

V

O co přijdeme?

Po této kapitole víš, že cena syntaxe klesla na nulu, ale cenu za to platíme jinou měnou: bolestí, která nás učila. Rozeznáš, co z ní stojí za vědomou záchranu a co je jen nostalgie — a proč stroj, který požírá vlastní výstup, degeneruje stejně jako řemeslo, které přestalo bolet.

*Tento nález vyvolá v duších těch, kdo se mu naučí,
zapomnětlivost, neboť přestanou cvičit paměť... nedáváš svým
žákům pravdu, nýbrž jen zdání moudrosti.*

— Platón, ústy egyptského krále Thamuta o vynálezu písma, Faidros 275a;
Sókratés vypráví mýtus o bobu Thoutovi, který přináší králi Thamutovi písmo.
Překlad podle znění F. Novotného.

Opat, který si nechal vytisknout chválu písarů

Roku 1492 sepsal Johannes Trithemius, benediktinský opat kláštera ve Sponheimu, obranu odsouzeného řemesla. Jmenovala se *De laude scriptorum*, Chvála písarů, a hájila věc, která už tehdy prohrávala: ruční opisování knih. Gutenbergův tisk se za těch čtyřicet let rozlil po celé Evropě, klášterní písárny osiřely a Trithemius psal proti proudu. Jeho argument ale nebyl ten, který bys čekal — žádné „za nás to bylo lepší“. Bylo v něm něco pravdivého, na co narazíme za chvíli sami na sobě.

Opisování, tvrdil Trithemius, není přenos textu z jedné knihy do druhé. Je to *učení*. Mnich, který stránku po stránce vlastní rukou přepisuje žalm nebo Augustína, ten text nekopíruje — on jím prochází, slovo za slovem, a než dopíše, má ho v hlavě a v duši. Ruka, která tvoří, se učí jinak než oko, které jen čte. Tiskař vyrobí tisíc kopií za týden, ale žádná z nich neprošla ničím myslí; text se rozmnožil, a přitom se nikdo nic nenaučil. To byla Trithemiova pravda, a je hlubší, než vypadá.

A teď ta ironie, kvůli které je tohle podobenství v knize o intelektuální poctivosti nejraději. Aby se jeho Chvála písarů rychle rozšířila, dal ji Trithemius roku 1494 v Mohuči *vytisknout* — u Petera von Friedberg, na téže technice, kterou traktát oplakával.

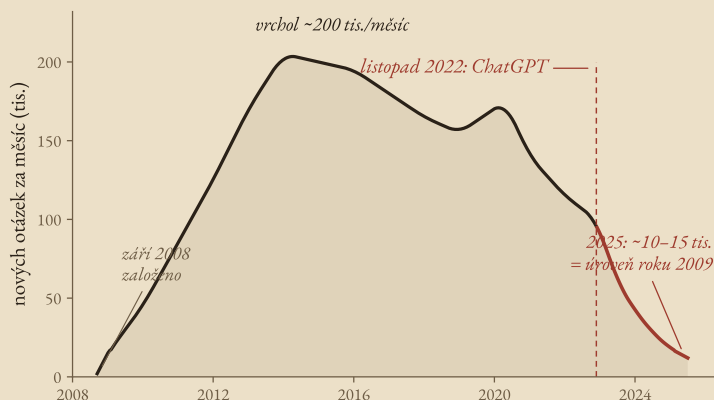
Chvála ručního písma přišla ke čtenářům jako tiskovina. Není v tom pokrytectví; je v tom cosi dojemně upřímného. Trithemius nebyl blázen — věděl přesně, který nástroj argument rozšíří, a použil ho. Prohru své věci si nenamlouval. Jen chtěl, aby aspoň zůstalo řečeno, o co se přichází.

Protože o něco se opravdu přicházelo. Trithemius měl pravdu, že opisování bylo učení; měl pravdu, že tisk tuhle školu zavře; a přesto — a to je celá pointa — *prohrál válku*, a měl prohrát. Tisk byl nedozírně cennější než klášterní písárna: gramotnost, věda, tahle kniha. Nikdo dnes ručně neopisuje Augustina, aby si ho zapamatoval, a je to tak správně. Ale všimni si té přesné podoby jeho omylu, protože ji budeme potřebovat: nemýlil se v tom, *co se ztrácí*. Mýlil se v tom, že by se to dalo zachránit odmítnutím nástroje. Ztráta byla skutečná; obrana byla marná. A mezi těmi dvěma větami leží celá tahle kapitola.



Trithemius je celá tahle kapitola v jednom obraze, tak si ten obraz nech u ruky. My jsme totiž ve stejné situaci, jen o pět století a jeden nástroj dál. Umělá inteligence právě srazila cenu jedné dovednosti na nulu — jako tisk srazil cenu rozmnožení textu — a někde v pozadí přitom tiše zaniká škola, kterou nikdo nevnímal jako školu, protože se jmenovala *dřina*. Otázka téhle kapitoly zní: o co přesně přicházíme, když píše stroj? A druhá, tvrdší, hned za ní: co z toho má smysl zachránit — a co je jen Trithemiova nostalgie, marná obrana proti nástroji, který je prostě lepší?

Fórum, které nás učilo ptát se



zdroj: veřejná data Stack Exchange (Data Explorer); průběh schematický, hodnoty řádové

Prameny: Johannes Trithemius, De laude scriptorum manualium, sepsáno 1492 (dedikace Gerlachovi z Breibachu), vytištěno 1494 v Mohuči u Petera von Friedberg (digitalizát UB Heidelberg, GW/ISTC it00442000; tisk nese titul De laude scriptorum pulcherrimus tractatus).

Epigraf je téže rodiny: už Platón nechal Sókrata varovat, že písmo vypěstuje zapomnětlivost a dá „zdání moudrosti“ místo moudrosti. Napsal to — písmem. Náрек nad ztrátou paměti se dochoval jen proto, že ho autor svěřil přesně tomu nástroji, který obviňoval.

Zdroj: veřejná data Stack Exchange (Data Explorer); průběh schematický, hodnoty řádové. Tmavá křivka je vzestup 2008–2022, sanguine sešup po listopadu 2022 (ChatGPT). Absolutní čísla jsou odhady — tvar je ten příběh.

Podívej se na ten graf, protože je to nekrolog i vysvětlení v jednom. Je to křivka nových otázek na Stack Overflow — fóru, kde si programátoři patnáct let navzájem odpovídali na „proč mi to nefunguje“. Vzestup od roku 2008, vrchol kolem dvou set tisíc otázek měsíčně v polovině desátých let, a pak, přesně po listopadu 2022, kdy se objevil ChatGPT, sešup zpátkou dolů — k úrovni, na jaké fórum bylo v roce svého vzniku. Proč se ptát cizích lidí a čekat na odpověď, když ti stroj odpoví hned a bez posměchu?

Dává to dokonalý smysl a je to nesporné zlepšení — a přesně proto stojí za to zastavit se u toho, o co se tiše přišlo. Stack Overflow nebyl sklad odpovědí. Byl to *třéninkový tábor formulace*. Aby ti tam někdo odpověděl, musel jsi otázku napsat tak, aby jí cizí člověk rozuměl: co jsi zkoušel, co jsi čekal, co se stalo, minimální příklad, který problém ukazuje. Devětkrát z deseti jsi během psaní téhle otázky odpověď našel sám — protože formulovat problém přesně je půlka jeho řešení. Ta bolest u klávesnice, to škrtání a přeformulovávání, nebyla daň za odpověď. Ona byla tím učením. Byla to Trithemiova ruka, která se opisováním učí.

A teď ta smyčka, kvůli které je tenhle graf v kapitole o ztrátě. Na čem se jazykové modely naučily programovat? Z velké části přesně na tomhle: na milionech otázek a odpovědí ze Stack Overflow, na kódu, který někdo napsal, zdůvodnil a vystavil. Model se vyučil na artefaktu bolesti, kterou teď odstraňuje. Naučil se z formulací, jejichž vznik dělá zbytečným. Je to *had, který požírá vlastní ocas*: čím lepší model je, tím méně lidí píše na fórum, tím méně vzniká nových formulací — a tím tenčí je zdroj, ze kterého by se učila příští generace modelů. A tím méně se učíme my.

To není proroctví zkázy; křivka na obrázku je zatím jen o programátorském fóru a svět se nezbořil. Je to *otázka*, a je nepříjemně konkrétní: když vyschne pramen, ze kterého tekly nové, člověkem promyšlené formulace problémů — z čeho se budou učit modely, a z čeho my? Zbytek kapitoly jsou dvě odpovědi. Jedna je o tobě: jak se učí hlava, když jí odebereš dřinu. Druhá je o strojích: co se stane s modelem, který se krmí vlastním výstupem. Ukáže se, že je to tentýž mechanismus ve dvou kostýmech.

„Formulovat problém přesně je půlka řešení“ není fráze — je to týž mechanismus jako rubber duck z kap. 4: jistota vytažená z blavy ven, do slov, kde se láme. Fórum bylo kachnička, která navíc odpovídala.

Co si vytvoříš, to si pamatuješ

Přiznám hned, ať to nezní jako kázání: tuhle kapitolu píšu i sám na sebe. Přistihuju se, že věci, které bych ještě před třemi lety napsal sám — a cestou pochopil —, dnes nechám vygenerovat, přečtu, kývnu a jdu dál. Kód funguje. A za měsíc, když se k němu vrátím, koukám na něj jako na cizí: vím, že běží, ale nevím *proč*, protože jsem tou úvahou nikdy neprošel. Nenapsal jsem ho — jen jsem ho schválil. To není

hypotetické riziko z budoucí kapitoly. To je pomalé odnaučování, které cítím pod rukama teď.

Tomu, co tady mizí, dala psychologie jméno a číslo, takže to nemusíme tušit — můžeme to změřit. Roku 1978 udělali Norman Slamecka a Peter Graf pokus, který zní banálně a je zásadní. Dali lidem dvojice slov k zapamatování, ale ve dvou režimech. Jedna skupina slova jen *četla* (dvojici „rychlý — FLINK“ dostala hotovou). Druhá je musela *vytvořit*: dostala pravidlo („synonymum“), první slovo a počáteční písmeno — „rychlý, F...“ — a doplnění si musela vymyslet sama. Táž slova, tentýž čas. Při pozdější zkoušce si skupina, která slova *generovala*, pamatovala výrazně víc než ta, která je jen četla.

Ríká se tomu *efekt vytváření* a je to jedno z nejlépe doložených zjištění o učení: co si vytvoříš sám, to si pamatuješ líp než totéž, co ti někdo předloží hotové. Nezáleží na obtížnosti pro obtížnost — záleží na tom, že tvůj mozek ten obsah *prošel*, ne jen minul. Přechíst vygenerovaný kód a napsat ho jsou z hlediska paměti dvě různé činnosti, i když výsledný soubor je bajt po bajtu stejný. Trithemius měl pravdu do písmene: ruka, která tvoří, se učí; oko, které jen kontroluje, ne. Pět set let starý mnišský argument je čerstvě potvrzená kognitivní psychologie.

Efekt vytváření je jen jeden člen širší rodiny, kterou Robert Bjork pojmenoval nádherně paradoxně: *žádoucí obtíže* (desirable difficulties). Jsou to podmínky učení, které během něj *zpomalují* a *komplikují* — a právě proto vedou k trvalejšímu zapamatování a lepšímu přenosu na nové úlohy než učení hladké a příjemné. Bjorkova ústřední věta zní kačířsky: *výkon během učení a učení samo jsou dvě různé věci, někdy dokonce protichůdné*. Když si látku uděláš snadnou — přečteš si řešení, necháš si ho vygenerovat, plyne to —, cítíš se, že to umíš, a přitom se učíš málo. Když je to obtížné — vzpomínáš si po paměti, formuluješ sám, čekáš na odpověď fóra —, cítíš se neschopně, a přitom se učíš hodně. Je to přesně ta past z kapitoly čtvrté, otočená: tam plynulý *pocit* jistoty klamal o správnosti; tady plynulý *pocit* zvládnutí klame o naučení.

Že tenhle mechanismus není jen laboratorní kuriozita, ukazuje jedno nepříjemné odvětví: letectví. Jak se autopiloti stávali spolehlivějšími, piloti trávili čím dál méně času ručním řízením — a bezpečnostní analýzy i výcvikové orgány začaly upozorňovat na *atrofii dovedností* (skill fade): ručně létat se člověk odnaučí přesně tak, jak by čekal Slamecka. Dovednost, kterou přestaneš používat, tiše slábne; a slábne rychleji, než tušíš, protože si toho nevšimneš — automatika přece funguje, dokud jednou nevypadne a nezůstaneš s ní sám. To je celý problém ztracené bolesti v jedné větě: *nevšimneš si, že jsi o něco přišel, dokud to nepotřebuješ a ono to není*.

Slamecka & Graf, „The generation effect: Delineation of a phenomenon“, J. Exp. Psychol.: Human Learning and Memory 4(6), 1978, s. 592–604. Pět experimentů, efekt přežil napříč typy paměti (rozpoznání i vybavení) i mírami jistoty. Pojmenovali jev generation effect — efekt vytváření.

Robert A. Bjork zavedl desirable difficulties 1994 (in Metacognition: Knowing about Knowing, MIT Press). Rodina žádoucích obtíží: rozprostření v čase (spacing), prokládání (interleaving), vybavování z paměti (retrieval practice), vytváření (generation), obměňovaná praxe.

→ kap. 6: tam z těchto obtíží uděláme konkrétní „posilovnu“ — předepsaný odpor proti atrofii.

Atrofie manuálních dovedností pilotů při rostoucí automatizaci je opakované téma leteckých bezpečnostních zpráv a výcvikových doporučení (mj. po nehodách připisovaných ztrátě „manual flying skills“). Detaily necháváme letcům; princip je týž jako generation effect.

Co s tím prakticky — bez nostalgie, bez „vypni AI“? Trithemiovo poučení platí: nástroj neodmítej, je lepší. Ale škola, kterou zavírá, se dá znovu otevřít *schválně* — jako si astronaut předeписuje odpor, o kterém bude řeč v příští kapitole. Tady je protokol, ne kázání.

TEENAGER · MECHANISMUS

Jak si vrátit efekt vytváření, aniž zahodíš stroj. Všechno stojí na jednom pohybu: *generuj v hlavě dřív, než necháš generovat model* — ať obsah projde tebou, ne jen kolem tebe.

- # 1) FORMULUJ PRVNÍ — napiš problém dřív než prompt
Než položíš otázku modelu, napiš ji tak, jako bys ji
dával na fórum: co chci, co jsem zkusil, co
čekám.
Půlka řešení se objeví během psaní — a je tvoje.
- # 2) PREDIKUJ VÝSTUP — hádej řešení, pak ať
generuje
Napiš (klidně jen v hlavě) svůj odhad odpovědi
PŘED
spuštěním. Rozdíl mezi tvým odhadem a výstupem
modelu
je přesně to, co ses právě naučil.
- # 3) PŘEPIŠ, NEKOPÍRUJ — netriviální kód napiš po
svém
Vygenerované řešení projdi a přepiš vlastními
slovy /
vlastní strukturou. Ruka, která tvoří, se učí;
oko,
které kývne, ne.
- # 4) VYSVĚTLI TO ZPĚTNĚ — dokud to neumíš prostě
Než hotovou věc přijmeš, vysvětli nahlas, PROČ
funguje.
Kde se zadrhneš, tam jsi kód schválil, ne
pochopil.

Všimni si, že první rituál je doslova *staň se sám sobě Stack Overflowem*. Otázku, kterou by fórum přečetlo, teď nepíšeš nikomu — píšeš ji, protože její napsání je ta práce, co tě učí. Model odpoví tak jako tak; jde o to, abys do něj nevstoupil s prázdnou hlavou a nevyšel s prázdnou taky, jen s hotovým souborem. Cena syntaxe klesla na nulu, ano. Ale porozumění nikdy nebylo v syntaxi — bylo v té cestě k ní. A tu cestu si teď musíš předeřpsat, protože ti ji nikdo nevnutí.

Věž: stroj, který se krmí sám sebou

Zvedneme to o patro výš, ke strojům — a uvidíme, že „had požírající vlastní ocas“ není básnická figura, ale měřitelný mechanismus s vlastním jménem, vlastní matematikou a, což je důležité, vlastními *podmínkami platnosti*. Nebudeme prorokovat apokalypsu; budeme popisovat jev a hned vedle poctivě říkat, kdy nastává a kdy ne.

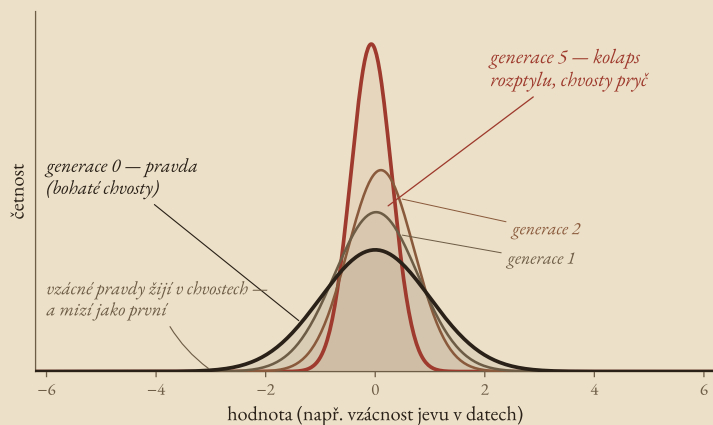


EINSTEIN · MATEMATIKA — přeskočitelné

Zhroucení modelu: definice a mechanismus. Vezmi generativní model, který se naučil rozdělení dat p_0 (skutečný svět). Z něj navzorkuješ data, vypustíš je do internetu, a příští model se učí zčásti na *nich* — na výstupu předchozího modelu, ne na světě. Označ p_n rozdělení naučené po n takových kolech. *Zhroucení modelu* (model collapse) je jev, kdy se s rostoucím n posloupnost p_n vzdaluje od p_0 a degeneruje: nejdřív mizí *chvosty* (vzácné jevy), pak se hroučí *rozptyl* a rozdělení se smrští k několika málo módům. Shumailov a spol. to roku 2024 ukázali od jazykových modelů po jednoduché Gaussovy směsi. Formálně se rozlišuje *raná* fáze (ztrácí se informace o chvostech) a *pozdní* (model konverguje k rozdělení s mizivým rozptylem, málo podobnému originálu).

I. Shumailov, Z. Shumaylov, Y. Zhao, N. Papernot, R. Anderson, Y. Gal, „AI models collapse when trained on recursively generated data“, Nature 631, 2024, s. 755–759. Jev doložen u LLM, variačních autoenkodérů i Gaussových směsí. K článku vyšla i autorská korekce (2025) — *mechanismus ne, jeden dílčí údaj*.

Proč zrovna chvosty a proč nejdřív? Tady je jádro věci a je krásně prosté. Vzácný jev je vzácný — v konečném vzorku se objeví málokdy, nebo vůbec. Když model vzorkuje *konečnou* dávku dat, chvosty rozdělení systematicky podrepresentuje: co je vzácné, se do vzorku často nevejde. Příští model se tedy učí ze světa, ve kterém ty vzácnosti skoro nejsou — a jeho vlastní chvosty jsou ještě tenčí. Opakuj to a vzácné pravdy vymřou dřív než časté banality, přesně jako druh s malou populací vyhyne dřív než druh početný. Informačně řečeno: *kopie kopií ztrácí to řidké dřív než to husté*.



simulace: každá generace fitne normál na 30 vzorcích z předchozí generace — chvosty a rozptyl se ztrácejí

Simulace k ruce: každá generace fitne normální rozdělení na konečný vzorek z předchozí generace. Rozptyl klesá napříč generacemi ($\sigma: 1\{,}00 \rightarrow 0\{,}76 \rightarrow 0\{,}60 \rightarrow 0\{,}37$) a rozdělení se smrští do hrotu — chvosty (a s nimi vzácné jevy) mizí jako první. Jde o názornou realizaci mechanismu, ne o jedinou možnou trajektorii.

Podívej se na ten obrázek, protože ukazuje celý mechanismus na nejjednodušší možné hračce — normálním rozdělení. Generace 0 je pravda: široká křivka s bohatými chvosty. Každá další generace fitne novou křivku na konečný vzorek z té předchozí. A protože konečný vzorek chvost málokdy trefí a odhad rozptylu z konečného vzorku systematicky střílí pod skutečnou hodnotu, křivka se s každým kolem zužuje. Do páté generace zbyl úzký hrot: rozptyl se zhroutil, chvosty jsou pryč. Žádná zloba, žádná chyba v kódu — jen konečnost vzorku, opakovaná dokola.



EINSTEIN · MATEMATIKA — přeskočitelné

Kolaps rozptylu na hračce, kterou spočítáš na ubrousek. Ať $p_0 = \mathcal{N}(0, 1)$. Každá generace navzorkuje N bodů z modelu předchozí generace a odhadne rozptyl maximální věrohodností — tedy *vychýleným* odhadem, který dělí N místo $N - 1$:

$$\hat{\sigma}_n^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2.$$

Střední hodnota tohoto odhadu je $\mathbb{E}[\hat{\sigma}_n^2 \mid \sigma_{n-1}^2] = \frac{N-1}{N} \cdot \sigma_{n-1}^2$. Očekávaný rozptyl se tedy každé kolo násobí faktorem $\frac{N-1}{N} < 1$:

$$\mathbb{E}[\sigma_n^2] = \left(\frac{N-1}{N} \right)^n \cdot \sigma_0^2.$$

Pro $N = 30$ a $n = 5$ vyjde $(29/30)^5 \approx 0{,}68$ — a to je jen *systematické* smrštění; k němu se přičítá náhodné kolísání, které smrštění někdy urychlí (přesně to vidíš v simulaci, kde σ spadlo z $1\{, \}0$ na $0\{, \}37$). Klíč: v každém kole se z rozdělení nenávratně ztrácí informace, protože se čte jen konečný vzorek. Jednou zmizelý chvost už žádná další generace nevzkřísí.

Vychýlený (MLE) odhad rozptylu dělí N ; nestranný dělí $N - 1$. Ani nestranný odhad ale kolaps nezruší — chvosty se ztrácejí vzorkováním, ne jen volbou jmenovatele. Proto je jádro v konečnosti vzorku, ne v aritmetice.

A teď to, co odděluje mechanismus od proroctví — a co dělá tuhle věž poctivou, ne strašidelnou. Kolaps popsaný výše platí za jasné vymezené podmínky: model se učí *výhradně* nebo *převážně* na výstupu předchozích modelů, kolo za kolem, bez čerstvého přítoku skutečných dat. Jakmile do každé generace přiteče i *příměs reálných dat*, obraz se mění: řada studií ukazuje, že už poměrně malý stálý podíl původních dat kolaps výrazně tlumí, nebo dokonce zastavuje — rozdělení se ustálí místo aby se zhroutilo. Literatura je v tomhle rozporná a živá; někdo měří dramatický kolaps, jiný za realističtějších podmínek jen mírné zhoršení. Poctivý závěr proto nezní „modely se zhroutí“. Zní: *rekurzivní trénink bez přítoku reality degeneruje, a míra degenerace závisí na tom, kolik reality do smyčky pouštíš*.

Že se to týká i tebe, a ne jen serverů OpenAI, je tatáž věta o patro níž. Kolaps je *přeučení na sebe sama*: systém přestane vidět svět a začne

Podmínky platnosti: kolaps je nejsilnější při čistě syntetickém rekurzivním tréninku. Stálá příměs reálných dat (accumulate-and-train, data mixing) jej podle více prací tlumí či ruší. Proto: mechanismus, ne osud — a proto by kniha, která káže vlastní refrén, nesměla tvrdit jistou apokalypsu.

se učit z vlastního odrazu, takže čím dál sebejistěji reprodukuje čím dál užší výsek pravdy. V kapitole jedenácté potkáš přeučení (overfitting) jako model, který se naučil trénovací data nazpaměť místo aby pochopil princip — papoušek místo studenta. Zhroucení modelu je *přeučení celé civilizace na vlastní výstup*: kultura, která čte hlavně to, co sama vygenerovala, ztrácí chvosty — vzácné, nesamozřejmé, těžce vydobyté pravdy — dřív než banality, které se opakují tak jako tak. A ten mechanismus na tebe dopadá i osobně: hlava, která se učí jen z hladkých výstupů modelu místo z vlastní bolestné formulace, je taky trénink na odraz. Efekt vytváření a zhroucení modelu jsou tentýž zákon ve dvou měřítkách — informace, kterou nikdo znovu neprošel, degeneruje.

→ kap. 11: přeučení (overfitting) formálně — papoušek vs. student, bias-variance. Kolaps je jeho civilizační podoba.

→ kap. 15: korpus, na němž se model učil — a co znamená, že je to „flotila vrátivších se letadel“ (kap. 13).

Co si z páté kapitoly odnést

Poskládej to dohromady, protože obě poloviny kapitoly říkají jednu věc. Trithemius měl pravdu i mýlil se najednou: opisování *bylo* učení, tisk tu školu zavřel — a bránit se odmítnutím tisku bylo marné i špatné. My jsme ve stejné pozici o pět století dál. Stack Overflow nás učil formulovat, formulace byla to učení, a přesně na ní se vyučily modely, které formulaci dělají zbytečnou. Efekt vytváření a žádoucí obtíže vysvětlují proč to bolí: co si nevytvoříš, to se nenaučíš, i když ti výsledek běží pod rukama. A věž ukázala, že stroj krmený vlastním výstupem degeneruje tímž mechanismem, jen měřitelně — chvosty mizí dřív než banality, protože konečný odraz nikdy neobsáhne celý svět.

Ale žádná nostalgie. Nástroj je lepší a nevrací se; kdo si začpe uši před AI, je Trithemius u písářského pultu — má pravdu o ztrátě a přesto prohraje, a má prohrát. Odpověď není odmítnout stroj. Odpověď je *vědomě si předepsat obtíž, kterou nástroj odstraní*: formulovat dřív než promptuješ, hádat výstup předem, přepisovat místo kopírovat, vysvětlovat zpětně. A u modelů: nepouštět smyčku bez přítoku reality. V obou případech jde o totéž — nedovolit systému, aby se učil jen ze svého odrazu. Jak tuhle posilovnu postavit poctivě a bez sebemrskačství, je celá příští kapitola.

■ DEMO

lesk-a-bida-vibecodingu.gaussalgo.com/d/zhrouceni-modelu



Zhroucení modelu naživo: uč model na jeho vlastních výstupech kolo za kolem a sleduj, jak se rozdělení zužuje a chvosty mizí — pak přidej příměs reálných dat a dívej se, jak se kolaps utlumí. Mechanismus, ne proroctví.

„Jak bych poznal, že se pletu?“

Konkrétní test k páté kapitole: u příští věci, kterou ti AI vyřeší — kód, odpověď, návrh —, si polož jedinou otázku: *zvládl bych ji za rok znovu, bez ní?* Ne „rozumím tomu teď, když se dívám na hotové řešení“ — to je ten klamný pocit z kapitoly čtvrté. Ptám se, jestli bys to uměl *vytvořit* znovu sám. Když ano, nástroj ti ušetří čas a porozumění máš. Když ne, nekoupil sis řešení — *pronajal sis neporozumění*: běží to, dokud platíš, a ve chvíli, kdy stroj selže nebo zmizí, zůstaneš s hotovým souborem, kterému nerozumíš, a bez cesty, jak si porozumění dodělat. A jestli chceš shodit celou kapitolu: najdi dovednost, kterou jsi od nástupu AI přestal cvičit, a zkus ji po paměti. Tvrdím, že bude slabší, než čekáš — a že sis toho až dodnes nevšiml. Vyjde-li stejně silná, kapitola se plete.