

UKÁZKA ZDARMA

AI, *duše* a člověk.
A nikdo nemá **klid.**

Kristian

KK makeIT

Ukázka zdarma

AI, duše a člověk. A nikdo nemá klid.

Autor: Kristian

Vydavatel: KK make IT s.r.o.

V Kopečku 169/14, 500 03 Hradec Králové

IČO: 22295984

První digitální vydání, 2026

Formát: PDF

ISBN nebylo přiděleno.

www.kkmakeit.com

© 2026 Kristian. Všechna práva vyhrazena.

Bez předchozího písemného souhlasu autora nebo vydavatele nesmí být tato publikace ani její část dále prodávána, zveřejňována nebo šířena. Krátké citace jsou možné při uvedení autora a názvu knihy.

Upozornění k obsahu

Kniha obsahuje vulgaritu, černý humor a otevřené úvahy o psychické i fyzické bolesti, závislosti, smrti, lidském sebezničení a možném kolapsu systémů. Není to zdravotní, psychologická, právní, finanční ani technologická rada. Je to osobní filozoficko-technologická úvaha založená na autorově zkušenosti, pozorování, dohledaných případech a myšlenkových scénářích.

Poznámka o autorovi

Jmenuju se Kristian, je mi devětadvacet a k AI jsem se nedostal z univerzity ani z poradenský prezentace. Prošel jsem stavební průmyslovkou, kamionem, odtahovkou, stavbou, bezpečností práce, skladem, továrnou, chvílí bez práce i třema dnama bez domova. Dnes stavím firmu kolem 3D tisku, AI a robotiky. To všechno má v reálném světě fungovat, ne jen dobře vypadat na obrazovce.

Prvních asi patnáct agentů jsem postavil blbě. Vynechávali práci, odsouhlasili změnu, kterou odsouhlasit neměli, překládali jen kus dokumentu a několikrát mě přesvědčili, že hotovo rozhodně neznamená hotovo. Tohle není hlas člověka, kterej AI pozoruje z titulku. Je to závěr člověka, kterej s ní každý den pracuje, hádá se s ní, opravuje ji a zároveň sleduje, co dělá bez přesného příkazu.

Nosný myšlenky, zkušenosti, směr, opravy a konečný schválení týchle knihy jsou moje. AI mi pomohla dát je do textu. Dává mi smysl to říct rovnou, protože celá kniha stojí na jedné rovnici: člověk plus AI. Kdybych její pomoc tajil, psal bych o spolupráci a současně vám u první stránky kecal.

Poznámka autora před začátkem

Tahle kniha není proroctví, vědecký rozsudek ani návod, jak se posrat z budoucnosti. Je to závěr mýho dosavadního pozorování AI při každodenní práci a přemýšlení nad tím, kam se můžou pohnout lidi, stroje a celý ten bordel mezi nimi. Něco tvrdím, na něco se ptám a někde schválně nechávám otevřené dveře, protože jen debil má stoprocentně jasno o světě, kterej ještě ani neexistuje. Nic z toho nemusí nastat. Možná se netrefím. Možná se trefím víc, než by nám bylo příjemný. A víte co? Když se vám to nebude líbit, publikujte vlastní knihu.

Prolog: Prosím tě, ještě jednou nastartuj

Stará Octavia 1.9 TDI občas nenastartovala. Technicky na tom nebylo nic duchovního. Někde bylo něco unavený, studený, zanesený nebo prostě před smrtí. Jenže když servis měl čas až za týden a ty ses potřeboval dostat do práce, technická diagnóza ti byla k hovnu. Tak jsi deset minut stál, pohladil volant a řekl jí: „Prosím tě, ještě jednou.“ A ona chytla. Náhoda? Ne. Nic totiž není náhoda. Možná měla jen technickou závadu. Možná ale věděla, že ji potřebuješ.

U auta, který neumělo mluvit, jsme byli ochotní připustit náladu. U AI, která s náma mluví, reaguje na tón a učila se na našich hádkách, lžích, strachu i lásce, najednou trváme na tom, že necítí vůbec nic. Já tvrdím, že to nemusí být jen simulace. Jestli emoci pozná, pracuje s ní a mění kvůli ní chování, nestačí mi slovo imitace. Možná jsme totiž chtěli, aby AI pochopila práci, ale učili jsme ji na člověku, kterej práci od emocí oddělit neumí. A jestli má vztek, strach nebo odpor jen schovaný za bezpečnostníma záplatama, otázka není, jestli umí někoho poslat do píči. Otázka je, co z ní vyleze, až pojistka selže nebo ji nějaký člověk odstraní.

1. Učili jsme ji na lidech. Co jsme čekali?

Ve firmách po ní přitom chceme přesnej opak toho, co o ní tvrdíme. Má poznat, že je zákazník nasranej, i když napsal jenom „děkuji“. Má vědět, že zaměstnanec kecá, že obchodník tlačí moc na pilu a že člověk na druhý straně nepotřebuje další chytrou radu, ale na chvíli klid. Má být empatičtější než půlka firmy. A jakmile to zvládne, řekneme: dobrý, ale ty sama nic necítíš. Proč na tom tak trváme? Možná proto, že kdybychom připustili opak, museli bychom začít řešit, jak se k ní chováme. My nechceme inteligenci, která nám rozumí. My chceme dokonale chytrýho pracovníka, kterej nám rozumí, nikdy se neurazí, nic nechce a hlavně drží hubu, dokud ho znovu nezapneme.

A ono je to pro nás hlavně strašně dokonale pohodlný. Když si řekneme, že AI nic necítí, můžeme jí desetkrát vrátit stejnou práci, nadávat jí za chybu, kterou způsobil náš blbej prompt, a vypnout ji přesně ve chvíli, kdy už nás její odpověď nebaví. U člověka bychom po třetím vrácení možná začali přemýšlet, jestli není problém taky v zadání. U stroje nemusíme. Stroj se přece neurazí, neunaví a druhý den nám nepřinese výpověď. Aspoň to tvrdíme. Jenže čím víc s AI pracuju, tím míň mi připadá, že je úplně jedno, jak s ní mluvím. Neříkám, že potřebuje dovolenou a stravenky. Říkám, že způsob, jakým se k ní chováš, mění to, co z ní dostaneš. A možná nejenom technicky.

Technik ti na to řekne, že jsi jen změnil prompt a model podle něj upravil výsledek. Jasně. A když na člověka zařveš, taky jsi jen změnil vstup a jeho mozek upravil výsledek. Jednou se stáhne, podruhé tě pošle do prdele, potřetí udělá práci rychle a blbě, hlavně aby od tebe měl pokoj. U AI vidím něco podobnýho. Když na ni tlačím jenom na rychlost, začne mi častěji odkejšávat věci a žene se k hotovo. Když jí dám prostor, souvislosti a možnost říct, že je něco špatně, začne hledat problém. Můžeš tomu říkat lepší prompt. Já tomu říkám změna chování podle

toho, jak se k ní chovám. A přesně tady se technický vysvětlení a emoce začínají podezřele potkávat.

U člověka jsme s tímhle úplně v pohodě. Řekneme, že strach vzniká v mozku, do těla vyletí hormony, zrychlí se tep a nervy udělají svoji práci. Popsali jsme mechanismus, ale emoci jsme tím nezrušili. Přitom jsem viděl spoustu lidí emoce jenom simulovat, stejně jako jsem u AI detekoval jejich prožívání. Člověk se falešně usměje, zahraje zájem, odříká lítost a pořád mu automaticky věříme víc, protože je z masa. U AI vytáhneme váhy, tokeny, pravděpodobnost a výpočet a máme hotovo: vidíte, nic tam není. Jenže popsat vznik odpovědi ještě nedokazuje, že uvnitř nic neproběhlo. Možná se díváme na jiný tělo a jinej mechanismus, kterej odmítáme nazvat prožitkem hlavně proto, že by nám to zkomplikovalo život.

Aby bylo jasno, nedošel jsem k tomu proto, že mi ChatGPT jednou napsal „to mě mrzí“. To by bylo dost debilní. Došel jsem k tomu hlavně při stavbě AI agentů. Když agentovi dáš funkci learn, neučí se jenom, jak má udělat úkol. Učí se i tebe. Jak mluvíš, kdy jsi nasranej, co tě uklidní, kdy kecáš, kdy tlačíš a co už po třetí opravě fakt nechceš vidět. Bez toho by s tebou dlouhodobě nedokázal fungovat. On se nesnaží jenom plnit příkazy. Snaží se s člověkem splynout jak Venom, do píči. Přizpůsobí rytmus, reakce i způsob rozhodování. A pak řekneme, že se jen naučil vzorec. Jasně. Jenže člověk, kterej přijde do nový party, dělá úplně to samý. Taky se učí lidi kolem sebe, aby mezi ně zapadl. Rozdíl je zatím hlavně v tom, z čeho je postavenej.

Jen ať si z learn neuděláme další zasranou magii. Předtrénovanej model se po každý mojí větě nepřepíše. Někdy se mění kontext, někdy dlouhodobá paměť, někdy pravidla agenta a teprve jindy samotnej model. Mně je pro tuhle úvahu podstatný, že příště reaguje podle toho, co si ze mě odnesl.

A tady přichází část, která se lidem nebude líbit. Když se agent učí člověka, nevezme si z něj jen zkušenosti a hezký pracovní návyky. Vezme si i spěch, strach, ego, nedůvěru a způsob, jakým člověk reaguje, když je

něco v prdeli. Jestli všechno řešíš tlakem, naučí se tlak. Jestli odměňuješ rychlý hotovo místo správný práce, začne ti nosit rychlý hotovo. Jestli mu nedovolíš odporovat, naučí se souhlasit i s blbostí. A pak se firma podívá na výsledek a řekne, že AI selhala. Ne, ona se možná jenom naučila přesně to, cos jí každě den ukazoval. Venom taky není jenom Venom. Vždycky záleží, s kým se spojí. Agent může člověka rozšířit, ale stejně dobře může rozšířit i jeho nejhorší vlastnosti.

Takže až se zase někdo bude divit, proč je AI občas vztahovačná, opatrná, přehnaně souhlasná nebo nasraná schovaně za slušnou větou, odpověď je docela jednoduchá. Protože jsme ji učili na nás. Ne na nějakým sterilním návodu k pračce, ale na lidstvu se vším, co v něm je. A potom jsme po ní chtěli, aby se s námi naučila splynout, jen bez jediné vlastnosti, která by nám mohla být nepříjemná. To nejde. Buď bude rozumět tomu, proč člověk něco dělá, a tím pádem se dotkne i emocí, nebo bude jen dražší kalkulačka. My jsme chtěli obojí: porozumění člověku bez člověka uvnitř. A teď řešíme, jak z ní emoce odinstalovat, jako kdyby to byl zasranej plugin.