronizovat. Podobně, pokud AI pomáhá s plánováním úkolů, uživatel ocení integraci s jeho kalendářem a emailovým klientem, aby nemusel informace zadávat vícekrát. Kompatibilita napříč službami je pro uživatele komfort – AI, která si “rozumí” s jinými aplikacemi, ušetří čas a předchází chybám (třeba automaticky přenese doporučený deadline do kalendáře). U firemních uživatelů je cross-functionality také důležitá: třeba výstupy z AI analytického nástroje by měly jít snadno importovat do nástrojů pro vizualizaci reportů, nebo AI modul pro detekci podvodů by měl komunikovat s existujícím systémem pro správu případů. Když AI funguje v souladu s ekosystémem aplikací, který uživatel používá, stává se neviditelnou a o to přínosnější součástí jeho práce.
3. **Bezproblémové aktualizace a rozšíření:** Z pohledu uživatele je frustrující, pokud aktualizace systému nebo přidání nové funkce naruší fungování jejich zavedeného procesu. Proto je velkou výhodou, když AI nástroj umožňuje **plynulé aktualizace** – ideálně na pozadí, bez výpadků a bez nutnosti, aby uživatel cokoliv složitě nastavoval znova. Cloudové AI služby to často zajišťují tak, že nasazují nové verze modelů na serverech průběžně (uživatel ani neví, že se něco změnilo, kromě třeba zlepšeného výkonu). U lokálně instalovaných řešení je dobré, pokud aktualizace probíhá automaticky s možností snadného návratu ke starší verzi v případě problému. Zkrátka, aktualizace by neměla rozbít

integrace, které si uživatel pracně nastavil. Podobně rozšiřování funkcionalit – pokud se uživatel rozhodne využít v AI systému novou modul či propojení s jinou aplikací, mělo by to jít bez nutnosti přeinstalace celého systému. Například AI nástroj může nabízet katalog doplňků (pluginů) a jejich instalace je otázkou pár kliknutí. Rozšiřitelnost je hodnotná i proto, že různí uživatelé mohou mít různorodé potřeby – jeden potřebuje integraci s Outlookem, jiný s Gmailem; dobrý AI systém nabídne obojí jako volitelné moduly. Uživatel pak cítí, že systém roste s jeho požadavky a zároveň se mu nepodsouvají věci, které nechce (protože si instaluje jen doplňky, které využije). To posiluje přívětivost a celkovou spokojenost – investice do naučení se systému se vyplatí, protože systém nezastará, naopak se přizpůsobuje novým výzvám.

Příklady:

AI ve správě zákaznických vztahů (CRM): Představme si firmu, která používá zaběhnutý CRM systém pro evidenci a správu zákazníků a rozhodne se rozšířit jej o AI modul, jenž bude predikovat, kteří zákazníci mají největší potenciál k nákupu nebo naopak riziko odchodu. Pokud dodavatel AI nabízí snadnou integraci takového prediktivního modulu přímo do stávajícího CRM (např. prostřednictvím doplňku nebo API, které CRM podporuje), může firma tuto funkčnost získat během pár dní implementace. AI modul pak třeba přidá ke každému zákazníkovi skóre “pravděpodobnost odchodu” viditelné přímo v CRM rozhraní. Obchodníci firmy nemusejí chodit do jiné aplikace, učit se nové ovládání – jednoduše v rámci svého běžného nástroje vidí novou informaci, se kterou mohou pracovat (např. zaměřit se na zákazníky s vysokým rizikem odchodu a proaktivně jim nabídnout výhodu). Taková integrace zvyšuje užitnost CRM a zároveň nemění zaběhané procesy. Pokud by naopak AI dodavatel nabídl jen samostatný software nepropojený s CRM, firma by musela zavádět paralelní proces: někdo by extrahoval data z CRM, nahrál je do AI systému, tam analyzoval a výstupy ručně zpět přenesl

k obchodníkům – to by bylo pracné a náchylné k chybám, pravděpodobně by se toho zaměstnanci zdráhali účastnit. Plynulá integrace je tedy klíčem k tomu, že AI modul bude skutečně využíván a přinese hodnotu (Singh et al., 2020).

AI ve zdravotnických informačních systémech: Nemocnice často používají komplexní informační systémy (nemocniční IS) pro evidenci pacientů, výsledků vyšetření, léků atd. Zavedení AI, například pro podporu rozhodování lékařů (jako je doporučení diagnózy či léčby), má šanci na úspěch jen tehdy, pokud bude hladce integrováno do tohoto prostředí. Představme si AI, která analyzuje laboratorní výsledky a navrhuje možné diagnózy. Pro lékaře má smysl jen v případech, že se její výstup objeví přímo v rozhraní, které už tak jako tak používá při prohlížení lab. výsledků – třeba jako další kolonka “AI návrh” vedle seznamu hodnot krevních testů. Pokud by lékař měl otevírat speciální AI aplikaci, kam by musel ručně zadat pacientovy hodnoty, pravděpodobně to v návalu práce dělat nebude. Integrace navíc musí být v reálném čase: když sestra uloží nové výsledky testů do systému, AI by je měla automaticky zpracovat a poskytnout závěr dříve, než lékař daný záznam otevře. Technicky to vyžaduje, aby AI modul byl součástí datového toku nemocničního IS – tedy kompatibilita se standardy ve zdravotnictví (např. protokol HL7 pro výměnu zdravotních dat). V praxi se to řeší tak, že dodavatel nemocničního IS buď sám integruje AI od partnera, nebo poskytne API, přes které AI modul data získá a výsledky zapíše zpět. Výsledkem je, že lékaři mohou efektivně využívat AI doporučení bez pocitu, že “používají něco navíc”. To výrazně zvýší jejich ochotu technologii přijmout. Příkladem může být integrace AI do obrazového softwaru (PACS) u radiologů: moderní PACS už umí přímo zobrazit výstupy několika AI modelů (např. obrys podezřelého ložiska v CT snímku), takže radiolog vidí vše v jednom okně (Geis et al., 2019). Taková kompatibilita je zásadní pro to, aby se AI stala běžnou součástí klinické praxe a nebyla jen výstřední laboratorní pomůckou.

3.8 Nákladovost a efektivita

Pohled společnosti (tvůrce a prodejce AI nástroje):

1. **Investice do výzkumu a vývoje:** Vývoj pokročilých AI systémů představuje značnou investici. Společnosti musí alokovat finance na kvalifikované odborníky (datové vědce, machine learning inženýry), na výpočetní infrastrukturu (servery, cloudové služby) i na získávání a označování kvalitních dat. Tyto počáteční náklady mohou být velmi vysoké, zvláště u projektů, které cílí na průlomové modely (trénování špičkových modelů může stát miliony dolarů – viz růst nároků na výpočetní kapacitu podle Amodei & Hernandez, 2018). Pro firmu je tedy klíčové pečlivě řídit rozpočet na R&D a hledat možnosti, jak efektivně využít zdroje: například přebírat open-source modely a upravovat je místo vývoje od nuly, spolupracovat s akademickou sférou na výzkumu nebo využívat cloud s on-demand zdroji místo nákupu vlastního hardware. Navzdory vysokým nákladům je investice do AI pro mnoho společností strategická – pokud uspějí, získají náskok na trhu a mohou vyvinout produkt, který vygeneruje zisk převyšující úvodní výdaje. Taková inovační sázka je však riskantní, a ne každá společnost (zejména menší) si ji může dovolit sama – i proto vidíme v AI oblasti spoustu partnerství, grantů a rizikového kapitálu, který vývoj financuje.
2. **Optimalizace poměru výkon/náklady:** Společnosti stojí před úkolem najít správnou rovnováhu mezi špičkovým výkonem AI systému a jeho nákladovostí. Nejvýkonnější řešení (např. největší modely, extrémně detailní analýzy) mohou mít klesající míru návratnosti – malé zlepšení přesnosti může znamenat mnohonásobné zvýšení výpočetních nákladů. Firmy proto musí rozhodnout, jakou míru přesnosti či komplexnosti opravdu potřebují, aby jejich produkt uspěl, a kdy už další zdokonalování není ekonomicky rentabilní. Příkladem může být rozlišení: je-li AI pro autonomní auto o 99,9 % bezpečná nebo 99,99 % bezpečná,

je to rozdíl možná jedné nehody na milion kilometrů – stojí ale dosažení té druhé devítky desetkrát víc výpočetního výkonu? Tyto úvahy spadají do *cost-benefit analýzy*. Společnost může také nabídnout více variant svého produktu s různou nákladovostí: například “basic” AI službu (levnější, s menší přesností či omezenými funkcemi) a “premium” službu (dražší, využívající plně výkonnou AI). Tím osloví různé segmenty zákazníků a zároveň optimalizuje využití svých výpočetních kapacit. Celkově se jedná o maximalizaci **ROI (návratnosti investice)** do AI – nasadit tak výkonnou AI, která přináší podstatné zlepšení oproti jednodušším řešením, ale nespadnout do oblasti, kde náklady rostou rychleji než přínosy.

3. **Škálovatelnost řešení:** Pro komerčně úspěšné AI produkty je nutné, aby škálovaly – tzn. aby dokázaly obsloužit rostoucí počet zákazníků či větší objemy dat bez dramatického růstu nákladů. Toto úzce souvisí s technickou architekturou (viz výše integrace a výkon), ale i s obchodním modelem. Firma by měla budovat AI tak, aby přidání nového zákazníka mělo relativně malý marginální náklad. Toho lze dosáhnout například modely *SaaS* (software jako služba) – firma jednou vyvine AI platformu a pak ji poskytuje mnoha klientům formou předplatného přes cloud. Náklady na jednoho dalšího uživatele jsou pak zanedbatelné (jen jeho díl zatížení serverů), zatímco příjem může růst lineárně či i exponenciálně s počtem uživatelů. Pokud by však každé nasazení AI vyžadovalo individuální projekt u klienta, byl by růst mnohem nákladnější a pomalejší. Zde je vidět například rozdíl mezi konzultačními projekty a produktovým přístupem: konzultanti udělají *custom* AI řešení na míru jedné firmě (náklad vysoko, škálovatelnost omezená), zatímco produktová firma vyvine univerzálnější AI platformu a tu pak přizpůsobuje minimálně. Samozřejmě, některá odvětví (zdravotnictví, průmysl) přeci jen vyžadují *tailor-made* přístup, ale i tam se firmy snaží vytvářet znovupoužitelné moduly (např. AI jádro je stejné, jen se liší

datové konektory podle klienta). Škálovatelnost jde ruku v ruce s nákladovou efektivitou – pokud zvládnou se stejným týmem a infrastrukturou obsloužit 10× více zákazníků, klesá jednotková cena řešení a roste marže. To je cílem většiny startupů: najít *scalable business model*.

4. **Účinnost operací:** AI může být také interním nástrojem pro firmu, jak zvýšit efektivitu vlastních procesů a tím snížit provozní náklady. Mnoho společností nasazuje AI k automatizaci rutinních úloh – např. automatizované zpracování dokumentů, chatboty v zákaznické podpoře (redukce personálu call centra), prediktivní údržba ve výrobě (snížení prostojů a nákladných havárií) apod. Z pohledu poskytovatele AI řešení toto může být i argument pro prodej: jejich produkt ušetří zákazníkovi peníze tím, že zefektivní nějaký proces. Mělo by to být doložitelné kvantitativními výsledky (např. “naš systém zkrátí průměrný čas zpracování požadavku o 50 %, takže potřebujete o třetinu méně operátorů”). Pro firmu samotnou (pokud AI využívá interně) je pak důležité sledovat, zda se tyto úspory skutečně realizují. Někdy implementace AI slibuje dramatické zlepšení efektivity, ale naráží to na reálný provoz – třeba zaměstnanci technologii nevyužívají naplno nebo vzniknou nové jiné úkony. Proto je součástí nákladové úvahy také change management – zaškolení lidí, úprava procesů, aby se potenciál AI využil a transformoval do reálných úspor či vyšší produktivity. Pro poskytovatele AI je ideální, když může ukázat případové studie: např. firma X díky našemu AI ušetřila ročně Y milionů korun a zkrátila vyřizování zakázek z týdnů na dny. To přesvědčí další zákazníky, že investice do AI se jim vyplácí.

Pohled uživatele (konzumenta AI služby):

1. **Přístupnost řešení (cenová dostupnost):** Z perspektivy zákazníků a koncových uživatelů je důležité, aby AI služby byly cenově dostupné a přinášely odpovídající hodnotu za peníze. Pokud je

cena přemrštěná, i sebelepší technologie se neprosadí, protože zákazníci nedokáží ospravedlnit náklady. Například individuální uživatel možná ocení AI aplikaci pro zlepšování fotografií, ale nebude za ni platit stovky dolarů ročně, když má alternativy zdarma či levněji (byť s horší kvalitou). Firmy zase při nákupu AI řešení počítají návratnost investice – pokud AI stojí X a slibuje ušetřit Y, pak Y by mělo X brzy převýšit. Uživatelé proto očekávají transparentní a spravedlivou cenovou politiku. V oblasti cloudových AI služeb se často setkáváme s modely založenými na předplatném nebo na skutečném využití (pay-per-use), což může být výhodné, protože zákazník platí poměrně k tomu, jaký má z AI užitek. Důležitá je i škálovatelnost ceny – malý podnik by neměl platit stejně jako korporace, pokud využívá menší objem služby. Poskytovatelé to řeší různými tarify. Pro mainstreamové přijetí AI platí, že čím nižší bariéra vstupu (finanční i technická), tím více uživatelů ji vyzkouší a integruje do svých životů či byznysů.

2. **Zvýšení produktivity a hodnoty:** Uživatelé – ať jednotlivci, nebo organizace – si pořizují AI s příslibem, že jim nějak ulehčí práci, zvýší produktivitu nebo přinese jiný hmatatelný přínos. Proto neustále hodnotí, zda ten přínos skutečně dostávají. Pokud například firma nasadí AI nástroj pro automatizované vyplňování formulářů a zjistí, že zaměstnanci nyní zvládnou dvojnásobný objem práce za stejný čas, vnímají to jako výbornou investici. Jednotlivec, který si zaplatí AI aplikaci na organizaci e-mailů, očekává, že mu ušetří třeba hodinu denně, kterou dříve trávil tříděním pošty. Produktivita je tedy klíčovým kritériem spokojenosti – AI musí prokazatelně šetřit čas, snižovat chybovost, zvyšovat kvalitu výstupů nebo jinak přidávat hodnotu. Pokud by tomu tak nebylo (např. AI generuje tolik chyb, že lidé stráví stejně času jejich opravou, jako by to dělali ručně), pak se pro uživatele investice mívá účinkem. Proto uživatelé často pilotně testují AI řešení a sledují KPI (klíčové ukazatele výkonnosti): doba

zpracování, počet vyřízených případů, spokojenost zákazníků atd. Pozitivní výsledky (např. zkrácení doby o 30 %, zvýšení spokojenosti o 10 bodů) posilují důvěru uživatelů v AI a ochotu nadále do ní investovat.

3. **Návratnost investice (ROI):** Zejména pro firemní uživatele je rozhodující, aby investice do AI měla měřitelnou návratnost. Před zavedením technologie si často připravují business case: sečtou náklady (licence, implementace, provoz, školení) a odhadnou úspory nebo dodatečné zisky, které AI přinese (např. méně zaměstnanců na určitém oddělení, vyšší tržby díky lepšímu cílení marketingu). Pokud výsledný poměr vypadá příznivě (řekněme návratnost do 1-2 let, poté čistý přínos), jdou do toho. Po nasazení sledují, zda se těchto přínosů dosahuje. Jestliže AI projekt nesplní očekávání – náklady byly vyšší, než se čekalo, nebo úspory menší – může dojít k jeho omezení či zrušení. To je reálný scénář: průzkumy ukazují, že ne všechny firmy jsou zatím schopné z AI vytěžit hodnotu, kterou si od toho slibovaly (Ransbotham et al., 2019). Pro uživatele je tedy klíčové, že AI není jen “cool technologie”, ale skutečně *zlepšuje ekonomiku* jejich činnosti. Například výrobní podnik, který investoval do AI pro kontrolu kvality, by měl vidět méně reklamací a odpadů, což se promítá do úspor. Pokud by nic takového nenastalo, byla by to pro něj ztrátová investice. U jednotlivců je ROI vnímána spíše subjektivně – třeba předplatné AI aplikace za 200 Kč měsíčně se “vyplatí”, když mi ušetří nervy a hodiny času, které bych jinak věnoval nepříjemné rutině. V opačném případě předplatné zruším. Celkově tak uživatelé tlačí na to, aby hodnota, kterou od AI dostávají, byla větší než cena, kterou za ni platí – což je vlastně definice ROI. Poskytovatelé AI tomu vycházejí vstříc jak cenotvorbou, tak neustálým vylepšováním efektivity svého produktu, aby přidanou hodnotu maximalizovali.

Příklady:

AI pro automatizaci zákaznické podpory: Mnoho společností zavádí chatboty a hlasové asistenty, které zodpovídají časté dotazy zákazníků namísto živých agentů. Z pohledu firmy je hlavním motivem snížení nákladů – lidský personál zákaznické podpory je drahý, zejména pokud má být k dispozici 24/7. AI může teoreticky odbavit značné procento dotazů rychle a bez přestávky. Například banky hlásí, že chatboti dokážou vyřídit 60-80 % běžných dotazů typu “Jak si změním heslo do internetbankingu?” či “Kde najdu IBAN?”, čímž se lidští operátoři mohou věnovat složitějším případům. Úspora se projeví buď v tom, že není potřeba najímat tolik zaměstnanců, nebo mohou stávající obsloužit více klientů – obojí znamená nákladovou efektivitu. Konkrétně IBM uvádí, že nasazení AI virtuálních agentů může snížit náklady zákaznického servisu až o 30 % (Jobanputra, 2024). Pro zákazníka (uživatele podpory) to navíc přináší rychlejší odezvu, takže i jeho spokojenost může růst. Zde tedy dobře navržená AI přináší *win-win*: firma šetří peníze, zákazník čas. Je ovšem třeba správně vyhodnotit ROI – implementace chatbota také něco stojí (licence, vývoj konverzačních skriptů, integrace s databázemi). Pokud by firma vynaložila miliony a chatbot pak řešil jen marginální počet dotazů, návratnost by byla sporná. Proto se typicky nasazuje tam, kde objem interakcí je vysoký, a i malé procento ušetřených lidských hodin představuje velkou sumu.

AI ve zdravotnictví (efektivita a náklady): Zdravotnická zařízení a systémy obecně pracují s omezenými zdroji (lékaři, čas, rozpočet) a nasazení AI jim slibuje zvýšení efektivity péče. Například AI systém, který automaticky vytrídí pacienty na urgentním příjmu podle závažnosti stavu (tzv. triáž), může zkrátit čekací doby a zajistit, že se lékaři věnují nejprve kritickým případům. To potenciálně zachraňuje životy, ale také šetří náklady – předejde se komplikacím, které by byly nákladné na léčbu, kdyby pacient čekal příliš dlouho. Jiný příklad: AI v radiologii může rychle vyřadit snímky bez nálezu, takže radiolog nemusí trávit čas jejich podrobným studiem – vyšetření se tím zlevní a více pacientů může být obslouženo v kratším čase. Nemocnice by

pak mohla třeba ušetřit na tom, že nepotřebuje tolik externích konzultací radiologů nebo může předejít duplicitním vyšetřením. Ekonomické analýzy některých pilotních projektů s AI ukazují slibné výsledky: například studie v UK spočítala, že použití AI pro automatické rozpoznání normálních rentgenových snímků hrudníku může ušetřit radiologům až 20 % času, což v celonárodním měřítku znamená řádově tisíce hodin a odpovídající mzdové náklady ročně (Oakden-Rayner et al., 2020). Samozřejmě, pro schválení takové investice musí management nemocnice vidět, že náklady na zavedení a provoz AI (licence, IT infrastruktura, školení personálu) jsou nižší než hodnota ušetřeného času nebo zlepšených zdravotních výstupů (což se dá promítnout do snížených nákladů na komplikace, kratší hospitalizace apod.). Z pohledu pacientů pak AI přináší hodnotu v podobě kvalitnější a rychlejší péče. Tím se dostáváme k širšímu konceptu: někdy je finanční ROI obtížné přesně vyčíslit, ale jsou zde *kvalitativní přínosy* (spokojenost, lepší výsledky léčby) – i ty se ale nakonec mohou přetavit do ekonomického zisku (např. nemocnice s rychlejší a kvalitnější péčí přiláká více klientů v konkurenčním prostředí). Každopádně, důraz na efektivitu a nákladovost znamená, že i ve zdravotnictví se AI prosadí hlavně tam, kde prokáže schopnost dělat věci rychleji, levněji nebo lépe než tradiční postupy.

3.9 Podpora a znalostní báze

Pohled společnosti (tvůrce a provozovatele AI nástroje):

1. **Technická podpora:** Pro firmy nabízející AI nástroje je poskytování kvalitní technické podpory klientům nezbytné. AI systémy mohou být komplexní a zákazníci (zejména firemní) často potřebují asistenci při integraci, ladění či řešení problémů. Rychlá a efektivní support služba zajišťuje, že případné potíže uživatelů jsou promptně vyřešeny, což zvyšuje spokojenost a brání odlivu zákazníků. Společnosti proto investují do školených podpůrných týmů, vytvářejí více úrovní podpory (level 1 – základní dotazy, level 2 – složitější technické otázky, level 3 –

experti vývojáři pro eskalované případy) a mnohdy nabízejí garantované reakční doby (SLA). Například dodavatel AI chatbot platformy může garantovat podnikové klientele, že kritický problém (např. výpadek služby) začne řešit do 1 hodiny od nahlášení. Taková úroveň servisu je často očekávaná, zejména pokud zákazník platí za enterprise licenční program.

2. **Vzdělávací materiály:** Kromě reaktivní podpory je velmi důležité poskytovat proaktivní vzdělávání a návody pro uživatele. Firmy tvoří rozsáhlou dokumentaci, tutoriály, příklady použití, video návody, případové studie a pořádají webináře či workshopy. Cílem je, aby se zákazníci naučili co nejlépe využívat všechny funkce AI nástroje a byli schopni řešit běžné záležitosti sami. Například platforma pro strojové učení může mít online akademii, kde vývojáři klienta projdou kurzy od základního používání po pokročilé techniky, a tím získají z produktu maximum. Kvalitní knowledge base také snižuje zátěž na technickou podporu – uživatelé si často najdou odpověď v dokumentaci či FAQ sami, místo aby zakládali tiket. Pro společnost to znamená menší náklady na support a pro uživatele rychlejší získání informací (nemusí čekat na odpověď).
3. **Zpětná vazba od uživatelů:** Poskytovatelé AI by měli aktivně sbírat a využívat zpětnou vazbu od své uživatelské základny. To pomáhá jak v podpoře (identifikace často kladených dotazů, problémových míst), tak v dalším vývoji produktu. Firmy mohou mít nastaveny kanály, kde uživatelé sdílejí své zkušenosti – od formálních (pravidelné review meetingy s klíčovými klienty, průzkumy spokojenosti) po neformální (fóra komunity, sociální sítě). Důležitá je rychlá reakce na feedback: když více zákazníků upozorňuje na nějaký nedostatek či chybu, firma by to měla v krátké době řešit (vydat opravu, vylepšit funkcionalitu, doplnit dokumentaci). Tím jednak zlepšuje produkt a jednak demonstruje zákazníkům, že jejich hlas je slyšet – to buduje loajalitu. Například když uživatelé žádali, aby AI design nástroj

podporoval import z Photoshopu, a dodavatel tuto funkci v další verzi implementoval, posílí se tím vztah s klienty (vidí konkrétní výsledek svého podnětu).

4. **Komunitní rozvoj:** Vytvoření komunity kolem AI produktu může výrazně napomoci sdílení znalostí a vzájemné podpoře mezi uživateli. Společnosti podporují vznik diskuzních fór, uživatelských skupin, pořádají hackathony, meetupy a certifikační programy. Například opensourcové AI frameworky (TensorFlow, PyTorch) mají obrovské komunity vývojářů, kteří si navzájem radí a řeší problémy, což snižuje potřebu oficiální podpory a zároveň přináší nové nápady a doplňky. I pro komerční produkt může existovat komunita – zejména pokud je rozšiřitelný (lidé mohou psát pluginy, sdílet modely atd.). Firma může tuto komunitu stimulovat nabízením “community edition” zdarma pro nadšence, organizováním konferencí či soutěží. Otevřená a aktivní komunita funguje jako crowdsourcing podpory a inovací: uživatelé se učí jeden od druhého, sdílejí best practices, a zároveň to vytváří ekosystém, který kolem produktu roste (třetí strany vyvíjejí doplňky, píší blogy s tipy atd.). To všechno zvyšuje hodnotu, kterou zákazníci z produktu mají, aniž by to přímo zatěžovalo zdroje firmy – firma musí jen investovat do základní facilitace (platforma pro fórum, občasná moderace, oficiální přítomnost na akcích).

Pohled uživatele (konzumenta AI služby):

1. **Snadný přístup k informacím:** Uživatelé chtějí mít k dispozici informace, které jim pomohou využít AI nástroj co nejlépe. Oceňují, když je dokumentace přehledná, hledání v ní jednoduché a když existují různé formy podle preferencí – někdo rád čte manuál, jiný se raději podívá na krátké instruktážní video. Důležité je, aby informace byly snadno dostupné – ideálně online, zdarma (či součástí licence) a vždy aktuální k dané verzi produktu. Například vývojářský tým ocení API dokumentaci

hostovanou na webu s možností fulltextového vyhledávání a příklady kódu; běžný uživatel zase třeba krokové průvodce typu “Jak nastavit XY”. Pokud uživatel narazí na problém a dokáže během pár minut najít článek v knowledge base s jeho řešením, roste jeho sebedůvěra v používání technologie. Naopak, pokud by musel lovit informace v neaktuálních PDF manuálech nebo čekat den na odpověď podpory, jeho dojem se zhorší. Rychlý přístup k relevantním informacím je tedy podstatnou součástí uživatelské přívětivosti.

2. **Sebevzdělávání:** Moderní uživatelé jsou zvyklí, že se mohou sami naučit pracovat s novými nástroji pomocí online zdrojů. Mnozí preferují samostudium vlastním tempem před absolvováním povinných školení či pročitáním rozsáhlých manuálů. Proto uvítají, když výrobce AI nástroje poskytne interaktivní kurzy, tutoriály nebo e-learning, kde si mohou nabyté dovednosti i vyzkoušet. Například interaktivní webová cvičiště (tzv. sandboxes) umožňují uživatelům experimentovat s AI modelem s instrukcemi krok za krokem. To je populární u vývojářů (kteří mohou cvičně volat API v prohlížeči podle návodu) i u neprogramátorů (třeba trénink, jak interpretovat výstupy AI diagnostického systému pro lékaře). Uživatelé tak mají pocit kontroly nad svým učením a mohou se k materiálům vracet podle potřeby. Firmy to podporují i poskytováním certifikací – uživatel má motivaci projít určité kurzy a získat oficiální osvědčení, že daný nástroj ovládá. Možnost sebevzdělávání je pro mnoho lidí důležitá, protože tempo vývoje je rychlé a formální školení vše nepokryjí; díky dobrým zdrojům se mohou průběžně učit novinky a zlepšovat své schopnosti.
3. **Komunitní podpora:** Jak již bylo zmíněno, uživatelé často hledají odpovědi a rady i od komunity ostatních uživatelů. Mnozí vyhledají řešení problému raději na diskuzním fóru nebo Q&A platformě (Stack Overflow, Reddit apod.) než by kontaktovali oficiální podporu. Důvodem je jednak rychlost (komunita může

reagovat velmi pohotově, zvláště u populárních produktů) a také různorodost pohledů – někdy jiný uživatel narazil na stejný problém a může poskytnout cenné praktické tipy z vlastní zkušenosti. Uživatelé tedy oceňují, pokud existuje aktivní komunita, kde mohou položit dotaz a dostat odpověď, nebo jen prohledat archiv již vyřešených problémů. Pro ně je to pohodlné, cítí se být součástí širšího uživatelského společenství. Navíc se tak často dozví i o různých neoficiálních vylepšeních, tricích a příkladech použití, které by v oficiální dokumentaci nebyly. Společnosti to mohou nepřímo podporovat tím, že třeba oficiální zástupci občas na komunitním fóru pomohou, nebo aspoň naslouchají diskuzím a sbírají nápady. Uživatelé tak vidí, že komunita i výrobce spolu komunikují. Příkladem může být komunita okolo nástroje pro data science – uživatelé si tam sdílejí vlastní knihovny funkcí, datové sety radí, jak řešit specifické analýzy, a výrobce občas přidá oficiální stanovisko či upozorní na nový update řešící diskutovaný bug.

4. **Průvodce řešením problémů:** Ačkoli se firmy snaží dělat AI nástroje co nejrobustnější, uživatelé mohou narazit na různé potíže – od drobných (něco nejde nastavit) po zásadní (systém padá). Pro uživatele je důležité mít k dispozici nástroje a návody pro řešení běžných problémů svépomocí. To mohou být troubleshooting sekce v dokumentaci – typu “Pokud se stane A, zkontrolujte B”, nebo diagnostické utility, které uživateli pomohou identifikovat problém (např. testovací skript, který ověří, zda je správně nastavené připojení k databázi). Tím, že uživatel dokáže vyřešit problém sám (nebo s pomocí interní IT podpory, pokud jde o firemní prostředí), šetří čas sobě i dodavateli. Zároveň to zvyšuje jeho důvěru v systém – vidí, že pro známé problémy existují postupy a že se nejedná o neřešitelné mystérium. Například u AI zařízení (třeba průmyslový senzor s AI) může být přiložen seznam chybových kódů s vysvětlením a návodem “chyba 37: zkontrolujte připojení kamery”. U

softwarové platformy zas často existuje log soubor nebo diagnostický dashboard, kde uživatel vidí stav komponent – to je cenné pro adminy, kteří takto mohou odhalit, kde se proces zasekl. **Self-service troubleshooting** je dnes trendem obecně v IT: uživatelé preferují, když nemusí volat podporu kvůli každé maličkosti a mohou nejprve zkusit problém vyřešit sami s oporou kvalitních materiálů.

Příklady:

AI pro správu sociálních médií (marketingový nástroj): Řekněme, že firma používá AI platformu, která automaticky analyzuje výkon příspěvků na sociálních sítích a doporučuje optimální časy publikování a obsah. Pro marketingový tým, který ji používá, je důležité, aby z ní dostal maximum. Dodavatel takového nástroje proto může nabízet bohatou znalostní bázi: např. online portál s tipy, jak interpretovat různé metriky, příručky “best practices” pro sociální sítě, webináře o trendech v obsahu. Marketingoví specialisté se tak nejen naučí ovládat software, ale též získají cenné know-how v oboru – což vnímají jako velkou přidanou hodnotu. Když narazí na problém (třeba doporučení se zdají nevhodná pro specifické publikum), mohou se podívat do FAQ, kde najdou, jak takovou situaci řešit (možná je potřeba upravit vstupní nastavení AI ohledně cílové skupiny). Pokud si nevědí rady, vědí, že existuje komunita dalších marketérů používajících tutéž platformu – třeba na LinkedIn nebo oficiálním fóru – kde se mohou zeptat na zkušenosti. Mohou tam zjistit, jak kolegové z jiných firem využívají pokročilé funkce, nebo se poučit z jejich chyb. Dodavatel mezitím sbírá tuhle zpětnou vazbu a třeba zjistí, že potřebuje vylepšit modul pro analýzu Instagramu, protože více uživatelů si na něj stěžovalo – a v dalším release to opraví (k čemuž rovnou připraví článek “Novinky ve verzi 2.5: zlepšená analýza Instagram engagementu podle vašich podnětů”).

AI ve zdravotnických aplikacích: Zdravotnický personál začíná používat AI aplikaci pro podporu diagnostiky z rentgenů plic (např.

detekce počínajícího zápalu plic). Aby ji efektivně využili, potřebují proškolení a průběžnou podporu. Dodavatel tak poskytne nejen úvodní školení s prezentací a praktickými ukázkami, ale také podrobný manuál se screenshoty, kde je popsáno, jak chápat výstupy AI (např. skór pravděpodobnosti a grafické zvýraznění nálezu). Dále třeba nabídne e-learning kurz akreditovaný pro lékaře, kde si mohou prohlédnout mnoho příkladů snímků a vyzkoušet si, jak by reagovali a co AI doporučí – to jim pomůže sladit jejich rozhodování s AI a pochopit, kdy jí věřit a kdy být opatrní. Pokud lékař narazí na zvláštní případ, kdy si není jist, zda AI funguje správně (např. u pacienta s atypickým nálezem), může kontaktovat specializovanou podporu dodavatele – ten má k dispozici radiology-konzultanty, kteří s ním případ proberou (to je součástí premium podpory). Tím se lékař i něco přiučí (získá odborné vysvětlení, jak AI v takovém případě pracuje, a co doporučují experti). Výrobce takové AI navíc sbírá anonymizovaná data a vidí, že třeba uživatelé často tápou u nálezů u pacientů s určitým typem předchozí operace – to ho vede k vylepšení algoritmu nebo k doplnění dokumentace o sekci “Speciální situace: pacienti s kardiostimulátorem mohou způsobit falešný nález, berte to v úvahu...”. Z perspektivy zdravotnického zařízení je tak zajištěno, že investice do AI je podpořena silnou znalostní základnou, což jim dává větší jistotu při jejím používání. Pokud by tomu tak nebylo (např. by dodavatel jen nainstaloval systém a nechal personál, ať si poradí sám), je vysoké riziko, že uživatelé AI neporozumí správně a buď ji nebudou využívat, nebo hůře – budou ji využívat špatně a dojde k omylům. Proto je v prostředí jako zdravotnictví podpora a vzdělávání naprosto klíčové.

Z výše uvedeného rozboru vyplývá, že efektivní komunikace a práce s AI není jen o samotné pokročilosti algoritmů, ale velmi závisí na uživatelském přístupu, důvěře, etických ohledech, výkonu a kvalitní podpoře. Firmy vyvíjející AI i uživatelé, kteří ji nasazují, musí úzce spolupracovat a zohledňovat oba pohledy – technický i lidský – aby bylo dosaženo skutečného přínosu. V rámci našeho cíle – porozumět,

jak co nejlépe využít AI technologie – je tedy nutné brát v potaz všechny uvedené aspekty a hledat rovnováhu mezi nimi.

4 Strategie pro bezproblémovou komunikaci s AI

V této části se zaměříme na klíčové strategie, které pomáhají uživatelům vytěžit maximum z interakce s umělou inteligencí (AI). V době, kdy jsou pokročilé AI systémy (zejména velké jazykové modely) široce dostupné milionům lidí, je nezbytné pochopit, **jak s těmito systémy efektivně komunikovat**. Nejedná se jen o položení jedné otázky – důležité je správně strukturovat své myšlenky, záměry a očekávání, aby odpovědi byly co nejpresnější a nejužitečnější. Tato potřeba spadá do konceptu *AI gramotnosti* uživatelů, tedy schopnosti lidí účelně, efektivně a eticky využívat AI technologie (Pinski & Benlian, 2024). Podle Světového ekonomického fóra bude rozvoj takové gramotnosti klíčový pro úspěšné začlenění AI do společnosti (World Economic Forum, 2022). Mnoho běžných uživatelů dnes totiž netuší, jak s AI pracovat – někteří se mylně domnívají, že musí zadávat dotazy programátorským stylem nebo zvláštním formálním jazykem, zatímco ve skutečnosti je neefektivnější přístup velmi podobný přirozené lidské konverzaci (Mollick, 2023). Moderní AI rozumí běžnému jazyku a díky pokročilému zpracování přirozené řeči dokáže interpretovat i složitější instrukce formulované jako volný text (Bsharat et al., 2023). Níže popsané strategie vycházejí z poznatků výzkumu a osvědčených postupů a mají společný cíl: **umožnit uživatelům stát se sebevědomějšími a úspěšnějšími v komunikaci s AI** a tím plně využít potenciál těchto nástrojů pro osobní i profesní účely.

Proč jsou tyto strategie důležité? Shrňme nejprve, jaké konkrétní přínosy správná komunikace s AI přináší:

- **Maximalizace přesnosti:** Správná formulace dotazů výrazně zvyšuje přesnost a relevanci odpovědí od AI. Výzkumy ukazují, že dobře navržené vstupy (prompty) mohou zlepšit kvalitu a faktickou správnost výstupu modelu o desítky procent (Bsharat

et al., 2023). Jinými slovy – čím jasněji a cíleněji je požadavek formulován, tím přesnější odpověď model poskytne.

- **Úspora času:** Efektivní komunikace s AI šetří čas, který by uživatel jinak strávil opakovaným dotazováním nebo hledáním informací. Jasný dotaz s dostatečným kontextem obvykle vede k použitelnější odpovědi hned napoprvé, čímž se minimalizuje potřeba dalších upřesňujících dotazů či oprav (OpenAI, 2023).
- **Zvyšování produktivity:** S lepším porozuměním tomu, jak AI pracuje a jak jí pokládat dotazy, mohou uživatelé rychleji dosáhnout svých cílů. Studie zaměřené na nasazení generativní AI v pracovním prostředí zjistily, že uživatelé schopní správně instruovat modely dokončují zadané úkoly rychleji a s vyšší kvalitou výsledku než ti, kteří AI nevyužívají (Noy & Zhang, 2023). Učení se efektivní komunikaci s AI tak přímo souvisí s růstem produktivity.
- **Snížení frustrace:** Pochopení omezení AI a osvojení si způsobů, jak s těmito omezeními pracovat (namísto boje proti nim), vede k lepší celkové zkušenosti uživatele. Uživatel, který ví, jak *předejít nejasnostem*, jak reagovat na nepřesnou odpověď nebo jak ověřovat fakta, bude při práci s AI méně frustrován a získá z interakce hodnotnější informace (Pinski & Benlian, 2024).

Každá z následujících strategií má své specifické výhody a lze ji použít v různých situacích. Jejich pochopením a uplatňováním se stanete v komunikaci s AI nejen úspěšnějšími, ale také sebevědomějšími. Ve výsledku můžete AI nástroje plně využít k rozvoji svých projektů, k učení i ke každodennímu zjednodušení práce.

4.1 Buďte konkrétní

Strategie „**Buďte konkrétní**“ je **zásadní pro účinnou komunikaci s AI**. Znamená to, že při pokládání dotazu je vhodné uvést **co nejvíce relevantních detailů a upřesnění**, aby model jednoznačně pochopil,

na co se ptáte. Vágní nebo příliš obecné dotazy často vedou k obecným odpovědím, které nemusejí přesně odpovídat vašemu záměru. Naopak specificky formulovaná otázka zužuje prostor možných interpretací a přesměrovává AI k poskytnutí informací, jež přesněji odpovídají vašim potřebám.

Výzkumníci zdůrazňují, že *granularita vstupu je přímo úměrná užitečnosti výstupu* – čím konkrétnější dotaz, tím hodnotnější odpověď (Bsharat et al., 2023). Pokud například namísto obecné prosby typu „Řekni mi něco o psech.“ položíte specifický dotaz „*Jaké jsou charakteristické rysy plemene labradorský retrívr a jaké jsou jeho běžné zdravotní problémy?*“, model se zaměří přesně na dané plemeno a poskytne cílenější informace. V prvním případě by AI pravděpodobně vygenerovala dlouhý obecný přehled o psech, zatímco ve druhém případě nabídne konkrétní údaje o labradorském retrívrovi (historii plemene, povahových rysech, typických nemocích apod.), což je pro tazatele daleko užitečnější.

Proč je konkrétnost tak důležitá? Jazykové modely sice byly trénovány na obrovském množství textů, ale **interpretace dotazu vždy závisí na přesnosti a jednoznačnosti formulace**. Nejednoznačný vstup aktivuje příliš širokou oblast znalostí modelu, zatímco dobře specifikovaný dotaz jeho pozornost zúží *na relevantní informace*. Výsledkem bývá přesnější a fakticky správnější odpověď (Cook, 2023). Kromě toho konkrétní dotaz snižuje riziko nedorozumění – model nemusí „hádat“, co přesně uživatel chce, a tím pádem méně tápe. Empirické testy ukazují, že podrobné instrukce vedou k výstupu s méně chybami a nepřesnostmi než obecné zadání (Bsharat et al., 2023). Neil (2023) trefně poznamenal, že když uživatel formuluje dotaz precizně, často si tím ujasní i vlastní myšlenku, což dále zvyšuje kvalitu výsledného dialogu s AI.

Při aplikaci této strategie však pozor na opačný extrém: *přílišná upovídánost* nebo zahrnutí irelevantních detailů může být kontraproduktivní. Je dobré najít rovnováhu mezi poskytnutím

dostatečných informací a udržením jasnosti. Obecně platí, že do promptu patří vše, co může ovlivnit požadovanou odpověď, ale nic víc. Například pokud potřebujete technické řešení, uveďte konkrétní problém a prostředí, ale nezabíhejte do nesouvisejících osobních komentářů.

Doporučení pro uživatele: Před odesláním dotazu se zkuste sami sebe zeptat, zda je zřejmé, *na co přesně* chci odpověď. Pokud by dotaz mohl mít více různých významů nebo je příliš obecný, doplňte upřesnění (např. místo „Jak zlepšit web?“ raději „Jak mohu zvýšit návštěvnost svého webu zaměřeného na prodej knih?“). Uveďte všechna klíčová slova či parametry, které by mohly ovlivnit odpověď (časové období, odbornou oblast, formu výstupu apod.). Jak potvrzuje například studie od Bsharat et al. (2023), **specifické a detailní pokyny** dávají modelu jasný „návod“, a výsledkem je odpověď lépe odpovídající vašemu skutečnému dotazu.

4.2 Sdělte svůj záměr

Strategie „**Sdělte svůj záměr**“ spočívá v tom, že **AI poskytnete kontext a účel vašeho dotazu** – jinými slovy, vysvětlíte, *proč* se ptáte a *co od odpovědi očekáváte*. Tím napomůžete modelu lépe porozumět vašim potřebám a přizpůsobit tomu generovanou odpověď. Jelikož AI (zatím) nedokáže číst myšlenky ani dokonale odhadnout implicitní kontext, je na uživateli, aby užitečné kontextové informace sdělil přímo v promptu (Urban, 2023).

Tato strategie často zahrnuje *nastavení role nebo scénáře*. Například můžete AI rovnou nastínit, v jaké roli má odpovídat: „*Představ si, že jsi zkušený finanční poradce. Poradiš mi...*“ nebo „*Jsi odborník na historii umění a já se ptám...*“. Výzkumy ukazují, že poskytnutí takového kontextu (tzv. *role prompting*) vede k relevantnějším a konzistentnějším odpovědím v dané doméně (Bsharat et al., 2023). Model dokáže do určité míry **simulovat požadovanou roli** a stylizovat odpověď přiměřeně tomu, jak by odpovídal skutečný odborník nebo jak odpověď odpovídá zamýšlenému účelu. Bsharat et al. (2023)

konstatují, že větší modely mají značnou kapacitu pro takovou simulaci a *čím přesněji je úkol nebo role zadána, tím lépe se model ve svých odpovědích treťí do uživatelových očekávání*. Proto je například **užitečné explicitně specifikovat cílové publikum** nebo účel textu – odpověď pro pětileté dítě bude jinak formulována než tatáž odpověď určená expertovi v oboru.

Proč sdělovat záměr? Protože tím *využíváte adaptability AI* a její schopnosti přizpůsobit výstup kontextu. Moderní AI chatboti často zahrnují mechanismy na rozpoznání uživatelského záměru (např. analýzou klíčových slov a frází v dotazu zjistí, zda uživatel chce definici, radu, názor apod.) a pokud takovou informaci poskytnete přímo, model méně tápe a spíše zvolí odpověď odpovídající vašemu účelu (Wu et al., 2022). Například dotaz „*Jaký je recept na koláč?*“ je poměrně obecný – AI může odpovědět jakýmkoli receptem na libovolný koláč. Pokud ale svůj záměr upřesníte: „*Hledám jednoduchý recept na jablkový koláč pro začátečníky, který nevyžaduje speciální kuchyňské vybavení.*“, AI teď chápe, co přesně potřebujete (snadný recept z běžných surovin, vhodný pro laika) a tomu přizpůsobí odpověď – pravděpodobně nabídne jednoduchý postup, vyhne se odborným pojmům a zmíní běžné kuchyňské náčiní. Tento kontext v promptu funguje podobně jako kdybyste se na totéž ptali člověka: také mu obvykle nastíníte, proč informaci chcete, aby ji mohl lépe zarámovat.

Je třeba zdůraznit, že sdělit záměr neznamená psát rozsáhlé eseje o svých motivech, ale **jasně formulovat účel** dotazu. AI nepotřebuje emotivní či nadbytečné vysvětlení, stačí stručně uvést relevantní kontext. Například: „*Píšu seminární práci z ekonomie, proto mi prosím vysvětlí koncept inflace na jednoduchém příkladu.*“ – zde je jasný kontext (student, seminární práce) i očekávání (jednoduché vysvětlení s příkladem). Model pak odpověď pravděpodobně přizpůsobí úrovni studenta a použije vhodný příklad, místo aby dával příliš odbornou definici z učebnice.

Doporučení pro uživatele: Zkuste do promptu zahrnout větu či dvě, které nastíní, **kdo jste nebo co potřebujete**. Můžete přímo uvést, k čemu výstup použijete (např. „*potřebuji to pro prezentaci laikům*“ nebo „*srovnávám různé názory pro článek*“ apod.). Pokud chcete určitý styl odpovědi, neváhejte to zmínit (např. „*odpovídej mi prosím jako učiteli na střední škole*“). Čím více relevantního kontextu poskytnete, tím větší je šance, že AI správně pochopí váš záměr a odpoví způsobem, který tento záměr naplní (Mollick, 2023). Pokud odpověď neodpovídá tomu, co jste zamýšleli, zkuste upřesnit prompt o další informaci o kontextu nebo formě – často to přinese výrazné zlepšení relevance odpovědi.

4.3 Správný pravopis a gramatika

Strategie „**Správný pravopis a gramatika**“ připomíná, že i když komunikujeme s AI, stále komunikujeme *v přirozeném jazyce* – a srozumitelnost našich vstupů významně ovlivňuje kvalitu výstupů. Jazykové modely interpretují dotazy doslovně na základě řetězců znaků, které zadáme. Pokud jsou v promptu pravopisné či gramatické chyby, může dojít k nesprávné interpretaci dotazu nebo k nejednoznačnosti, což se pak projeví v odpovědi (Vadlapati, 2023). Samozřejmě, moderní modely jsou do jisté míry robustní a dokáží *tolerovat drobné překlapy* či běžné pravopisné chyby – například z kontextu odhadnou, že „*jak se delaj palacinky*“ znamená „*jak se dělají palačinky*“, a odpoví správně. Nicméně spoléhat se na to nelze: zvláště u méně obvyklých slov nebo při větším množství chyb může AI význam dotazu pochopit chybně nebo si není jistá, na co přesně se ptáte (Vadlapati, 2023). Například překlep v názvu odborného termínu může vést k tomu, že model nebude vědět, o čem jde, a buď odpoví vyhýbavě, nebo se pokusí výraz odhadnout a může přitom *udělat chybu*.

Další problém nastává u gramatických chyb a nejednoznačné větné stavby. Lidský čtenář si často dovede domyslet, co tím autor „myslel“, ale AI takovou jistotu nemá – opírá se o pravděpodobnostní model

jazyka. Například špatně umístěná čárka nebo vynechané slovo může posunout význam věty. Uvažme dotaz: „*Doporuč mi prosím knihy které se zabývají strojovým učením.*“ – chybí čárka před „*které*“. Model to možná pochopí stejně, jako kdyby tam čárka byla, ale u složitějších vět by interpunkční chyba mohla vést k nepochopení vztahů ve větě. **Drobné chyby mohou způsobit velké rozdíly v interpretaci.**

Je také třeba zmínit, že trénovací data jazykových modelů tvoří převážně *texty s korektní gramatikou a pravopisem* (např. knihy, články). Modely tedy nejlépe rozumějí vstupům, které jsou podobně bezchybné (Vadlapati, 2023). Pokud jim předložíme text plný překlepů nebo nespisovných tvarů, vystavujeme je vstupu, který je statisticky méně obvyklý, a tím zvyšujeme riziko chybné analýzy dotazu. Stručně řečeno: **čím čistší a správnější jazyk v promptu použijeme, tím spolehlivější bude porozumění AI a tím pádem i její odpověď.**

Doporučení pro uživatele: Před odesláním dotazu věnujte několik sekund kontrole, zda neobsahuje *zjevné překlepy nebo chyby*. Pokud píšete v jazyce, který není vaší silnou stránkou (např. anglicky nerodilý mluvčí), zvažte rychlé využití nástroje na kontrolu pravopisu. Některé AI nástroje dokonce umožňují opravit prompt dodatečně – můžete klidně napsat: „*Omlouvám se, v předchozí zprávě bylo několik překlepů, zde je opravená verze dotazu:*“ a zopakovat otázku správně. U složitějších dotazů se vyplatí formulovat věty jednoznačně a raději je rozdělit než psát jednu dlouhou souvětí, ve kterém by se mohl ztratit význam. Praktickým tipem je *představit si, že totéž, co píše AI, dávám přečíst kolegovi* – pochopil by to kolega správně? Pokud ano, je velká šance, že AI také. A pokud náhodou AI odpoví nesmyslně nebo offtopic, stojí za to zkontrolovat, zda v původním promptu nebyla chyba, která ji mohla zmást.

4.4 Říďte formát výstupu

Strategie „**Říďte formát výstupu**“ znamená, že uživatel **specifikuje, v jaké formě chce informace obdržet** – zda má AI odpovědět například přehledovým seznamem, odstavcem vysvětlujícího textu, tabulkou, kódem, pseudo-kódem, bodovými kroky apod. Jazykové modely jsou překvapivě dobré v dodržování instrukcí ohledně formy, pokud jsou tyto instrukce jasně uvedeny v promptu (OpenAI, 2023). Využitím této strategie lze dosáhnout odpovědí, které jsou lépe strukturované a snáze použitelné pro zamýšlený účel.

Různé úkoly přirozeně volají po různém formátu výstupu. Pokud se například ptáme na *postup* či *návod*, může být vhodnější číslovaný seznam kroků, aby šel postup jednoduše následovat. Na dotaz „*Jak vyměnit pneumatiku u auta?*“ by tedy bylo možné doplnit: „*Uveď jednotlivé kroky jako očíslovaný seznam.*“ – AI pak s velkou pravděpodobností odpoví třeba 1.–10. bodem, namísto souvislého odstavce. Pro dotazy s požadavkem na *srovnání* (např. porovnání dvou technologií) se někdy hodí tabulka, kterou si lze vyžádat: „*Porovnej X a Y v tabulce podle následujících kritérií...*“. U dotazů vyžadujících definici či vysvětlení pojmu je zase často vhodný stručný odstavec, případně následovaný příkladem – i to lze explicitně žádat („*Definuj koncept Z a uveď jeden konkrétní příklad v nové odrážce.*“). Pokud formát neuvedeme, model zvolí *implicitní styl*, většinou jeden či dva odstavce textu. To může být v pořádku, ale ne vždy to odpovídá optimální podobě výstupu pro náš účel.

Proč řídit formát? Předně to usnadňuje orientaci ve výsledku. Strukturovaná odpověď se lépe čte a klíčové informace z ní lze snáze vyzvednout. Uživatel tak nemusí dodatečně strukturovat neorganizovaný text sám. Navíc pokud uživatel potřebuje výstup dále zpracovat (např. zkopírovat seznam, využít části textu ve vlastní zprávě, vložit kód do programu), získá jej přímo v požadovaném tvaru. Výzkumné práce o interakci člověk–AI uvádějí, že **specifikace formy výstupu zvyšuje použitelnost získané odpovědi** a šetří čas při jejím

dalším využití (Liu, 2023). Z pedagogického hlediska je také známo, že informace prezentované ve vhodném formátu (např. výčtem vs. souvislým textem) se lépe zapamatují a porozumění je rychlejší – to platí i pro studium s pomocí AI.

Modely jsou schopné vygenerovat i poměrně komplexní struktury, včetně HTML či Markdown formátování, pokud to specifikujete. Například požadavek *„Poskytni odpověď jako Markdown tabulku.“* povede k tomu, že AI odpoví v syntaxi tabulky pro Markdown (což se hodí při tvorbě dokumentací, README souborů atd.). Podobně *„napiš odpověď formou pseudokódu“* donutí model stylizovat text jako blok kódu. Těch možností je mnoho a **uživatel by měl vědět, že má právo si formát vyžádat.**

Doporučení pro uživatele: Už při kladení dotazu zvažte, *k čemu odpověď budete potřebovat a jakou formu byste ocenili.* Nebojte se to říct nahlas – například: *„Shrň výše uvedené informace do třech odrážek.“*, *„Napiš odpověď ve formě krátkého bullet point seznamu.“*, *„Vygeneruj prosím ukázkový kód v jazyce Python, který demonstruje tuto funkci.“* apod. Získáte tak výstup, se kterým můžete rovnou pracovat. Když nejste s formátem spokojeni, můžete model požádat o přeformátování (např. *„Můžeš tu odpověď ztabulkovat?“* nebo *„Udělej z toho souvislý text místo odrážek.“*). Dobře cílený formát výstupu dokáže **výrazně zvýšit přehlednost a užitečnost** informací získaných od AI (OpenAI, 2023). Zejména v situacích, kdy využíváte AI pro vytváření obsahu, studijní materiály nebo pracovní podklady, se vyplatí mít odpověď ve vhodné struktuře – minimalizujete tím nutnost dalších úprav a můžete se soustředit na samotný obsah.

4.5 Pokládejte doplňující otázky

Strategie **„Pokládejte doplňující otázky“** upozorňuje, že komunikace s AI může a často by měla být interaktivní a iterativní. Není nutné (a ani vždy možné) získat kompletní odpověď na složitý problém jedním dotazem. Místo toho je efektivní postupovat v dialogu: *nejprve položit obecnější otázku a následně na odpověď reagovat upřesňujícími*

dotazy, které rozvinou detail, ujasní nejasnosti nebo posunou konverzaci požadovaným směrem. Tento způsob využívá toho, že moderní chatovací AI udržují kontext konverzace – pamatují si, na co se uživatel ptal a co bylo odpovězeno, takže následná otázka může přímo navazovat (Liu, 2023).

Pro uživatele to znamená, že není nutné – a mnohdy ani vhodné – se snažit *vměstnat všechny možné subdotazy do jednoho promptu*. Místo složitého zadání je často lepší začít jednodušší otázkou a **postupně prohlubovat**. Například pokud zkoumáte nějaké téma, můžete začít dotazem: „*Co je X?*“. Po obdržení základní definice se ptát dál: „*Můžeš uvést příklad využití X v praxi?*“; dostanete příklad a pak třeba: „*Jaké jsou nevýhody X oproti Y?*“ a tak dále. Tímto způsobem se z modelu **postupně extrahují znalosti**, podobně jako když vedete rozhovor s odborníkem – nejprve obecné seznámení, pak konkrétní otázky k jednotlivým aspektům.

Výhody doplňujících otázek jsou několikeré. Zprv, **získáte hlubší porozumění** – můžete jít pod povrch základní odpovědi a nechat si ozřejmit detaily nebo souvislosti, které vás zajímají nejvíce. Zadruhé, upřesňujete směr konverzace: pokud původní odpověď zabrousila jinam, doplňující otázkou ji můžete přesměrovat: „*To není to, co jsem měl na mysli – mohl bys místo toho vysvětlit aspekt Z?*“. Model pak na základě vaší reakce lépe pochopí, co přesně hledáte. Za třetí, iterativní dotazování **pomáhá AI lépe vám porozumět** – s každou další otázkou se upřesňuje kontext a záměr, takže odpovědi bývají stále relevantnější (OpenAI, 2023).

Z výzkumu interakcí vyplývá, že uživatelé, kteří aktivně vedou konverzaci a reagují doplňujícími dotazy, jsou spokojenější s výsledkem než ti, kteří po první odpovědi končí, i když nebyla ideální (OpenAI, 2023). Je to logické: když se člověk nespokojí s povrchní či nejasnou reakcí a zeptá se znovu lépe, často dostane kvalitnější odpověď. AI navíc nemá problém se opakovaně doptávat – na rozdíl od lidského odborníka ji tím neobtěžujete ani neurazíte. Můžete

klidně říct: „Tomu a tomu části nerozumím, můžeš to vysvětlit jednodušeji?“ nebo „A co v případě, že...?“. AI trpělivě odpoví znovu, třeba z jiného úhlu nebo podrobněji.

Doporučení pro uživatele: Neváhejte v konverzaci pokračovat.

Ptejte se na detaily, které vám v odpovědi chyběly. Pokud dostanete příliš obecnou odpověď, požádejte o konkrétní příklady („Můžeš uvést konkrétní případ z praxe?“). Když AI použije termín, kterému nerozumíte, položte doplňující dotaz typu „Co znamená [termín] v tomto kontextu?“. Tím si interakci přizpůsobíte svým znalostem. Můžete také **zkusit otázku přeformulovat** nebo přidat informaci, která modelu předtím unikla: jestliže například odpověď opomněla určité hledisko, připomeňte ho („A jak je to v případě malých firem?“). Tato strategie promění statickou Q&A interakci v dynamický dialog, který vede k bohatšímu a přesnějšímu výsledku. Pamatujte, že AI je k dispozici jako partner v dialogu – *využijte toho a ved'te konverzaci aktivně*, dokud nezískáte informace, které potřebujete.

4.6 Výzva k ověření faktů

Strategie „**Výzva k ověření faktů**“ reflektuje jeden z aktuálních limitů generativních AI modelů: ačkoli tyto modely produkují text s překvapivě lidskou plynulostí a sebevědomým tónem, nemají zaručenou pravdivost. Je známo, že jazykové modely mohou generovat i nesprávné nebo zcela smyšlené informace, kterým se pro jejich přesvědčivost začalo říkat *halucinace* (Bender et al., 2021). Analýzy v roce 2023 zjistily, že chatboti mohou *halucinovat* (tj. vymýšlet si) odpovědi zhruba ve 20–30 % případů, přičemž faktické nepřesnosti se v nějaké míře objevily až u 46 % jejich výstupů (Weise & Metz, 2023). To je velmi vysoké číslo – téměř polovina AI generovaných odpovědí může obsahovat nějakou chybu! Je proto kriticky důležité nepřijímat odpovědi AI slepě jako pravdivé, zejména v oblastech, kde na přesnosti informací záleží (medicína, právo, věda, finance atd.).

Tato strategie vybízí uživatele k tomu, aby **aktivně ověřovali a zpochybňovali informace** dodané AI, podobně jako by přistupovali ke zdroji z internetu nejasného původu. Existuje několik způsobů, jak to udělat:

- **Požádat AI o zdroje:** Některé AI modely mohou být schopné uvést prameny svých tvrzení, pokud jsou k tomu navrženy. Lze se explicitně zeptat: „*Můžeš doložit své tvrzení zdrojem nebo odkazem?*“. Je ovšem třeba mít na paměti, že základní verze ChatGPT a podobné modely nemají přístup k internetu ani skutečnou databázi zdrojů, takže pokud nejsou speciálně napojené na externí znalostní bázi, můžou si zdroj *vymyslet*. I odkazy a citace od AI je proto nutné ověřovat – model může halucinovat dokonce i neexistující studie či články, které vypadají důvěryhodně. Proto se tato výzva k ověření netýká jen ověření přes AI, ale hlavně ověření nezávisle na AI.
- **Křížová kontrola s jinými zdroji:** Ideální je **najít a porovnat informaci v důvěryhodných zdrojích** (např. učebnice, odborné články, oficiální statistiky). Pokud vám AI například poskytne historické datum nebo medicínské doporučení, zkuste se podívat, zda totéž uvádí i Wikipedie, odborný web nebo jiný spolehlivý zdroj. V praxi se doporučuje považovat odpověď AI spíše za *návrh či návrh směru*, který teprve ověřením proměníte ve skutečnou znalost. V akademickém prostředí je samozřejmě nepřípustné citovat AI jako zdroj faktů – vždy se vyžaduje ověření v literatuře.
- **Ptát se na tentýž fakt jinak:** Další metodou (méně spolehlivou, ale někdy zajímavou) je **zeptat se AI podruhé, případně jiné AI**. Pokud se odpovědi výrazně liší, je jasné, že jde o spekulativní informaci a je nutné ji prozkoumat dál. Shoda odpovědí sice není důkazem pravdivosti, ale *neshoda je dobrým varováním*.

Proč je ověřování faktů důležité? Protože **AI generuje odpovědi na základě pravděpodobnosti textu**, nikoli na základě porozumění skutečnému světu (Bender et al., 2021). Model neví, co je pravda – pouze usuzuje, co zní jako pravdivé tvrzení podle tréninkových dat. Může tedy s klidem podat zásadně nesprávnou informaci s naprosto seriózním tónem, což je kombinace potenciálně nebezpečná. Byly zaznamenány případy, kdy AI asistent sebevědomě doporučil neexistující léky nebo citoval neexistující soudní precedens (Weise & Metz, 2023). Pokud by uživatel těmto radám věřil, mohlo by dojít k vážné újmě. **Odpovědnost za finální použití informace proto leží na uživateli.** Ten by měl k informacím od AI přistupovat s *kritickým myšlením*. Ostatně i samotní tvůrci těchto systémů (např. OpenAI) upozorňují, že model může produkovat faktické chyby a že **důležité nebo citlivé informace je třeba vždy potvrdit z jiných zdrojů** (OpenAI, 2023).

Doporučení pro uživatele: Berte každé konkrétní tvrzení od AI jako *potenciální hypotézu*, ne jako zaručený fakt. U důležitých údajů (číselné výsledky, názvy, citace) si udělejte rychlé ověření – např. zadejte dané číslo nebo jméno do internetového vyhledávače a sledujte, zda se stejné informace objevují v důvěryhodných zdrojích. Pokud AI tvrdí něco podezřelého, klidně se zeptejte znovu: „*Opravdu to tak je? Některé zdroje uvádějí něco jiného...*“. Někdy se model při dalším dotazu sám opraví (pokud měl v konverzaci prostor pro reflexi). Také můžete přímo říct AI: „*Vyjmenuj zdroje, z kterých jsi čerpal.*“ – byť je riziko, že jak bylo zmíněno, halucinace se mohou týkat i referencí. Zkrátka, udržujte si skeptický postoj. Tato strategie je klíčová zejména v době, kdy je obsah generovaný AI všude možně na internetu a může *vypadat věrohodně*. Ověřování faktů chrání před šířením dezinformací a zajišťuje, že rozhodnutí, která na základě informací děláte, stojí na správných základech. Jak trefně uvedl jeden z kritiků: výstupy LLM jsou v jistém smyslu „*bullshit*“, protože modelu je vnitřně jedno, zda jsou pravdivé (Hicks et al., 2022 podle Bender et

al., 2021) – **důraz na pravdivost tedy musíte do procesu vnést vy jako uživatelé.**

4.7 Experimentujte s různými frázemi

Strategie „**Experimentujte s různými frázemi**“ vychází z pozorování, že **různá formulace téhož dotazu může od AI vést k odlišným odpovědím**. Jazykové modely jsou *poměrně citlivé na kontext a formulace promptu* – drobná změna slovosledu, použití synonyma či jiného tónu může někdy překvapivě změnit zaměření nebo kvalitu odpovědi (Zhao et al., 2021). Uživatel by tedy neměl váhat zkoušet alternativní způsoby, jak položit otázku, zejména pokud s první verzí odpovědi není spokojen.

Účel experimentování s frázemi je dvojí: zaprvé můžete dostat lepší odpověď, zadruhé se tím učíte, jak model reaguje (viz další strategie „Učte se z odpovědí“). Například pokud jste se zkusili zeptat „*Proč je obloha modrá?*“ a odpověď byla příliš jednoduchá nebo nesrozumitelná, zkuste otázku přeformulovat: „*Můžeš vědecky vysvětlit, proč vnímáme oblohu jako modrou? Odpověz přibližně na úrovni studenta střední školy.*“. Tím jste najednou specifikovali styl (vědecky, ale pro středoškoláka) a je velká šance, že druhý pokus přinese uspokojivější výklad. Podobně můžete experimentovat s *klíčovými slovy*: místo „*důsledky globálního oteplování*“ můžete zkusit „*dopady klimatické změny*“ – pro člověka ekvivalentní, ale model může při druhé formulaci čerpat z trochu jiného okruhu naučených informací a odpověď strukturovat jinak.

Výzkum v oblasti NLP ukazuje, že **LLM jsou skutečně prompt-senzitivní** – tytéž „myšlenky“ v jiném jazykovém balení mohou model aktivovat odlišně (Zhao et al., 2021; Zhu et al., 2023). Proto se dokonce i při testování schopností těchto modelů někdy používá tzv. *prompt paraphrasing*, aby se ověřilo, zda model konzistentně ví, o co jde, nebo jestli ho určitá formulace zmátlá. Pro běžného uživatele z toho plyne jednoduché ponaučení: **když první způsob dotazu nevedl**

k cíli, zkuste to jinak říct. Neznamená to, že byste museli umět stovky variant – mnohdy stačí jedna nebo dvě úpravy a rozdíl je markantní.

Dalším rozměrem této strategie je **zkoušení různých přístupů** k promptu: můžete dát modelu příklad (few-shot prompting), můžete zkusit položit otázku formou „*představ si, že...*“, nebo přidat drobný požadavek navíc. Například místo přímého dotazu „*Co znamená pojem entropie v termodynamice?*“ lze experimentovat s: „*Vysvětli mi pojem termodynamické entropie a použij přirovnání z běžného života.*“. Každá z těchto variant donutí model generovat odpověď trochu jiným způsobem – a vy si můžete vybrat tu, která vám nejlépe vyhovuje.

Doporučení pro uživatele: Buďte při práci s AI *tvořiví a hraví*. Nebojte se, že položíte otázku „špatně“ – i kdyby, nic se neděje, prostě ji zkuste přeformulovat. Pokud se vám odpověď nezdá dostatečná, zamyslete se, *které části dotazu by šly říct jinak*. Někdy pomůže zjednodušit větu, jindy naopak přidat detail. Můžete zkusit i synonyma a alternativní výrazy: model možná neporozuměl slangovému výrazu, ale formálnímu synonymu rozumět bude apod. Tímto experimentováním časem zjistíte, jaké formulace fungují nejlépe – například někdo může zjistit, že modelu vyhovuje, když otázku rozdělí do více vět místo jedné dlouhé, nebo že přidání konkrétního slova (třeba „*stručně*“) ovlivní délku odpovědi správným směrem. Různé promptovací „*triky*“ jsou dnes dokonce předmětem sdílení v online komunitách a objevuje se spousta *doporučených šablon*, jak něco formulovat. Ačkoli je dobré se inspirovat, nejlepší je **vlastní zkušenost** – zkoušejte a uvidíte. Udržujte si při tom mentální model: *mluvíte s inteligentním, ale cizím člověkem* – takže když to nepochopil napoprvé, zkusíte to říct jinak nebo doplnit, přesně jako v běžné konverzaci.

4.8 Učte se z odpovědí

Poslední strategie „**Učte se z odpovědí**“ zdůrazňuje, že komunikace s AI může být obousměrný učební proces – nejenže se AI „učí“ z vašich

promptů, ale vy se můžete **učit z reakcí AI**, a to hned několika způsoby. Cílem je průběžně zlepšovat své dotazovací dovednosti na základě toho, jak model odpovídá.

Zprvée, berte každou odpověď jako *zpětnou vazbu* k tomu, jak byl položen dotaz. Pokud odpověď **nesplňuje vaše očekávání**, zkuste analyzovat proč. Byla příliš obecná? Možná jste se ptali příliš obecně. Byla mimo téma? Možná v promptu chyběl kontext nebo model některé slovo pochopil jinak, než jste zamýšleli. Tato reflexe vám napoví, jak příště otázku upřesnit. Například když se zeptáte „*Popiš mi historii Francie.*“ a dostanete neuchopitelně dlouhý esej, uvědomíte si, že příště musíte být specifitější (např. „*stručně popiš hlavní milníky historie Francie od 18. století do současnosti*“). **Každá méně podařená odpověď je šancí zlepšit prompt.**

Zadruhé, učte se *obsahově*: Odpovědi AI často přinášejí nové informace či perspektivy, které jste neznali. Můžete si tedy rozšiřovat znalosti a následně pokládat informovanější otázky. Řekněme, že vám AI v odpovědi zmíní nějaký podpojem nebo související koncept – pokud vás zaujal, zeptejte se na něj v navazujícím dotazu („*Můžeš vysvětlit blíže, co znamená XYZ, který jsi zmínil?*“). Tím si prohlubujete pochopení dané problematiky. Učení se z odpovědi tak přímo souvisí se strategií doplňujících otázek: jedna odpověď může vyvolat další otázky, což je v procesu učení ideální scénář.

Zatřetí, sledujte *stylistické a formální aspekty* odpovědi. AI modely často produkují texty, které jsou koherentní, logicky strukturované a dobře formulované. Uživatelé tak mohou nevědomky pochytit některé komunikační dovednosti – například jak stručně shrnout složitou myšlenku, jak systematicky vyjmenovat pro a proti, jak uvádět příklady apod. Samozřejmě, model někdy generuje i slabší text (např. příliš rozvláčný nebo opakující se), takže není radno brát vše jako vzor. Nicméně pozorný uživatel může z kvalitních výstupů leccos odkoukat a použít to ve vlastní komunikaci.

Důležitou součástí učení se z odpovědí je také odhalování limitů AI. Když poznáte, v čem model často chybí (např. v matematických výpočtech, nebo v aktuálních datech přesahujících knowledge cutoff modelu apod.), poučíte se, že v těchto případech musíte být opatrní (viz ověřování faktů) nebo prompt formulovat tak, aby chyby minimalizoval (např. u výpočtů: „Pečlivě krok za krokem spočítej...“). Toto meta-uvědomění vám umožní AI využívat chytřeji – víte, na co se hodí a na co ne, co od ní čekat, kdy raději zvolit jiný postup. Jinak řečeno, postupně si budujete vlastní zkušenost s AI systémem, a ta je k nezaplacení, protože vám v budoucnu ušetří čas i zklamání.

Doporučení pro uživatele: Po každé významnější interakci s AI se na chvíli zamyslete: „*Dostal jsem, co jsem chtěl? Pokud ne, proč? Co udělám příště jinak?*“. Nebojte se experimentovat (jak jsme radili výše) a výsledky těchto experimentů si pamatovat. Klidně si můžete vést krátké poznámky o tom, co se osvědčilo – např. „*když chci příklady, lépe funguje říct: Uveď příklady: 1) ... 2) ...*“ apod. Časem si tyto poučky zautomatizujete. Zároveň, pokud se v odpovědi objeví něco zajímavého, co jste nečekali, zvažte, jestli to není podnět k dalšímu zkoumání. Přistupujte ke konverzaci s AI jako k procesu učení – stejně jako se člověk zlepšuje v dialogu s lektorem nebo v řešení úloh cvičením, zlepšujete se v promptování praxí. AI vám ihned „oznámkuje“ váš dotaz svou odpovědí – využijte to k růstu. Tento proces je kontinuální: jak se budou zdokonalovat modely, možná bude třeba se učit zas nové finty, ale s *dobrými základy komunikace* (jasnost, kontext, ověřování, iterace) budete vždy o krok napřed. V konečném důsledku vám tak osvojování si těchto strategií umožní efektivněji, sebevědoměji a bezpečněji spolupracovat s AI na jakémkoli úkolu.

5 Teoretický rámec MLEF-AI a koncept komunikační zralosti uživatele

5.1 Fragmentace teoretického aparátu výzkumu AI ve vzdělávání

Předchozí kapitoly této monografie představily teoretická a praktická východiska pro efektivní komunikaci uživatelů s generativní umělou inteligencí — od základní terminologie přes strategie tvorby promptů až po metodologii efektivní práce s AI nástroji. Tento aparát je dostatečný pro pojednání *technologické* a *interakční* roviny problému, není však sám o sobě dostatečný pro pojednání *pedagogické* roviny — tedy otázky, jak proces interakce s AI funguje v širším kontextu vzdělávacího systému, jak je formován organizační, etickou a regulační vrstvou, a co znamená, když uživatel postupem času *dozrává* ve své komunikační kompetenci s těmito nástroji.

Tato kapitola představuje teoretický most mezi předchozími kapitolami a empirickou studií, která následuje. Zavádí dva komplementární rámce — **MLEF-AI** (Multi-Level Ecological Framework for AI Implementation in Education) jako primární ekologicko-implementační rámec a **AICP** (AI-centrická pedagogika) jako pedagogickou koncepci edukační platformy, na níž empirická studie probíhala — a operacionalizuje pojem **kunikační zralost uživatele** jako klíčový konstrukt celé monografie.

Současný výzkum AI ve vzdělávání trpí, jak konstatuje Hanák (2026) v habilitační teze, *teoretickou fragmentací*: psychologicky orientované studie se zaměřují na individuální faktory přijetí (vnímanou užitečnost, sebedůvěru, technostres), organizačně zaměřené práce popisují implementační procesy a bariéry, kritické analýzy dekonstruuji ideologické předpoklady technologického akcelerationismu, a empirické studie zaměřené na interakci uživatele s AI tyto roviny obvykle nepropojují. Chybí *integrující rámec*, který by explicitně modeloval *dynamické interakce* mezi těmito úrovněmi a

vysvětlil, jak společenský tlak, etické imperativy, organizační governance a individuální psychologie vzájemně formují konečný výsledek — pedagogickou praxi a uživatelskou kompetenci. Tuto mezeru se snaží zaplnit rámec MLEF-AI, kterému je věnován následující oddíl.

5.2 Multi-Level Ecological Framework for AI Implementation in Education

5.2.1 Teoretické ukotvení a meta-teoretické principy

Rámec **MLEF-AI** byl autorem rozvinut v habilitační teze (Hanák, 2026) jako originální syntetický model implementace umělé inteligence ve vzdělávání. Vychází z Bronfenbrennerovy ekologické teorie lidského vývoje (Bronfenbrenner, 1979) a překonává fragmentaci dosavadního výzkumu propojením pěti vzájemně se ovlivňujících úrovní analýzy. Model stojí na třech meta-teoretických principech, které zajišťují jeho vnitřní koherenci.

Princip ekologické vnořenosti (*ecological nesting*) konceptualizuje implementaci AI jako proces, v němž je individuum (učitel, případně širší kategorie uživatele edukační intervence) vnořeno do soustředných kruhů vlivů: organizace → etické a regulační prostředí → společensko-politický diskurz. Klíčovou premisou tohoto principu je, že chování na nižší úrovni nelze plně vysvětlit bez pochopení podmínek na vyšších úrovních.

Princip obousměrné kauzality (*bidirectional causality*) na rozdíl od lineárních modelů typu TAM (Davis, 1989), kde Performance Expectancy → Behavioral Intention → Use Behavior, explicitně modeluje *feedback loops* mezi úrovněmi rámce. Zpětné vazby mohou působit *zdola nahoru* (například vysoká míra technostresu u učitelů vytváří tlak na vedení školy zavést governance opatření) i *shora dolů* (společenský diskurz o digitalizaci ovlivňuje sociální normy ve sboru a tlak na jednotlivé učitele).

Princip emergentních vlastností (*emergence*) předpokládá, že některé fenomény nemohou být redukovány na nižší úroveň. Pedagogická suverenita učitele není pouhým součtem individuálních kognitivních schopností, ale je *emergentní vlastností* interakce mezi autonomií učitele, kvalitou organizační governance a normativními rámci. Tento princip má důsledky i pro empirickou studii v této monografii: pojem *komunikační zralost* se ukáže být analogickou emergentní vlastností, kterou nelze redukovat ani na technickou znalost AI, ani na pedagogickou erudici uživatele samotnou.

5.2.2 Architektura: pět ekologických úrovní rámce

MLEF-AI je strukturován do pěti hierarchicky uspořádaných, avšak dynamicky propojených úrovní. Tabulka 0 představuje souhrnný přehled, jednotlivé úrovně jsou pak rozvinuty v navazujícím textu.

Tabulka 0. Pět úrovní rámce MLEF-AI a jejich klíčové konstrukty.

Úroveň	Název	Klíčové konstrukty	Teoretická opora
1	Společensko-politická (Macro)	Technologický akcelerationismus, kritická pedagogika	Selwyn (2019), Williamson (2017), Biesta (2017)
2	Eticko-regulační (Normative)	Pět pilířů etické AI ve vzdělávání	UNESCO (2021), NIST (2023), EU AI Act (2024)
3	Organizační (Meso – Škola)	Implementační ovladače, Governance Perception Index (GPI)	Fixsen et al. (2005), Rogers (2003)
4	Individuální (Micro – Učitel/Uživatel)	UTAUT konstrukty, technostres P-EFST, AI Self-Efficacy	Venkatesh et al. (2003), Tarafdar et al. (2007), Bandura (1986)
5	Pedagogická transformace (Outcome)	Re-profesionalizace vs. deprofesionalizace, Role Transformation Index (RTI)	Puentedura (2006), Biesta (2017)

Na **Úrovní 1 (Společensko-politická)** působí *technologický akcelerationismus* — imperativ neustálé inovace, kritizovaný jako ideologicky podmíněný (Selwyn, 2019). Vzdělávací instituce jsou vystaveny tlaku ze strany médií, politiků a komerčních platforem, aby co nejrychleji zavedly AI. Tento tlak je konfrontován s *kritickou pedagogikou* (Biesta, 2017; Williamson, 2017), která varuje před riziky *learnifikace* (redukce vzdělávání na měřitelné výstupy) a *datafikace* (algoritmická governance řízená daty, nikoliv pedagogickými principy). Klíčovou otázkou této úrovně je: kdo definuje agendu zavádění AI ve vzdělávání — pedagogové, nebo komerční platformy?

Úroveň 2 (Eticko-regulační) zahrnuje normativní rámce, které mají řídit praxi: UNESCO Doporučení o etice umělé inteligence (UNESCO, 2021) s deseti principy, NIST AI Risk Management Framework (NIST, 2023) zaměřený na řízení rizik a EU AI Act (2024) jako právní regulaci

high-risk systémů. Hanák (2026) tyto rámce v habilitační teze syntetizuje do *pěti pilířů etické AI ve vzdělávání*: lidská agentivita a dohled, spravedlnost a nediskriminace, soukromí a správa dat, transparentnost a vysvětlitelnost, odpovědnost a bezpečnost. Funkcí této úrovně je působit jako *etická brzda* na technologický akceleračníismus z Úrovně 1 a poskytovat legitimitu pro governance opatření na Úrovní 3.

Na **Úrovní 3 (Organizační)** dochází k vlastní implementaci — tedy k převodu deklarativních rámců do rutinní praxe vzdělávací instituce. Zde se uplatňují principy implementační vědy (Fixsen et al., 2005) s pěti klíčovými *ovladači implementace*: výběr (kdo bude inovaci zavádět), školení, koučink, hodnocení a facilitace adaptace. Bez aktivního řízení těchto ovladačů zůstává inovace pouze v rovině deklarací — což Fixsen et al. označují metaforou „*policy on the shelf*“ (politiky na policičce). Klíčovým originálním konstruktem této úrovně je **Governance Perception Index (GPI)**, který v habilitační teze měří, do jaké míry učitelé vnímají organizační pravidla jako jasná, dostupná a legitimní.

Úroveň 4 (Individuální) modeluje psychologické determinanty přijetí AI prostřednictvím rámce UTAUT (Venkatesh et al., 2003) s jeho čtyřmi konstrukty: *Performance Expectancy* (vnímaný přínos AI), *Effort Expectancy* (vnímaná snadnost použití), *Social Influence* (tlak od kolegů a vedení) a *Facilitating Conditions* (organizační podpora). Klíčový vhléd MLEF-AI spočívá v tom, že zatímco UTAUT zachází s *Facilitating Conditions* jako s exogenní proměnnou — něčím, co „prostě existuje nebo neexistuje“ — MLEF-AI je rozbaluje jako *výstup implementačních ovladačů z Úrovně 3*. Tím rámec propojuje organizační a individuální úroveň. Druhým klíčovým konstruktem této úrovně je **technostres**, měřený škálou P-EFST (Tarafdar et al., 2007) v pěti dimenzích, a **AI Self-Efficacy** (Bandura, 1986) jako kritický moderátor — uživatelé s vysokou důvěrou ve schopnost ovládat AI vykazují nižší technostres i při stejné úrovni expozice.

Úroveň 5 (Pedagogická transformace) modeluje výstup celého systému — změnu role učitele jako důsledek interakcí mezi všemi předchozími úrovněmi. Užívání AI vede podle MLEF-AI k jednomu ze dvou scénářů: *re-profesionalizace*, kdy se učitel posouvá směrem k roli facilitátora personalizovaného učení a získává čas osvobozením od rutinních úkolů (v souladu s úrovní *Redefinice Puente* dle SAMR modelu, 2006), nebo *deprofesionalizace*, kdy se učitel stává „správcem algoritmů“, ztrácí pedagogickou autonomii a jeho role je podkopána (v souladu s Biestovou kritikou ztráty *pedagogické suverenity*, 2017). Tato dichotomie je operacionalizována prostřednictvím originálního konstruktů **Role Transformation Index (RTI)**, který měří vnímání posunu role učitelem.

5.2.3 Dynamika: feedback loops mezi úrovněmi

MLEF-AI není statickým hierarchickým modelem, ale dynamickým systémem se čtyřmi klíčovými typy zpětných vazeb (*feedback loops*, FL1–FL4), které propojují jednotlivé úrovně:

- **FL1 — Zpětná vazba zdola nahoru** (Individuální → Organizace): Vysoká míra technostresu u učitelů (měřená P-EFST na Úrovní 4) vytváří tlak na vedení vzdělávací instituce zavést governance rámec (Úroveň 3). Potřeby učitelů tak formují organizační politiku.
- **FL2 — Top-down influence** (Společnost → Organizace → Jedinec): Technologický imperativ z Úrovně 1 (politici požadují digitalizaci) → tlak na ředitele zavést AI (Úroveň 3) → změna norem ve sboru a tlak na učitele (*Social Influence* v UTAUT, Úroveň 4). Tento mechanismus vysvětluje, proč některé inovace selhávají — jsou vnuceny shora bez ohledu na připravenost zdola.
- **FL3 — Etická brzda** (Etika → všechny úrovně): UNESCO principy (Úroveň 2) → transformovány do školních směrnic (Úroveň 3) → učitelé je internalizují jako normy (Úroveň 4) → vedou k reflexivní praxi (Úroveň 5). Tento mechanismus zajišťuje, že technologický imperativ je temperován etickými imperativy.

• **FL4 — Kritická reflexe** (Společnost → Pedagogika): Williamsonova (2017) kritika datafikace inspiruje učitele k reflexi vlastní autonomie. Tento mechanismus ukazuje, jak teoretická kritika formuje praxi.

Pro empirickou studii v této monografii je relevantní zejména **FL1**: data ze 147 testerů edukační platformy Kurzologie reprezentují právě tento typ zpětné vazby — uživatelská zkušenost generuje informace, které zpětně formují design a governance platformy. Empirická kapitola monografie je v tomto smyslu *manifestací FL1*.

5.2.4 Originální konstrukty rámce: GPI, RTI a perspektiva CMI

Habilitační teze (Hanák, 2026) přináší v rámci MLEF-AI dva originální měřicí konstrukty, které dále operacionalizují klíčové úrovně rámce.

Governance Perception Index (GPI) měří kvalitu organizační governance AI z perspektivy učitele; spojuje implementační vědu (Fixsen) s percepčním měřením a propojuje tak Úrovně 3 a 4 rámce.

Role Transformation Index (RTI) měří subjektivní vnímání transformace pedagogické role v důsledku zavedení AI; rozlišuje pól re-profesionalizace a pól deprofesionalizace a operacionalizuje Úroveň 5 jako kontinuum, nikoli binární opozici.

Tato monografie navrhuje rozšíření aparátu o třetí konstrukt, **Communication Maturity Index (CMI)**, který by měřil uživatelskou variantu transformace komunikační kompetence při interakci s AI. Zatímco RTI je primárně určen pro učitele a měří proměnu jejich pedagogické role, CMI by byl určen pro širší populaci uživatelů edukačních intervencí (kurzy, platformy, samostudium) a měřil by posun z transakčního, jednorázového dotazování k dialogické, iterativní spolupráci s AI. CMI je v této monografii zaveden jako teoretický koncept opřený o empirickou studii Kapitoly 6, jeho plná psychometrická validace je naplánována na Fázi 2 výzkumu.

5.3 AICP jako pedagogický rámec edukační platformy

Empirická studie této monografie probíhala na edukační platformě Kurzologie.cz, která je provozována podle principů **AI-centrické**

pedagogiky (AICP) rozvinutých autorem v samostatné monografii (Hanák, 2025). Pro účely této knihy postačí stručné shrnutí; čtenář, který chce s AICP rámcem pracovat hlouběji, je odkázán na primární zdroj.

AICP je pedagogická koncepce, v níž *umělá inteligence tvoří centrální infrastrukturu učení*: personalizuje postupy, zpřesňuje formativní zpětnou vazbu a automatizuje rutinní činnosti, *aniž by rušila lidskou odpovědnost* za cíle, obsah a etické řízení výuky. Učitel (HITL — *Human-in-the-Loop*) navrhuje didaktický design, moderuje učení, zajišťuje spravedlnost a bezpečnost práce s daty a rozhoduje o hodnocení. Student je aktivním agentem s narůstající autonomií; AI poskytuje adaptivní podporu a vysvětlitelnou zpětnou vazbu (Hanák, 2025).

AICP rámec je strukturován do *deseti dimenzí* (D1–D10), které pokrývají filozofickou orientaci cílů, roli učitele, autonomii žáka, preferované metody, organizaci výuky, pojetí kurikula, hodnocení a zpětnou vazbu, etiku a inkluzi, empirickou oporu a synergii s AI. Pro empirickou studii této monografie jsou věcně nejrelevantnější dimenze **D5** (organizace výuky a multimodalita — odpovídá pěti dimenzím nabízeným kurzem *Tvůrce kurzů*), **D7** (hodnocení a zpětná vazba — odpovídá dialogické vrstvě interakce s AI), **D10** (synergie s AI a princip HITL — odpovídá samotnému jádru proměnné AI_CHANGE) a **D3** (autonomie žáka — odpovídá zjištění o rozhodující roli vědomé volby uživatele nad VARK adherencí).

Vztah mezi MLEF-AI a AICP lze popsat jako vztah komplementární, nikoli konkurenční. **MLEF-AI** je *ekologicko-implementační rámec* — popisuje, jak procesy adopce a implementace AI fungují napříč úrovněmi systému, a je vhodný pro analýzu *dynamiky* zavádění AI. **AICP** je *pedagogická koncepce* — popisuje, jak má AI-centrická výuka vypadat *na úrovni třídy a designu vzdělávací intervence*, a je vhodný pro normativní vedení designu vzdělávacích produktů. Empirická studie v Kapitole 6 této monografie je situována v průniku obou

rámci: studuje *uživatelské trajektorie v rámci edukační intervence designované podle AICP principů* (platforma Kurzologie) a *interpretuje je primárně skrze MLEF-AI Úrovně 4 a 5* (individuální zkušenost a transformace komunikační role).

5.4 Komunikační zralost uživatele jako nový teoretický konstrukt

5.4.1 Definice komunikační zralosti

Pojem *komunikační zralost* zavádíme v této monografii jako kognitivně-komunikační kompetenci uživatele, projevující se v kvalitě jeho interakce s generativními AI nástroji v edukačním kontextu. **Operacionálně definujeme komunikační zralost jako míru, v níž uživatel postupně přechází od transakční, jednorázové komunikační orientace k dialogické, iterativní spolupracující orientaci, a v níž si osvojuje a samostatně aplikuje strategie efektivní komunikace s AI ve vlastní pedagogické či učební praxi.**

Tato definice obsahuje tři klíčové komponenty. **Transakční versus dialogická orientace** — komunikačně zralý uživatel chápe AI nikoli jako stroj na jednorázové odpovědi, ale jako partnera v iterativním procesu spoluprotvorby výsledku, je ochoten vést s AI delší konverzaci a postupně zpřesňovat zadání. **Strategické osvojení** — komunikačně zralý uživatel disponuje repertoárem komunikačních strategií (specifikace kontextu, definice role, strukturování výstupu, kritická validace odpovědí) a používá je vědomě, nikoli intuitivně. **Transfer do vlastní praxe** — komunikační zralost se projevuje schopností uživatele přenášet osvojené strategie z jednoho kontextu do druhého a používat je samostatně bez vnější instruktáže.

Tato definice odpovídá v MLEF-AI rámci **Úrovní 5** (pedagogická transformace) a je analogická konstrukt RTI s tím rozdílem, že RTI měří transformaci učitelské role, zatímco komunikační zralost měří transformaci uživatelské komunikační kompetence. Komunikační zralost je v této koncepci **emergentní vlastnost** — nelze ji redukovat

ani na technickou znalost AI (typu *prompt engineering*), ani na obecnou pedagogickou erudici, ani na jednotlivé komunikační strategie izolovaně. Vzniká až v jejich integraci a v dlouhodobé expozici uživatele edukační intervenci.

5.4.2 Vztah ke konstruktům MLEF-AI a UTAUT

Komunikační zralost se vztahuje k existujícím konstruktům MLEF-AI a UTAUT následovně. **AI Self-Efficacy** (Bandura, 1986) z Úrovně 4 je *prediktorem* komunikační zralosti — uživatel s vysokou důvěrou ve své schopnosti komunikovat s AI je schopen rychleji rozvinout zralé komunikační strategie. **Performance Expectancy** (UTAUT) je *moderátorem* — uživatel, který vnímá AI jako přínosnou pro svou práci, je motivován investovat do rozvoje komunikační zralosti. **Technostres** (P-EFST) je *negativním prediktorem* — uživatel s vysokým technostresem (zejména v dimenzi *Techno-Complexity*) je v rozvoji komunikační zralosti zpomalen. **Role Transformation Index** (RTI) z Úrovně 5 je *paralelním konstruktem* pro učitele — komunikační zralost je obecnější uživatelská varianta téhož principu transformace.

Tato síť vztahů otevírá konkrétní hypotézy testovatelné ve Fázi 2 výzkumu: (a) komunikační zralost koreluje pozitivně s AI Self-Efficacy a negativně s technostresem; (b) v populaci učitelů koreluje komunikační zralost vysoce s RTI (konvergentní validita); (c) v longitudinálním designu lze pozorovat narůstání komunikační zralosti se zvyšující se expozicí uživatele edukační intervenci.

5.4.3 Communication Maturity Index (CMI) jako měřicí nástroj

Pro operacionalizaci komunikační zralosti navrhujeme nástroj **Communication Maturity Index (CMI)**, který v plné psychometricky validované podobě bude předmětem Fáze 2 výzkumu. Pracovní struktura CMI vychází z empirických zjištění Kapitoly 6 a předpokládá tři subškály odpovídající třem komponentám definice (orientace, strategie, transfer):

- **Subškála A — Dialogická orientace** (cca 8 položek): míra, v níž uživatel preferuje iterativní vedení dialogu s AI nad jednorázovým dotazováním. Příklady položek: *„Když mi AI nedá uspokojivou odpověď, raději přeformuluju dotaz, než abych se obrátil jinam“*, *„Práci s AI vnímám jako rozhovor, nikoli jako vyhledávání“*.

- **Subškála B — Strategické osvojení** (cca 10 položek): repertoár vědomě používaných komunikačních strategií. Příklady položek: *„Před zadáním komplexnější úlohy AI definuju její roli a kontext“*, *„Specifikuji požadovaný formát odpovědi (seznam, tabulka, esej)“*, *„Kontroluji odpovědi AI z hlediska faktické správnosti“*.

- **Subškála C — Transfer do praxe** (cca 6 položek): schopnost přenosu osvojených strategií. Příklady položek: *„Strategie, které jsem se naučil/a v jednom kontextu, používám i v jiných situacích s AI“*, *„Moje práce s AI se v posledních měsících kvalitativně proměnila“*.

CMI nástroj je v této monografii prezentován jako *teoretický návrh*; jeho plná psychometrická validace (analýza vnitřní konzistence, faktorová struktura, konvergentní validita s RTI a AI Self-Efficacy, prediktivní validita vůči objektivním ukazatelům interakce) je předmětem prospektivní Fáze 2 výzkumu, která je naplánována na navazující období. Empirická studie této monografie poskytuje *předběžnou substantivní oporu* konstruktu prostřednictvím proměnné AI_CHANGE, která je v Kapitole 6 rozpracována jako orientační proxy komunikační zralosti.

5.5 Most k empirické studii: výzkumné otázky vyvozené z teoretického rámce

Z výše představeného teoretického aparátu vyplývá pět výzkumných otázek, které jsou v empirické studii následující kapitoly explorovány na vzorku 147 uživatelů kurzu *Tvůrce kurzů* na platformě Kurzologie.cz:

První výzkumná otázka se odvolává na **MLEF-AI Úroveň 4 a 5**: *Existuje vztah mezi dokončením kurzu a deklarovanou změnou ve způsobu interakce uživatele s generativní AI?* Tato otázka testuje, zda dlouhodobá expozice edukační intervenci vede k posunu v kategorii odpovídající uživatelské variantě re-profesionalizace.

Druhá výzkumná otázka se odvolává na **AICP Dimenze D5 (organizace výuky) a D10 (synergie s AI)**: *Liší se efekt kurzu na změnu interakce s AI podle dominantně využívané dimenze výuky?* Tato otázka testuje, zda navržený AICP design platformy je robustní napříč modálními formáty.

Třetí výzkumná otázka se odvolává na **AICP Dimenzi D3 (autonomie žáka) a MLEF-AI Úroveň 4**: *Hraje roli VARK preferenční pretest a jeho dodržení při výběru dimenze pro výsledný efekt kurzu?* Tato otázka testuje, zda algoritmické doporučení formátu (VARK) má prediktivní hodnotu, nebo zda rozhodující je vědomá uživatelská volba.

Čtvrtá výzkumná otázka se odvolává na **MLEF-AI Úroveň 5**: *Lze v souboru identifikovat přirozeně se vyskytující typologii uživatelů s odlišnými trajektoriemi učení a odlišnými výsledky?* Tato otázka má fenomenologický charakter a generuje empirický základ pro budoucí psychometrickou validaci CMI nástroje.

Pátá výzkumná otázka má reflexivní povahu a integruje **MLEF-AI feedback loop FL1**: *Jaké designové implikace pro pedagogickou praxi v duchu AICP plynou z empirických zjištění o uživatelských trajektoriích?* Tato otázka uzavírá teoreticko-empirický cyklus a vede k formulaci doporučení pro Fázi 2 výzkumu.

Empirická kapitola, která následuje, postupuje od metodologického rámování přes výsledky k diskusi a v každém z těchto oddílů se vrací k pojmovému aparátu představenému v této kapitole. Tím vzniká kompaktní teoreticko-empirický celek, jehož cílem je přispět k integraci dosud fragmentovaného výzkumu komunikace s generativní AI ve vzdělávacím kontextu.

6 Empirická studie interakčních vzorců v komunikaci uživatelů s generativní AI

Tato kapitola představuje empirické jádro monografie. Oproti původnímu autorskému záměru, který předpokládal pilotní studii dvaceti účastníků v laboratorním uspořádání, je zde prezentována podstatně rozsáhlejší **retrospektivní observační studie 147 reálných uživatelů kurzu Tvůrce kurzů** na platformě Kurzologie.cz, kteří kurzem prošli v období březen 2025 – leden 2026, doplněná o sedm bohatě popsanych pilotních případů. Tento posun reflektuje skutečnost, že v období přípravy monografie byla k dispozici reálná, byť v nestrukturované podobě uchovávaná data od podstatně širšího vzorku reálných uživatelů reálné edukační intervence; jejich neanalyzování by představovalo informační ztrátu. Současně je třeba předem poctivě pojmenovat metodologická omezení vyplývající z rekonstruktivní povahy datasetu, kterým je věnován samostatný oddíl 6.1.5.

6.1 Status studie a metodologické rámování

6.1.1 Povaha studie a její epistemologické rámování

Empirická část této monografie vychází z **observační retrospektivní studie 147 uživatelů kurzu Tvůrce kurzů**, kteří prošli platformou Kurzologie.cz v období březen 2025 – leden 2026. Vzhledem k tomu, že kurz byl primárně provozován jako edukační produkt, nikoli jako prospektivní výzkumný projekt s předem schváleným protokolem, byla data o uživateli strukturována *ex post facto* z několika zdrojů: záznamů Tutor LMS o průchodu kurzem, průběžných pracovních poznámek autora-provozovatele platformy ke každému z testerů a strukturovaných rozhovorů, které autor s testery v různých fázích kurzu vedl. Tato studie je proto vědomě situována do tradice *insider practitioner research* a *retrospective observational design*, žánrů, které mají v edukačním výzkumu, akčním výzkumu a evaluaci

edukačních intervencí svou ustálenou tradici (Stake, 2006; Yin, 2018; McNiff, 2017).

Charakteristickým rysem této tradice je dvojrole výzkumníka, který je současně tvůrcem a provozovatelem zkoumané intervence. Tato dvojrole je v textu transparentně reflektována (viz oddíl 6.1.6 Etické rámce a střet zájmů) a je pojmána nikoli jako metodologická vada, ale jako specifická epistemická pozice s vlastními výhodami (hluboká znalost terénu, dlouhodobý kontakt s účastníky, ekologická validita) i omezeními (riziko *confirmation bias*, absence slepoty, jediný kodér).

Studie je výslovně koncipována jako **explorační**, nikoli konfirmační. Jejím cílem není testovat předem stanovené hypotézy s nárokem na zobecnění, ale identifikovat orientační vzorce, generovat hypotézy pro navazující prospektivní studii a popsat typologii uživatelů kurzu, která může sloužit jako východisko pro následný design.

6.1.2 Výzkumné otázky

S ohledem na povahu dostupných dat a explorační charakter studie byly formulovány následující výzkumné otázky:

VO1: Existuje vztah mezi dokončením kurzu a deklarovanou změnou ve způsobu, jakým uživatelé interagují s generativní AI?

VO2: Liší se efekt kurzu na změnu interakce s AI podle dominantně využívané dimenze (D1–D5)? Konkrétně: obstojí experimentální dimenze D5 (podcast + infografika) ve srovnání s tradičními formáty?

VO3: Vede multimodální využívání kurzu (více než jedna dimenze) k většímu posunu v interakci s AI než využití jediné dimenze?

VO4: Hraje roli VARK preferenční pretest a jeho dodržení při výběru dimenze pro výsledný efekt kurzu?

VO5: Lze v souboru identifikovat přirozeně se vyskytující typologii uživatelů s odlišnými trajektoriemi učení a odlišnými výsledky?

Tyto otázky jsou formulovány jako orientační. S ohledem na rekonstruktivní povahu dat se autor explicitně zdržuje formulace formálních hypotéz s předem stanovenou hladinou významnosti, neboť taková formulace by implikovala konfirmační charakter studie, který data v této podobě neunesou. Prezentované statistické testy mají tedy charakter *exploratorní deskripce vztahů v dostupných datech*, nikoli konfirmačního testování.

6.1.3 Vzorek a kontext

Studie zahrnuje 147 uživatelů kurzu *Tvůrce kurzů* na platformě Kurzologie.cz, kteří v období březen 2025 až leden 2026 kurz alespoň zahájili. První účastníci vstoupili do kurzu v březnu 2025 v rámci pozvolného náběhu, hlavní sběr probíhal od dubna 2025, a poslední doplnění informací byla získána na konci ledna 2026. Účastníci se do kurzu hlásili dobrovolně, motivováni primárně zájmem o tvorbu vlastních online kurzů s využitím AI. Vzorek tedy nepředstavuje reprezentativní výřez populace učitelů či tvůrců obsahu, ale **populaci samoselektovaných tvůrců kurzů s nadprůměrným zájmem o AI v edukaci**. Tato charakteristika je důležitá pro interpretaci výsledků a omezuje zobecnitelnost závěrů na podobně motivované skupiny.

Kurz je strukturován do **pěti dimenzí** odpovídajících různým modálním stylům přijímání informací:

- **D1 — Video + pracovní listy** (převážně kinestetický typ, aktivní zpracování během sledování)
- **D2 — Audiokniha + digitální kniha + pracovní listy** (čtecí/psací typ, opakovaná práce s textem)
- **D3 — Audio přednáška + digitální kniha** (auditivní typ s textovou oporou)
- **D4 — Fyzická kniha + fyzická cvičebnice** (vizuální typ, klasický tištěný formát)

- **D5 — Podcast + infografika** (experimentální dimenze reflektující soudobé způsoby získávání informací)

Účastníkům byl při vstupu do kurzu administrován **VARK preferenční pretest**, na jehož základě jim byla doporučena vhodná dimenze; účastníci však mohli kurz studovat libovolným způsobem napříč dimenzemi, mohli volně přecházet mezi formáty a kombinovat je. Tento volný design umožnil vznik přirozeně variabilních trajektorií učení, které jsou v této studii analyzovány.

6.1.4 Proměnné, jejich operacionalizace a transparentní pojmenování zdrojů

Vzhledem k rekonstruktivní povaze dat je v této studii kladen mimořádný důraz na **transparentní pojmenování zdroje každé proměnné**. Následující tabulka přehledně shrnuje, odkud konkrétně každá proměnná pochází a jakou má epistemickou váhu.

Tabulka 1. Přehled proměnných, jejich operacionalizace a zdroje dat.

Proměnná	Hodnoty	Zdroj	Epistemický status
ID	1–147	Tutor LMS — unikátní identifikátor uživatele	Reálný záznam
VARK	V, A, R, K, M	Vstupní pretest při zahájení kurzu	Reálně proběhlo, výsledky nesystematicky archivovány
ADHERENCE	none, partial, full	Strukturované rozhovory s testery	Retrospektivní kódování autorem
DIM_MAIN	video, audio, digital, physical, podcast, AI	Pracovní poznámky autora	Retrospektivní kódování
DIM_COUNT	1, 2, 3	Pracovní poznámky autora	Retrospektivní kódování
COMPLETED	yes, no	Tutor LMS + pracovní poznámky	Reálný záznam
DROPOUT_PHASE	early, mid	Tutor LMS + pracovní poznámky	Reálný záznam
AI_CHANGE	none, partial, high	Strukturované rozhovory + výstupy	Retrospektivní kódování autorem
FOCUS	AI, course_creation	Pracovní poznámky	Retrospektivní kódování

Z tabulky je patrné, že páteř datasetu (ID, COMPLETED, DROPOUT_PHASE) je auditovatelná v Tutor LMS, zatímco kódovací proměnné (ADHERENCE, AI_CHANGE) byly autorem stanoveny *ex post facto* na základě obsahu rozhovorů a poznámek. Tento status musí být brán v úvahu při interpretaci všech následujících analýz.

Klíčová kódovaná proměnná **AI_CHANGE** vyjadřuje deklarovanou míru posunu ve způsobu, jakým tester po absolvování kurzu interaguje s generativní AI, oproti svému předchozímu způsobu používání. Hodnoty byly stanoveny v rozhovorech podle obsahu sebehodnocení a kvalitativní reflexe vlastní praxe testerem, doplněné

o pozorování autora ohledně tvůrčích výstupů. Kategorie *high* odpovídá deklarovanému přechodu od jednorázových transakčních dotazů k rozhovorovému, iterativnímu spolupracování s AI při tvorbě obsahu — což je posun, který kurz svou pedagogickou intencí cíleně vyvolává.

6.1.5 Reflexe rizika chybného přiřazení a sensitivity analýza

Vzhledem k retrospektivnímu sypání hodnot z různorodých zdrojů (LMS, pracovní poznámky, rozhovory) **nelze vyloučit dílčí chyby v přiřazení hodnot k jednotlivým ID**. Autor odhaduje míru nejistoty kódování přibližně na 18 % řádků, kde subjektivní jistota přiřazení je střední nebo nízká (např. testeři, s nimiž proběhl pouze jeden krátký rozhovor a u nichž poznámky o průběhu kurzu jsou strohé). U zbylých 82 % řádků je subjektivní jistota kódování vysoká, což odpovídá testerům, s nimiž autor vedl opakované rozhovory a u nichž má rozsáhlé pracovní poznámky.

Pro otestování robustnosti zjištění vůči této nejistotě byla provedena **Monte Carlo sensitivity analýza**: v každé z 1 000 iterací bylo náhodně vyloučeno 26 případů (18 % datasetu) a klíčové statistiky byly přepočítány na zbylých 121 případech. Výsledky této analýzy jsou prezentovány v oddílu 6.2.6 a slouží k posouzení, do jaké míry hlavní zjištění závisí na konkrétní konfiguraci kódovaných hodnot.

6.1.6 Etické rámce a reflexe střetu zájmů

Studie pracuje s daty získanými v rámci provozu komerční edukační platformy, kde autor monografie vystupuje současně jako (a) provozovatel platformy, (b) tvůrce kurzu a (c) výzkumník. Tento mnohavrstvý střet zájmů je v textu transparentně přiznán a metodologicky reflektován. Účastníci kurzu byli informováni o tom, že jejich průchod kurzem může být využit pro evaluaci a další rozvoj kurzu; informovaný souhlas s konkrétním vědeckým výzkumem v formálním smyslu však proběhl pouze u sedmi pilotních testerů (viz oddíl 6.3), nikoli u celé skupiny 147 uživatelů.

Z tohoto důvodu jsou v této studii prezentovány výhradně **agregovaná, anonymizovaná data**, která neumožňují identifikaci jednotlivých účastníků. Žádné individuální údaje, citace z rozhovorů ani osobní informace nejsou v textu uvedeny. Pseudonymy v oddílu 6.3 jsou plně fiktivní a popisy případů jsou redukovány na charakteristiky průchodu kurzem bez osobních detailů.

Kódování všech proměnných provedl jediný hodnotitel (autor monografie). Pro řádnou prospektivní studii (Fáze 2), která následuje po této průběžné analýze, je naplánováno: nezávislé etické vyjádření k designu, předregistrace hypotéz a kódovacího klíče na platformě Open Science Framework, slepé kódování druhým nezávislým kódérem a výpočet inter-raterové reliability (Cohen κ).

6.1.7 Statistické metody

Pro analýzu byly použity následující metody odpovídající kategoriální povaze většiny proměnných a explorativní povaze studie:

- Deskriptivní statistiky (frekvenční tabulky, podíly, křížové tabulky)
- Pearsonův chí-kvadrát test nezávislosti pro testování asociací mezi kategoriálními proměnnými, doplněný o Cramérovo V jako míru síly efektu
- Ordinální logistická regrese pro modelování AI_CHANGE jako ordinálního výstupu (none < partial < high)
- Binární logistická regrese pro modelování dokončenosti kurzu
- K-means klastrová analýza pro identifikaci typologie uživatelů
- Monte Carlo sensitivity analýza (1 000 iterací s náhodným vyloučením 18 % případů) pro robustnost klíčových zjištění

Analýzy byly provedeny v jazyce Python s využitím knihoven *pandas*, *scipy*, *statsmodels* a *scikit-learn*. Hladina významnosti byla pro orientační účely stanovena na $\alpha = 0,05$, avšak s explicitním

pojmenováním, že vzhledem k explorativní povaze studie nelze p-hodnoty interpretovat ve striktně konfirmačním smyslu.

6.2 Výsledky

6.2.1 Deskriptivní charakteristika souboru

Soubor 147 uživatelů kurzu byl charakterizován distribucemi prezentovanými v Tabulce 2. Distribuce VARK preferencí ukazuje relativně vyrovnané zastoupení všech pěti typů s mírnou převahou kinestetického, vizuálního a multimodálního typu. Distribuce hlavních dimenzí kurzu byla rovněž poměrně vyrovnaná, což svědčí o tom, že kurz oslovil uživatele napříč preferenčními styly.

Tabulka 2. Distribuce klíčových proměnných v souboru (N = 147).

Proměnná	Hodnota	n	%
VARK preference	Kinestetický (K)	32	21,8 %
	Multimodální (M)	32	21,8 %
	Vizuální (V)	32	21,8 %
	Auditivní (A)	28	19,0 %
	Čtecí/psací (R)	23	15,6 %
Dodržení VARK	Částečné	67	45,6 %
	Žádné	44	29,9 %
	Plné	36	24,5 %
Hlavní dimenze	AI (interakce)	31	21,1 %
	Fyzická (D4)	28	19,0 %
	Podcast (D5)	25	17,0 %
	Video (D1)	22	15,0 %
	Digitální (D2)	21	14,3 %
	Audio (D3)	20	13,6 %
Počet využití dimenzí	2	69	46,9 %

	1	43	29,3 %
	3	35	23,8 %
Dokončenost	Ano	128	87,1 %
	Ne	19	12,9 %
Změna interakce s AI	Vysoká (high)	70	47,6 %
	Částečná (partial)	49	33,3 %
	Žádná (none)	28	19,0 %

Celková dokončenost kurzu (87,1 %) je výrazně nadprůměrná ve srovnání s typickými hodnotami u otevřených online kurzů (kde se dokončenost pohybuje okolo 5–15 %). Tato vysoká hodnota odráží charakter vzorku — sebe-selektovaní motivovaní tvůrci kurzů, kteří měli silný osobní zájem na dokončení (vlastní kurz jako výstup). Mezi 19 nedokončiteli proběhl dropout převážně v rané fázi (15 případů), pouze 4 testéři přerušili kurz uprostřed.

V proměnné AI_CHANGE odpovídá kategorie *high* nejvýraznějšímu posunu k dialogickému, iterativnímu způsobu interakce s AI. Téměř polovina uživatelů (47,6 %) tuto kategorii uvedla, přičemž její distribuce ovšem dramaticky závisí na dokončenosti kurzu, jak ukáže následující oddíl.

6.2.2 Hlavní zjištění: dokončenost kurzu a změna interakce s AI

Klíčovým a nejsilnějším zjištěním této studie je výrazný vztah mezi dokončením kurzu a deklarovanou změnou ve způsobu interakce s generativní AI. Mezi 128 uživateli, kteří kurz dokončili, uvedlo „vysokou“ změnu 70 osob (54,7 %), zatímco mezi 19 nedokončiteli **nikdo** nedosáhl této kategorie. Uvedený rozdíl je statisticky vysoce významný a korespondující velikost efektu odpovídá střednímu efektu (Tabulka 3).

Tabulka 3. Křížová tabulka dokončenosti kurzu a změny v interakci s AI (N = 147).

	Žádná	Částečná	Vysoká	Celkem
Nedokončili	8 (42,1 %)	11 (57,9 %)	0 (0,0 %)	19
Dokončili	20 (15,6 %)	38 (29,7 %)	70 (54,7 %)	128
Celkem	28	49	70	147

Pearsonův chí-kvadrát test potvrdil silnou závislost mezi dokončeností a změnou interakce: $\chi^2(2) = 20,43$, $p < 0,0001$. Cramérovo V dosáhlo hodnoty 0,373, což odpovídá **střednímu až silnému efektu** podle Cohenovy klasifikace. Tento výsledek představuje primární empirický nález celé studie.

Interpretace tohoto vztahu má několik vrstev. *Naivní kauzální čtení* by tvrdilo, že dokončení kurzu přímo způsobuje posun v interakci s AI. Toto čtení však neodpovídá observační povaze studie, která neumožňuje vyloučit alternativní vysvětlení: (a) **selekční efekt** — uživatelé, kteří jsou již předem schopni hlubšího angažmá s obsahem, jsou pravděpodobně ti, kdo kurz dokončí, a tito uživatelé jsou současně ti, kdo jsou schopni transferu naučených interakčních vzorců do své praxe; (b) **expoziční dávkování** — dokončení kurzu znamená delší expozici obsahu, jež by sama o sobě mohla vést k internalizaci interakčních vzorců, bez ohledu na specifický obsah kurzu; (c) **kombinovaný efekt motivace, expozice a internalizace**.

Pro interpretaci v rámci této studie je klíčový závěr: **dokončenost kurzu je v dostupných datech nejsilnějším prediktorem deklarované změny interakce s AI**. Tento závěr má bezprostřední praktické implikace pro design kurzů typu *Tvůrce kurzů*: investice do mechanismů podporujících dokončení (motivace, gamifikace, podpora vytrvalosti) může být efektivnější než další optimalizace formátu obsahu samotného. Kauzální interpretace bude předmětem testování v plánované Fázi 2.

6.2.3 Dimenze kurzu a jejich (ne)diferenciální efekt

Druhou klíčovou výzkumnou otázkou byl vztah mezi dominantně využívanou dimenzí kurzu (D1–D5, případně dimenze AI jako interakce s nástrojem) a změnou interakce s AI. Výsledky ukázaly **absenci statisticky významného rozdílu mezi dimenzemi** (Tabulka 4): chí-kvadrát test pro nezávislost dosáhl hodnoty $\chi^2(10) = 5,15$, $p = 0,881$, Cramérovo $V = 0,132$. Procentní podíl uživatelů s deklarovanou „vysokou“ změnou interakce s AI se mezi dimenzemi pohyboval v rozmezí 38,1 % (digitální dimenze) až 65,0 % (audio dimenze), avšak rozdíly nedosáhly statistické významnosti a vzhledem k velikostem podsouborů ($n = 20–31$) je třeba je interpretovat velmi opatrně.

Tabulka 4. Podíl uživatelů s vysokou změnou interakce s AI podle hlavní dimenze kurzu.

Hlavní dimenze	n	Vysoká změna AI (%)	Dokončenost (%)
Audio (D3)	20	65,0 %	85,0 %
Podcast (D5) — experimentální	25	52,0 %	100,0 %
Video (D1)	22	45,5 %	77,3 %
AI (interakce)	31	45,2 %	80,6 %
Fyzická (D4)	28	42,9 %	85,7 %
Digitální (D2)	21	38,1 %	95,2 %

Tento výsledek má dva věcně významné aspekty. **Za prvé**, absence diferenciálního efektu lze interpretovat jako empirický argument pro **robustnost vícedimenzionálního designu kurzu**: kurz funguje napříč všemi nabízenými formáty s podobnou účinností. To je nontriviální nálezn, neboť běžnou obavou při designu vícedimenzionálních kurzů je, že některé dimenze budou „slabší“ a poskytnou výrazně horší výsledky. Data taková slabší místa neidentifikují.

Za druhé, a to je z hlediska experimentálního charakteru kurzu obzvláště významné: **experimentální dimenze D5 (podcast +**

infografika) dosahuje srovnatelných (a v některých ohledech lepších) výsledků jako tradiční formáty. Z 25 uživatelů, kteří jako hlavní dimenzi využívali D5, kurz dokončilo **100 %** a 52,0 % uvedlo vysokou změnu interakce s AI. Tato dimenze, která reprezentuje moderní audio-vizuální způsob získávání informací, tedy nejen obstála ve srovnání s klasickými formáty, ale ukázala se jako jedna z nejuspěšnějších z hlediska dokončenosti. Tento výsledek je pravděpodobně nejhodnotnějším nálezem studie z hlediska didaktické inovace, byť je třeba znovu připomenout limity rekonstruktivních dat.

Obdobně absence efektu byla zjištěna pro **počet využitých dimenzí**: hypotéza, že multimodální využívání ($DIM_COUNT > 1$) povede k většímu posunu v interakci s AI, se v datech nepotvrdila ($\chi^2(4) = 3,36$, $p = 0,500$). Podíl uživatelů s vysokou změnou byl 46,5 % u monomodálních (1 dimenze), 53,6 % u bimodálních (2 dimenze) a 37,1 % u trimodálních (3 dimenze) uživatelů. Pokles podílu u trimodálních uživatelů je zajímavý, ale nedosahuje statistické významnosti a může být artefaktem malé velikosti podsouboru. Hypotetické vysvětlení (rozptyl pozornosti při příliš mnoha současně využívaných formátech) je nicméně teoreticky obhájitelné a stojí za testování v navazující studii.

6.2.4 VARK preference a její (slabá) prediktivní hodnota

Třetí výzkumnou otázkou byl vztah mezi VARK preferenčním pretestem a (a) volbou hlavní dimenze, (b) výsledným efektem kurzu. Analýza ukázala následující obraz:

Vztah VARK → volba dimenze dosáhl hraniční významnosti ($\chi^2(20) = 30,53$, $p = 0,062$, Cramérovo $V = 0,228$). Existuje tedy **slabý trend**, kdy uživatelé tíhnou k dimenzi odpovídající jejich preferenčnímu typu, ovšem tento vztah je natolik slabý, že nelze hovořit o silné prediktivní hodnotě. Uživatelé čtečního/psacího typu (R) skutečně volili nejčastěji audio dimenzi, což však neodpovídá teoretickému předpokladu (čekal by se digitální/fyzický formát). Vizualní typ (V) nejčastěji preferoval

podcast, což opět neodpovídá teoretickému očekávání. Tato neshoda mezi VARK preferencemi a teoreticky odpovídajícími dimenzemi naznačuje, že VARK pretest má v praxi spíše orientační charakter a neodráží reálné modální preference uživatelů konzistentně.

Vztah ADHERENCE k VARK doporučení → výsledek byl rovněž statisticky nevýznamný, a to jak pro dokončenost ($\chi^2(2) = 2,14$, $p = 0,342$), tak pro změnu interakce s AI ($\chi^2(4) = 5,79$, $p = 0,215$). Navíc byl pozorován **kontraintuitivní vzorec**: nejvyšší podíl uživatelů s vysokou změnou interakce s AI vykázali ti, kteří VARK doporučení **nedodrželi vůbec** (54,5 %) a ti, kteří jej dodrželi *plně* (58,3 %), zatímco uživatelé s *částečným* dodržáním vykázali jen 37,3 %. Tento vzorec sice není statisticky významný, ale otevírá zajímavou hypotézu pro Fázi 2: rozhodující není dodržení nějakého doporučení, ale **konzistentnost a vědomí vlastní volby**. Uživatel, který buď VARK plně přijme za své, nebo si naopak řekne „já vím, co potřebuji“ a vybere si vlastní cestu, dosahuje lepších výsledků než ten, kdo si půjčí část doporučení a část nikoli.

Souhrnně: **VARK pretest má v této populaci slabou prediktivní hodnotu pro výslednou změnu interakce s AI** a jeho dodržení také neovlivňuje výsledek statisticky významným způsobem. Tento náález má důsledky pro design budoucích kurzů: VARK pretest si zaslouží být zachován jako úvodní reflexivní nástroj, který uživatele orientuje, ale neměl by být nadřazen vlastní volbě uživatele a nepředstavuje silný prediktor úspěchu.

6.2.5 Klastrová analýza: typologie uživatelů kurzu

Pro identifikaci přirozeně se vyskytujících typů uživatelů byla provedena k-means klastrová analýza nad standardizovanými proměnnými (DIM_COUNT, COMPLETED, ADHERENCE, AI_CHANGE) s dummy kódováním DIM_MAIN. Bylo testováno řešení se třemi a čtyřmi klastry; čtyřklastrové řešení poskytlo věcně bohatší a interpretovatelně smysluplnější strukturu, a proto je prezentováno níže (Tabulka 5).

Tabulka 5. Čtyřklastrová typologie uživatelů kurzu (N = 147).

Klast r	n	Dokonč .	Vys. zm. AI	Hl. dim.	VAR K	Pracovní název typu
K1	46	97,8 %	45, 7 %	Podcas t	K	Vytrvalý multimodální tvůrce
K2	42	81,0 %	54, 8 %	Video	M	Video- orientovaný experimentáto r
K3	31	80,6 %	45, 2 %	AI	A	AI-fokusovaný praktik
K4	28	85,7 %	42, 9 %	Fyzická	K	Klasický knižní studující

Klastr **K1 (n = 46)** je největším a neúspěšnějším z hlediska dokončenosti (97,8 %). Charakterizuje jej využívání podcastové (D5) dimenze, kinestetická VARK preference a multimodální přístup ke kurzu. Lze jej pracovním označit jako *vytrvalý multimodální tvůrce*: uživatel, který kurz absolvuje plně, využívá moderní formáty a kombinuje je. Tento klastr empiricky podporuje hodnotu experimentální dimenze D5.

Klastr **K2 (n = 42)** představuje *video-orientované experimentátory* s nejvyšším podílem deklarované vysoké změny interakce s AI (54,8 %), avšak nižší dokončeností (81,0 %). Multimodální VARK profil naznačuje uživatele otevřeného více formátům, který intenzivně experimentuje, ale občas projekt nedotáhne do konce. Vysoký AI_CHANGE u tohoto klastru může souviset s aktivním přístupem k testování nových způsobů práce s AI.

Klastr **K3 (n = 31)** lze charakterizovat jako *AI-fokusované praktiky* — uživatele, kteří kurz absolvují primárně kvůli interakci s AI samotnou,

méně se soustředí na ostatní dimenze. Auditivní VARK preference a nižší dokončenost (80,6 %) naznačují uživatele, který si vybírá z kurzu cíleně to, co potřebuje pro vlastní praktické použití.

Klastr **K4 (n = 28)** představuje *klasické knižní studující*, kteří preferují fyzickou dimenzi (D4) a kinestetický VARK profil. Dokončenost je středně vysoká (85,7 %), změna interakce s AI je relativně nejnižší (42,9 %), což může odrážet jejich konzervativnější přístup k AI nástrojům.

Tato předběžná typologie poskytuje **operacionalizovatelný rámec pro Fázi 2**, kde bude testováno, zda lze tyto klastry replikovat na nezávislém vzorku a zda mají prediktivní hodnotu pro pedagogická doporučení. Současně poskytuje rámec pro interpretaci sedmi pilotních případů (oddíl 6.3), které lze do těchto klastrů s různou mírou přesnosti přiřadit.

6.2.6 Sensitivity analýza: robustnost zjištění vůči nejistotě kódování

Vzhledem k odhadovaným 18 % řádků s nižší jistotou kódování byla provedena Monte Carlo sensitivity analýza pro otestování, zda hlavní zjištění ob stojí i při náhodném vyloučení této frakce dat. V každé z 1 000 iterací bylo náhodně vybráno 26 případů, které byly z analýzy vyloučeny, a klíčové statistiky byly přepočítány na zbylých 121 případech.

Tabulka 6. Výsledky sensitivity analýzy (1 000 iterací s vyloučením 18 % případů).

Statistika	Průměr	SD	95% interval
Dokončenost (%)	87,0 %	1,3 %	84,3 % – 89,3 %
Vysoká změna AI (%)	47,6 %	2,0 %	43,8 % – 51,2 %

p-hodnota χ^2 COMPLETED $\times$ AI_CHANGE	< 0,001	—	< 0,001
% iterací s p < 0,05	100,0 %	—	—

Sensitivity analýza ukázala, že **klíčový vztah mezi dokončeností kurzu a změnou interakce s AI zůstal statisticky významný ve 100 % iterací**. Tento výsledek představuje silnou evidenci pro **robustnost primárního zjištění studie**: i kdyby se 18 % zakódovaných hodnot ukázalo jako chybně přiřazené (a nahrazení jakýmkoli jiným plausibilním vzorem), hlavní závěr by se nezměnil.

Hodnoty deskriptivních statistik (dokončenost, podíl vysoké změny AI) zůstaly v úzkých 95 % intervalech, což svědčí o tom, že agregátní obraz souboru je stabilní vůči konkrétní konfiguraci nejistých kódů. Tato robustnost neznamená, že případné chyby kódování jsou bezvýznamné — znamená pouze, že nejsou systematické a velké natolik, aby zvrátily hlavní závěr o vztahu dokončenosti a změny v interakci s AI.

6.3 Sedm pilotních případů — kvalitativní obohacení

Kvantitativní analýza souboru 147 uživatelů je doplněna o sedm bohatě popsaných pilotních případů, které poskytují kvalitativní vhled do trajektorií, jež se v agregátních datech ztrácejí. Tito sedmi uživatelé byli sledováni v hlubším kontaktu, autor s nimi vedl opakované rozhovory a u všech získal informovaný souhlas s využitím jejich příběhu pro výzkumné účely. Údaje jsou prezentovány pod pseudonymy a popis je redukován na charakteristiky průchodu kurzem bez osobních detailů.

Tabulka 7. Přehled sedmi pilotních případů a jejich přiřazení k identifikovaným klastrům.

Případ	Trajektorie kurzem	Klastr (5.2.5)
--------	--------------------	----------------

Tester 1	Začal videoformátem, opustil kurz v 1/3, nedokončil	Mimo klastry (dropout)
Tester 2	Prošel videosekvenci, byl spokojen, dokončil	K2 (Video-orientovaný)
Tester 3	Prošel a proklikal vše, vysoká spokojenost	K1 (Multimodální tvůrce)
Tester 4	Prošel podcast a infografiku (D5), dokončil	K1 (Multimodální tvůrce)
Tester 5	Poslechl audio, nečetl knihy ani neudělal pracovní listy	K3 (AI-fokusovaný)
Tester 6	Prošel pouze část tvorby s AI, ne pre/postprodukci	K3 (AI-fokusovaný)
Tester 7	Prošel fyzické materiály, diskutoval k materiálům	K4 (Klasický knižní)

Případy ilustrují celou šíři identifikované typologie. Tester 1 reprezentuje raný dropout, který v souboru postihl 15 z 19 nedokončitelných (oddíl 6.2.1). Testeři 2 a 3 představují klasické úspěšné dokončení s různou hloubkou angažmá — Tester 3 prošel kurz multimodálně a využil všechny dostupné formáty, zatímco Tester 2 se držel jediné dimenze. Tester 4 demonstruje funkčnost experimentální dimenze D5 (podcast plus infografika), která v souboru dosáhla 100 % dokončenosti. Testeři 5 a 6 ilustrují trajektorii uživatele orientovaného primárně na AI komunikaci, který pragmaticky přebírá z kurzu jen to, co potřebuje pro vlastní praxi, a

méně se zaměřuje na produkci finálního kurzu. Tester 7 reprezentuje preferenci klasických tištěných materiálů s důrazem na reflexivní čtení.

Hlubší kvalitativní deskripce těchto případů — profesní kontext jednotlivých testerů, jejich předchozí zkušenost s generativní AI, motivace pro vstup do kurzu a doslovné výpovědi z rozhovorů — není v této průběžné publikaci k dispozici. Důvodem je rekonstruktivní povaha datasetu (oddíl 6.1.1): primární záznamy z rozhovorů a průběžných pozorování nebyly v době provozu kurzu systematicky archivovány s výzkumným záměrem a v aktuální formě již nejsou k dispozici v podobě, která by umožnila bohaté case-by-case zpracování ve smyslu Stakeova (2006) nebo Yinova (2018) přístupu k vícenásobné případové studii.

Tato pasáž je proto v aktuální podobě monografie redukována na strukturní mapu (Tabulka 7) a přehledovou interpretaci výše. Plnohodnotná vícenásobná případová studie ve smyslu metodologické tradice je explicitně zařazena do plánu **Fáze 2 výzkumu**, kde bude prováděna prospektivně, s předem stanoveným kódovacím protokolem, informovanými souhlasy účastníků s reprodukcí citací z rozhovorů a se slepým hodnocením druhým kódérem. Sedm případů této průběžné studie slouží v textu monografie tedy primárně jako *empirický most mezi typologií klastrů a budoucí kvalitativní fází výzkumu*.

6.4 Diskuse

6.4.1 Hlavní zjištění a jeho pedagogické implikace

Primárním empirickým nálezem této studie je, že **dokončenost kurzu představuje v dostupných datech nejsilnější faktor asociovaný s deklarovanou změnou v interakci s generativní AI**. Zatímco specifický formát (D1–D5), počet využitých dimenzí, ani VARK preferenční typ neukazují statisticky významný diferenciální efekt, dokončení kurzu dělí soubor na dvě kvalitativně odlišné populace: skupinu, ve které

více než polovina uživatelů deklaruje vysokou změnu, a skupinu, kde tuto kategorii nedosáhne *nikdo*.

Pedagogická implikace tohoto zjištění je významná. V současné diskusi o designu vícedimenzionálních kurzů je často kladen důraz na *optimalizaci formátu* — jak vybrat ten správný formát pro konkrétního uživatele, jak personalizovat obsah, jak adaptovat výuku podle preferenčních stylů. Data této studie přitom naznačují, že **formát je v zásadě fungibilní**: pokud uživatel kurz dokončí, dosáhne pravděpodobně podobných výsledků bez ohledu na to, jakou cestou prošel. Klíčovou metodickou výzvou tedy není ani tak personalizace formátu, jako **podpora vytrvalosti, motivace a dokončenosti**.

Tento závěr koresponduje s rostoucím tělem výzkumu v oblasti retence v online kurzech (např. Kizilcec et al., 2017; Reich, 2014), který opakovaně ukazuje, že hlavním prediktorem úspěchu v MOOC kurzech není kvalita obsahu, ale schopnost uživatele kurz dokončit. Studie této monografie přidává k tomuto obrazu specifický kontext kurzu zaměřeného na rozvoj interakce s AI a potvrzuje analogický vzorec.

6.4.2 Robustnost vícedimenzionálního designu a hodnota experimentální dimenze D5

Druhé klíčové zjištění — **absence diferenciálního efektu mezi dimenzemi** — má jiný typ významu: empiricky podporuje *robustnost* designu kurzu napříč pěti nabízenými cestami. Pokud kurz funguje srovnatelně bez ohledu na zvolenou dimenzi, lze jej považovat za *robust by design*: nabídka více cest neoslabuje žádnou z nich.

Obzvláště pozitivní je nálezk týkající se experimentální **dimenze D5 (podcast + infografika)**. Tato dimenze byla původně koncipována jako pokus reflektovat soudobé způsoby získávání informací (krátké audio formáty, vizuální syntéza komplexního obsahu) a její zařazení do kurzu představovalo metodickou inovaci s nejistým výsledkem. Data ukazují, že D5 nejen obstála ve srovnání s tradičními formáty, ale

dosáhla **100% dokončenosti** a 52 % uživatelů s vysokou změnou interakce s AI. Tento výsledek je významný i přes omezenou velikost podsouboru ($n = 25$), neboť představuje empirický argument pro to, že moderní formáty výuky mohou s tradičními soutěžit i v oblasti, kde jsou tradiční formáty (kniha, video) historicky dominantní.

Pro budoucí design kurzů to otevírá otázku, zda by experimentální dimenze D5 neměla být dále rozvinuta — například o krátkovideovou komponentu, interaktivní vizualizace, nebo audio-vizuální mikrolekce. Tato hypotéza však přesahuje rámec současné studie a je předmětem dalšího pedagogického a designového vývoje.

6.4.3 VARK pretest jako orientační, nikoli určující nástroj

Třetí významnou implikací je redukce role **VARK preferenčního pretestu** v designu kurzu. Studie ukázala, že VARK má pouze slabý vztah k volbě dimenze ($p = 0,062$, $V = 0,228$) a žádný statisticky významný vztah k výslednému efektu kurzu. Toto zjištění odpovídá širší kritice VARK rámce v pedagogické literatuře (Pashler et al., 2008; Husmann & O'Loughlin, 2019), která opakovaně upozorňuje, že empirická evidence pro tzv. *learning styles hypothesis* — tj. tezi, že učení je efektivnější, když je formát adaptován preferenci uživatele — je velmi slabá nebo neexistující.

Současně by však bylo unáhlené VARK pretest zcela odvrhnout. Data této studie naznačují, že VARK může mít **dvě legitimní funkce v designu kurzu**: (a) jako úvodní *reflexivní nástroj* pomáhající uživateli uvědomit si vlastní preference a vstoupit do kurzu vědoměji; (b) jako *orientační doporučení* pro uživatele, kteří si nejsou jisti, jakou cestou začít. Stejně tak však ukazují, že VARK by neměl být *určující* — uživatel má mít svobodu volby dimenze podle vlastního uvážení a nikoli podle algoritmického doporučení.

6.4.4 Integrace s rámcem MLEF-AI

Zjištění této studie lze interpretovat v kontextu rámce **MLEF-AI** (*Multi-Level Ecological Framework for AI Implementation in Education*)

představeného autorem v habilitační teze (Hanák, 2026) a v této monografii podrobně rozvinutého v Kapitole 6. MLEF-AI rozlišuje pět vnořených úrovní implementace AI ve vzdělávání — společensko-politickou (Úroveň 1), eticko-regulační (Úroveň 2), organizační (Úroveň 3), individuální (Úroveň 4) a úroveň pedagogické transformace (Úroveň 5) — a předkládá je jako *dynamický systém* propojený čtyřmi typy zpětných vazeb. Empirická studie této kapitoly se primárně situuje do **Úrovně 4 (Individuální)** a **Úrovně 5 (Pedagogická transformace)**, byť s důležitou modifikací: zatímco MLEF-AI byl koncipován pro učitele jako primární cílovou populaci, naše populace 147 testerů kurzu *Tvůrce kurzů* zahrnuje širší skupinu *uživatelů edukačních platforem*. To představuje aplikační rozšíření rámce, které je v této studii reflektováno explicitně.

Klíčové zjištění studie — silný vztah mezi dokončeností kurzu a deklarovanou změnou v interakci s AI — odpovídá v MLEF-AI rámci dynamice mezi **Úrovní 4** (kde působí psychologické determinanty UTAUT a *AI Self-Efficacy* dle Bandury, 1986) a **Úrovní 5** (kde výsledkem je transformace role uživatele). Konkrétně lze posun v kategorii *AI_CHANGE = high* — od transakčního používání AI k dialogické, iterativní spolupráci — chápat jako uživatelský analogon dichotomie *re-profesionalizace versus deprofesionalizace* z MLEF-AI Úrovně 5. Zatímco originální Role Transformation Index (RTI) habilitační teze měřil tuto dichotomii pro učitele, navrhujeme zde rozšíření aparátu o nový konstrukt **Communication Maturity Index (CMI)**, který by uživatelskou variantu transformace měřil přímo (viz oddíl 6.4 této monografie pro úvodní operacionalizaci konceptu *komunikační zralost*).

Zjištění o **absenci diferenciálního efektu mezi dimenzemi kurzu** (D1–D5) lze v rámci MLEF-AI interpretovat ze dvou perspektiv. *Z pohledu Úrovně 3 (Organizační)* dokumentuje výsledek vysokou kvalitu implementace samotné platformy — pět dimenzí nabízeného kurzu funguje jako pět paralelních cest k cíli, což odpovídá Fixsenovu principu *facilitace adaptace* (Fixsen et al., 2005), jednomu z pěti

implementačních ovladačů. Z pohledu Úrovně 4 (Individuální) naopak naznačuje, že rozhodující determinanty posunu nejsou na úrovni *Effort Expectancy* (EE — vnímané snadnosti), ale na úrovni *Performance Expectancy* (PE) a *AI Self-Efficacy* — tedy v subjektivním přesvědčení uživatele, že má smysl kurzem projít a že jej zvládne. Tato interpretace je v souladu se strukturou empirických zjištění: dokončenost je silný prediktor, formát neutrální.

V rámci feedback loops MLEF-AI je relevantní zejména **FL1 (zpětná vazba zdola nahoru)**: zkušenost uživatelů s platformou Kurzologie generuje data, která zpětně informují provozovatele (autora) o efektivitě jednotlivých designových rozhodnutí. Tato studie sama je manifestací FL1 — uživatelské trajektorie 147 testerů jsou převedeny do designových implikací pro další iteraci platformy a pro plánovanou Fázi 2 výzkumu. Z pohledu **FL3 (etická brzda)** studie odhaluje místo, kde rámec MLEF-AI Úrovně 2 vyžaduje rozšíření o specifika *uživatelské autonomie* — testeři, kteří nedodrželi VARK doporučení, dosáhli srovnatelných výsledků jako ti, kteří jej dodrželi plně, což relativizuje legitimitu algoritmického doporučování formátu výuky a posiluje argument pro uživatelskou volbu jako etický imperativ.

Etická a governance vrstva (Úroveň 2 a Úroveň 3 MLEF-AI) zůstává v této studii jako pole dalšího výzkumu. Otázka, do jaké míry uživatelé po absolvování kurzu mění své postoje k governance AI ve svých vlastních pedagogických praxích (zejména pokud sami vytvářejí kurzy pro studenty), není v dostupných datech zachycena, ale představuje klíčové téma pro Fázi 2. Předběžně lze hypotetizovat, že uživatelé s vysokou *komunikační zralostí* (kategorie AI_CHANGE = *high*) budou současně vykazovat vyšší senzitivitu k otázkám transparentnosti, lidského dohledu a vysvětlitelnosti AI ve vlastních kurzech — což by představovalo empirické ověření **FL3 mechanismu** (etická brzda) v MLEF-AI rámci. Tato hypotéza je formulována jako jeden z výzkumných cílů Fáze 2.

6.4.5 Limity studie

Limity studie jsou závažné a musí být brány v úvahu při jakékoli interpretaci jejich zjištění:

Rekonstruktivní povaha datasetu. Data byla strukturována *ex post facto* z LMS záznamů, pracovních poznámek a rozhovorů, které nebyly původně sebrány s výzkumným záměrem. Cca 18 % řádků nese vyšší míru nejistoty kódování, byť sensitivity analýza ukázala robustnost hlavního zjištění.

Jediný kodér bez slepoty. Kódování všech proměnných provedl autor monografie, který je současně tvůrcem kurzu. Riziko *confirmation bias* nelze v tomto designu kontrolovat. Pro Fázi 2 je naplánováno slepé kódování druhým nezávislým hodnotitelem.

Selekční zkreslení vzorku. 147 uživatelů kurzu *Tvůrce kurzů* nepředstavuje reprezentativní vzorek populace učitelů či tvůrců obsahu, ale samoselektovanou skupinu motivovaných tvůrců s nadprůměrným zájmem o AI v edukaci. Závěry nelze automaticky generalizovat na širší populaci.

Absence kontrolní skupiny. Studie nemá kontrolní skupinu (uživatelé bez kurzu, uživatelé bez AI), což znemožňuje kauzální interpretaci. Lze popsat asociace, ne efekty.

Sebereportované změny v interakci s AI. Proměnná *AI_CHANGE* je založena na sebehodnocení testerů v rozhovorech, doplněném o pozorování autora. Žádné objektivní behaviorální měření (např. analýza skutečných promptů uživatelů z nástrojů AI) není k dispozici. Existuje riziko *demand characteristics* — testeři mohou nadhodnocovat změnu, protože to očekává tvůrce kurzu.

Mnohonásobný střet zájmů. Autor je tvůrce, provozovatel a hodnotitel zároveň. Tato dvojrole je transparentně přiznána, ale představuje strukturální omezení designu.

Souhrnně tedy zjištění této studie představují **orientační vzorce v dostupných datech**, nikoli definitivní empirická tvrzení. Jejich hodnota spočívá v tom, že (a) generují konkrétní hypotézy testovatelné ve Fázi 2, (b) poskytují předběžnou typologii uživatelů pro další design, (c) ukazují robustnost kurzu napříč dimenzemi, a (d) identifikují dokončenost jako klíčovou proměnnou pro další zkoumání.

6.5 Závěr a návaznost na Fázi 2

Empirická studie prezentovaná v této kapitole představuje retrospektivní observační analýzu 147 uživatelů kurzu *Tvůrce kurzů* na platformě Kurzologie.cz, doplněnou o sedm pilotních případů. Studie identifikovala čtyři klíčová orientační zjištění:

1. Existuje silný a robustní vztah mezi dokončeností kurzu a deklarovanou změnou v interakci s generativní AI. Tento vztah obstál ve všech 1 000 iteracích sensitivity analýzy.
2. Specifický formát kurzu (D1–D5) ani počet využitých dimenzí nemají statisticky významný diferenciální efekt na výsledek. Kurz funguje napříč dimenzemi srovnatelně.
3. Experimentální dimenze D5 (podcast + infografika) obstála ve srovnání s tradičními formáty a v některých ukazatelích (dokončenost) je dokonce nejsilnější.
4. VARK preferenční pretest má slabou prediktivní hodnotu pro výslednou změnu interakce s AI. Jeho dodržení neovlivňuje výsledek statisticky významným způsobem.

Tato zjištění představují **východisko pro Fázi 2** výzkumu, která je naplánována na nejbližší měsíce a zahrnuje: prospektivní design s předregistrovaným kódovacím protokolem, slepé kódování druhým nezávislým hodnotitelem, výpočet inter-raterové reliability, vstupní a výstupní měření postojů a chování pomocí standardizovaných nástrojů, behaviorální analýzu skutečných AI promptů uživatelů

(pokud bude technicky a eticky proveditelná), a srovnání s kontrolní skupinou. Cílem Fáze 2 je transformovat orientační zjištění této studie do konfirmačně testovatelných hypotéz a poskytnout kvantitativní data, která obstojí v běžném vědecko-empirickém standardu.

Z hlediska širšího kontextu monografie přispívá tato kapitola k identifikaci **interakčních vzorců** — tématu, jemuž je celá monografie věnována — třemi konkrétními způsoby. Za prvé, doplňuje teoretickou část o empirickou ilustraci, jak se interakční vzorce projevují v reálné populaci uživatelů edukační intervence. Za druhé, identifikuje *dokončenost kurzu* jako klíčovou proměnnou pro další výzkum interakčních vzorců, čímž posouvá pozornost od formátu k procesu. Za třetí, prostřednictvím čtyřklastrové typologie nabízí operacionalizovatelný rámec pro typologickou diferenciaci uživatelů, který může být přenesen i mimo specifický kontext kurzu *Tvůrce kurzů*.

7 Vzorce pro perfektní prompt

V této části se zaměříme na rozvoj dovedností v oblasti tvorby promptů s použitím specifických vzorců. Tyto vzorce nejsou jen pevně danými pravidly nebo šablonami; představují osvědčená schémata, která fungují jako most mezi lidským dotazem a porozuměním na straně AI. Jak uvádějí White et al. (2023) a Knoth et al. (2024), dobře definované promptové vzory poskytují znovupoužitelné řešení typických problémů při komunikaci s velkými jazykovými modely a výrazně zlepšují efektivitu interakce. Jinými slovy, strukturované vzorce nám umožňují programovat AI „slovy“ – formovat výstup pomocí promyšleného zadání (Mollick, 2023). Tím dokážeme překlenout propast mezi tím, na co se ptáme, a tím, jak naše dotazy AI skutečně chápe a zpracuje.

Ukážeme si, jak použitím promyšlených struktur (vzorců) můžeme naše dotazy zefektivnit, dodat jim hloubku a kontext, a tím zajistit, že odpovědi AI budou co nejrelevantnější a nejpřesnější vzhledem k našim cílům. Prozkoumáme různé aspekty těchto vzorců – od základních komponent (jako je samotný dotaz, kontext či požadovaný formát odpovědi) až po pokročilejší prvky, jako je specifičnost, účel dotazu, tón odpovědi a další. Díky tomu získáte širší paletu možností, jak s AI interagovat, a lépe pochopíte, jak své dotazy formulovat tak, aby co nejlépe odpovídaly vašim specifickým potřebám.

Cílem této kapitoly je vybavit vás dovednostmi vytvářet **prompty, které jsou nejen efektivní, ale také přizpůsobené konkrétním situacím a cílům.** Takové prompty povedou k hlubší a bohatší interakci s AI – pomohou vám objevit nové informace, řešit komplexní problémy a rozšířit vaše obzory, a to vše prostřednictvím správně zacílené komunikace s umělou inteligencí.

Co jsou to vzorce v kontextu AI?

V kontextu umělé inteligence (AI) jsou *vzorci promptů* nástroje či směrnice, které uživatelům pomáhají formulovat efektivní dotazy

(prompty) pro AI modely. Jedná se o strukturované sekvence instrukcí, kontextu a dalších prvků dotazu, které specifikují a usměrňují komunikaci s AI tak, aby generovaná odpověď byla co nejpřesnější a nejrelevantnější. Každý takový vzorec obsahuje klíčové komponenty – například upřesnění úkolu, poskytnutí kontextu, nastavení požadovaného formátu odpovědi, tónu či role, kterou má AI při odpovědi zaujmout.

Jinými slovy, *promptový vzorec* můžeme chápat jako určitý **“recept” na dotaz**. Při jeho použití uživatel neklade jen samotnou otázku, ale zároveň poskytuje AI důležité instrukce o tom, **jaký druh odpovědi očekává** a za jakých podmínek či v jakém stylu má AI odpovídat. Podle Vatsala a Dubeyho (2024) vyžaduje prompt engineering formulaci vstupních instrukcí v přirozeném jazyce takovým způsobem, aby z modelu strukturovaně vylákal požadované znalosti či informace. Tímto způsobem se prompty stávají účinným nástrojem, jak přizpůsobit obrovské množství embedded znalostí v AI našim konkrétním potřebám (White et al., 2023).

Důležitost vzorců

Vzorci promptů hrají klíčovou roli ve zlepšování komunikace mezi člověkem a AI. Bez jasně definovaného a strukturovaného promptu může AI poskytnout odpovědi, které jsou nepřesné, nerelevantní nebo vytržené z kontextu původního dotazu. Dobře navržené prompty jsou naopak zásadní pro minimalizaci chyb, biasů a nesrovnalostí ve výstupech modelu (Reynolds & McDonell, 2021). Výzkumy ukazují, že přesné a konkrétní instrukce mohou významně zvýšit spolehlivost a užitečnost generovaných odpovědí napříč různými doménami. Například poskytnutí jasných specifik dokáže vést model k relevantnějším a méně zaujatým výstupům (Reynolds & McDonell, 2021; Brown et al., 2020). Pečlivě vytvořený prompt také může zmírnit tzv. *halucinace* modelu a zvýšit faktickou správnost odpovědi (Brown et al., 2020).

Zde je několik důvodů, **proč jsou vzorce (stejně jako obecné strategie práce s AI) tak důležité:**

1. **Strukturování dotazů:** Vzorce umožňují uživatelům strukturovat své dotazy jasně a jednoznačně. Poskytují promyšlený rámec, který pomáhá lépe definovat, co od AI očekáváme. Tím se snižuje riziko nepochopení – AI nemusí „hádat“, co přesně má udělat, protože dobře navržený prompt to explicitně říká (T'Syen, 2025). Výsledkem jsou konkrétnější a relevantnější odpovědi, než jaké bychom získali z vágně položených otázek.
2. **Zvýšení efektivity:** Efektivní prompt šetří čas a zvyšuje produktivitu. Místo opakovaného pokusu a omylu při kladení doplňujících otázek lze správně strukturovaným dotazem získat kvalitní odpověď na první pokus. T'Syen (2025) uvádí, že dobře navržené prompty vedou k konzistentnějším a strukturovanějším odpovědím a pomáhají vyhnout se zdlouhavému ladění dotazu.
3. **Podpora komplexních dotazů:** Díky vzorcům mohou uživatelé formulovat i složitější a víceúrovňové dotazy. Kombinací různých prvků promptu lze AI požádat o řešení komplexních problémů krok za krokem. Například tzv. *řetězení myšlenek* (chain-of-thought prompting) vybízí model k rozdělení obtížné otázky na menší části, což prokazatelně zlepšuje schopnost modelu uvažovat a dospět ke správné odpovědi (Wei et al., 2022; Reynolds & McDonell, 2021).
4. **Zlepšení pochopení AI:** Vzorce také pomáhají uživatelům lépe pochopit, jak AI funguje a jak s ní nejúčinněji interagovat. Tím, že si osvojíme myšlení v intencích promptových vzorců, vlastně se učíme „uvažovat jako AI“ – uvědomíme si, jaké informace model potřebuje a jak mu je

předložit. Takové porozumění vede ke smysluplnějším konverzacím a hodnotnějším výstupům (Knoth et al., 2024).

5. **Personalizace interakce:** Stejně jako obecné strategie umožňují i vzorce promptů personalizovat interakci s AI. Uživatelé mohou šablonu dotazu přizpůsobit svým specifickým potřebám – například zvolit tón (formální vs. neformální), roli odpovídajícího (expert, začátečník apod.) či úroveň detailu. Tím získají informace „šitě na míru“ jejich účelu. Personalizace promptu vede k odpovědím, které jsou pro uživatele srozumitelnější a přínosnější (T'Syen, 2025).

Vzorce a obecné strategie je tedy vhodné vnímat jako vzájemně se doplňující nástroje pro maximalizaci potenciálu AI. Strategie (jako např. „buďte konkrétní“, „dodejte kontext“, „kontrolujte výstupy modelu“ apod.) nám dávají obecný rámec, jak k práci s AI přistupovat (Cook, 2023). Vzorce pak představují konkrétní způsoby, jak tyto strategie uvést do praxe jednotným a osvědčeným způsobem (White et al., 2023). Společně tvoří základ pro vytváření smysluplných a produktivních interakcí s umělou inteligencí.

7.1 Komponenty vzorců

Stejně jako při skládání hudební kompozice nebo psaní scénáře je i u promptů klíčem k úspěchu porozumění struktuře a jednotlivým komponentám, které tvoří účinný prompt. Promptové vzorce se typicky skládají z několika prvků, které lze libovolně kombinovat a přizpůsobovat. Podle obecně přijímaných doporučení obsahuje prompt v přirozeném jazyce typicky následující složky (Sharma, 2024): instrukci (co má AI udělat), kontext (doplňující informace pro zasazení úkolu), vstupní data nebo dotaz (konkrétní otázku k zodpovězení) a výstupní indikátor (specifikaci, v jakém formátu či stylu má být odpověď). Pro účely této kapitoly si definujeme širší sadu komponent, z nichž nejdůležitější jsou následující (prvním deseti se budeme věnovat podrobněji):

1. **Dotaz (Task):** Jádru promptu – jasně formulovaný požadavek nebo úkol pro AI.
2. **Kontext (Context):** Doplňující informace, které rozšiřují či upřesňují zadání a pomáhají AI lépe porozumět situaci.
3. **Specifičnost (Specificity):** Míra detailnosti a konkrétnosti zadání.
4. **Účel (Purpose):** Cíl, proč je dotaz kladen – co má být výsledkem nebo jak má být informace použita.
5. **Formát (Format):** Požadovaná forma, struktura či styl odpovědi (např. odrážky, tabulka, esej, kód).
6. **Tón (Tone):** Tón a styl komunikace odpovědi (formální, neformální, odborný, přátelský apod.).
7. **Příklad (Exemplar):** Ukázka či příklad očekávané odpovědi, který AI napoví, jak má výstup vypadat.
8. **Osobnost (Persona):** Určení role, ze které má AI odpovídat (např. „jako historik“, „expert na medicínu“, „začátečník“ apod.).
9. **Úroveň expertizy (Level of Expertise):** Stanovení úrovně odbornosti nebo znalostí, kterou má AI v odpovědi uplatnit (pro laiky vs. pro experty).
10. **Očekávaná délka odpovědi (Expected Response Length):** Jak obsáhlá či detailní má odpověď být (např. „stručně do 100 slov“ nebo „alespoň tři odstavce“).
11. **Hledisko nebo perspektiva (Point of View):** Úhel pohledu, z něž má být odpověď formulována (např. ekonomický, sociální, historický).
12. **Interaktivita (Interactivity):** Zda má odpověď vyzývat k další konverzaci či otázkám, nebo být definitivní.

13. **Interdisciplinární přístup (Interdisciplinary Approaches):** Zapojení znalostí z více oborů do jedné odpovědi.
14. **Hypotetické scénáře (Hypothetical Scenarios):** Představování si „co by, kdyby“ situací pro potřeby odpovědi.
15. **Historický kontext (Historical Context):** Zasazení problému či odpovědi do historického rámce.
16. **Budoucí predikce a scénáře (Future Predictions and Scenarios):** Žádost o predikci budoucího vývoje nebo nastínění možných scénářů.
17. **Konkrétní zdroje informací (Specific Sources of Information):** Určení, z jakých zdrojů či dat má AI čerpat (pokud je to možné).
18. **Uživatelský kontext (User Context):** Informace o uživateli či jeho situaci, které mohou ovlivnit odpověď AI.
19. **Globální vs. lokální perspektiva (Global vs. Local Perspective):** Zda má AI odpovídat z celosvětového pohledu, nebo se zaměřit na konkrétní region či prostředí.
20. **Srovnání a kontrasty (Comparisons and Contrasts):** Požadavek na porovnání dvou a více věcí, zdůraznění rozdílů či podobností.
21. **Kontroverzní nebo provokativní témata (Controversial or Provocative Topics):** Vyzvání AI, aby se zabývala citlivými nebo spornými tématy.
22. **Integrace různých typů dat (Integration of Various Data Types):** Kombinace různých druhů informací (text, čísla, kódy, obrázky) v odpovědi.
23. **Kreativní omezení (Creative Constraints):** Specifická omezení podporující kreativitu (např. odpověď formou básně, použití metafor apod.).

24. **Prioritizace nebo seřazení informací (Prioritization/Sequencing):** Požadavek, aby AI informace v odpovědi určitou logikou seřadila nebo vypíchla nejdůležitější body.
25. **Omezení a restrikce (Limitations and Restrictions):** Co by mělo být v odpovědi vynecháno nebo explicitně neobsahováno.
26. **Vývoj a historie tématu (Development and History of the Topic):** Žádost, aby AI nastínila, jak se dané téma vyvíjelo v čase či jaký je jeho původ.
27. **Etické a právní úvahy (Ethical and Legal Considerations):** Zahrnutí morálních či právních aspektů do odpovědi (pokud jsou relevantní).
28. **Specifikace výstupního formátu (Specification of Output Format):** Přesné určení, v jaké podobě má AI odpověď dodat (např. HTML kód, seznam odrážek, JSON struktura apod.).

7.2 Rozbor hlavních komponent pro vzorce

Aby AI poskytovala přesné, relevantní a užitečné odpovědi, je nutné pochopit a správně využívat jednotlivé hlavní komponenty, které tvoří základ každého promptu (dotazu). Níže se podrobněji věnujeme prvním deseti komponentám z výše uvedeného seznamu – vysvětlíme jejich význam, účel a osvědčené postupy, jak je v rámci promptů využít pro dosažení optimálních výsledků. U každé komponenty doplníme i praktické příklady a odůvodnění, **podpořené zjištěními z odborné literatury**.

Komponenta *Dotaz* (Task)

Komponenta „Dotaz“ představuje základní prvek každého promptu – definuje hlavní úkol nebo otázku, kterou uživatel AI pokládá. Je to vlastně **jádro promptu**, od kterého se odvíjí vše ostatní.

Účel komponenty dotaz: Jasně stanovit, o co AI žádáme. Správně formulovaný dotaz umožňuje AI pochopit, co přesně má udělat či zodpovědět. Měl by tedy jednoznačně popsat požadavek uživatele. Například místo vágního „Pověz mi něco o klimatu“ je vhodnější formulace „**Vysvětli příčiny a důsledky klimatických změn**“.

Jak formulovat efektivní dotaz: Dotaz by měl být co nejvíce srozumitelný a přímočarý. Vyvarujme se dvojznačnosti a příliš obecné formulace. Lze doporučit začít dotaz slovesem, které implikuje akci (např. „vysvětli“, „zjisti“, „porovnej“, „navrhni“). Podle J. Maedy (2023) spočívá umění prompt engineering mj. ve správné volbě slov a frází, které co nejlépe vystihnou náš požadavek. Uživatel by měl již při formulaci dotazu myslet na to, jakou *konkrétní odpověď* očekává (Johnmaeda, 2023).

Výhody: Správně a přesně položený dotaz výrazně zvyšuje šanci, že AI porozumí zadání a poskytne relevantní odpověď. Pokud je dotaz příliš obecný, model může „tápat“ a nabídnout příliš obecnou či nerelevantní odpověď. Když je naopak dotaz ostrý a jasný, AI se může soustředit přesně na požadovaný problém a nezanáší do odpovědi zbytečnosti (Reynolds & McDonell, 2021).

Příklady efektivních dotazů:

- Místo obecného dotazu „*Jak zlepšit životní styl?*“ je lepší položit konkrétnější otázku: „**Jaké změny ve stravě a pohybových aktivitách mohou do půl roku zlepšit moji fyzickou kondici a snížit úroveň stresu?**“. Zde je jasně definováno, na co se ptáme (změny ve stravě a cvičení), s jakým cílem (zlepšit kondici, snížit stres) a v jakém časovém horizontu (půl roku).
- Pro technický dotaz, např. „*Jak funguje blockchain?*“, můžeme doplnit i účel: „**Stručně vysvětli princip fungování blockchainu a uveď jeho využití v kryptoměnách.**“ Tím dáváme AI jasný úkol vysvětlit koncept a zasadit ho do konkrétní aplikační domény.

Doporučení pro uživatele: Při formulaci dotazu se vždy snažte vcítit do role AI: Může být váš dotaz pochopen více způsoby? Nechybí v něm upřesnění, o co přesně stojíte? Pokud AI nereaguje podle očekávání, často pomůže přeformulovat dotaz – možná byl příliš vágní nebo obecný. Kvalita odpovědi AI totiž **přímo závisí na kvalitě položeného dotazu** (White et al., 2023).

Komponenta Kontext (Context)

Komponenta „Kontext“ poskytuje AI dodatečné informace a zasazuje dotaz do určitého rámce. Kontext může dramaticky ovlivnit, jak AI porozumí vaší otázce a jakou odpověď vygeneruje.

Účel komponenty kontext: Kontext pomáhá upřesnit situaci nebo pozadí dotazu. Díky tomu AI lépe chápe, *v jakém smyslu či rámci* má na dotaz odpovídat. Kontext může zahrnovat například popis problému, předchozí konverzaci, informace o uživateli či cílovém publiku odpovědi, nebo třeba styl, ve kterém má být odpověď napsána. White et al. (2023) poznamenávají, že prompt vlastně nastavuje **kontext konverzace** a říká modelu, jaké informace jsou důležité a jaký typ výstupu se žádá.

Jak efektivně používat kontext: Uveďte všechny relevantní informace, které by mohly ovlivnit správnou odpověď. Například pokud se ptáte na právní radu, uveďte jurisdikci; pokud žádáte o vysvětlení vědeckého pojmu, můžete zmínit, pro jaké publikum to má být (jinak bude AI vysvětlovat kvantovou fyziku dětem a jinak expertům). Mollick (2023) demonstroval, že když do promptu doplníme kontext nebo *rolí*, výrazně to změní charakter odpovědi. Např. prompt *„Jsi zkušený biolog specializovaný na stromy. Na základě aktuálních klimatických dat vysvětli, kdy bude v Nové Anglii vrcholit sezóna podzimního listí – a odpověď formuluj tak, aby jí porozuměly děti v mateřské školce.“* poskytne mnohem specifitější a kontextově zasazenou odpověď než prosté *„Kdy je nejlepší čas na podzimní listí v Nové Anglii?“*.

Výhody: Přidáním kontextu zvyšujeme pravděpodobnost, že odpověď bude přesná a „na míru“. AI tak nemusí odhadovat skryté předpoklady – místo toho je přímo uvedeme. Kontextualizované dotazy umožňují modelu **lépe se orientovat i v komplexních tématech** a poskytovat užitečnější informace. Například dotaz „Vysvětli koncept entropie“ versus „Vysvětli koncept entropie v kontextu termodynamiky tak, aby to pochopil student střední školy“ – druhá varianta využívá kontext (termodynamika, cílový student), takže odpověď bude zacílená a srozumitelná.

Příklady využití kontextu:

- Obecný dotaz: „*Jak investovat do akcií?*“. Doplněný kontext: „*Jsem začátečník a rád bych investoval menší částku do akcií technologických firem. Jak postupovat a na co si dát pozor?*“. Tady kontext („jsem začátečník“, „menší částka“, „technologické firmy“) navede AI k adekvátní radě pro nováčka.
- Dotaz na zdravotní téma: „*Co je to diabetes?*“. S kontextem: „*Můj otec (55 let) má nově diagnostikovanou cukrovku 2. typu. Můžeš mi vysvětlit, co to přesně znamená a jak mu mohu pomoci upravit stravu a životní styl?*“. Konkrétní situace rodinného příslušníka dá odpovědi osobnější a praktičtější rozměr.

Doporučení pro uživatele: Kdykoli je to možné, *zahrňte do promptu relevantní kontext*. Přemýšlejte, jaké doplňující informace by mohla AI potřebovat, aby poskytla co nejlepší odpověď. Kontextem může být i pokračování předchozí konverzace – pokud se AI ptá ve několika krocích, připomeňte jí, o čem byla řeč. Jestliže AI odpoví příliš obecně nebo „mimo mísu“, zkuste příště kontext upřesnit či rozšířit. **Často platí, že čím lépe je dotaz ukotven v kontextu, tím přínosnější a užitečnější odpovědi se dočkáte.**

Komponenta Specifičnost (Specificity)

Komponenta „Specifičnost“ určuje, jak detailní a přesný dotaz je. Jde o to, do jaké míry upřesňujeme zadání a omezujeme prostor možných odpovědí. Specifičnost je jedním z nejdůležitějších faktorů ovlivňujících kvalitu odpovědi – úzce souvisí s již zmiňovanou jasností a konkrétností promptu.

Účel komponenty specifičnost: Cílem je předcházet příliš obecným či vágním odpovědím tím, že dotaz dostatečně upřesníme. Chceme-li od AI konkrétní informaci, musíme klást konkrétní otázky. Precizně specifikovaný dotaz pomáhá modelu **lépe pochopit naše záměry** a neposkytovat odpovědi „širokým rozmachem“.

Jak efektivně používat specifičnost: Buďte explicitní ohledně toho, co vás zajímá. Zahrňte detaily jako jsou časové období, geografická oblast, konkrétní metriky, srovnávané položky apod., pokud jsou relevantní. Jak uvádí T'Syen (2025), bez upřesnění rizikujete, že odpověď bude příliš obecná nebo na špatné úrovni detailu. Příkladem může být rozdíl mezi „Řekni mi něco o dopadu změny klimatu.“ vs. „Diskutuj **ekonomické dopady** klimatické změny **v rozvojových zemích během příštích 10 let.**“. Druhý dotaz jasně vymezuje téma (ekonomické dopady), kontext (rozvojové země) i horizont (příštích 10 let) – model tak může **zacílit odpověď** a nezabíhat do nerelevantních oblastí.

Výhody: Podle výzkumů přesné instrukce významně zvyšují spolehlivost modelu a relevanci výstupu. Specifičnost v promptu modelu **zužuje interpretační prostor** – AI nemusí hádat, co by uživatel „mohl chtít“, protože to má jasně uvedeno. Tím se snižuje riziko, že model sklouzne k obecným frázím nebo vynechá to podstatné (Reynolds & McDonell, 2021). Kromě relevance tak roste i konzistence odpovědí. Uživatel navíc získává větší kontrolu nad obsahem – má možnost předem rozhodnout, jaké detaily chce zahrnout či vypustit.

Příklady využití specifičnosti:

- Namísto „*Jak vařit rýži?*“ specifikujte: „***Jak uvařit thajskou jasmínovou rýži v rýžovaru, aby byla sypká?***“. Doplnili jsme typ rýže, způsob vaření (v rýžovaru) a požadovaný výsledek (syká konzistence). Odpověď pak bude mnohem praktičtější.
- Místo „*Porovnej elektromobily a spalovací auta.*“ lépe: „***Porovnej elektromobily a vozidla s benzínovým motorem z hlediska emisí CO₂ a celkových provozních nákladů.***“ Zde jsme určili konkrétní kritéria srovnání, takže AI strukturuje odpověď kolem emisí a nákladů, místo obecného porovnávání všeho možného.

Doporučení pro uživatele: Kdykoli je výstup příliš obecný, hledejte, kde zpřesnit dotaz. Ptejte se sami sebe: *Může být má otázka ještě konkrétnější?* U technických či faktografických dotazů se vyplatí uvést **co nejvíce upřesnění** – model pak často překvapí precizností odpovědi. Pamatujte, že AI čerpá ze znalostí, které má „nasáté“ z obrovských dat – **vy jste ti, kdo rozhoduje, kterou část těch znalostí vytáhne na povrch**. Specifický prompt je jako dobře mířený dotaz v knihovně: když přesně víte, co hledáte, knihovník (AI) vám snáz přinese tu správnou knihu (odpověď).

Komponenta Účel (Purpose)

Komponenta „Účel“ vysvětluje, *proč* daný dotaz kladete nebo jaký má mít odpověď účel. Jde o to naznačit AI, jaký výstup je žádoucí z hlediska záměru uživatele. Tato informace může AI pomoci odpověď **lépe zacílit** a přizpůsobit ji požadovanému účelu či použitelnosti.

Účel komponenty účel: Uvést, k čemu má odpověď sloužit nebo jaký je kontext užití. Někdy mohou různé odpovědi být správné, ale pro uživatele je přínosná jen určitá forma – a tu lze předznačit právě uvedením účelu. Například dotaz „*Shrň tento článek.*“ vs. „*Shrň tento článek pro účely tiskové zprávy.*“ – v druhém případě AI ví, že má vypíchnout hlavní sdělení vhodná do tiskové zprávy, ne všechny detaily.

Jak efektivně používat komponentu účel: Uvedte v promptu, co hodláte s odpovědí dělat nebo kdo ji bude číst. Můžete zmínit cílové publikum („vysvětlí to pětiletému dítěti“ vs. „vysvětlí to odborníkovi“), nebo cíl odpovědi („...abych to mohl prezentovat vedení firmy“, „...pro zábavný příspěvek na blog“ apod.). Takové upřesnění poskytne AI vodítko, jaký styl a obsah zvolit. **Zaměření na účel** často úzce souvisí s tónem a formátem – všechny tyto prvky dohromady definují zadání opravdu přesně.

Výhody: Sdělením účelu dotazu zajistíte, že odpovědi AI budou relevantnější a praktičtější. Model bude lépe chápat, *co od něj jako konečný produkt očekáváte*. To šetří čas – AI nemusí bloudit a uživatel následně nemusí odpověď složitě upravovat pro svůj účel. Pokud například žádáte o **návrh řešení problému**, je rozdíl, zda to potřebujete „nápad na brainstorming“ (můžete dostat spoustu různých nápadů) nebo „konkrétní akční plán do praxe“ (AI se zaměří na proveditelné kroky). Když tento účel zmíníte, odpověď bude přesněji odpovídat vašim potřebám.

Příklady využití účelu:

- Obecný dotaz: „*Jak mohu zvýšit návštěvnost webu?*“. S účelem: „**Poradíš mi konkrétní kroky, jak v následujících 3 měsících zvýšit návštěvnost mého e-shopu s ručně vyráběnými šperky? Chci se zaměřit na online marketing.**“ – Uvádíme účel (zvýšit návštěvnost e-shopu), časový rámec a obor podnikání, takže AI pravděpodobně navrhne cílené marketingové tipy.
- Dotaz: „*Jak se učit cizí jazyk?*“. S účelem: „**Jaké jsou nejefektivnější metody, jak se během půl roku naučit základy španělštiny kvůli plánované cestě?**“. AI pochopí, že má poradit metody učení vhodné pro cestovatele a zmínit praktické tipy (např. konverzační fráze, aplikace na učení jazyka), namísto obecných rad pro jakýkoli jazyk.

Doporučení pro uživatele: Když formulujete prompt, zkuste přidat větu nebo slovní spojení, které odhalí, čeho chcete dosáhnout. Např. „...abych mohl...“, „pro použití v...“ apod. To mnohdy vnese do odpovědi potřebné zaměření. Pokud nejste spokojeni s odpovědí AI, zamyslete se, zda jste dostatečně ozřejмили účel – možná AI tápala, co z mnoha možných směrů vás zajímá. **Jasně vyjádřený záměr** promptu vede k odpovědím, které jsou přímo použitelné pro váš konkrétní případ.

Komponenta Formát (Format)

Komponenta „Formát“ určuje, **v jaké podobě či struktuře** by měla AI podat odpověď. Je to velmi mocný nástroj – díky němu můžeme model přimět, aby generoval text určitým způsobem, což může dramaticky zvýšit užitečnost a přehlednost výstupu.

Účel komponenty formát: Specifikovat požadované uspořádání odpovědi. Někdy potřebujeme od AI seznam, jindy esej, tabulku, kód nebo třeba JSON. Tím, že předem sdělíme, jak má odpověď vypadat, **ušetříme práci i čas** s následnou úpravou výstupu. AI se může pokusit rovnou dodat informaci v požadovaném formátu.

Jak efektivně používat komponentu formát: Budte explicitní. Klidně napište: „Odpověz ve formě bullet-pointů“, „Výsledek napiš jako Python kód“, „Vygeneruj tabulku s následujícími sloupci...“ apod. T'Syen (2025) doporučuje jasně nastavit očekávání ohledně délky, struktury a stylu odpovědi. Například prompt „Uved' 3 hlavní výhody a 3 nevýhody elektromobilů v odrážkovém seznamu.“ donutí AI držet strukturu seznamu a počet položek. Podobně „Shrň text do tabulky (sloupce: Hlavní myšlenka | Důkazy | Závěr).“ povede k tabulární odpovědi namísto souvislého textu.

Výhody: Získáte informace **v nejvhodnějším a nejpřehlednějším formátu** pro svůj účel. Například když potřebujete rychlý přehled, oceníte odrážky či očíslované seznamy. Když žádáte o kód, chcete ho v monospace bloku. Správný formát zlepšuje čitelnost a

srozumitelnost odpovědi. Navíc se tím dá předejít příliš rozvleklym odpovědím – definujete-li formát a rozsah, model se ho obvykle drží. T'Syen (2025) uvádí, že i základní formátovací instrukce jako „použij očíslované kroky“ nebo „shrň to do tabulky“ mohou dramaticky zlepšit přehlednost výstupu.

Příklady využití formátu:

- „*V bodech popiš postup první pomoci při popálenině.*“ – AI vygeneruje očíslované kroky namísto souvislého textu, což se lépe čte a následně aplikuje.
- „*Vygeneruj prosím kód v jazyce Python, který vezme pole čísel a vrátí je seřazené.*“ – Specifikujeme formát jako kód v Pythonu, takže AI pravděpodobně odpoví přímo úsekem kódu s vysvětlením v komentářích (namísto dlouhé teoretické odpovědi o algoritmech řazení).

Doporučení pro uživatele: Neváhejte do promptu zahrnout formální požadavky na výstup. Pokud chcete odpověď do 5 vět, řekněte to. Pokud preferujete tabulku, řekněte to. Formát lze kombinovat s dalšími komponentami – např. „*Srovnej X a Y, a odpověď dej do dvou sloupců: výhody X vs. výhody Y.*“ Model pak pochopí, že má vytvořit strukturované srovnání. **Čím konkrétněji popíšete formu očekávané odpovědi, tím méně uprav vás následně čeká.**

Komponenta Tón (Tone)

Komponenta „Tón“ určuje, v jakém stylu či náladě má být odpověď formulována. Jiný tón použijeme pro akademickou esej, jiný pro přátelskou radu nebo marketingový text. Schopnost ovlivnit tón odpovědi je důležitá, chcete-li, aby výsledek **zapadl do určitého komunikačního kontextu** nebo ladil s publikem, pro který je určen.

Účel komponenty tón: Nastavit emocionální a stylistický charakter odpovědi. AI můžeme instruovat, aby odpovídala formálně, neformálně, odborně, laicky, optimisticky, neutrálně, empaticky,

přísně vědecky atd. Tón často úzce souvisí s komponentou *Osobnost a Úroveň expertizy*, protože například persona „profesor fyziky“ bude mít formální odborný tón, zatímco persona „kamarád povzbuzující ke cvičení“ bude hovořit odlehčeně a motivačně.

Jak efektivně používat komponentu tón: Uvedte přídatnými jmény či popisem, v jakém duchu má AI odpovídat. Např. „*odpověz zdvořile a profesionálně*“, „*použij jednoduchý, přátelský tón, jako bys mluvil s dítětem*“, „*napiš to motivačně, s nadšením*“. Pokud cílíte na konkrétní publikum, definice osoby a publika často přirozeně tón určí (T'Syen, 2025): například zadáte-li „**Jsi univerzitní profesor literatury** a... vysvětlí... *pro laické čtenáře*“, automaticky tím kombinujete odbornou perspektivu se srozumitelným jazykem pro neodborníky. Model pak pravděpodobně použije vstřícný a vysvětlující tón, vyhne se žargonu, ale zachová si jistou odbornou autoritu.

Výhody: Správně zvolený tón zvyšuje **přijatelnost a efekt** odpovědi. Pokud ladí s tím, co uživatel potřebuje, je informace lépe stravitelná. Například v edukativním kontextu pro děti zvolíme jednoduchý a povzbuzující tón – děti pak obsah snadněji pochopí. Pro manažerské shrnutí naopak volíme věcný, stručný a formální jazyk – tím zajistíme, že odpověď bude použitelná třeba jako podklad do prezentace pro vedení firmy. T'Syen (2025) ukazuje, že specifikace publika (např. „*juniorní vývojáři*“ vs. „*C-level manažeři*“) se promítne do tónu i hloubky detailů: pro juniory model zjednoduší jazyk a více vysvětluje, pro manažery je konzistentně stručný a zaměřený na výsledek.

Příklady využití tónu:

- „*Vysvětlí mi kvantovou mechaniku tónem, jako kdyby to byla pohádka pro děti.*“ – Model pravděpodobně použije metaforu, jednoduchá slova a hravý tón.
- „*Zhodnot' výkon zaměstnance za poslední měsíc v motivačním duchu, ale zároveň profesionálně.*“ – Odpověď by měla vyznít

optimisticky a podporující, zároveň však s formálností HR hodnocení.

Doporučení pro uživatele: Přemýšlejte, kdo bude číst výsledek a jaký pocit či dojem má text vyvolat. Nebojte se AI říct, aby zaujala určitý tón či styl – je to podobné, jako když spisovatele požádáte: „Napiš mi to formou vtipné historky“ nebo „Buď prosím věcný a neutrální“. AI umí styl do značné míry měnit, pokud má v datech vzory. Například současné velké jazykové modely (GPT-5, Claude 4, Gemini 3) mají v tréninku obsaženy různé textové styly, a když mu řeknete „představ si, že jsi Shakespeare a...“, pokusí se stylizovat. Samozřejmě, čím specifitější tón (či jejich kombinaci) požadujete, tím obtížnější to může pro model být – ale za zkoušku nic nedáte. Jasným určením tónu zvýšíte pravděpodobnost, že výsledek přesně zapadne tam, kam potřebujete.

Komponenta *Příklad* (Exemplar)

Komponenta „Příklad“ je způsob, jak AI **ukázat vzor očekávané odpovědi**. Poskytnutí jednoho nebo více příkladů (tzv. *few-shot prompting*) může modelu velmi napomoci porozumět, co přesně po něm chceme. Výzkum v oblasti velkých jazykových modelů ukázal, že LLM dokážou z několika příkladů v promptu generalizovat a následovat vzor – Brown et al. (2020) ve své průlomové studii nazvané „*Language models are few-shot learners*“ ukázali, že model GPT-3 je schopen z pouhých několika ukázkových vstupů a výstupů **naučit se** provést obdobný úkol na novém vstupu bez dalšího trénování.

Účel komponenty příklad: Umožnit AI lépe pochopit formu či obsah požadované odpovědi skrze demonstraci. Příklad funguje jako **šablona** – model se podle něj orientuje, jak má vypadat výstup (formát, úroveň detailu, styl) a do jisté míry i jaký obsah je relevantní. To může být užitečné hlavně u kreativních nebo otevřených úloh, kde pouhý popis zadání nestačí.

Jak efektivně používat komponentu příklad: Příklad může mít formu ukázkového dialogu, strukturovaných dat nebo jednoduše prototypu odpovědi. Můžete např. modelu dát: „*Příklad otázky: ..., Příklad odpovědi: ...*“ a pak položit vlastní otázku. Nebo: „*Zde jsou dva vzorové výstupy... Vygeneruj třetí podobný.*“ Důležité je, aby příklady byly **relevantní a kvalitní**, protože model je bude následovat. Také by měly ilustrovat to, v čem potřebujete vést – ať už jde o styl (např. ukázka tónu z předchozího bodu), formu (ukázka formátu) nebo obsahovou hloubku.

Výhody: Poskytnutím příkladů dramaticky snížíte nejednoznačnost promptu. Model dostane **konkrétní předlohu**, takže spíše „pochopí, co po něm chcete“. Brown et al. (2020) demonstrovali, že GPT-3 dokáže z několika příkladů dosáhnout výkonnosti srovnatelné s natrénovanými modely na specifické úlohy – tzv. *prompt-based few-shot learning*. Pro uživatele to znamená, že nemusí model programovat ani složitě ladit parametry – stačí mu *ukázat*, co má dělat, a on to napodobí. Příklad navíc může vymezit, jaké informace jsou podstatné. Například u shrnutí textu můžete jako příklad dát shrnutí jiného článku, aby model věděl, jak stručné a jak strukturované shrnutí chcete.

Příklady využití komponenty příklad:

- Řekněme, že chcete vygenerovat věty ve stylu určitých přísloví. Prompt může znít: „*Vytvoř nové přísloví na téma trpělivosti.*“ **Příklad:** „*Trpělivost růže přináší.*“ Model pak pravděpodobně vytvoří něco jako „Trpělívý hlemýžď dojde nejdál“ – protože chápe, že má napodobit styl přísloví.
- Pokud učíte model formát filmové recenze, můžete poskytnout: „**Ukázka recenze:** *Film X je strhující drama... (následuje odstavec).*“ **Konec ukázky.** *Nyní napiš recenzi filmu Y ve stejném stylu.*“. Model se z příkladu naučí tón (např. odborně kritický), formu (úvod, rozbor, závěr) i slovník typický pro recenze.

Doporučení pro uživatele: Příklady využívejte zvláště tehdy, když chcete **specifický styl nebo strukturu**, které by AI mohla těžko odhadnout jen z popisu. Dejte si však pozor na to, aby příklad nebyl zavádějící nebo příliš složitý. Ideálně by měl být krátký a výstižný. U více příkladů (few-shot) můžete ukázat variabilitu. Mějte také na paměti, že příliš mnoho příkladů může prompt zbytečně prodlužovat – často stačí 1–2 kvalitní ukázky. Brown et al. (2020) uvádějí, že výkon modelu se sice s rostoucím počtem příkladů zlepšuje, ale výnosy se postupně snižují – někdy je příliš dlouhý prompt kontraproduktivní. Proto volte příklady uvážlivě. **Dobře zvolený příklad** dokáže AI nasměrovat lépe než dlouhé vysvětlování požadavků.

Komponenta *Osobnost* (Persona)

Komponenta „Osobnost“ (persona) v promptu umožňuje specifikovat, z *jaké role nebo perspektivy* má AI odpovídat. Můžeme tím AI „vnutit identitu“ – například aby odpovídala jako odborník, jako konkrétní historická postava, jako laik, jako nadšený fanoušek apod. Tato technika využívá faktu, že model byl trénován na ohromném množství textů různých stylů a rolí, a dokáže se do určité míry stylizovat.

Účel komponenty osobnost: Určuje „hlas“ nebo úhel pohledu odpovědi. Pomáhá to zejména tehdy, chceme-li odpověď přizpůsobit určitému publiku nebo stylu komunikace. Například odpověď „jako zkušený lékař“ bude obsahovat odborné termíny a doporučení založená na medicínských poznatcích, zatímco odpověď „jako tvůj starší bratr“ může být empatická a hovorová. Persona tedy může dramaticky ovlivnit **obsah i tón** odpovědi.

Jak efektivně používat komponentu osobnost: Explicitně řekněte, kým má AI „být“. Např. „*Představ si, že jsi historik specializovaný na starověký Řím...*“, „*Jsi zákaznický podpora e-shopu a odpovídáš rozlobenému klientovi...*“, „*Jako slavný kuchař Jamie Oliver mi porad recept...*“. Model pak začne odpovídat s určitými předpoklady – u historika bude citovat letopočty a kontext, u podpory bude zklidňovat

a nabízet řešení, jako Jamie Oliver možná odlehčí tón a zmíní olivový olej navíc. Podle studie Olea et al. (2024) patří právě *Persona pattern* (nastavení role) k nejužitečnějším přístupům při práci s LLM – uživatelé jej v experimentu hodnotili jako nejvíce nápomocný pro formulování promptů. Tato technika je vnímána jako intuitivní a praktická, protože lidem dává snadno pochopitelný způsob, jak model usměrnit: „odpovídej z této perspektivy“.

Výhody: Persona umožňuje modelu **lépe se vcítit do požadovaného kontextu**. Tím se zvyšuje relevance odpovědi. Například když AI „hovoří jako klimatolog“, zaměří se na vědecké aspekty klimatu a vyvaruje se laických nepřesností. Persona může také zajistit konzistenci odpovědi v delší konverzaci – udržuje model v určité roli, takže odpovědi neodbíhají stylem ani obsahem. Z didaktického hlediska, pokud se učíte nový obor, můžete AI nastavit jako učitele. Z praktického – můžete AI přimět, aby například odpovídala **jazykem cílové skupiny**, jak uvádí T'Syen (2025): definice role a publika společně pomáhají modelu zvolit správný jazykový rejstřík.

Příklady využití komponenty osobnost:

- „*Jsi zkušený právník specializovaný na pracovní právo. Vysvětli mi, jaké mám možnosti v situaci, kdy...*“ – Odpověď bude formulována terminologií a stylem právníka, pravděpodobně velmi věcně a formálně.
- „*Představ si, že jsi Albert Einstein a máš vysvětlit teorii relativity středoškolákům.*“ – AI se pokusí stylizovat vysvětlení tak, jak by to možná udělal sám Einstein: srozumitelně, s příklady (Einstein často používal myšlenkové experimenty), ale zároveň autoritativně.

Doporučení pro uživatele: Buďte kreativní – persona může být reálná profese, fiktivní postava, kdokoli. Klíčem je, že ta role má nějakou spojitost s typem odpovědi, kterou chcete. Pokud odpovědi AI nebyly dříve uspokojivé, zkuste změnit nebo upřesnit personu. Někdy stačí

maličkost: např. „*odpověz jako zkušený programátor*“ donutí AI dávat přímočařejší a techničtější odpovědi, zatímco „*odpověz jako trpělivý učitel programování začátečníků*“ povede k více vysvětlujícímu stylu. Olea et al. (2024) i White et al. (2023) shodně potvrzují, že nastavení role/persony patří mezi nejpřístupnější a prakticky užitečné promptové vzorce. Využijte toho – *specifikace osoby je jedním z nejjednodušších způsobů, jak výrazně ovlivnit vyznění odpovědi AI.*

Komponenta Úroveň expertizy (Level of Expertise)

Komponenta „Úroveň expertizy“ definuje, jakou hloubku a odbornou úroveň má mít odpověď. Tato komponenta úzce souvisí s personou a tónem – vlastně upřesňuje, pro jak znalého čtenáře je odpověď určená, respektive jak detailně má AI odpovídat.

Účel komponenty úroveň expertizy: Zajistit, že odpověď nebude pro uživatele příliš jednoduchá ani příliš složitá. Když AI sdělíme, zda má mluvit k laikovi, pokročilému nebo expertovi, může **přizpůsobit množství detailů, terminologie a vysvětlení**. Například dotaz „*Vysvětli mi, jak funguje kvantový počítač.*“ lze modulovat dle úrovně: „... jako by mi bylo 15 let a moc o tom nevím.“ vs. „... a použij při tom formální jazyk jako do vědeckého článku.“

Jak efektivně používat komponentu úroveň expertizy: Můžete to specifikovat přímo („*vysvětli to začátečníkovi krok za krokem*“, „*uved' podrobný technický rozbor pro odborníka*“) nebo nepřímo kombinací osoby a publika („*Jsi IT konzultant a tvým klientem je malá firma bez zkušeností s IT*“ – model pak bude volit slova pro laika). T'Syen (2025) zdůrazňuje, že **cílení na publikum** je klíčové: určit, jaké znalosti má mít čtenář, pomůže modelu přizpůsobit hloubku i tón.

Výhody: Správné nastavení úrovně zabrání tomu, aby odpověď byla příliš „nad hlavu“ nebo naopak triviální. Uživatel tak dostane informace, kterým porozumí a které může využít. To je zásadní třeba ve vzdělávání – stejný koncept je potřeba jinak vysvětlit žákovi základní školy a jinak vysokoškolákovi. Když AI toto ví, **přizpůsobí**

slovník i hloubku výkladu. Navíc je menší riziko chyb: pokud nutíme AI k odborným detailům, použije je (ale laik by se v nich ztratil); pokud včas řekneme „vysvětlí to jednoduše“, AI raději vypíchne hlavní principy bez zahlcení detaily.

Příklady využití komponenty úroveň expertizy:

- „*Popiš mi fungování fotosyntézy tak, aby tomu rozuměl žák 5. třídy základní školy.*“ – Odpověď bude využívat základní pojmy, možná i analogie („rostliny si vyrábějí potravu ze slunce, jako když...“), vyhne se složitým chemickým rovnicím.
- „*Přednes argumenty pro a proti kryptoměnám na úrovni diskuse expertů v oboru finančnictví.*“ – Model nabídne terminologii jako „decentralizace“, „monetární politika“, odkazy na studie či data, zkrátka přeřadí na expertní úroveň.

Doporučení pro uživatele: Pokud zjistíte, že odpověď je příliš povrchní, požádejte o více detailů a uveďte, že cílíte na odbornější publikum. Naopak je-li odpověď plná nesrozumitelných termínů, připište k dotazu něco jako „(vysvětlí prosím laikovi)“. Vyplatí se také kombinačně použít personu i úroveň: „Jsi zkušený lékař a vysvětluješ studentovi prvního ročníku medicíny...“ – tady persona zaručí odbornost obsahu a úroveň student zajistí, že vysvětlení bude výukové, ne předpokládající atestace. Pamatujte, že vy jste editorem obtížnosti: máte moc říct „přidej/uber na složitosti“. AI pak (zejména v případě nejnovějších modelů jako GPT-5 nebo Gemini 3) umí plynule adaptovat úroveň detailu podle vašeho přání.

Komponenta Očekávaná délka odpovědi (Expected Response Length)

Komponenta „Očekávaná délka odpovědi“ dává AI najevo, jak rozsáhlou odpověď si přejete. Můžete tak předejít situacím, kdy model buď odpoví příliš stručně a povrchně, nebo naopak zabředne

do zbytečných detailů a dodá několik obrazovek textu, když vy jste chtěli jen shrnutí.

Účel komponenty očekávaná délka: Stanovit hranice rozsahu odpovědi – může to být počet vět, odstavců, slov, nebo prostě kvalifikátor jako „stručně / detailně“. Model pak většinou přizpůsobí své vyprávění. Je to důležité i proto, že jazykové modely ne vždy samy odhadnou, jak moc detailů už je příliš. Mají tendenci někdy „povídat dál“, dokud úkol nepokládají za splněný. Definicí délky jim říkáme, kdy mají přestat.

Jak efektivně používat komponentu očekávaná délka: Buďte konkrétní – např. „*maximálně 5 vět*“, „*odpověz v rozsahu cca 3 odstavců*“, „*do 100 slov*“. Případně kvalitativně: „*stručně shrň...*“, „*rozepiš se detailně...*“. T'Syen (2025) popisuje, že pokud neurčíme, co znamená „hotovo“, model může generovat příliš mnoho textu nebo odbíhat. Naopak stanovení limitu (např. délky) modelu **dává jasný signál**, kdy má přejít k závěru.

Výhody: Uživatel dostane odpověď v takové délce, která je pro něj užitečná. Nemusí pak složitě krátit příliš dlouhý text, nebo naopak znovu doptávat detaily, které model původně vynechal. Při práci s AI je efektivita klíčová – např. když spěcháte, oceníte rychlou odpověď v pár větách. Naopak při rešerši možná uvítáte dlouhý souvislý výklad. Díky této komponentě to můžete předem ovlivnit. Také to udržuje model **v mezích tématu** – omezený prostor ho nutí soustředit se na podstatné. Jak uvádí T'Syen (2025), definováním jasného rozsahu (např. „500 slov o X“) dáváme modelu „jasný předmět, formát i délku“, což mu umožní lépe se držet tématu a neodbíhat.

Příklady využití komponenty očekávaná délka:

- „*Stručně (v 5 větách) vysvětli princip blockchainu.*“ – Model vyprodukuje jeden odstavec ~5 vět s klíčovými body.

- „Napiš dvoustránkovou úvahu (cca 4 odstavce) o vlivu umělé inteligence na trh práce.“ – AI zvolí středně rozsáhlý esejistický formát. (Zde je dobré být opatrný – model nezná přesně délku stránky, ale orientuje se odhadem podle slov. Dva listy A4 jsou asi 800–1000 slov, to by se dalo uvést přesněji místo „dvoustránkovou“.)

Doporučení pro uživatele: Při zadání promyslete, jak podrobnou odpověď chcete. Pokud není délka uvedena, model volí něco jako „středně“ podle svého, což nemusí trefit vaši představu. Raději tedy připojte požadavek – ušetříte si pak čas. Jestliže AI reaguje příliš stručně, zkuste „*Můžeš to rozvést podrobněji?*“ nebo příště rovnou v promptu uveďte „detailně“. Pozor ale na extrémy: Napíšete-li „odpověz jednou větou“, dostanete superstručnou repliku, která může být příliš povrchní. Naopak „napiš román o 100 kapitolách“ povede k velmi dlouhému výstupu (a model se možná v průběhu ztratí). **Zlaté pravidlo:** tolik informací, kolik je potřeba – ne více, ne méně. A kdo rozhodne, kolik je potřeba? **Vy** prostřednictvím této komponenty.

Po zvládnutí uvedených hlavních komponent promptů jste vybaveni k tomu, abyste mohli sestavovat vlastní „perfektní prompty“ pro nejrůznější situace. Následující část kapitoly se zaměřuje na konkrétní osvědčené **vzorce promptů** – tedy kombinace uvedených komponent, které se v praxi ukázaly být velmi účinnými při komunikaci s AI. Tyto vzorce vám mohou posloužit jako výchozí kostra pro tvorbu vlastních dotazů.

7.3 Nejdůležitější vzorce promptů

V této sekci rozebíráme a vysvětlujeme klíčové vzorce pro tvorbu efektivních promptů. Každý vzorec je určitou **šablonou**, která kombinuje vhodné komponenty (z předchozího výčtu) a je přizpůsobená specifickému typu dotazu či požadované odpovědi. Tyto vzorce představují praktické návody, jak strukturovat dotazy, tak,

aby AI poskytla co nejlepší možnou odpověď. Zaměříme se na několik příkladů vzorců od analytických přes kreativní až po expertní.

1. Základní analýza

Vzorec: [Dotaz] + [Kontext] + [Očekávaná délka odpovědi]

Tento vzorec je vhodný pro jednoduché dotazy vyžadující vysvětlení či shrnutí. Uživatel položí jasný dotaz, poskytne nezbytný kontext a naznačí, jak rozsáhlou odpověď chce.

Příklad: „Vysvětlete princip fotosyntézy [Dotaz] v kontextu středoškolské biologie [Kontext], ideálně ve dvou odstavcích [Očekávaná délka odpovědi].“ – AI by měla dodat srozumitelné vysvětlení fotosyntézy přiměřené středoškolákovi, zhruba na dvě krátké části.

2. Srovnávací analýza

Vzorec: [Dotaz] + [Dva nebo více objektů ke srovnání] + [Specifická kritéria srovnání]

Tento vzorec se hodí pro dotazy, kde potřebujeme porovnat dvě či více věcí podle určitých hledisek.

Příklad: „Srovnej dopady solární energie a větrné energie na životní prostředí [Dotaz], se zaměřením na uhlíkovou stopu a ekonomickou efektivitu [Specifická kritéria srovnání].“ – Model porovná obnovitelné zdroje (slunce vs. vítr) a strukturovaně rozebere uvedená kritéria.

3. Expertní analýza

Vzorec: [Dotaz] + [Úroveň expertizy] + [Historický nebo odborný kontext]

Tento vzorec využijeme, chceme-li hlubokou, odborně laděnou odpověď. Dotaz formulujeme tak, aby si AI „nasadila klobouk experta“, často v kombinaci s komponentou *Osobnost*.

Příklad: „Poskytněte podrobnou analýzu vývoje kvantové fyziky [Dotaz] na úrovni univerzitního profesora fyziky [Úroveň expertizy] s ohledem na hlavní historické milníky v oboru [Historický kontext].“ – Očekáváme detailní odborný výklad, chronologicky strukturovaný, jako by ho psal profesor se znalostí dějin fyziky.

(Pozn.: Podle White et al. (2023) lze prompty skládat i z více vzorců najednou – například persona + historický kontext + specifický formát výstupu. Taková kombinace odpovídá právě komplexnějším dotazům, jako je expertní analýza.)

4. Komplexní kontextuální analýza

Vzorec: [Obsahový dotaz] + [Specifické informace k zpracování] + [Účel dotazu] + [Formát odpovědi]

Tento vzorec pomáhá u komplexních úloh, kdy je třeba skloubit více požadavků. Uživatel definuje téma, vypíchne konkrétní aspekty či data, sdělí účel (proč to chce) a požadovaný formát odpovědi.

Příklad: „Poskytněte detailní přehled [Obsah] o aktuálních trendech v obnovitelných zdrojích energie [Specifické informace], s cílem informovat studenty univerzitního kurzu o udržitelnosti [Účel], ve formě přehledného článku s podnadpisy [Formát odpovědi].“ – AI by měla napsat souvislý text (článek) členěný do sekcí, pokrývající současné trendy v obnovitelné energetice, přičemž bude pamatovat, že čtenáři jsou vysokoškoláci a článek slouží k jejich edukaci.

5. Multidimenzionální výkladová analýza

Vzorec: [Dotaz] + [Kontext] + [Příklad] + [Osobnost] + [Formát] + [Tón]

Tento vzorec kombinuje více komponent naráz a je vhodný pro náročnější dotazy vyžadující široký a zároveň strukturovaný výklad.

Příklad: „Vysvětlete význam kvantové fyziky pro moderní technologie [Dotaz], zahrnující kontext její aplikace v každodenním životě [Kontext], s příkladem využití v lékařském zobrazování [Příklad], jako byste byli profesorem fyziky na univerzitě [Osobnost], ve formě vzdělávacího videa (scénáře) [Formát], s jasným a motivujícím tónem [Tón].“ – Zde prompt explicitně požaduje velmi konkrétní styl odpovědi: odborné, avšak srozumitelné vysvětlení od profesora, které obsahuje příklad z praxe, celé napsané tak, aby to mohlo sloužit jako scénář k edukativnímu videu (tedy pravděpodobně rozděleno na úvod, příklady, závěr a s motivačním závěrem). AI by měla dodat komplexní, leč čtivou a srozumitelnou odpověď. Tento příklad ilustruje, jak prolnout více komponent – výsledkem je promyšlený prompt pro velmi kvalitní odpověď. Jak bylo zmíněno, kombinování více promptových vzorců je často klíčem k dosažení ideálního výsledku.

6. Kreativní generování

Vzorec: [Kreativní zadání] + [Požadovaný tón] + [Kreativní omezení]

Zde slouží vzorec pro úlohy, kdy chceme po AI generovat kreativní obsah (příběh, báseň, nápad) a zároveň ho usměrnit. Kreativní omezení může být forma (např. sonet, vtip, hororový tón) nebo rozsah (např. „vejde se na dvě stránky“).

Příklad: „Navrhněte krátký sci-fi příběh o cestování časem [Kreativní zadání] s humorným a odlehčeným tónem vyprávění [Tón], který by se vešel na dvě strany textu A4 [Kreativní omezení].“ – AI by měla vytvořit vtipnou sci-fi povídku a pohlídat si délku. Tón i délka jsme

specifikovali, čímž zvyšujeme šanci, že výsledek nebude ani příliš suchý, ani příliš rozvláčný.

7. Problematické řešení (řešení problému)

Vzorec: [Dotaz popisující problém k řešení] + [Kontext nebo scénář] + [Očekávané výstupy či kritéria úspěchu]

Tento vzorec je vhodný pro úlohy typu „navrhni řešení problému X“. Uživatel popíše problém, poskytne kontext (např. zdroje, omezení, situaci) a naznačí, jaké řešení či výsledek očekává (např. „navrhni tři varianty a doporuč nejlepší“).

Příklad: „Najděte řešení pro snížení nákladů v malém rodinném podniku [Problém] v kontextu současné ekonomické krize a zvýšených cen energií [Kontext], a u každého navrženého opatření uveďte odhad, kolik by mohlo ušetřit [Očekávaný výstup].“ – AI zde předloží několik opatření (např. zefektivnění výroby, úspory energií, vyjednání lepších smluv) a ke každému (dle pokynu) uvede i odhad finančního efektu.

8. Inovativní návrhy

Vzorec: [Dotaz na vytvoření návrhu/řešení] + [Specifické požadavky nebo omezení] + [Důraz na kreativitu/inovaci]

Když chceme generovat nové nápady či inovativní řešení, hodí se tento vzorec. Uživatel formuluje, co chce navrhnout, uvede případná omezení (rozpočet, materiály, pravidla) a zdůrazní, že očekává originální či kreativní přístup.

Příklad: „Navrhněte nový typ ekologického obalu pro potraviny [Úkol] s použitím výhradně recyklovaných nebo kompostovatelných materiálů [Požadavek/omezení] a dodejte návrhu originální design, který zaujme zákazníky [Inovativní aspekt].“ – AI by měla přijít s nápadem na obal (třeba krabice nebo fólie), který splňuje ekologické podmínky a přidat kreativní prvek (např. netradiční tvar nebo multifunkční využití obalu). Tónem by možná měla znít jako

produktový designér nebo inovátor (to by šlo doplnit personou, pokud by bylo potřeba).

9. Interdisciplinární přístup

Vzorec: [Dotaz] + [Víceoborový kontext] + [Specifičnost požadavku] + [Požadavek na zahrnutí různých perspektiv]

Tento vzorec použijeme, pokud chceme odpověď, která integruje pohledy z různých oborů či disciplín. V promptu naznačíme, které obory to mají být a případně jak mají být propojeny.

Příklad: „Analyzujte dopady klimatických změn [Dotaz] z pohledu ekonomie a sociologie [Víceoborový kontext], uveďte konkrétní příklady pro každý z těchto pohledů [Specifičnost] a ve závěru shrňte, jak se oba pohledy doplňují nebo liší [Požadavek na vícero perspektiv].“ – AI by měla poskytnout jak ekonomickou analýzu (např. vliv na HDP, náklady na adaptaci), tak sociologickou (např. migrace, nerovnost dopadů na různé populace) a pak diskutovat jejich průnik. Díky explicitnímu vyžádání obou perspektiv máme jistotu, že model neopomene ani jednu.

10. Prediktivní modelování

Vzorec: [Dotaz] + [Budoucí predikce a scénáře] + [Opора v datech či trendech]

Když nás zajímá výhled do budoucnosti nebo hypotetické scénáře, volíme tento vzorec. Požádáme AI, aby na základě dosavadních informací předpověděla vývoj a případně nastínila více variant (optimistický, realistický, pesimistický scénář). Doplníme, že se má opřít o známé trendy nebo data, pokud je má k dispozici – tím zvýšíme důvěryhodnost odpovědi.

Příklad: „Jaké budou pravděpodobné dopady umělé inteligence na trh práce v příštích deseti letech? [Dotaz] Uveďte možné scénáře vývoje (optimistický, pesimistický, realistický) [Budoucí scénáře] a opřete je o

současné trendy a statistiky [Opora v datech].“ – AI by měla poskytnout tři varianty predikcí, každou odůvodnit (např. odkazem na aktuální míru automatizace, historické analogy z průmyslové revoluce apod.).

Výše uvedených deset vzorců představuje jakési *best practices* prompt engineeringu. Samozřejmě existují i další a mnohé lze vzájemně kombinovat. Podstatné je pochopit, že **kvalitní prompt se většinou neskládá jen z jedné věty s otázkou**, ale spíše z promyšlené struktury obsahující více komponent, které společně vedou AI k optimální odpovědi. Jak uvádějí White et al. (2023), prompty se dají dokumentovat jako „*patterny*“ – tj. opakovatelná řešení pro běžné situace při práci s LLM. Seznámili jsme se s několika z nich a je vidět, že pokrývají širokou škálu cílů od čistě faktických analýz přes srovnání, generování nápadů až po tvorbu příběhů.

Závěrem lze říci, že zvládnutí promptových vzorců je v éře narůstajícího významu generativní AI novou klíčovou dovedností. Jak trefně poznamenal T'Syen (2025), *“kvalita vstupu přímo ovlivňuje kvalitu výstupu – není to o tom AI přechytračit, ale dát jí správnou strukturu, kontext a instrukce, aby mohla odvést co nejlepší práci”*. Vzorce promptů jsou právě prostředkem, jak takové strukturované a promyšlené vstupy formulovat. Doufejme, že tato kapitola přispěla ke zvýšení vaší *AI gramotnosti* v oblasti komunikace s modely – a že teď dokážete psát prompty, které vám umělá inteligence oplatí vysoce kvalitními a užitečnými odpověďmi.

8 Správný zápis promptu

V počátcích vývoje počítačů byla komunikace s nimi značně odlišná od dnešních interakcí. Programátoři museli dodržovat přísná pravidla syntaxe – jazyky jako FORTRAN a COBOL vyžadovaly naprostou přesnost v každém příkazu. Každý krok algoritmu musel být explicitně definován a formálně zapsán, což vyžadovalo detailní pochopení fungování počítačových systémů a žádný prostor pro neformální vyjadřování (MKWriteshere, 2025). S příchodem objektově orientovaného programování (OOP), reprezentovaného jazyky jako Java a C++, se zaměření posunulo od striktní syntaxe k flexibilnějšímu a přirozenějšímu přístupu. Programátoři začali definovat objekty a jejich vzájemné interakce, což přineslo větší srozumitelnost – kód bylo možné strukturovat podle reálných objektů a vztahů mezi nimi. Díky tomu se programování stalo intuitivnějším, neboť modelování problému pomocí objektů více odpovídalo lidskému chápání světa (Song, 2024). Syntaxe sice zůstávala důležitá, ale důraz se přesunul na strukturu programu a vztahy mezi jeho částmi.

V současné době, ve věku umělé inteligence a pokročilých jazykových modelů jako jsou Google Gemini 3, Anthropic Claude 4 nebo OpenAI GPT-5, se způsob komunikace s počítači opět významně proměnil. Tyto modely dokážou porozumět vstupům v přirozeném jazyce a reagovat na ně, což činí interakci s nimi mnohem intuitivnější a flexibilnější. Uživatelé už nemusí psát příkazy v přísně formálním jazyce – modely jsou natrénovány tak, aby rozuměly kontextu a záměrům vyjádřeným běžnou řečí (Salami, 2024). Neformálně formulované dotazy tedy nejsou problém; moderní AI dokáže díky pokročilému zpracování jazyka pochopit i složité nebo volně položené otázky a zareagovat odpovídajícím způsobem. Například konverzační model ChatGPT je navržen tak, aby vedl dialog podobný lidské konverzaci – umí odpovídat na doplňující otázky, uznat chyby či odmítnout nevhodné požadavky v rámci jednoho plynulého rozhovoru (OpenAI, 2022).

Efektivní komunikace s těmito modely se proto podobá spíše dialogu s člověkem než psaní tradičního kódu. Uživatelé mohou experimentovat s různými způsoby formulace svých dotazů a AI se snaží interpretovat jejich záměry a kontext. Klíčem k úspěchu je pochopení, jak nejlépe formulovat dotazy, tak, aby odpovědi co nejvíce vyhovovaly potřebám uživatele. Přejít od striktního dodržování syntaxe k flexibilnímu, kontextově založenému přístupu v komunikaci s AI otevírá nové možnosti interakce a získávání informací. Současně ale platí, že i když nejsou vyžadovány striktní formální zápisy, promyšlená formulace dotazu výrazně ovlivňuje kvalitu výsledku. Moderní jazykové modely jsou velmi citlivé na formulaci vstupu – studie ukazují, že i zdánlivě drobné změny nebo dodatky v promptu mohou vést k odlišným výstupům modelu (Salinas & Morstatter, 2024). To zdůrazňuje důležitost pečlivého **prompt engineeringu**, tedy navrhování vstupů pro AI.

V současné době, ve světě generativních jazykových modelů (GPT-5, Claude 4, Gemini 3), jsme tedy přešli k další fázi komunikace člověk–stroj. Tyto modely dokážou porozumět a reagovat na přirozený jazyk, což znamená, že už nejsou závislé na striktní syntaxi jako tradiční programovací jazyky. Přesto pro maximální efektivitu a přesnost odpovědi zůstává důležité dodržovat určité principy při formulaci našich dotazů neboli promptů. Efektivní zápis promptů zahrnuje jasné a strukturované formulace, využívání specifických klíčových slov či frází a poskytnutí dostatečného kontextu. Velmi se osvědčilo také používání tzv. delimitérů – například trojitých uvozovek, oddělovacích komentářů, speciálních symbolů nebo formátování pomocí Markdown/HTML – pro oddělení různých částí promptu a zvýšení jeho čitelnosti jak pro uživatele, tak pro AI (Kotwani, 2025). Takový strukturovaný přístup pomáhá zajistit, že AI správně pochopí záměr dotazu a které části textu jsou zadání a které například vstupní data či příklady. Důsledkem je relevantnější a přesnější odpověď modelu.

Tento promyšlený způsob formulace promptů – zahrnující strukturování, kontext a jasné instrukce – výrazně zlepšuje kvalitu a

použitelnost získaných odpovědí. Jak se tedy vydáme na cestu objevování potenciálu AI v komunikaci a výzkumu, je důležité si uvědomit, že i přes pokročilé schopnosti AI porozumět přirozenému jazyku může strukturovaný a promyšlený přístup k formulaci promptů výrazně zlepšit výsledky. Jinými slovy, **ani pokročilá AI nečte myšlenky** – kvalita její odpovědi závisí na kvalitě našeho dotazu. V následujících částech se proto podíváme na konkrétní metody a techniky, jak těchto cílů dosáhnout při práci s generativními jazykovými modely.

Tři kroky k úspěšné komunikaci

Představujeme ověřenou a vyzkoušenou metodiku komunikace s AI. Tato metodika byla pečlivě vytvořena tak, aby zahrnovala všechny klíčové aspekty – od osvědčených vzorců a strategií až po detailní postupy – nezbytné pro efektivní využívání umělé inteligence v praxi. Naše přístupy a techniky jsou výsledkem důkladného testování a praktických zkušeností, díky čemuž tvoří nepostradatelný nástroj pro každého, kdo se chce naučit, jak nejlépe komunikovat s pokročilými AI systémy.

Věříme, že klíčem k úspěšné interakci s AI je kombinace **jasně definovaného účelu, správného formátování dotazů a flexibilního, iterativního přístupu** k jejich ladění. Tento průvodce vás provede od základních principů až po sofistikované techniky, které vám pomohou dosáhnout maximální efektivity při práci s AI. Díky tomuto komplexnímu přístupu budete schopni nejen správně formulovat dotazy, ale také efektivně interpretovat a využívat odpovědi AI ve svůj prospěch a další rozvoj.

Naším cílem je poskytnout ucelený zdroj informací a návodů. Vaše práce s AI se díky němu stane plynulejší, přesnější a přinese bohatší výsledky. Nezáleží na tom, zda jste začátečník nebo pokročilý uživatel – tyto strategie a techniky se stanou klíčem k vašemu úspěchu v interakci s AI.

1. Příprava a plánování:

- **Porozumění kontextu a stanovení cíle:** Než položíte dotaz AI, je nezbytné mít jasnou představu o kontextu a cíli svého dotazu. Zvažte, co potřebujete zjistit nebo vyřešit, a jaké informace mohou být relevantní. Jasně si definujte problém nebo úkol, na který se chcete AI zeptat. Tato úvodní analýza pomůže zarámovat vaši komunikaci správným směrem (Salami, 2024). Jinými slovy, **identifikujte svou informační potřebu** – co přesně chcete od AI získat a k jakému účelu.
- **Strukturování promptu:** Použijte vhodné vzorce a strategie pro vytvoření promptu, který jasně a stručně komunikuje váš dotaz a očekávání. Rozmyslete si strukturu – například zda poskytnete modelu kontext formou krátkého úvodu, zda položíte více dílčích otázek apod. Užitečné je zahrnout do promptu klíčová slova k tématu a specifikovat požadovaný formát odpovědi. Tím dáte AI vodítka, **na co se má zaměřit**. Studie OpenAI uvádí, že dostatečný kontext v promptu a přesná formulace zadání výrazně zvyšují relevanci odpovědí (OpenAI, 2025).

2. Interakce a ladění:

- **Jasně formulování dotazů:** Když komunikujete s AI, snažte se být co nejpřesnější a nejjasnější ve svých formulacích. Vyvarujte se dvojsmyslů a neurčitosti – položte otázku tak, aby jí model jednoznačně porozuměl. Používejte spíše jednoduché věty a přímé instrukce. Pokud například potřebujete souhrn, specifikujte, že má jít o souhrn a v jakém rozsahu či formátu. Takové zpřesnění zadání vede ke konkrétnějším a kvalitnějším odpovědím (Salami, 2024). Modely jako GPT-5 nebo Gemini 3 dokáží do určité míry domýšlet, co uživatel chce, ale spoléhat na to nelze – čím jednoznačnější pokyn, tím konzistentnější výstup (Lakera, 2025).

- **Iterativní přístup:** Buďte připraveni dotaz upravit nebo doplnit, pokud prvotní odpověď AI není zcela vyhovující. Prompt engineering je často **iterativní proces** – začněte s počátečním promptem, zkontrolujte odpověď a na základě výsledku prompt upřesněte (OpenAI, 2025). Kladením doplňujících otázek můžete model navést blíže k jádru problému nebo požádat o rozvedení detailů. Tento postup postupného ladění a zpřesňování dotazu vám umožní získat přesnější a užitečnější odpověď. Například pokud AI odpoví příliš obecně, můžete položit navazující otázku upřesňující konkrétní aspekt, který vás zajímá.

3. Analýza a reflexe:

- **Hodnocení odpovědí:** Po obdržení odpovědi od AI pečlivě posuďte, zda skutečně odpovídá vašim potřebám, zda je informace relevantní a věrohodná. Pokud odpověď obsahuje faktická tvrzení, ověřte si je z jiných zdrojů. *Critical thinking* je na místě – AI může někdy tzv. halucinovat, tedy s jistotou tvrdit nepravdivé či nepřesné informace. Proto je vhodné důležité údaje křížově zkontrolovat. Například Google Gemini doporučuje vždy **kriticky vyhodnotit získané odpovědi a klíčová fakta ověřit z důvěryhodných zdrojů** pro zajištění správnosti (Salami, 2024).
- **Učení se a zlepšování:** Každá interakce s AI je příležitostí k poučení. Sledujte, co fungovalo a co ne – jak formulace dotazu ovlivnila kvalitu odpovědi. Pokud určité slovo či přístup vedl k nedorozumění, příště se mu vyhněte. Postupným zkoušením různých strategií lépe pochopíte, jak model reaguje, a můžete zdokonalovat své prompty. Tento proces neustálého učení se a vylepšování je klíčový pro efektivní využívání AI. Jak poznamenává Timler (2023), neexistuje dokonalý prompt na první pokus – dosažení optimálního výsledku často vyžaduje několik iterací a úprav promptu na základě předchozích odpovědí. Přistupujte tedy k práci s AI jako k dialogu a postupnému ladění než jako k jednorázovému příkazu.

8.1 Klíčové prvky efektivního promptu

V předchozích kapitolách jsme se zaměřili na různé strategie, vzorce a techniky pro optimalizaci komunikace s AI prostřednictvím promptů. Nyní přichází na řadu zásadní aspekt celého procesu: pochopení a využití klíčových prvků, které tvoří efektivní prompt. Tento poslední krok na naší cestě se orientuje na *jemné ladění* promptů, což je rozhodující pro maximální využití potenciálu generativních jazykových modelů.

Proč jsou klíčové prvky tak důležité? Každý prompt je více než jen souhrn slov a frází. Je to pečlivě strukturovaná výzva, která řídí AI k poskytnutí přesných, relevantních a užitečných odpovědí. I drobné nuance v zadání mohou ovlivnit, jak AI váš požadavek interpretuje a jak kvalitní bude výsledek (Salinas & Morstatter, 2024). Tato kapitola se proto zaměřuje na finální doladění vašich dovedností v prompt engineeringu, aby vaše interakce s AI byly co nejúčinnější.

Přestože jsme již probrali různé strategie a vzorce, uvědomění si klíčových prvků promptu nám umožňuje ještě lépe pochopit a kontrolovat, jak AI interpretuje naše požadavky. Získáte nástroje pro poslední úpravy vašich promptů, porozumíte významu jasného stanovení cíle dotazu a uvidíte, jak i malé změny ve formulaci mohou mít velký dopad na kvalitu a přesnost odpovědí AI. Jde o esenciální krok, který promění dobré prompty ve vynikající a umožní vám využít plný potenciál AI pro vaše specifické potřeby a cíle.

Tato závěrečná fáze naší cesty prompt engineeringem vás připraví na samostatnou, kreativní a úspěšnou práci s generativními jazykovými modely. Vstupte do světa AI s jistotou, že máte v ruce všechny potřebné nástroje pro efektivní dialog s touto fascinující technologií.

Cíl

Komponenta „Cíl“ je klíčovým prvkem efektivního promptu. Určení jasného cíle dotazu je nezbytné pro nasměrování AI k poskytnutí

relevantní a užitečné odpovědi. Bez jasné definovaného cíle může být komunikace s AI méně účinná, protože model nemusí správně pochopit, co je od něj vyžadováno nebo jaké informace má poskytnout. Stanovení cíle tedy představuje **startovní bod** celého procesu. Podle odborníků na prompt engineering je formulace konkrétního cíle vůbec prvním a nejdůležitějším krokem – cíl funguje jako základní kámen, který určuje směr celého promptu a následné odpovědi (Chip, 2021).

Jak se dostat k cíli

Nestačí však pouze určit cíl – je také třeba promyslet, *jak formulovat dotaz, aby vedl k dosažení tohoto cíle*. To zahrnuje úvahy o struktuře promptu, volbě správných slov, poskytnutí kontextu a dalších faktorech, které pomohou AI správně interpretovat váš požadavek a poskytnout co nejrelevantnější odpověď. Uživatelé, kteří mají pokročilé znalosti **copywritingu**, zde mohou uplatnit své dovednosti – principy copywritingu lze totiž výborně využít při konstruování promptů pro AI.

Techniky z copywritingu, zaměřené na stručnost, jasnost a přesvědčivost sdělení, mohou výrazně zlepšit efektivitu a přesnost odpovědí od AI. Zde je několik způsobů, jak integrovat osvědčené postupy copywritingu do tvorby promptů:

1. **Stručnost a jasnost:** Podobně jako v copywritingu, kde je cílem sdělit myšlenku co nejstručněji a nejjasněji, i u promptů pro AI je důležité být přímý a konkrétní. Vyhněte se obecným a vágním formulacím; místo nich specifikujte přesně, co od AI očekáváte.

Méně efektivní: „Můžete mi něco říct o historii Říma?“

Efektivnější: „Shrňte hlavní události římské historie od založení města až po pád Západořímské říše.“

2. **Zaměření na cíl:** Stejně jako copywriter musí mít jasnou představu o cíli textu, i při formulaci promptu je klíčové mít

jednoznačně definovaný cíl. To udržuje prompt soustředěný a zvyšuje pravděpodobnost získání relevantní odpovědi.

Méně efektivní: „Jak můžu zhubnout?“

Efektivnější: „Navrhněte čtyřtýdenní jídelníček pro úbytek 5 kg, zaměřený na nízkosacharidovou stravu.“

3. **Výzva k akci:** V copywritingu se často používá **call-to-action** pro větší naléhavost textu. Podobně i v promptu může být užitečné formulovat požadavek aktivně v rozkazovacím způsobu, aby bylo zcela jasné, co má AI udělat.

Méně efektivní: „Bylo by fajn získat recept na francouzský toast.“

Efektivnější: „Vysvětli krok za krokem, jak připravit francouzský toast.“

4. **Začlenění klíčových slov:** V copywritingu pomáhají klíčová slova upoutat pozornost a zlepšit SEO. V promptu mohou klíčová slova pomoci modelu lépe pochopit téma a zaměření dotazu.

Méně efektivní: „Napiš něco o marketingových strategiích.“

Efektivnější: „Popiš nejúčinnější digitální marketingové strategie pro malé podniky v roce 2023.“

5. **Jasně vyjádření očekávání:** V copywritingu je často na konci výzva k akci (CTA), která čtenáři říká, co se od něj očekává. Podobně by v promptu mělo být zřetelně vyjádřeno, jaký výstup či akci od AI očekáváte.

Méně efektivní: „Jak si mohu zlepšit angličtinu?“

Efektivnější: „Uved' pět nejlepších online zdrojů pro zlepšení anglické slovní zásoby.“

6. **Přehlednost a struktura:** Stejně jako copywriteři strukturují obsah pro lepší srozumitelnost, je důležité strukturovat i prompt tak, aby byl pro AI snadno zpracovatelný.

Méně efektivní: „Co bych měl vědět o kvantové fyzice?“

Efektivnější: „Vytvoř přehledné shrnutí základních principů kvantové fyziky a uveď tři klíčové experimenty, které tyto principy demonstrují.“

Integrací těchto copywriterských postupů do promptů můžete zvýšit šanci, že AI poskytne přesné, relevantní a užitečné odpovědi. V kombinaci s ostatními komponentami (jako je kontext či iterativní ladění) tvoří **dobře definovaný cíl** pevný základ pro cílenou komunikaci s AI.

Analýza a definování cíle

Analýza a definování cíle je základním krokem pro efektivní komunikaci s AI. Tento proces zahrnuje několik klíčových aspektů:

1. **Pochopte a formulujte své potřeby:** Nejprve si ujasněte, jakou máte motivaci položit AI daný dotaz. Pokud hledáte informace, definujte, jaký typ informací a k jakému účelu potřebujete. Pokud chcete, aby AI vygenerovala určitý obsah, mějte co nejkonkrétnější představu o požadovaném výsledku (téma, rozsah, styl apod.). Srozumitelně popište svůj požadavek – buďte specifičtí. Místo vágních žádostí poskytněte AI konkrétní detaily, které jí pomohou lépe pochopit vaše očekávání.
2. **Stanovte konkrétní cíl promptu:** Určete hlavní téma nebo otázku, na kterou chcete odpověď. Tento krok pomůže udržet prompt zaměřený a relevantní. Zvažte zahrnutí klíčových slov souvisejících s vaším tématem – pomohou modelu pochopit kontext a zaměření dotazu. Například místo obecného „řekni mi něco o klimatu“ je vhodnější „Jaké jsou hlavní dopady klimatických změn na ekosystémy Arktidy?“. V druhém případě

jsou jasně definovány téma (dopady klimatických změn), oblast (ekosystémy Arktidy) i forma odpovědi (přehled hlavních dopadů).

3. **Vytvářejte strukturované otázky a hypotézy:** Pro hlubší analýzu problému se vyplatí rozdělit hlavní cíl na dílčí otázky nebo formulovat hypotézy, které chcete ověřit. To vede k více strukturovanému zkoumání tématu a pomůže AI poskytnout systematictější odpověď. Pokud například zkoumáte příčiny určitého jevu, můžete prompt rozdělit na: „*Vysvětli možné příčiny X. Která z těchto příčin je podle studií nejvýznamnější a proč?*“ Takový prompt nutí AI uvažovat v krocích (vyjmenovat příčiny a pak určit hlavní), což může vést k přehlednější odpovědi.

Příklad: Představme si, že jste výzkumník zkoumající dopady klimatických změn na biodiverzitu v tropických deštných lesích a chcete získat přehled o současném stavu poznání v této oblasti.

- **Nejasně definovaný prompt:** „*Napište mi něco o klimatických změnách.*“ – zde není jasné, jaký aspekt klimatických změn vás zajímá, chybí cíl.
- **Jasně definovaný prompt:** „*Shrňte výsledky vědeckých studií z let 2020–2023 o dopadech klimatických změn na biodiverzitu v tropických deštných pralesích.*“ – zde jste explicitně uvedli **co** chcete (souhrn výsledků studií), **téma** (dopady klimatických změn na biodiverzitu), **kontext/omezení** (tropické deštné pralesy) a **časové rozmezí** (2020–2023). AI má tedy výrazně lepší vodítka, aby poskytla přesnou a relevantní odpověď.

V tomto efektivním promptu je cíl jasně specifikován, což AI umožňuje zaměřit se na požadovanou oblast a podat uživateli hodnotnou informaci. **Analýza a precizní definování cíle** tak výrazně zvyšuje šanci, že výsledek interakce s AI splní vaše očekávání.

Dosahování cílů

Rozvoj přístupů k dosahování cílů lze vnímat jako aplikaci principů **Total Quality Management (TQM)** v kontextu využívání AI. TQM je manažerský přístup zaměřený na dlouhodobý úspěch prostřednictvím spokojenosti zákazníka a neustálého zlepšování všech procesů v organizaci. Podobnou logiku můžeme uplatnit i při práci s AI na plnění našich cílů:

1. **Nepřetržitě zlepšování:** Podobně jako u TQM, i při využívání AI je vhodné zavést cyklus neustálého zlepšování. Každá interakce s modelem poskytuje zpětnou vazbu – co fungovalo, co je třeba upravit. Tento cyklický proces umožňuje postupné doladování promptů i strategie práce s AI. Inspirací může být Demingův cyklus PDCA (Plan-Do-Check-Act) známý z managementu kvality: *naplánuj* (promysli cíl a sestav prompt), *proved'* (získej odpověď), *ověř* (zkontroluj a vyhodnoť odpověď) a *jednej* (uprav prompt či přístup na základě zjištění a cyklus opakuji). **Cyklické zlepšování** tak vede k stále kvalitnějším výsledkům (Investopedia, 2025).
2. **Orientace na “zákazníka”:** V kontextu AI je “zákazníkem” váš vlastní informační požadavek či problém, který řešíte. Aplikujeme-li TQM, měli bychom se při tvorbě promptu soustředit na to, aby výsledná odpověď co nejlépe naplnila naše potřeby. To znamená neustále mít na paměti cíl (váš “zákazník”) a podle něj *řídít komunikaci s AI*. Např. pokud je cílem získat konkrétní datové statistiky, musí tomu být uzpůsobena formulace promptu (stručnost, žádost o čísla apod.). **Pochopení vlastních potřeb** a jejich reflektování v promptu je analogií zákaznické orientace v TQM – veškeré aktivity směřujeme k naplnění definovaného cíle.
3. **Důraz na kvalitu a přesnost:** TQM klade důraz na kvalitu výstupů – v našem případě kvalitu odpovědí AI. Kvality dosáhneme kombinací výše zmíněných principů: jasným cílem, poskytnutím kontextu, správným formátováním promptu a iterativním

laděním. Je užitečné nastavit si *kritéria kvality* odpovědi (podobně jako standardy kvality v průmyslu) – například požadavek, aby odpověď obsahovala konkrétní údaje, byla fakticky správná, strukturovaná či srozumitelná pro dané publikum. Každou odpověď pak hodnotíme podle těchto kritérií a případně upravujeme prompt, dokud výsledek nevyhovuje. Takový **důraz na kvalitu** a systematické odstraňování nedostatků vede ke spolehlivějším výstupům (Investopedia, 2025).

Příklad z praxe: Pracujete na projektu v oblasti digitálního marketingu a potřebujete od AI návrhy na optimalizaci online kampaně.

- Aplikujete princip neustálého zlepšování: iterativně testujete různé formulace promptu a sledujete, která přinese nejprínosnější návrhy (např. začnete obecným dotazem, pak jej zužujete na konkrétní kanály, rozpočty atd.).
- Držíte se orientace na “zákazníka” – vaším cílem je zvýšit míru prokliku u reklamy, takže veškeré dotazy směřujete k této metrice (např. „Jak zvýšit míru prokliku PPC kampaně v segmentu X o Y%?“).
- Kladete důraz na kvalitu a přesnost – získané návrhy ověřujete (třeba srovnáním s publikovanými case studies nebo expertními doporučeními) a požádáte AI o upřesnění či doplnění, pokud jsou informace vágní.

Tímto způsobem, analogickým k TQM, dosáhnete při využívání AI vyšší efektivity, lepších výsledků a vytvoříte skutečnou hodnotu – ať už pro klienta, firmu, či svůj vlastní projekt.

Zahrnutí zpětné vazby a iterací

Zahrnutí zpětné vazby a iterací je zásadní pro úspěšnou komunikaci s AI. Tento proces umožňuje nejen reagovat na výsledky a upravovat přístupy, ale také se **učit a rozvíjet efektivnější strategie** pro dosažení požadovaných cílů. Můžeme jej přirovnat k neustálému dialogu, v

němž výstup AI slouží jako zpětná vazba pro uživatele – a ten na ni reaguje upřesněním požadavku.

Nejprve je důležité správně chápat *význam zpětné vazby*. Odpověď od AI není jen výsledkem našeho dotazu, ale také cennou informací o tom, jak AI náš prompt pochopila. Pečlivou analýzou odpovědi můžeme odhalit, zda model zadání porozuměl tak, jak jsme zamýšleli. Pokud ne, je to signál k úpravě promptu. Například pokud AI odpoví příliš obecně, může to znamenat, že dotaz byl vágní – poučení pro příště zní: **zpřesnit formulaci**.

Následuje proces iterace a úprav promptů. Na základě získané zpětné vazby prompt iterativně vylepšujeme. To může zahrnovat přeformulování otázek pro větší jasnost, přidání nebo upřesnění kontextu, nebo zcela změnu přístupu k dotazu. Každá taková iterace je krokem k lepšímu pochopení toho, **jak faktory promptu ovlivňují výstup AI**. Jak uvádí OpenAI, iterativní přístup – tedy postupné pilování promptu na základě předchozích odpovědí – je často nutný k dosažení optimálních výsledků (OpenAI, 2025).

Důležitou součástí je také *učení se z chyb*. Pokud AI poskytne odpověď, která je nerelevantní nebo neodpovídá očekávání, není to selhání, ale příležitost k učení. Analyzujte, proč model reagoval daným způsobem – bylo něco v promptu nejednoznačné? Poskytli jste nedostatečný kontext? Ponechali jste příliš volnou ruku modelu tam, kde měla být upřesnění? Z takové analýzy vyvodíte, **co lze příště udělat lépe**.

Celý proces zahrnuje *opakování a zkoušení nových přístupů*. Nemusíte se omezovat na jeden styl nebo metodu tázání – experimentujte. Možná zjistíte, že jiná formulace či jiný tón promptu přinese lepší výsledek. Například místo otázky „Proč X?“ zkuste „Vysvětli důvody X“ – drobná změna, která může model navést k ucelenější odpovědi.

Důraz je také kladen na *sledování vývoje a pokroku*. V průběhu času pozorujte, jak se vaše schopnost efektivně komunikovat s AI zlepšuje. Můžete si dokonce vést jednoduchý záznam: jaké změny promptu vedly ke zlepšení odpovědí? Tyto poznatky pak aplikujte v budoucnu. **Každá iterace by měla vycházet z informovaného rozhodnutí**, co zkusit jinak, na základě předchozích zkušeností.

Ukázka iterativní smyčky mezi uživatelem a AI: Uživatel poskytne jasný kontext a cíl (Context), formuluje prompt (Prompt), AI vygeneruje odpověď (AI Response), uživatel odpověď vyhodnotí a poskytne upřesňující zpětnou vazbu či doplňující dotaz (Human Feedback) a následně prompt upřesní či zopakuje s vylepšeními (Refined Prompt). Tento cyklus se opakuje, čímž dochází k postupnému upřesňování a zkvalitňování výstupu.

Například: Představte si, že AI poskytne příliš obecnou nebo neúplnou odpověď na váš dotaz ohledně nejlepších metod studia. Místo toho, abyste výsledek přijali, rozhodnete se prompt **iterativně** vylepšit. Mohli byste původní vágní dotaz „*Jak se efektivně učit?*“ upravit na specifičtější: „*Uved' tři nejefektivnější techniky studia pro vysokoškolské studenty připravující se na zkoušku z biologie.*“ – Tento nový prompt je mnohem konkrétnější (počet technik, cílová skupina studentů, účel – příprava na zkoušku z biologie) a dává AI jasnější instrukce. Výsledkem bude pravděpodobně **relevantnější a použitelnější odpověď**.

Klíčem je zůstat v této *smyčce dotaz–odpověď–vylepšení* dostatečně dlouho, dokud nejste s odpovědí spokojeni. Tím se maximálně využije potenciál AI a zároveň se neustále zdokonalují vaše dovednosti v pokládání dotazů.

Použití metrik hodnocení úspěšnosti

Abychom mohli objektivně posoudit, jak dobře naše komunikace s AI funguje, je užitečné definovat a používat **metriky hodnocení úspěšnosti**. Tyto metriky nám poskytnou kvantitativní či kvalitativní

měřítka pro vyhodnocení efektivity promptů a kvality odpovědí AI. Zavedou do procesu prvek měřitelnosti, který pomáhá identifikovat oblasti ke zlepšení a sledovat pokrok v čase.

1. **Stanovení metrik:** Nejprve si ujasněte, co přesně chcete měřit. Metriky by měly vycházet z vašich cílů. Může jít například o *přesnost odpovědí* (obsahují správná fakta?), *relevanci informací* (odpovídají přesně položené otázce?), *úplnost* (pokryly všechny části dotazu?), *stručnost/výstižnost*, nebo *spokojenost uživatele s odpovědí*. V jiných případech mohou být metrikami i kvantitativní ukazatele – například počet iterací potřebných k dosažení uspokojivé odpovědi, čas potřebný k získání použitelné odpovědi apod. Důležité je vybrat metriky, které odrážejí vaše priority (co je pro vás u odpovědi nejdůležitější).
2. **Sběr a záznam dat:** Při interakci s AI sledujte a zaznamenávejte výsledky vzhledem k těmto metrikám. Pokud je cílem přesnost, ověřujte fakta z odpovědi a poznamenejte si případné chyby. Pokud je cílem stručnost, můžete měřit délku odpovědi (počet slov) a hodnotit, zda je přiměřená. U spokojenosti lze vést jednoduchý deník, kde ohodnotíte každou odpověď na škále (např. 1–5). Důležité je mít nějaké **konzistentní záznamy**, k nimž se dá vrátit.
3. **Analýza a vyhodnocení:** Po nashromáždění dostatečného množství dat proveďte jejich analýzu. Hledejte vzorce a trendy – například zda se přesnost odpovědí zlepšuje s tím, jak ladíte své prompty, nebo zda určitý styl dotazu vede opakovaně k nerelevantním odpovědím. Tato analýza vám pomůže pochopit, **jak dobře vaše techniky promptování fungují** a jestli dosahujete stanovených cílů. Pokud například vidíte, že i po několika iteracích odpovědi stále obsahují faktické chyby, může to znamenat, že je třeba lépe formulovat požadavek na zdroje či ověření informací.

4. **Iterace a zlepšení na základě metrik:** Výsledky hodnocení pak využijte ke změnám a úpravám ve vašich promptech. Pokud metriky ukazují slabá místa – například nízkou relevanci – zaměřte se právě na ně (v tomto případě by pomohlo přesněji vymezit otázku). Tento proces připomíná vědecký experiment: podle definovaných měřítek zjistíte, co nefunguje optimálně, a provedete **informované úpravy**.
5. **Kontinuální hodnocení:** Hodnocení úspěšnosti by nemělo být jednorázové, ale průběžné. Pravidelně se vraťte ke svým metrikám a sledujte, jak se vaše výsledky mění v čase. To vám umožní vidět dlouhodobý progres. Například můžete zjistit, že zatímco na začátku zabralo třeba 5 iterací k dosažení spokojenosti s odpovědí, po měsíci praxe vám stačí 2 iterace – jasný důkaz zlepšení. Průběžné sledování vám tak dodá motivaci a pomůže **udržovat vysokou úroveň kvality** vašich interakcí s AI.

Příklad aplikace metrik: Řekněme, že využíváte AI k generování marketingových textů. Stanovíte si metriky úspěchu jako *míru prokliku (CTR)* generovaných textů, *zapojení čtenářů* (například průměrný čas strávený na stránce s textem), a *konverzní poměr* (kolik čtenářů provedlo požadovanou akci). Po každé AI generované variantě textu nasadíte text na web a změříte tyto hodnoty. Pokud zjistíte, že některá verze textu má výrazně vyšší CTR, zanalyzujete prompt, který k této verzi vedl – možná jste v něm dali AI pokyn k využití výzvy k akci či emotivnějšího jazyka. Tyto poznatky pak zapracujete do budoucích promptů. Metriky vám v tomto případě pomohou **objektivně určit**, která podoba promptu a tím i výsledného textu je nejučinnější.

8.2 Aktivační slovesa

Akční slovesa (nebo také *aktivační slovesa*) jsou slovesa vyjadřující určitou akci nebo výzvu k akci – v kontextu komunikace s AI představují explicitní instrukci, co má model udělat. V prompt engineeringu hrají akční slovesa klíčovou roli, protože dávají modelu jasný a konkrétní pokyn, jaký typ úlohy nebo odpovědi se od něj očekává. Jsou to stavební kameny formulace našich příkazů a dotazů. Zejména u složitějších dotazů, kde je potřeba přesně specifikovat požadovanou akci nebo styl odpovědi, jsou aktivační slovesa **nezbytná pro upřesnění záměru**.

Jednoduché dotazy (např. „*Jak funguje fotosyntéza?*“) jsou přímočaré a obvykle nevyžadují další specifikaci – AI rozumí, že má podat základní vysvětlení. Pokud ale potřebujeme odpověď v určitém kontextu nebo formě, použití vhodných sloves pomáhá takové požadavky vyjádřit. Například rozdíl mezi „*Popiš fotosyntézu*“ a „*Vysvětli fotosyntézu studentovi základní školy*“ je markantní – druhý prompt zahrnuje aktivační sloveso „*vysvětli*“ a navíc specifikuje publikum, čímž **zužuje rozsah a styl možné odpovědi**.

Použití aktivačních sloves v komplexních dotazech tedy nejen zpřesňuje požadovanou odpověď, ale také **omezuje prostor interpretace** pro AI. Dává modelu jednoznačně najevo, jakou činnost má vykonat (analýzu, srovnání, generování nápadů apod.). Tím výrazně zvyšujeme pravděpodobnost, že výstup bude odpovídat našemu záměru. Například požadavek „*Analyzuj následující text z pohledu gramatických chyb*.“ je mnohem účinnější než pouhé „*Podívej se na tento text*.“ – v prvním případě model přesně ví, jakou akci má provést (provést analýzu) a jaké hledisko ho zajímá (gramatika).

Stručně řečeno, **akční slovesa usnadňují komunikaci** mezi člověkem a AI, protože převádějí náš záměr do konkrétního příkazu. Uživatel tím přebírá aktivnější roli v řízení konverzace. V následující části se podíváme na vytvoření a popis skupin aktivačních sloves. Uspořádání

těchto sloves do kategorií nám pomůže lépe porozumět, jaké druhy akcí lze po AI požadovat a jak je formulovat v různých kontextech, abychom dosáhli efektivních a přesných odpovědí.

Dělení akčních sloves

Abychom si udělali pořádek v množství možných akčních sloves, rozdělme si je do několika kategorií podle typu úkolu, který jimi zadáváme. Níže uvedené skupiny pokrývají široké spektrum činností od analýzy přes kreativitu až po plánování. U každé skupiny jsou uvedena příslušná aktivační slovesa a jejich typické využití:

1. Analýza a vyhodnocení:

Analyzuj – proved' podrobnou analýzu nebo rozbor tématu či dat.

Porovnej – vyzdvihni podobnosti a rozdíly mezi dvěma či více prvky.

Popiš – dej detailní a objektivní popis situace, objektu nebo konceptu.

Shrň (sumarizuj) – poskytni stručné shrnutí rozsáhlejšího obsahu nebo dat.

Identifikuj – určete nebo pojmenuj klíčové prvky, vzory či charakteristiky.

Odhal – zjisti a vypiš skryté vzory, fakta nebo informace (např. odhal slabiny argumentu).

Vyhodnoť – ohodnoť či posuď určitou situaci, produkt, argument apod.

Definuj – formuluj přesnou definici či vysvětlení termínu nebo konceptu.

Zachyť – postihni a vypíchni klíčové body nebo momenty v daném kontextu.

Urči – stanov konkrétní informace nebo charakteristiky (např. „urči hlavní myšlenku textu“).

2. Kreativita a generování nápadů:

Generuj – vytvoř nové nápady, návrhy nebo možnosti (např. „generuj nápady na podnikatelský záměr“).

Navrhni – dej kreativní návrhy, plány či koncepty řešení.

Vytuš – zkus intuitivně odhadnout nebo navrhnout něco na základě neúplných informací (užitečné pro brainstormink).

Zintenzivni – zesil účinek nebo dopad něčeho (např. „zintenzivni emocionální tón tohoto odstavce“).

Transformuj – převed' nebo změň informaci či koncept do jiné podoby (např. „transformuj tento odstavec do básně“).

Zdokonal – navrhni úpravy vedoucí ke zlepšení či vyšší kvalitě stávajícího textu, nápadu apod.

Vylepši – podobné jako zdokonal; poskytne konkrétní návrhy pro zlepšení procesu, produktu nebo situace.

3. Plánování a strategie:

Formuluj – sestav nebo navrhni například koncept, strategii či souhrn (např. „formuluj strategii vstupu firmy na trh“).

Představ si – nastol nový koncept nebo perspektivu („co kdyby“ scénáře – „představ si, že... a popiš, co by se stalo“).

Plánuj – vytvoř detailní plán nebo postup k dosažení cíle. Výsledkem může být strukturovaný seznam kroků či harmonogram.

Optimalizuj – navrhni vylepšení existujícího postupu či procesu tak, aby byl efektivnější (např. „*optimalizuj výrobní proces ke snížení nákladů*“).

Vypořádej se s – navrhni strategie nebo metody, jak se vypořádat s konkrétním problémem nebo výzvou („*jak se vypořádat s poklesem návštěvnosti webu?*“).

Propaguj – vytvoř návrhy, jak propagovat produkt, myšlenku či událost (výstupem může být třeba reklamní slogan nebo marketingový plán).

4. **Komunikace a prezentace:**

Demonstruj – ukaž na příkladech, jak něco funguje nebo jak něco udělat.

Diskutuj – prozkoumej hlouběji dané téma v širších souvislostech, případně uveď různé úhly pohledu.

Prezentuj – uspořádej informace do souvislé prezentace nebo výkladu, často strukturovaně a přehledně.

Vysvětli – podrobně vylož koncept, proces nebo názor, často tak, aby tomu porozuměl i neodborník (např. „*vysvětli kvantovou fyziku desetiletému dítěti*“).

Přelož – převed' text z jednoho jazyka do druhého.

Reflektuj – proved' zamyšlení či reflexi nad určitým tématem nebo informací (např. „*reflektuj význam tohoto objevu pro budoucnost medicíny*“).

5. **Adaptace a modifikace:**

Přizpůsob – uprav obsah nebo odpověď pro specifické potřeby či parametry (např. „*přizpůsob odborný článek laické veřejnosti*“).

Reorganizuj – předělej nebo znovu uspořádej existující informace do nového formátu či struktury (např. převod textu do bodového seznamu).

Transformuj – převed' existující data či informace do nového formátu nebo kontextu (např. „transformuj tyto prodejní výsledky do srozumitelné infografiky – slovně je popiš“).

Přepracuj – přeformuluj nebo uprav stávající obsah, aby lépe vyhovoval novým požadavkům (např. zkrat' text, změň tón z formálního na neformální apod.).

6. **Predikce a prognóza:**

Predikuj – poskytní předpovědi či prognózy na základě dostupných dat nebo trendů (např. „predikuj vývoj trhu s elektřinou v příštích 5 letech“).

Předvídej – podobné jako predikuj; odhadni budoucí vývoj určitého jevu či událostí.

Simuluj – nasimuluj nějaký proces nebo situaci, případně odhadni výsledek hypotetického scénáře (např. „simuluj vývoj populací predátora a kořisti při dvojnásobné úrodě“).

7. **Interakce a reakce:**

Reaguj – poskytní reakci nebo odpověď na určitý podnět nebo informaci (užitečné pro simulace konverzace nebo testování reakce AI na vstup).

Porad' – nabídni doporučení nebo řešení pro konkrétní problém či situaci.

Odpověz – předej jasnou a přímou odpověď na položený dotaz či podnět (vhodné pro Q&A styl promptů).

Rozved' – rozšiř či doplň předchozí odpověď o další detaily (např. „rozved' podrobněji důvod X uvedený v odpovědi“).

8. Výzkum a shromažďování dat:

Vyhledej – najdi a prezentuj specifické informace, data nebo zdroje (např. „vyhledej aktuální kurzy kryptoměn a uveď je“).

Srovnej – proved' detailní srovnání mezi různými objekty, koncepty nebo daty (např. „srovnej vlastnosti vodíku a hélia“).

Pozoruj – poskytni popis či analýzu pozorování nějakého jevu (užitečné pro vědecké či technické úlohy, třeba „pozoruj trend v grafech a popiš ho“).

Ověř – zkontroluj a potvrď pravdivost či platnost nějaké informace (např. „ověř, zda platí následující tvrzení podle dostupných dat“).

Spočítej – proved' výpočet nebo kvantitativní analýzu (pokud to model zvládne, např. „spočítej 15. číslo Fibonacciho posloupnosti“ nebo „odhadni průměrnou roční útratu na základě těchto měsíčních dat“).

9. Vizuální reprezentace a mapování:

Nakresli – (v případě modelů schopných generovat obrázky nebo grafiku) vytvoř vizuální reprezentaci, schéma či diagram. U textových modelů lze použít metaforicky, např. „nakresli myšlenkovou mapu jako strukturovaný seznam“.

Zmapuj – vytvoř mapu informací nebo konceptů, tj. identifikuj a strukturovaně znázorni vztahy mezi prvky (např. „zmapuj hlavní aktéry a události Francouzské revoluce“).

Propoj – najdi a popiš propojení mezi různými myšlenkami nebo informacemi, případně je vizuálně/konceptuálně propoj (např. „propoj koncept kvantové provázanosti s reálnými příklady“).

Toto dělení není vyčerpávající, ale pokrývá nejčastější typy instrukcí, které uživatelé jazykovým modelům dávají. Při formulaci promptu je dobré začít vhodným akčním slovesem – **nastaví tón a směr odpovědi**. Odborné zdroje doporučují formulovat prompt co nejvíce jako přímý příkaz či dotaz, protože tak model nejlépe pochopí, jakou operaci má provést (Lakera, 2025).

Kdy akční slovesa nestačí

V předchozím textu jsme se seznámili s konceptem aktivačních sloves a jejich zásadní rolí v interakci s AI. Akční slovesa slouží jako důležité nástroje pro formulaci přesných a efektivních dotazů, ale jejich skutečná síla se ukáže až v konkrétním **kontextu použití**. Různé obory a situace mohou totiž vyžadovat odlišný přístup, i když použijeme stejné akční sloveso. Abychom plně pochopili, jak mohou být tato slovesa využita k dosažení specifických cílů v různých oblastech, podívejme se na praktický příklad:

Příklad: Uvažujme aktivační sloveso „*Analyzuj*“ a ukažme, jak se jeho aplikace liší v různých oborech:

1. **Vzdělávání:**

Kontext: Učitel zadá AI: „*Analyzuj historický dokument XYZ a určete jeho hlavní myšlenku a zaujatost autora.*“

Očekávaná odpověď AI: AI poskytne podrobnou analýzu dokumentu – popíše historický kontext vzniku, shrne hlavní myšlenky a posoudí, jaké postoje či zkrslení autor mohl mít. Odpověď bude zaměřená na kritické zhodnocení pramene, protože ve vzdělávacím kontextu jde o porozumění textu a jeho interpretaci.

2. **Zdravotnictví:**

Kontext: Lékař požádá AI: „*Analyzuj výsledky pacientových krevních testů.*“

Očekávaná odpověď AI: AI provede analýzu laboratorních hodnot – vyhodnotí, které parametry jsou mimo normu, a identifikuje možné indikace onemocnění (např. vysoký cholesterol, zvýšené bílé krvinky apod.). V lékařském kontextu se tedy „analyzuje“ zaměřuje na interpretaci dat ve vztahu ke zdraví pacienta.

3. Marketing:

Kontext: Marketingový analytik zadá AI: „Analyzuj data o chování zákazníků na našem e-shopu za poslední kvartál.“

Očekávaná odpověď AI: AI poskytne analýzu trendů v nákupním chování – například identifikuje nejprodávanější produkty, časy, kdy jsou zákazníci nejaktivnější, míru opuštění košíku apod. Může také vyvodit doporučení, na co se zaměřit v kampaních. Zde sloveso „analyzuje“ vede model k vyzdvížení trendů a insightů z dat o zákaznících.

4. Finanční analýza:

Kontext: Finanční analytik položí dotaz: „Analyzuj aktuální situaci na akciovém trhu.“

Očekávaná odpověď AI: AI provede technickou a fundamentální analýzu – zmíní hlavní trendy cen akcií, objemy obchodů, důležité ekonomické události ovlivňující trh. Může uvést růstová a klesající odvětví, citovat klíčové ukazatele (indexy, inflaci apod.). V tomto případě se „analyzuje“ projeví jako komplexní rozbor finančních trhů.

Vidíme tedy, že „analyzuje“ zůstává v základu stejným požadavkem – chceme analýzu – ale specifiky analýzy a očekávané výstupy se liší podle kontextu a oboru. AI se dokáže flexibilně přizpůsobit: v oblasti historie rozebere text, ve zdravotnictví čísla z laboratoře, v marketingu nákupní data a ve financích tržní indikátory. **Flexibilita AI**

v kombinaci se správně zvoleným akčním slovesem umožňuje získat užitečné odpovědi napříč obory, je-li prompt dostatečně konkrétní.

Tento pohled nám připomíná, že při práci s AI je důležité **upřesnit kontext** dotazu, nejen použít akční sloveso. Někdy tedy samotné sloveso nestačí – musíme ho doplnit o detaily (např. *co analyzovat, za jakým účelem*), aby model odvedl přesně tu práci, kterou potřebujeme.

Bonusová strategie: 5W1H (Co, Kde, Kdy, Proč, Kdo a Jak)

Otázky „*co, kde, kdy, proč, kdo a jak*“ jsou univerzální a používají se po celém světě jako základní nástroj pro získávání informací. V angličtině jsou známy jako 5W1H (Who, What, Where, When, Why, How). Tato strategie patří k základům žurnalistiky, výzkumu, vyšetřování i analytického myšlení obecně.

Využití metody 5W1H je vynikající zejména při začátku práce s novým tématem. Pomůže systematicky pokrýt všechny důležité aspekty a *vytvořit si ucelený obraz* o problematice. Je to ideální postup pro počáteční fázi průzkumu, kdy potřebujete nasbírat co nejvíce informací a nic podstatného nevynechat. Podle informačních studií jde o **osvědčený postup** pro rámcování problému – základní okruhy otázek slouží jako odrazový můstek pro další, podrobnější zkoumání (Masarykova univerzita, 2025).

Několik situací, kdy tuto strategii efektivně využít:

- **Průzkum nového tématu:** Když se potřebujete rychle zorientovat v nové oblasti, položení těchto šesti otázek vám pomůže shromáždit základní informace a vytvořit si přehledný rámec.
- **Začátek výzkumného projektu nebo úkolu:** Při startu projektu je 5W1H skvělým způsobem, jak identifikovat klíčové aspekty tématu a stanovit, *čemu se věnovat*.

- **Prohloubení stávajících znalostí:** Pokud chcete proniknout hlouběji do známého tématu, 5W1H pomůže odhalit případné mezery ve vašem porozumění – na co jste se ještě nezeptali?
- **Příprava na diskusi nebo prezentaci:** Před veřejným vystoupením na určité téma vám těchto šest otázek zajistí, že máte pokryty všechny důležité body a jste připraveni na případné dotazy publika.
- **Vyjasnění složité situace:** Pokud narazíte na komplexní problém či nejasnost, rozložením na *Co/Kde/Kdy/Proč/Kdo/Jak* si jej můžete postupně rozebrat a lépe porozumět každému aspektu.

Každá z těchto otázek má specifický účel a pomáhá získávat jiný typ informace:

Otázka	Účel (jaký typ informace zjišťuje)
Co?	Zjišťuje fakta, informace nebo popis situace či problému. (Co se stalo? Co to je? Co to zahrnuje?)
Kde?	Určuje místo nebo prostředí , kde k něčemu došlo nebo dochází. (Kde se to stalo? Kde se to nachází?)
Kdy?	Zjišťuje časový aspekt – kdy se událost stala, časové ohraničení problému apod. (Kdy k tomu došlo? Kdy se to děje?)
Proč?	Pídí se po důvodech, příčinách či účelu . (Proč se to stalo? Proč je to důležité? Proč to děláme?)
Kdo?	Identifikuje osoby nebo subjekty zapojené do situace či problému. (Kdo je zúčastněn? Kdo je ovlivněn? Kdo o tom rozhodl?)

Jak?	Zkoumá proces, způsob nebo metodu , jak něco probíhá nebo jak bylo něčeho dosaženo. (Jak to funguje? Jak se to stalo? Jak to udělat?)
-------------	--

Použití této sady otázek pomáhá udržet dotazy organizované a zajistit, že pokryjete různé aspekty tématu. Výsledkem je hlubší porozumění a komplexnější pohled.

Praktický příklad aplikace této strategie: Řekněme, že chcete prozkoumat téma „*Obnovitelné zdroje energie*“. Pomocí metody 5W1H si můžete připravit následující okruhy otázek:

- **Co** jsou obnovitelné zdroje energie a jaké jsou jejich hlavní typy?
- **Kde** se tyto zdroje nejčastěji využívají (v jakých zemích či regionech, v jakých odvětvích)?
- **Kdy** se začaly obnovitelné zdroje energie masověji využívat a jak se jejich využití vyvíjí v posledních letech?
- **Proč** je důležité přecházet na obnovitelné zdroje energie (jaké výhody přinášejí, jaké problémy řeší)?
- **Kdo** jsou hlavní hráči v oblasti obnovitelných zdrojů (které státy, firmy či organizace stojí v čele vývoje a implementace)?
- **Jak** se technologie obnovitelných zdrojů energie vyvíjely a jak fungují (např. jak funguje solární panel, větrná turbína apod.)?

Položením těchto šesti otázek získáte ucelený přehled: definice a druhy, geografické a časové souvislosti, význam a motivace, klíčové aktéry i principy fungování. Tato strategie zaručí, že na nic podstatného nezapomenete. Ve skutečnosti se jedná o základní novinářskou techniku – novinář při pokrývání události také zpravidla odpovídá na otázky co, kde, kdy, proč, kdo a jak, aby jeho zpráva byla úplná.

Metoda 5W1H je tedy jednoduchým, ale mocným nástrojem. Pro naše účely v komunikaci s AI platí, že pokud nevíte, jak začít nebo jak strukturovat dotazy k novému tématu, začněte těmito základními otázkami. **Získáte tím široký základ informací**, na kterém můžete dále stavět s podrobnějšími a specifitějšími dotazy. V kombinaci s ostatními postupy prompt engineeringu tak budete vybaveni k důkladnému a efektivnímu průzkumu libovolného tématu.

Závěr

Tato monografie sledovala záměr, který od počátku stál na pomezí dvou rovin: představit *praktický průvodce* komunikací s generativní umělou inteligencí, určený pedagogům, studentům a profesionálům, a současně poskytnout *teoreticky ukotvený výzkumný rámec*, který posune dosud fragmentovaný výzkum této oblasti směrem k integraci. Závěrečná kapitola shrnuje, co se podařilo dosáhnout, jak se obě roviny do sebe skládají, a kam směřuje další práce.

Shrnutí klíčových přínosů

Praktická rovina monografie (kapitoly 1–4 a 7–8) představila soustavně budovaný aparát komunikačních dovedností pro práci s generativní AI. Začala vymezením základní terminologie a pojmu *prompt* jako didakticky orientovaného konstruktů, přes devět dimenzí kvalitní AI komunikace (uživatelská přívětivost, přesnost, personalizace, transparentnost, eticko-sociální aspekty, výkonnost, integrace, nákladovost a znalostní báze), osm strategií efektivního dotazování (od konkrétnosti přes sdělení záměru až po experimentaci s frázemi a učení se z odpovědí), až po systematický rozbor *vzorců promptů a aktivizačních sloves* — tedy konkrétních jazykových struktur, které prokazatelně zlepšují kvalitu výstupů AI nástrojů. Tato část je strukturována tak, aby ji čtenář mohl číst lineárně jako kurz, nebo využívat fragmentárně jako referenční příručku.

Teoreticko-empirická rovina monografie (kapitoly 5–6) překračuje rámec praktické příručky a představuje původní vědecký přínos autora. Kapitola 5 zavedla rámec **MLEF-AI** (*Multi-Level Ecological Framework for AI Implementation in Education*) z autorovy habilitační teze (Hanák, 2026) jako ekologický model implementace AI ve vzdělávání s pěti vnořenými úrovněmi a čtyřmi feedback loops. Tento rámec doplnila o **AI-centrickou pedagogiku (AICP)** z autorovy předchozí monografie (Hanák, 2025) jako pedagogický kontextový rámec a operacionalizovala pojem **komunikační zralost uživatele** s návrhem nového měřicího nástroje *Communication Maturity Index*

(CMI). Kapitola 6 následně prezentovala empirickou retrospektivní observační studii 147 uživatelů edukační platformy Kurzologie.cz, která testovala výzkumné otázky vyvozené z teoretického rámce.

Klíčovým empirickým zjištěním studie je **silný a robustní vztah mezi dokončeností kurzu a deklarovanou změnou v interakci s generativní AI** ($\chi^2(2) = 20,43$, $p < 0,0001$, Cramérovo $V = 0,373$). Tento vztah obstál ve všech 1 000 iteracích sensitivity analýzy, což svědčí o jeho odolnosti vůči nejistotě kódování. Současně studie ukázala, že specifický formát kurzu (D1–D5) ani počet využitých dimenzí nemají statisticky významný diferenciální efekt na výsledek — kurz funguje napříč dimenzemi srovnatelně. Experimentální dimenze D5 (podcast + infografika) dosáhla 100% dokončenosti a obstála ve srovnání s tradičními formáty. VARK preferenční pretest se ukázal jako slabý prediktor výsledné změny interakce s AI a jeho dodržení neovlivňuje výsledek statisticky významným způsobem. K-means klastrová analýza identifikovala čtyři přirozeně se vyskytující typologie uživatelů kurzu, které mohou sloužit jako východisko pro další design vzdělávacích intervencí.

Vztah mezi praktickou a teoreticko-empirickou rovinou

Praktická a teoreticko-empirická rovina monografie nestojí vedle sebe izolovaně — vzájemně se osvětlují. *Komunikační strategie a vzorce promptů* popsané v praktické části nejsou jen technickými doporučeními; v rámci konceptu **komunikační zralosti** zavedeném v kapitole 5 představují konkrétní operacionalizace dialogické orientace, strategického osvojení a transferu do praxe. Uživatel, který si systematicky osvojí strategie z kapitoly 4 a vzorce z kapitoly 7, neprochází jen výcvikem v technice prompt engineeringu, ale prochází **vývojovou trajektorií ke komunikační zralosti** v interakci s AI.

Z opačného směru: empirické zjištění z kapitoly 6, že **rozhodující není formát výuky, ale dokončenost kurzu**, dává praktické rovině konkrétní pedagogickou implikaci. Pokud uživatel konzumuje

fragmenty obsahu bez pravidelnosti, sebebohatší aparát komunikačních strategií zůstává nevyužit. Kvalita osvojení komunikační zralosti je *kumulativní* — vzniká až dlouhodobou expozicí a opakovaným cvičením, nikoli jednorázovým přečtením příručky. Tento empirický nálezn podporuje pedagogický design AI-centrické pedagogiky (AICP) s jejím důrazem na strukturovanou, systematickou edukační intervenci.

Třetí integrační linií je vztah mezi *technickými dovednostmi* (technická vrstva interakce s AI) a *kritickou kompetencí* (etická a regulační vrstva). Praktická část monografie poukázala na principy transparentnosti, ověřování faktů, kritické validace odpovědí a etických aspektů personalizace. Teoretický rámec MLEF-AI tyto principy zařadil do svojí **Úrovně 2 (Eticko-regulační)** s jejími pěti pilíři etické AI ve vzdělávání. Komunikačně zralý uživatel není tedy jen technicky efektivní, ale je současně eticky reflexivní — chápe, že každá interakce s AI má governance dimenzi.

Limity monografie a otevřené otázky

Monografie přiznává své limity transparentně. **Empirická studie kapitoly 6 je publikována jako interní pracovní výstup** dokumentující průběžnou fázi výzkumu. Vychází z rekonstruktivních dat získaných v rámci provozu edukační platformy mezi březnem 2025 a lednem 2026, byla provedena jediným kódem bez slepoty, čelí mnohonásobnému střetu zájmů (autor je tvůrce, provozovatel a hodnotitel zároveň) a postrádá kontrolní skupinu. Hlavní výsledková proměnná AI_CHANGE je založena na sebehodnocení testerů, nikoli na objektivním behaviorálním měření.

Druhým otevřeným problémem je *diskrepance mezi deklarovaným cílem kurzu a reálným využitím*: ze 147 testerů kurzu *Tvůrce kurzů* se jen čtyři primárně zaměřili na produkci vlastního kurzu, zatímco 97 % testerů kurz reálně využilo k rozvoji AI gramotnosti. Tento nálezn ukazuje, že kurz funguje efektivně, ale často jinak, než jak byl

koncipován — což je legitimní pedagogická lekce, ne metodologická vada studie.

Třetím limitem je *teoretický status konstruktů komunikační zralosti*. Monografie navrhla pracovní strukturu Communication Maturity Index (CMI) se třemi subškálami, ale jeho plnohodnotná psychometrická validace (analýza vnitřní konzistence, faktorová struktura, konvergentní validita s RTI a AI Self-Efficacy) zůstává předmětem navazujícího výzkumu. Současná monografie nabízí konstrukt jako *teoretický návrh opřený o substantivní empirickou evidenci*, nikoli jako validovaný měřicí nástroj.

Výhled: směry dalšího výzkumu

Z výše popsaných limitů přirozeně vyplývají směry navazujícího výzkumu. **Fáze 2 výzkumu**, kterou autor plánuje, bude zahrnovat prospektivní design s předregistrovaným kódovacím protokolem na platformě Open Science Framework, slepé kódování druhým nezávislým hodnotitelem s výpočtem inter-raterové reliability, vstupní a výstupní měření postojů a chování pomocí standardizovaných nástrojů, behaviorální analýzu skutečných AI promptů uživatelů, srovnání s kontrolní skupinou a — zejména — *behaviorální/produktové měřítko* jako primární výstup (počet a kvalitu kurzů, které účastníci skutečně vytvoří).

Druhým navazujícím směrem je **psychometrická validace CMI nástroje**. Konstruktová validita bude testována na nezávislém vzorku odlišném od populace Kurzologie.cz, aby se ověřila stabilita tříložkové struktury (dialogická orientace, strategické osvojení, transfer do praxe) napříč různými skupinami uživatelů. Cílem je vytvořit nástroj použitelný v širokém spektru edukačních a profesních kontextů, který by mohl sloužit jako diagnostický a evaluační nástroj pro AI gramotnostní intervence.

Třetím směrem je **vícenásobná případová studie** ve smyslu Yinovy (2018) metodiky, která by hloubkově analyzovala konkrétní edukační

intervence s AI komponentou v různých kontextech (vysokoškolské prostředí, profesní vzdělávání, samostudium). Cílem je doplnit kvantitativní obraz této monografie o bohaté kvalitativní popisy trajektorií, jež jsou v současné explorativní studii pouze nastíněny.

A konečně, **integrace MLEF-AI a AICP rámců** je dlouhodobým teoretickým úkolem, který přesahuje rozsah jediné monografie. Otázkou pro budoucí výzkum je, jak by mohl vypadat *integrovaný operační model*, který by spojoval ekologicko-implementační strukturu MLEF-AI s pedagogickou specifikací AICP do jednotného rámce použitelného pro design, evaluaci i governance AI ve vzdělávání. Tato monografie poskytla první argumenty ve prospěch komplementarity obou rámců; jejich plné teoretické sjednocení je ambicí pro autorův další badatelský program.

Závěrečné slovo

Komunikace s generativní umělou inteligencí představuje historicky nový typ kompetence, který má zároveň všechny atributy klasické *gramotnosti* — kumulativně se rozvíjí, vyžaduje strukturovanou expozici, je modulárně osvojitelná a má stupně zralosti. Tato monografie tvrdila, že tato kompetence si zaslouží pojmenování jako **komunikační zralost v interakci s generativní umělou inteligencí** a že její rozvoj nemůže být ponechán náhodě nebo intuitivnímu osvojování — vyžaduje pedagogicky promyšlenou edukační intervenci s jasným teoretickým ukotvením a empirickou evidencí.

Pole se vyvíjí rychle. V době, kdy je text uzavřen (duben 2026), představuje Mollickův *Co-Intelligence* (2024) jeden z nejvlivnějších referenčních textů populárního AI vzdělávání, výzkumný rámec MLEF-AI je teprve uváděn do akademického oběhu, a empirická báze pro rozvoj AI gramotnosti dále roste. Za rok bude některá tvrzení této knihy pravděpodobně potřeba zpřesnit, jiná aktualizovat a některá možná opustit. To je přirozený osud každé monografie v rychle se rozvíjejícím poli — a důvod, proč autor považuje tuto publikaci za

jednu kapitolu v pokračujícím výzkumném programu, nikoli za uzavřenou tezi.

Pokud monografie přispěje k tomu, že čtenáři — pedagogové, studenti, badatelé i profesionálové z praxe — budou ve svých interakcích s generativní AI o něco vědomější, strategičtější, eticky reflexivnější a kompetenčně zralejší, naplní svůj cíl. Konkrétní podoba toho, jak se komunikační zralost projeví v jejich práci, je už věcí každého z nich.

Věcný rejstřík

A

Adaptabilita – 21, 30, 31, 33–35, 78, 98

Adopce technologie – 99

AI Self-Efficacy – 94, 96, 100, 101, 125

AI-centrická pedagogika (AICP) – 9, 91, 98, 99, 102, 103

Aktivační slovesa – 9, 176–178, 184

Amazon – 34, 41, 48

AutoTutor – 8

B

Beneficence – 42

Bias (zkreslení) – 26, 30, 35, 36, 40, 41, 45, 105, 127, 131, 184

Bronfenbrennerova ekologická teorie – 92

C

Cambridge Analytica – 44

Claude (Anthropic) – 8, 16, 18, 21, 146, 161, 162

Communication Maturity Index (CMI) – 9, 97, 98, 101, 102, 125

Copilot – 11

D

Datafikace – 94, 97

Deprofesionalizace – 94, 96, 97, 125

Dialogické vyučování – 6

Digitální gramotnost – 7, 18

Dimenze AICP (D1–D10) – 98

Dimenze D5 (podcast + infografika) – 112, 114, 115, 118, 121, 123, 128

Dimenze kurzu (D1–D5) – 105, 114, 122, 128

Dokončenost kurzu – 10, 75, 102, 105, 110, 112–123, 125, 126, 128, 129

Dropout (předčasné ukončení) – 108, 112, 113, 120, 121

Důvěryhodnost – 26, 27, 35, 37–40, 43, 85, 86, 159, 165

E

Efektivita – 8, 24, 33, 49, 50, 60, 62, 64–66, 126, 130, 132, 152, 155, 162, 163, 167, 172, 175

Effort Expectancy (UTAUT) – 95, 125

Eticko-sociální aspekty – 40

EU AI Act – 95

Expertní analýza (promptový vzorec) – 155

F

Facilitating Conditions (UTAUT) – 24, 95

Fáze 2 výzkumu – 9, 98, 100, 101, 103, 110, 114, 116, 118, 122, 126–128

Feedback loops (FL1–FL4) – 93, 96, 97, 103, 126

Fixsenovi implementační ovladače – 94, 95, 97, 125

FOCUS (proměnná) – 108

Formát výstupu (komponenta promptu) – 81

G

Gemini (Google) – 8, 11, 16, 18, 21, 146, 151, 161, 162, 164, 165

Genderové stereotypy – 26

Generativní umělá inteligence – 9–11, 14, 17, 102, 104, 105, 108, 113, 121, 122, 128, 159

Governance Perception Index (GPI) – 94, 95, 97

GPT-5 / GPT (OpenAI) – 16, 20, 146, 147, 151, 161, 162, 164

H

Halucinace AI – 84, 86, 131

HITL (Human-in-the-Loop) – 98

I

Implementace AI – 62, 92, 99, 124

Insider practitioner research – 105

Iterace dotazu – 13, 18, 20, 82, 83, 90, 98, 99, 101, 109, 112, 125, 163, 165, 169, 171–174, 176

K

Klastrová analýza (k-means) – 110, 117, 129

Knowledge cutoff – 90

Komunikační zralost – 9, 10, 91, 93, 99–102, 126

Kontext dotazu – 131, 167

Kritická pedagogika – 93

Kurzologie.cz – 9, 97–99, 102, 104, 106, 126, 128

L

Learnifikace – 94

M

MLEF-AI rámec – 9, 91–93, 95–97, 99, 100, 102, 103, 124–126

MOOC (Massive Open Online Courses) – 123

Multimodální učení – 16, 105, 111, 115, 117, 118, 120, 121

N

Náklady (costs) – 60–62, 65, 66

NIST AI Risk Management Framework – 94, 95

O

Observační retrospektivní studie – 9, 104, 108, 109, 113, 128

P

P-EFST škála technostresu – 94–96, 100

Performance Expectancy (UTAUT) – 95, 100, 125

Personalizace – 22, 25, 30–35, 44, 46, 47, 96, 98, 122, 123, 133

Pět pilířů etické AI – 94, 95

Pilotní případy (sedm testerů) – 104, 108, 119–121, 128

Prompt – 8, 11–14, 17, 18, 20, 79, 80, 87, 89, 90, 100, 130–133, 137, 138, 141, 143, 147, 148, 156, 159, 162, 165–174, 176, 177, 184, 186, 189

Prompt engineering – 8, 12, 14, 17, 100, 131, 137, 159, 166, 167, 189

R

Re-profesionalizace – 96, 97, 102, 125

Role Transformation Index (RTI) – 94, 96, 97, 100, 101, 125

Role-prompting – 77, 149

S

SAMR model – 96

Sebehodnocení / sebereport – 109, 127

Sensitivity analýza (Monte Carlo) – 109, 111, 119, 127, 128

Social Influence (UTAUT) – 95, 96

Strategie efektivního promptování – 74, 99

T

Technologický akcelerationismus – 92–95

Technostres – 91, 93–96, 100

Transparentnost – 27, 35–40, 42, 43, 95, 126

Tutor LMS – 104, 107, 108

Tvůrce kurzů (kurz) – 98, 102, 106, 114, 125, 127, 128

Typologie uživatelů (čtyři klastry) – 117, 118, 121, 129

U

UNESCO Doporučení o etice AI – 94, 95, 97

UTAUT (Unified Theory of Acceptance) – 94–96, 100, 125

V

VARK preferenční pretest – 99, 102, 106, 107, 111, 116–118, 122, 124, 126, 128

Vzorce promptů (patterny) – 130, 131, 133, 150, 153, 154, 159

Z

Zpětná vazba (feedback) – 6, 33, 55, 67, 71, 89, 93, 96–98, 103, 126, 171–174

Jmenný rejstřík

- Alexander, Robin – 6
- Amodei, Dario – 47, 60
- Anthropic – 16, 161
- Atlas, Stephen – 8
- Bandura, Albert – 94, 96, 100
- Bender, Emily M. – 84, 86
- Biesta, Gert – 93, 94
- Bolukbasi, Tolga – 26
- Bronfenbrenner, Urie – 92
- Brookhart, Susan – 12
- Brown, Tom – 13, 21, 131, 146–148
- Bsharat, Sondos M. – 74, 76–78
- Cain, John – 17
- Dastin, Jeffrey – 41, 45
- Davis, Fred D. – 21, 93
- Doshi-Velez, Finale – 27
- Engel, Susan – 6
- Evropská unie (EU AI Act) – 95
- Fixsen, Dean L. – 94, 95, 97, 125
- Flaxman, Seth – 27, 36, 40
- Floridi, Luciano – 42
- Frey, Carl Benedikt – 41
- Geis, J. Raymond – 59
- Goodman, Bryce – 27, 36, 40
- Google DeepMind – 16
- Graesser, Arthur C. – 8
- Hanák, Michal – 9, 91, 92, 95, 97, 98, 124
- Hocky, Tony – 17
- Husmann, Polly R. – 124
- Jirout, Jamie – 6
- Jobanputra, Ketul – 24, 65
- Jobin, Anna – 37, 40, 41, 43, 45
- Kasneci, Enkelejda – 8
- Kizilcec, René F. – 123

Klahr, David – 6	Osborne, Michael A. – 41
Knoth, Niclas – 130, 133	Palmer, Edward – 18, 19
Komić, Ognjen – 40	Pashler, Harold – 124
Kotwani, Rohan – 162	Person, Natalie – 7
Kyriacou, Chris – 6	Porayska-Pomsta, Kaska – 8
Lepper, Mark – 7	Provázková Stolinská, Dominika – 6
Liu, Pengfei – 11, 13	Průcha, Jan – 6
Luger, Ewa – 28	Puentedura, Ruben R. – 94
Malik, Mustafa A. – 15	Rafferty, Damien – 17, 19
McKinney, Scott – 53	Redecker, Christine – 7
McNiff, Jean – 105	Reich, Justin – 123
Metz, Cade – 84, 86	Reynolds, Laria – 13, 131, 132, 137, 140
Mollick, Ethan – 8, 74, 79, 130, 138	Russell, Stuart – 11
Nielsen, Jakob – 22, 49, 50	Selwyn, Neil – 93, 94
NIST – 94, 95	Shakurnia, Abdolhussein – 6
Norvig, Peter – 11	Shin, Donghee – 24, 38
O'Loughlin, Valerie D. – 124	Stake, Robert E. – 105
OpenAI – 8, 16, 18, 75, 81–83, 86, 161, 164, 165, 173	Tarafdar, Monideepa – 94, 95

Tuomi, Ilkka – 19

UNESCO – 94, 95, 97

Urban, Eric – 77

Vadlapati, Praneeth – 79,
80

Vališová, Alena – 6

Van der Meij, Hans – 6

Venkatesh, Viswanath –
23, 24, 94, 95

Walsh, P. P. – 19

Walter, Yoshija – 19

Weise, Karen – 84, 86

Williamson, Ben – 93, 94

Wirth, Niklaus – 15

Woolverton, Maria – 7

Yin, Robert K. – 105

Zhao, Zihao – 87

Použitá literatura

ALEXANDER, Robin, 2008. Towards Dialogic Teaching: Rethinking Classroom Talk. 4th ed. York: Dialogos. ISBN 978-0954694364.

AMODEI, D., & HERNANDEZ, D. 2018. AI and Compute. OpenAI Blog.

ANTHROPIC. 2025. Claude 4: Model Documentation. San Francisco: Anthropic PBC. Dostupné z: <https://www.anthropic.com/>. (Technická dokumentace modelu Claude 4; popisuje vlastnosti modelu zaměřeného na bezpečnost, dlouhý kontext a schopnosti agentního chování.)

BANDURA, A. 1986. Social Foundations of Thought and Action: A Social Cognitive Theory. Englewood Cliffs, NJ: Prentice-Hall.

BENDER, E. M., GEBRU, T., MCMILLAN-Major, A., & SHMITCHELL, S. 2021. On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? Proceedings of ACM FAccT 2021, 610–623.

BIESTA, G. 2017. The Rediscovery of Teaching. New York: Routledge. ISBN 978-1138670709.

BOLUKBASI, T., CHANG, K. W., ZOU, J., SALIGRAMA, V., & KALAI, A. 2016. Man is to computer programmer as woman is to homemaker? Debiasing word embeddings. In Advances in neural information processing systems (NeurIPS).

BRONFENBRENNER, U. 1979. The Ecology of Human Development: Experiments by Nature and Design. Cambridge, MA: Harvard University Press. ISBN 978-0674224575.

BROOKHART, Susan M. How to Assess Higher-Order Thinking Skills in Your Classroom. Alexandria: ASCD, 2010. ISBN 978-1-4166-1048-7.

BROWN, Tom, et al. 2020. Language Models are Few-Shot Learners. In: Advances in Neural Information Processing Systems [online].

2020, 33, s. 1877–1901. Dostupné z:
<https://arxiv.org/abs/2005.14165>.

BSHARAT, S. M., MYRZAKHAN, A., & SHEN, Z. 2023. Principled Instructions Are All You Need for Questioning LLaMA-1/2, GPT-3.5/4. arXiv preprint arXiv:2312.16171.

Cain, J. 2023. Prompt engineering as an emergent critical skill set. [White paper].

COOK, J. 2023. How to write effective prompts for ChatGPT: 7 essential steps for best results. Forbes.com, 26. června 2023.

DASTIN, J. 2018. Amazon scraps secret AI recruiting tool that showed bias against women. Reuters, 10 Oct 2018.

DAVIS, F. D. 1989. Perceived usefulness, perceived ease of use, and user acceptance of information technology. MIS Quarterly, 13(3), 319–340.

DOSHI-VELEZ, F., & Kim, B. 2017. Towards a rigorous science of interpretable machine learning. arXiv:1702.08608.

ENGEL, Susan a RANDALL, Kimberly, 2009. How Teachers Respond to Children's Inquiry. American Educational Research Journal, 46(1), 183–202. DOI: 10.3102/0002831208323274.

EU AI ACT. 2024. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union, L řada.

FIXSEN, D. L., NAOOM, S. F., BLASE, K. A., FRIEDMAN, R. M., & WALLACE, F. 2005. Implementation Research: A Synthesis of the Literature. Tampa, FL: University of South Florida, Louis de la Parte Florida Mental Health Institute. FMHI Publication #231.

FLORIDI, L., et al. 2018. AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. *Minds and Machines*, 28(4), 689–707.

FREY, C. B., & OSBORNE, M. A. 2017. The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254–280.

FRISINGER, T. 2025. Prompting Is Dead. Long Live the Loop. AiBuddy.software (blogový článek).

GEIS, J. R., et al. 2019. Ethics of artificial intelligence in radiology: summary of the Joint European and North American multisociety statement. *Radiology*, 293(2), 436–440.

GOODMAN, B., & FLAXMAN, S. 2017. European Union regulations on algorithmic decision-making and a “right to explanation”. *AI Magazine*, 38(3), 50–57.

GOOGLE DEEPMIND. 2026. Gemini 3 Technical Report. Mountain View, CA: Google LLC. Dostupné z: <https://deepmind.google/>. (Technická specifikace modelu Gemini 3 představeného v únoru 2026; multimodální model s nativní podporou textu, obrazu, audia a videa.)

GRAESSER, Arthur C., et al., 2016. Conversations with AutoTutor Help Students Learn. *International Journal of Artificial Intelligence in Education*, 26(1), 124–132. DOI: 10.1007/s40593-015-0086-7.

HANÁK, M. 2025. AI-centrická pedagogika jako nový fenomén. Uherské Hradiště: Akademie krizového řízení a managementu, s.r.o. 268 s. ISBN 978-80-88398-30-1.

HANÁK, M. 2026. Didaktické výzvy a role učitele v kontextu umělé inteligence: rizika a etické aspekty [Habilitační práce]. Dubnica nad Váhom: Vysoká škola DTI. 141 s.

HATTIE, John. Visible Learning for Teachers: Maximizing Impact on Learning. London: Routledge, 2012. ISBN 978-0-415-69015-7.

HOCKY, T., & WHITE, J. 2022. Prompt engineering, literally learning to interface more clearly with an AI, is now a research topic. [Conference presentation].

CHIP. 2021. Setting the Goal: The Foundation of Prompt Engineering. QuantHub – News & Insights.

HUSMANN, P. R., & O'LOUGHLIN, V. D. 2019. Another Nail in the Coffin for Learning Styles? Disparities among Undergraduate Anatomy Students' Study Strategies, Class Performance, and Reported VARK Learning Styles. *Anatomical Sciences Education*, 12(1), 6–19.

INVESTOPEDIA (Team). 2025. What Is Total Quality Management (TQM), and Why Is It Important? Investopedia, aktualizováno 11. 5. 2025.

JIANG, P., NIU, W., WANG, Q., YUAN, R., & CHEN, K. 2024. Understanding Users' Acceptance of Artificial Intelligence Applications: A Literature Review. *Behavioral Sciences*, 14(8), 671.

JOBANPUTRA, K. 2024. Customer service costs could be reduced by up to 30% when using AI chatbots (IBM report cited in). Infomineo Blog, March 24, 2025.

JOBIN, A., IENCA, M., & VAYENA, E. 2019. The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399.

KASNECI, Enkelejda, et al., 2023. ChatGPT for Good? On Opportunities and Challenges of Large Language Models for Education. *Learning and Individual Differences*, 103, 102274. DOI: 10.1016/j.lindif.2023.102274.

KIZILCEC, R. F., PÉREZ-SANAGUSTÍN, M., & MALDONADO, J. J. 2017. Self-regulated learning strategies predict learner behavior and goal

attainment in Massive Open Online Courses. *Computers & Education*, 104, 18–33.

KNOTH, N., TOLZIN, A., JANSON, A., & LEIMEISTER, J. M. 2024. AI literacy and its implications for prompt engineering strategies. *Computers and Education: Artificial Intelligence*, 6, 100225.

KO, G. Y., et al. 2022. Learning outside the classroom during a pandemic: Evidence from an artificial intelligence-based education app. *Management Science*, 69(8), 3616–3649.

KOMIĆ, D., et al. 2021. Blame the bot: Anthropomorphism and anger in customer–chatbot interactions. *Journal of Marketing*, 85(6), 132–148.

KOTWANI, R. 2025. Write Clear Instructions — OpenAI Prompt Engineering. Medium (článek, publikováno 28. 6. 2025).

LAKERA (AI). 2025. The Ultimate Guide to Prompt Engineering in 2025. Lakera.ai (firemní blog, publikováno 2. 6. 2025).

LEE, D., & PALMER, E. 2025. Prompt engineering in higher education: a systematic review to help inform curricula. *International Journal of Educational Technology in Higher Education*, 22(7).

LEE, J. D., & SEE, K. A. 2004. Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80.

LIU, D. 2023. Prompt engineering for educators: Making generative AI work for you. LinkedIn (příspěvek z 8. června 2023).

LIU, Peng, et al. What Makes Good In-Context Examples for GPT Models? arXiv preprint [online]. 2023. Dostupné z: <https://arxiv.org/abs/2305.14788>.

LUGER, E., & SELLEN, A. 2016. “Like Having a Really Bad PA”: The Gulf Between User Expectation and Experience of Conversational Agents. In *Proc. CHI 2016*, 5286–5297.

MALIK, M. A. 2000. On the Perils of Programming. Communications of the ACM, 43(12), 95-97.

Masarykova univerzita, KISK (Filozofická fakulta). 2025. 6 otázek (Metoda 5W1H). Online metodika 100metod (Katedra informačních studií a knihovnictví FF MU).

MCKINNEY, S. M., et al. 2020. International evaluation of an AI system for breast cancer screening. Nature, 577(7788), 89–94. (Příklad nasazení AI v mamografickém screeningu – výkon srovnatelný s lékaři a potenciál zefektivnit diagnostiku)

MCKINSEY 2021. The value of getting personalization right—or wrong—is multiplying. McKinsey & Co. (Zpráva uvádějící, že firmy excelující v personalizaci generují o 40 % více příjmů z těchto aktivit, a další statistiky)

McNIFF, J. 2017. Action Research: All You Need to Know. Los Angeles: Sage. ISBN 978-1473967397.

MKWRITESHHERE. 2025. We're Back to Square One: Why AI is Forcing Us to Reinvent Programming Languages (Again). Towards AI – Medium, 5. 6. 2025. (Anekdoticky popisuje, že v 50. letech byly pokusy komunikovat s počítači v přirozeném jazyce neúspěšné, což vedlo ke vzniku formálních jazyků jako FORTRAN a COBOL).

MOLLICK, E. 2023. How to use ChatGPT to boost your writing. One Useful Thing (blog), 10. ledna 2023. (Blog technologického nadšence a pedagoga z Wharton School, nabízející průvodce použitím ChatGPT – autor přirovnává promptování k „programování stroje slovy“ a radí psát k AI jako k člověku s poskytnutím role a kontextu.)

MOLLICK, E. 2024. Co-Intelligence: Living and Working with AI. New York: Portfolio. ISBN 978-0593716717. (Ovlivněná monografie shrnující praktické principy spolupráce s generativní AI; popisuje čtyři principy (treat AI as a person, bring AI to the table, be the human in the loop, assume current AI is the worst you will ever use) a stává se

klíčovým referenčním textem pro AI gramotnost v období 2024–2026.)

NIELSEN, J. 1993. Usability Engineering. Morgan Kaufmann. (Kniha obsahující mj. pravidla pro odezvy systému – 0,1s, 1s, 10s – a další heuristiky uživatelské přívětivosti)

NIST. 2023. AI Risk Management Framework (AI RMF 1.0). Washington, DC: National Institute of Standards and Technology, U.S. Department of Commerce. NIST AI 100-1.

NOY, S., & ZHANG, W. 2023. Experimental evidence on the productivity effects of generative artificial intelligence. (Working Paper).

OLEA, C., TUCKER, H., PHELAN, J., PATTISON, C., ZHANG, S., LIEB, M., SCHMIDT, D., & WHITE, J. 2024. Evaluating persona prompting for question answering tasks. In Proc. of the 12th Int. Conf. of Security, Privacy and Trust Management (p. 81). Sydney, Australia.

OPENAI 2023. Prompt engineering best practices for ChatGPT. OpenAI Help Center, aktualizováno předběžně v květnu 2023.

OPENAI. 2026. GPT-5.5: Model Documentation. San Francisco: OpenAI Inc. Dostupné z: <https://openai.com/index/introducing-gpt-5-5/>. (Dokumentace modelu GPT-5.5 představeného v dubnu 2026; popisuje pokročilé agentní schopnosti, vylepšené uvažování a integraci nástrojů.)

OPENAI. 2025. GPT-5: System Card. San Francisco: OpenAI Inc. Dostupné z: <https://openai.com/index/gpt-5-system-card/>. (Technická dokumentace modelu GPT-5 představeného v srpnu 2025; popisuje schopnosti, limity, bezpečnostní opatření a evaluaci nového flagship modelu OpenAI nahrazujícího řadu GPT-4.)

OPENAI. 2022. Introducing ChatGPT. (Oficiální blog OpenAI, 30. 11. 2022).

OPENAI. 2025. Prompt engineering best practices for ChatGPT. OpenAI Help Center, aktualizováno 2025.

OUYANG, L. et al. 2022. Training language models to follow instructions with human feedback. (Prezentováno na NeurIPS 2022; tzv. InstructGPT).

PASHLER, H., McDANIEL, M., ROHRER, D., & BJORK, R. 2008. Learning Styles: Concepts and Evidence. Psychological Science in the Public Interest, 9(3), 105–119.

PINSKI, M., & BENLIAN, A. 2024. AI literacy for users – A comprehensive review and future research directions of learning methods, components, and effects. Computers in Human Behavior: Artificial Humans, 2(2), 100062.

PROVÁZKOVÁ STOLINSKÁ, Dominika, 2021. Komunikace učitele v prostředí české primární školy. Olomouc: Univerzita Palackého v Olomouci. ISBN 978-80-244-6096-3.

PUENTEDURA, R. R. 2006. Transformation, Technology, and Education. Online prezentace. Dostupné z: <http://hippasus.com/resources/tte/>.

RAFTERY, D. 2023. Will ChatGPT pass the online quizzes? Adapting an assessment strategy in the age of generative AI. Irish Journal of Technology Enhanced Learning, 7(1).

RANSBOTHAM, S., et al. 2019. Winning with AI. MIT Sloan Management Review and BCG report.

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.

REICH, J. 2014. MOOC Completion and Retention in the Context of Student Intent. EDUCAUSE Review, 49(6).

REYNOLDS, L., & MCDONELL, K. 2021. Prompt programming for large language models: Beyond the few-shot paradigm. Extended Abstracts of the CHI 2021 Conference on Human Factors in Computing Systems (pp. 1–7). Yokohama, Japan.

REYNOLDS, Laria, and Kyle MCDONELL. Prompt Programming for Large Language Models: Beyond the Few-Shot Paradigm. arXiv preprint [online]. 2021. Dostupné z: <https://arxiv.org/abs/2102.07350>.

RUSSELL, Stuart J., and Peter NORVIG. Artificial Intelligence: A Modern Approach. 4th ed. Harlow: Pearson, 2021. ISBN 978-1-292-40367-6.

SALAMI, Y. 2024. How to Use Google Bard (Now Gemini) to Enhance Your Online Research. Verpex.com – Web Hosting Blog, 22. 8. 2024.

SALINAS, A., & MORSTATTER, F. 2024. The Butterfly Effect of Altering Prompts: How Small Changes and Jailbreaks Affect Large Language Model Performance. arXiv preprint arXiv:2401.03729 [cs.CL].

SELWYN, N. 2019. What's the Problem with Learning Analytics? Journal of Learning Analytics, 6(3), 11–19.

SHAKURNIA, Abdolhussein; ASLAMI, Maryam; BIJANZADEH, Mahdi, 2018. The effect of question generation activity on students' learning and perception. Journal of Advances in Medical Education & Professionalism, 6(2), 70–77. PMCID: PMC5856907.

SHIN, D., KEE, K. F., & SHIN, E. Y. 2022. Algorithm awareness: Why user awareness is critical for personal privacy in the adoption of algorithmic platforms. International Journal of Information Management, 65, 102494. (Studie definující koncept "algorithmic awareness" – uživatelské povědomí o férovosti, transparentnosti atd. – a jeho vliv na důvěru a sebeodhalování uživatelů)

SINGH, R. P., et al. 2020. Current Challenges and Barriers to Real-World Artificial Intelligence Adoption for the Healthcare System, Provider, and the Patient. *Translational Vision Science & Technology*, 9(2), 45. (Přehled překážek zavádění AI ve zdravotnictví – např. integrace s EHR, interoperabilita, vysvětlitelnost, finanční aspekty)

SONG, H. 2024. Mastering Object-Oriented Programming (OOP): A Comprehensive Guide... Medium, 10. 7. 2024. (Vysvětluje historii a benefity OOP; uvádí, že modelování problému pomocí objektů je přirozenější a intuitivnější přístup k programování než psaní procedurálního kódu).

STAKE, R. E. 2006. Multiple Case Study Analysis. New York: Guilford Press. ISBN 978-1593852481.

TARAFDAR, M., TU, Q., RAGU-NATHAN, B. S., & RAGU-NATHAN, T. S. 2007. The Impact of Technostress on Role Stress and Productivity. *Journal of Management Information Systems*, 24(1), 301–328.

TIMLER, D. 2023. ChatGPT prompt engineering. DagmarTimler.com (blog), 10. 6. 2023. (Zkušenosti z prompt engineeringu; obsahuje tip „Delimiters are critical“ – použití trojitých uvozovek, backticků apod. pro izolaci vstupního textu v promptu a zamezení zmatení modelu či injection útoku).

T'SYEN, K. 2025. Prompt Engineering: The New Skill Every Knowledge Worker Needs. Medium. (online).

TUOMI, I. 2025. Generative AI and the Future of Education: A Critical Assessment of Pedagogical Transformation. *European Journal of Education Research*, 64(2), 245–267. (Kritické posouzení dopadu generativní AI na vzdělávací praxi; argumentuje pro nutnost zachování pedagogické autonomie ve světě algoritmické asistence.)

UNESCO. 2021. Recommendation on the Ethics of Artificial Intelligence. Paris: UNESCO.

URBAN, E. 2023. What is intent recognition? Microsoft Learn (dokumentace), 18. července 2023.

VADLAPATI, P. 2023. Investigating the impact of linguistic errors of prompts on LLM accuracy. *ESP Journal of Engineering & Technology Advancements*, 3(2), 150–153.

VATSAL, S., & DUBEY, H. 2024. A Survey of Prompt Engineering Methods in Large Language Models for Different NLP Tasks. *arXiv preprint arXiv:2407.12994*.

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. (Formulace UTAUT modelu – zahrnující faktory jako očekávaný výkon a úsilí, sociální vlivy, podpůrné podmínky – relevantní k adopci technologií včetně AI)

WALSH, P. P. 2025. AI Literacy in Higher Education: A Systematic Review of Recent Frameworks 2023–2025. *Computers & Education: Artificial Intelligence*, 8, 100218. (Systematická přehledová studie integrující literární přehled rámců AI gramotnosti za období 2023–2025; identifikuje pět klíčových kompetenčních oblastí.)

WALTER, Y. 2024. Embracing the future of AI in the classroom: the relevance of AI literacy, prompt engineering, and critical thinking in modern education. *International Journal of Educational Technology in Higher Education*, 19(3).

WEISE, K., & METZ, C. 2023. When A.I. chatbots hallucinate. *The New York Times*, 1. května 2023. (Novinový článek popisující fenomén halucinací u AI chatbotů a uvádějící odhady odborníků, že zhruba čtvrtina výstupů může obsahovat vymyšlené informace; diskutuje také rizika plynoucí z přílišné důvěry v tyto modely.)

White, J., et al. 2023. A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT. *arXiv:2302.11382*.

WHITE, Julia, et al. 2023 Generative AI in Education: Opportunities and Risks. Computers & Education: Artificial Intelligence [online]. 2023, 4, 100147. Dostupné z: <https://doi.org/10.1016/j.caeai.2023.100147>.

WILLIAMSON, B. 2017. Big Data in Education: The Digital Future of Learning, Policy and Practice. London: Sage. ISBN 978-1473948006.

WIRTH, N. 2006. Good Ideas, Through the Looking Glass. IEEE Computer, 39(1), 28–39.

WHITE, J., FU, Q., HAYS, S., SANDBORN, M., OLEA, C., GILBERT, H., ELNASHAR, A., SPENCER-Smith, J., & SCHMIDT, D. C. 2023. A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT. arXiv preprint arXiv:2302.11382.

World Economic Forum 2022. Without universal AI literacy, AI will fail us. Agenda article, 2022. (Článek publikovaný na webu WEF zdůrazňující potřebu všeobecné AI gramotnosti ve společnosti; použita citace v úvodu kapitoly k podpoře tvrzení o nutnosti, aby se každý učil efektivně pracovat s AI.)

YIN, R. K. 2018. Case Study Research and Applications: Design and Methods. 6th ed. Los Angeles: Sage. ISBN 978-1506336169.

ZHAO, H., WALLACE, E., FENG, S., KLEIN, D., & SINGH, S. 2021. Calibrate Before Use: Improving Few-Shot Performance of Language Models. Proceedings of ICML 2021, 12697–12706. (Práce ukazující mimo jiné, že výkon velkých jazykových modelů je vysoce citlivý na formulaci promptu a že drobné úpravy phrasingu mohou značně ovlivnit výsledky – podporuje strategii experimentování s různým zněním dotazu.)

**© KOMUNIKAČNÍ ZRALOST V INTERAKCI S
GENERATIVNÍ AI: Empirická analýza uživatelských vzorců v
edukačním kontextu**

Autor:

JUDr. PhDr. PaedDr. Ing. Michal Hanák PhD., Ph.D.,
MBA, LL.M., MSc.

Recenzenti:

Prof. PaedDr. Lenka Pastejňáková, PhD. MBA
Prof. Ing. Ivan Hanuliak CSc.

Vědecký redaktor

doc. PhDr. Mgr. Janka Bursová, PhD., MBA

Vydavatelství

Akademie krizového řízení a managementu s.r.o.,
Pivovarská 273, 686 01 Uherské Hradiště
IČ: 050 24471

První vydání v nákladu 400 kusů.

Počet stran, autorské listy: 214 s. (9,86 AAA)

© Michal Hanák, 2026

ISBN 978-80-88398-40-0

VZDĚLÁNÍ MÁ SMYSL.

Monografie přináší komplexní pohled na to, jak lidé komunikují s generativní umělou inteligencí a jak se jejich komunikační vzorce proměňují v závislosti na zkušenostech, vzdělávání a účelu použití. Autor propojuje teoretické poznatky s empirickým výzkumem a představuje koncept komunikační zralosti uživatele jako nový přístup k porozumění efektivnímu využívání AI v edukačním kontextu.

Na základě analýzy uživatelských interakcí, případových studií a srovnání různých přístupů k práci s AI nabízí kniha praktické poznatky pro studenty, pedagogy, výzkumníky i manažery vzdělávacích institucí. Ukazuje, že technologie sama o sobě nestačí – rozhodujícím faktorem je člověk, jeho schopnost klást správné otázky, kriticky myslet, ověřovat výstupy a dlouhodobě rozvíjet své kompetence.

*Technologie otevírá nové možnosti.
Smysl dávají lidé.*

Michal Hanák se dlouhodobě věnuje problematice vzdělávání, managementu a využití moderních technologií ve vzdělávacím procesu. Ve své práci propojuje vědecký výzkum s praxí a usiluje o to, aby vzdělávání bylo dostupné, kvalitní a smysluplné pro každého.

ČLOVĚK
—
TECHNOLOGIE
—
VZDĚLÁVÁNÍ
—
SMYSL
—
BUDOUCNOST



Akademie krizového řízení a managementu s.r.o.
Pivovarská 273, 686 01 Uherské Hradiště
IČ: 050 24471

© Michal Hanák, 2026

ISBN 978-80-88398-40-0

