

II

To už
tu bylo?

Po této kapitole poznáš hype cyklus v přímém přenosu, umíš vyjmenovat obě velké AI zimy a hlavně přesně říct, co je na dnešní vlně stejné jako na všech předchozích a co je doopravdy jiné. A dostaneš do ruky nástroj, kterým každý příští „průlom“ prosvítíš na dřevě: detektor Whitehouse/Kelvin.

Kdo si nepamatuje minulost, je odsouzen k tomu, aby si ji zopakoval.

— George Santayana, „Those who cannot remember the past are condemned to repeat it.“ · *The Life of Reason*, sv. I, 1905, kap. XII

Kabel, který zabily oslavy

Pátého srpna 1858 se stalo něco, co do té chvíle patřilo do říše zázraků: dva světadíly si promluvíly v reálném čase. Po dně Atlantiku ležel měděný kabel dlouhý přes tři tisíce kilometrů a šestnáctého srpna po něm královna Viktorie poslala prezidentu Buchananovi blahopřání. Osmadevadesát slov — a přenos té krátké zprávy trval bezmála šestnáct hodin. Nikomu to nevadilo. Amerika se zbláznila radostí: zvonily zvony, střílelo se z děl, konaly se průvody. V New Yorku slavili tak divoce, že se oslavný ohňostroj vymkl kontrole a v noci ze sedmnáctého na osmnáctého srpna *skutečně zapálil radnici* — shořela kupole, střecha i socha Spravedlnosti na vršku newyorské City Hall. Kabel spojil dva světy a v přímém přenosu podpálil vlastní oslavu. Lepší předzvěst si dějiny nemohly vymyslet.

Protože ten kabel byl už tehdy nemocný. Vyrobili ho a skladovali nedbale, gutaperčová izolace popraskala, signál po něm dorazil slabý a rozmazaný. A tady vstupuje na scénu člověk, kterého tahle kapitola potřebuje vyprávět poctivě, ne jako padoucha: Wildman Whitehouse, hlavní elektrikář projektu. Byl to lékař a nadšený amatér, žádný šarlatán — dělal, co uměl nejlíp, a věřil své metodě. A jeho metoda zněla logicky: je-li signál slabý, přidej sílu.

Zapojil velké indukční cívky a hnal do slábnoucího kabelu napětí kolem dvou tisíc voltů, podle některých odhadů i vyšší. Kabel to nezachránilo — dorazilo to už tak nemocné vedení. Do měsíce oněmělo nadobro.

Následoval pád tvrdý jako vzestup. Tisk, který kabel ještě před pár týdny velebil, teď psal o humbuku a podvodu; objevily se dokonce hlasy, že žádné zprávy z Ameriky nikdy nepřišly a všechno to byl švindl. Investoři utekli, projekt na léta zamrzl. Přišla zima.

Roztála až o osm let později, a roztál ji rozdíl v *metodě*. Roku 1866 položila obří loď Great Eastern nový, poctivěji vyrobený kabel — a u přístrojů stál William Thomson, fyzik, který o telegrafii uvažoval jinak než Whitehouse. (Rytířem ho královna pasovala právě za tehle úspěch v roce 1866; slavný titul lorda Kelvina dostal až mnohem později, roku 1892 — tady je pořád ještě prostě Thomson.) Thomson nevěřil na křik vyšších voltů. Vsadil na opak: na *šepot*, který dokázal přechytit. Zkonstruoval zrcadlový galvanometr — přístroj tak citlivý, že nepatrný proud v cívice pohnul zrcátkem a to vrhlo světelný paprsek na vzdálenou stupnici, kde i tichounký signál povyroste do čitelného výkyvu. Whitehouse chtěl slabý signál překřičet. Thomson ho chtěl líp *slyšet*. Vyhrál ten druhý — a od té chvíle spolu oba světadíly mluví bez přestání.



Ten příběh není o telegrafu. Je o *tvaru*, který uvidíš znovu a znovu, jakmile ho jednou rozpoznáš: zázrak, který předběhne své vlastní pochopení. Nejdřív průlom a oslava, pak přehnané sliby, pak náraz na realitu, pak zima a řeči o podvodu — a teprve po ní tichá, nezáživná práce, která věc konečně rozchodí. Těhle tvar má jméno, *hype cyklus*, a AI ho za svou krátkou historii prošla už nejmíní dvakrát. Tahle kapitola je o tom naučit se ho číst — abys u příští oslavy poznal, jestli stojíš u kabelu z roku 1858, nebo u toho z roku 1866.

Jaro, zima, jaro — v přímém přenosu

Umělá inteligence nezačala ChatGPT. Začala jako sen a ten sen má přesné datum zrození. Roku 1950 napsal Alan Turing stať, ve které se ptal „Mohou stroje myslet?“ — a protože ta otázka je beznadějně mlhavá, nahradil ji hrou: *brou na napodobování*. Když u dálkopisu nepoznáš, jestli ti odpovídá člověk, nebo stroj, přestává podle Turinga mít smysl upírat stroji myšlení. Všimni si, co udělal: nevyřešil filozofii, obešel ji. Nezodpověditelnou otázku „myslí to?“ proměnil v měřitelnou „poznáš rozdíl?“. To je pohyb, který se v téhle knize bude opakovat — a je to zároveň první drobná trhlinka, kterou dnešek nasvítí: Turingova hra

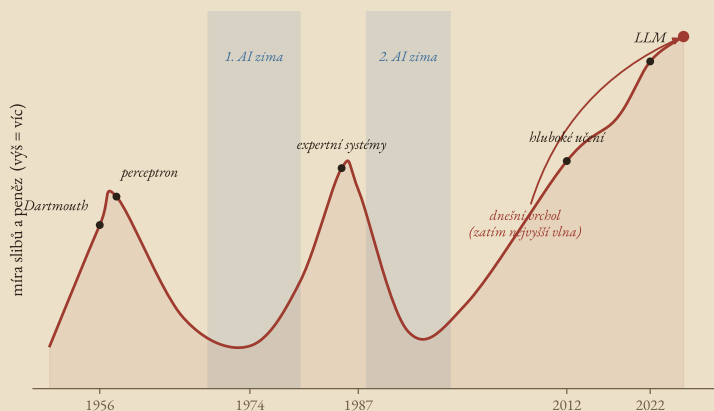
Zrcadlový galvanometr Thomson zkoušel už na kabelu 1858 (na irské straně) — roku 1866 to nebyla novinka, ale vítězná metoda. Zpráva Viktorie z roku 1858 měla 98 slov a její přenos trval bezmála 16 hodin; přes kabel z roku 1866 by proletěla v minutách.

Whitehouse vs. Thomson není příběh blupáka a génia. Oba byli poctiví, oba dřeli. Lišila se metoda: jeden přidával sílu do systému, druhý citlivost do měření. Tenhle rozdíl si zapamatuj — je to nástroj, který z téhle kapitoly odneseš.

měří *přesvědčivost*, ne porozumění, a to jsou, jak uvidíš, dvě úplně různé věci.

Za semínkem přišlo první jaro. V létě 1956 se v Dartmouthu sešla hrstka matematiků a v žádosti o grant napsali větu, která se stala legendární svou sebejistotou: věřili, že *podstatný pokrok* v otázkách jazyka, abstrakce a sebezdokonalování strojů lze udělat, „když na tom pečlivě vybraná skupina vědců bude společně pracovat jedno léto“. Jedno léto. Ta věta je dokonalá miniatura celé kapitoly: nadšení, které si spletlo *vidím cíl za jsem u cíle*. Cesta k tomu cíli trvá dodnes.

Vrchol prvního jara přišel roku 1958 a vypadá skoro k neuvěření podobně jako newyorská oslava kabelu. Psycholog Frank Rosenblatt představil *perceptron* — jednoduchou síť, která se učila třídit obrázky. Osmého července o něm *New York Times* napsal pod titulkem „Nové námořní zařízení se učí činem“, že jde o zárodek elektronického počítače, který *bude umět chodit, mluvit, vidět, psát, rozmnožovat se a být si vědom své existence*. Chodit, mluvit, být si vědom — z jedné vrstvy umělých neuronů. Ohňostroj u radnice.



Právě v Dartmouthu se poprvé napsal termín umělá inteligence (John McCarthy). Mezi organizátory byl i Claude Shannon z kapitoly 15 — týž, který o pět let dřív měřil s manželkou entropii angličtiny.

Křivka není předpověď budoucnosti — je to paměť minulosti. Každý vrchol slibil myslící stroj; každá zima přišla, když sliby narazily na realitu. Vlny sílí (víc peněz, víc pozornosti), ale tvar se opakuje. Dnešní vrchol je zatím nejvyšší — což neznamená „konec dějin“, jen „zatím nejvyšší vlna“.

Pak přišla první zima. Roku 1969 dokázali Marvin Minsky a Seymour Papert, že Rosenblattův jednovrstvý perceptron má tvrdý formální strop — a věž na konci kapitoly ukáže přesně jaký. Roku 1973 shrnul britský matematik James Lighthill pro tamní vládu roky nesplněných slibů do zprávy, po níž financování oboru vyschlo. Nadšení, které slibovalo vědomí z jedné vrstvy neuronů, narazilo na to, co jedna vrstva neuronů opravdu umí. Zima trvala roky.

A „zima“ tu není básnická licence — je to popis skutečnosti. Vysychají granty, mizí laboratoře, chytří lidé z oboru odcházejí, samo slovo se stává málem sprostým, takže i ti, kdo v něm dál pracují, radši píší „strojové zpracování dat“ než „umělá inteligence“. Zima neznamená, že se přestalo bádát; znamená, že se přestalo *věřit* — a s vírou odeče kapital i pozornost. Tohle si drž v hlavě, až budeš vážit dnešní vlnu:

sestup nepřichází proto, že by nástroje přestaly fungovat, ale proto, že sliby doběhly dál než výsledky a rozdíl si někdo nakonec spočítá.

Roztálo druhé jaro, na jiném příslibu. V osmdesátých letech dozrály *expertní systémy* — programy, do kterých experti ručně vlévali tisíce pravidel typu „když teplota a tlak, pak porucha“. Fungovaly, prodávaly se, firmy do nich lily peníze. A pak narazily i ony: pravidla se nedala udržet, křehla, každá výjimka znamenala další ruční záplatu. Kolem roku 1987 přišla *druhá AI zima*. Zase sliby, zase náraz, zase mráz.

Stojí za to zastavit se u toho, *proč* expertní systémy narazily, protože v tom je schovaný zárodek dneška. Vědění do nich vkládal člověk, pravidlo po pravidle, ručně — a svět je bohatší, než kolik pravidel stihne kdokoli napsat; na každou napsanou výjimku číhá deset nenapsaných. Dnešní modely jdou opačnou cestou: nikdo do nich pravidla nekládá, ony si je *samy vytáhnou* z obrovského množství příkladů (to je celá kapitola patnáctá). Kde expertní systém dostal pravidla darem a došel mu dech u výjimek, dnešní model se pravidla učí sám — a přesně to mu dovolila škála, první z těch tří věcí, které jsou jiné.

Vidíš ten vzorec? To je celý smysl grafu nahoře. Neříká, že AI je podvod — to by byla stejně hloupá chyba jako pohřbívat obor po každé zimě. Říká něco strážlivějšího: nadšení kolem myslících strojů *nikdy nebylo přímka*. Bylo to vlnobítí a my zrovna stojíme na jeho zatím nejvyšším hřebeni. A na hřebeni vlny se špatně soudí, jestli je to konečně příliv, nebo jen další vlna, po níž přijde odliv. Přesně proto potřebuješ znát tvar celé pláže.

Poctivost velí přiznat, že „AI je za rohem“ je nejstarší předpověď oboru. Slibovala se roku 1956 na jedno léto, roku 1958 z jedné vrstvy neuronů a od té doby prakticky každý rok. Občas dorazí věci ohromující — jen skoro nikdy ta, co byla slíbená, a skoro nikdy podle jízdního řádu.

Co je stejné a co je jiné

Kdyby tahle kapitola skončila u „všechno se opakuje, buď skeptik“, byla by levná a hlavně by lhala. Protože něco je tentokrát opravdu jiné, a poznat *co* přesně, je celé umění. Rozděl to natvrdo na dvě hromádky.

Na hromádce *stejně* leží věc, kterou už znáš: křivka slibů a peněz. Přehnaná očekávání, titulky o vědomí, obří investice hnané spíš pocitem než měřením — to všechno jsme viděli v roce 1958 i 1985 a vidíme to dnes. Kdo tvrdí, že tentokrát hype neexistuje, nečetl graf. Sem patří i to, že nejhlasitější sliby zase mívají dál, než kam nástroj dohlédne.

Na hromádce *jiné* leží tři věci, a stojí za to je pojmenovat přesně, protože právě ony dělají dnešní vlnu nejvyšší v historii. Nejde o rozdíl

ve stupni nadšení — nadšení bylo obří i v roce 1958. Jde o rozdíl v tom, *co ta vlna doopravdy nese*.

☹ TEENAGER · MECHANISMUS

Tři věci, které jsou tentokrát doopravdy jiné:

- | | |
|---|---|
| 1. ŠKÁLA
obrázků. | Perceptron 1958 se učil na hrstce
obrázků. |
| nesou | Dnešní modely přečtou biliony slov a
stovky miliard parametrů. Ne „víc
téhož“ —
o deset řádů víc. Řád mění povahu
věci. |
| 2. PŘESVĚDČIVOST | Expertní systém vyplivl strohé
„PORUCHA“. |
| vtipně, a zní | Dnešní model mluví plynně, česky,
jistě i když se mýlí. Přesvědčivost
je NOVÁ |
| veličina — a nezávisí na pravdě. (→
kap. 3) | |
| 3. DOSTUPNOST | Perceptron byl skříň v laboratoři
námořnictva. |
| Dnešní model je zadarmo v každé
kapse, hned, | pro každého. Ne exponát — nástroj v
ruce miliard. |

Whitehousův kabel byl jeden a v ruce hrstky inženýrů. Dnešní „kabel“ drží v ruce každý — a každý u něj dělá tutéž volbu jako Whitehouse: víc síly, nebo lepší měření?

Tyhle tři rozdíly nejsou důvod k naivitě, ale ani k jejímu opaku. Škála je skutečná a dělá věci, které dřív nešly. Přesvědčivost je skutečná — a je to past, protože náš mozek bere plynulou řeč jako důkaz mysli za slovy (o tom je celá příští kapitola). A dostupnost je skutečná: ať už vlna opadne, nebo ne, nástroj v kapse ti zůstane. Právě proto nemá smysl ptát se „je to bublina, ano nebo ne“. Užitečnější otázka zní jinak — a tady se rodí nástroj, kvůli kterému tahle kapitola existuje.

Detektor Whitehouse/Kelvin

Vrať se ke kabelu, protože v něm je schovaný přenosný nástroj myšlení. Whitehouse i Thomson stáli před tímž problémem — slabý, nespolehlivý signál — a každý ho řešil opačným směrem. Whitehouse přidal *sílu*

do systému: víc voltů, hlasitěji, překřič to. Thomson přidal *citlivost do měření*: lepší přístroj, který i tichý signál přečte. Jeden bojoval hlukem, druhý sluchem. A vyhrál sluch.

Z toho rozdílu se dá udělat otázka, kterou přiložíš k *jakémukoli* řešení, nástroji nebo ohlášenému průlomů — vlastnímu i cizímu.

☹ TEENAGER · MECHANISMUS

Detektor Whitehouse/Kelvin. U každého nového „průlomů“ nebo návrhu řešení se zeptej:

Přidává tohle SÍLU do systému,
nebo CITLIVOST do měření?

víc síly (Whitehouse) (Thomson/Kelvin)	lepší měření
---	--------------

víc voltů	citlivější přístroj
víc parametrů „na sílu“	umím poznat, kdy to selže
hlasitější tvrzení	umím to nezávisle ověřit
„přidáme výkon“	„změříme, jestli to funguje“

Síla není nutně špatně a měření není samospásné — často potřebuješ obojí. Ale jakmile řešitel sahá *jen* po síle a nemá čím výsledek změřit, díváš se na Whitehouse: na někoho, kdo do už tak nemocného kabelu žene další volty.

Přilož ten detektor na dnešek a začne pracovat okamžitě. Firma ohlásí model s desetkrát větším počtem parametrů — to je síla; ptej se, o kolik líp *měří* svou vlastní chybu, jestli vůbec. Někdo ti ukáže oslnivé demo — ptej se, jestli změřil, kdy selže, nebo jen přidal hlasitosti. A hlavně, obrať ten detektor na sebe: když ti model vyhodí výsledek a ty s ním nejsi spokojený, sáhneš po delším, důraznějším promptu (víc voltů), nebo si postavíš způsob, jak nezávisle poznat, jestli je výstup správně (citlivější měření)? První cesta je Whitehousova a vede do němoty. Druhá je Thomsonova a mluví přes oceán.

Tím se dostáváme k druhé polovině čtenářského řemesla téhle kapitoly: umět odlišit *demo* od *produkce*. Protože přesně tady se svádí dnešní bitva slibů.

Jak číst benchmark a demo proti produkční realitě. Než uvěříš ohlášenému „umí to“, projdi si čtyři otázky:

1. VIDĚL JSEM PRODUKCI, NEBO DEMO?

Demo je vybraný nejlepší běh na vybraném příkladu.

Produkce je tisíc běhů na vstupech, které nikdo nevybral.

2. NA ČEM SE TO MĚŘILO?

Benchmark je trať. Zvládnout trať ≠ zvládnout terén.

Ptej se, jestli trať náhodou nebyla v tréninkových datech.

3. KDO VYBÍRAL PŘÍKLAD?

Když příklad vybral prodejce, vybral ten, na kterém to září.

4. JAK POZNÁM, ŽE TO SELHALO?

Když na to neumíš odpovědět, neviděl jsi funkci – viděl jsi ohňostroj.

Tenhle inventář není cynismus. Je to Thomsonův přístroj přeložený do práce s dnešními stroji: místo abys hlasitost dema bral jako důkaz, postavíš si měření, které dohlédne až tam, kam ohňostroj nesvítí — na produkci, kde věc běží ve tmě a bez potlesku.

Věž: čím se tahle vlna liší, spočteno

Zatím jsme mluvili obrazem. Teď vylezeme na věž a řekneme přesně, *proč* dnešní vlna dokáže to, co perceptron nedokázal — a kde má naopak tvrdé stropy, které žádné nadšení neposune. Tři výhledy z věže: proč jedna vrstva neuronů selhala, proč škála funguje kvantitativně, a proč se každá technologie nakonec ohne do téhož tvaru.

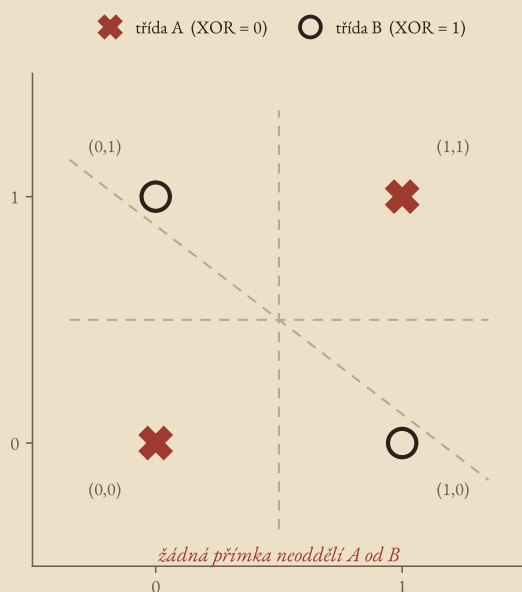


Proč přímka neoddělí XOR. Rosenblattův jednovrstvý perceptron počítá jedinou věc: zváží vstupy, sečte je a porovná s prahem. Geometricky to znamená, že do prostoru vstupů umí nakreslit jedinou *přímku* (obecně nadrovinu) a rozdělit jí body na dvě třídy. Cokoli, co jde takhle oddělit, se naučí. Cokoli jiného ne.

A hned nejjednodušší zajímavá funkce se takhle oddělit *neda*: XOR, tedy „právě jedno ze dvou“. Její tabulka je

$$0 \oplus 0 = 0, \quad 0 \oplus 1 = 1, \quad 1 \oplus 0 = 1, \quad 1 \oplus 1 = 0$$

a když ty čtyři body zakreslíš do roviny, dvě jedničky leží v protilehlých rozích a dvě nuly v druhých dvou. Žádná přímka je neoddělí — ať ji nakloníš jakkoli, jeden bod vždycky zůstane na špatné straně.



Minsky & Papert, Perceptrons (1969). Pozor na rozšířený mem: netvrdili, že limit platí i pro vícevrstvé sítě — naopak věděli, že sít s vrstvou navíc XOR zvládne. Zabíla obor domněnka, že se takové sítě nedají trénovat. Trvalo do 80. let, než algoritmus zpětného šíření chyby ukázal, že dají — a druhé jaro moblo začít.



Lekce XOR přežila padesát let a platí i pro dnešní obří modely: síla jedné vrstvy je omezená *tvar*em, ne velikostí. Přidat víc neuronů do jedné vrstvy XOR nevyřeší; musí přibýt vrstva, která skládá jednoduché dělicí čáry do složitějších hranic. Odtud slovo *hluboké* v „hlubokém učení“: hloubka, vrstva na vrstvě, je ta věc, která dělá rozdíl — ne pouhá šířka. Zapamatuj si to jako protijed na hype: „větší“ samo o sobě není „chytřejší“, dokud nevíš, jaký *tvar* problému ta velikost zvládne.

Tím se dostáváme k druhému výhledu z věže — a k jedinému místu, kde dnešní vlna stojí na tvrdé kvantitativní půdě, jakou žádná předchozí neměla.



EINSTEIN · MATEMATIKA — přeskočitelné

Škálovací zákony. Roku 2020 změřila skupina kolem Jareda Kaplana něco pozoruhodného: chyba jazykového modelu neklesá s velikostí nahodile, ale po *mocninném zákoně*. Označ L ztrátovou funkci (křížovou entropii z kapitoly 15 — průměrné překvapení na token). Pak zhruba platí

$$L(N) \approx (N_c/N)^{\alpha_N}$$

kde N je počet parametrů modelu, N_c a α_N konstanty. Slovy: zvětší model desetkrát a chyba klesne o předvídatelný kousek — a tatáž pravidelnost platí zvlášť pro množství dat i pro výpočetní výkon. Poprvé v dějinách oboru šlo *dopředu spočítat*, o kolik bude větší model lepší. Nadšení dostalo pod nohy křivku, ne jen pocit.

Kaplan et al., „Scaling Laws for Neural Language Models“ (2020). Pozoruhodné je, že *mocninný vztah drží přes sedm řádů výpočetního výkonu — vzácná pravidelnost pro tak mladý obor.*



EINSTEIN · MATEMATIKA — přeskočitelné

Roku 2022 tým DeepMind (Hoffmann et al., model *Chinchilla*) ukázal, že se dosavadní vlna škálovala *špatně*: honila počet parametrů a šetřila na datech. GPT-3 nesl 175 miliard parametrů, ale přečetl jen zhruba 300 miliard tokenů — necelé dva tokeny na parametr. Chinchilla dokázal, že při daném rozpočtu se parametry a data mají zvětšovat *úměrně*, zhruba

$$D_{\text{opt}} \approx 20 \cdot N$$

tedy kolem dvaceti tokenů na každý parametr. Menší model nakrmený víc daty porazil většího hladovce. Všimni si, co je tohle za druh objevu: není to „přidali jsme sílu“ (víc parametrů), je to „změřili jsme, co doopravdy pomáhá“. Chinchilla je Thomson téhle kapitoly — vyhrál lepším měřením, ne hlasitějším modelem.

Hoffmann et al., „Training Compute-Optimal Large Language Models“ (2022).
→ kap. 15: ta klesající L je táž křížová entropie, kterou si v Shannonově hře naměříš na sobě. Pokrok modelů = pokles průměrného překvapení, měřený na textech, které model neviděl.

Ten rozdíl mezi Whitehousem a Thomsonem se ve věži vrací v čisté podobě. Kaplan přinesl *sílu s křivkou*: ukázal, že přidávání parametrů a dat funguje a jde předpovědět. Chinchilla přinesl *měření*: ukázal, že dosavadní přidávání síly bylo špatně vyvážené, a spočítal správný poměr. Obojí je potřeba. Ale všimni si, komu patří poslední slovo — tomu, kdo lépe změřil, ne tomu, kdo hlasitěji přidával. To není náhoda téhle kapitoly; to je její teze převlečená za dějiny strojového učení.

Zbývá poslední výhled — ten, který ti brání spadnout z věže do opačného extrému, totiž věřit, že když křivka roste po mocninném zákoně, poroste tak navěky.



S-křivka difuze. Žádná technologie neroste donekonečna. Šíření sleduje *logistickou křivku* ve tvaru písmene S:

$$f(t) = \frac{K}{1 + e^{-r(t-t_0)}}$$

Zkraje pomalu (málokdo o věci ví), pak prudce vzhůru (všichni ji berou), a nakonec do nasycení u stropu K (trh je plný). Ta prostřední, strmá část se *zevnitř* tváří jako exponenciála bez konce — a přesně tam se dělají přehnané předpovědi, protože nikdo neví, kde leží K . Perceptron i expertní systémy vypadaly na své strmé části jako začátek nekonečna; obě to byly esíčka mířící ke svému stropu. Škálovací zákony jsou tvrdý fakt *uvnitř* dnešní křivky; kde má tahle křivka svoje K , dnes nikdo nezměřil. Proto věž nekončí věštbou, ale detektorem: až uvidíš příští čáru mířící k nebi, ptej se, jestli ti někdo *změřil* její strop — nebo ti jen ukazuje strmou část a doufá.

Co si z druhé kapitoly odnést

Poskládej to dohromady. Kabel z roku 1858 ti dal tvar, který se opakuje: zázrak, sliby, náraz, zima, a teprve pak tichá práce, co věc rozhodí. AI ten tvar prošla nejmíň dvakrát — perceptron a jeho zima, expertní systémy a jejich zima — a dnes stojíme na zatím nejvyšším hřebeni. Rozdělili jsme, co je *stejné* (křivka slibů a peněz) od toho, co je *jiné* (škála, přesvědčivost, dostupnost). A z Whitehouse a Thomsona jsme vydestilovali nástroj, který ti zůstane, i kdybys zapomněl všechno ostatní: ptej se, jestli řešení přidává *sílu*, nebo *citlivost měření*.

Věž pak ukázala, že dnešní vlna stojí na tvrdší půdě než ty minulé — má změřené škálovací zákony — ale i ona je nejspíš esíčko s neznámým stropem, a limity tvaru (XOR!) žádná velikost neobejde.

Tři nitě si ponesem dál. Přesvědčivost, kterou jsme sotva pojmenovali, je celá příští kapitola: proč nás plynulá řeč strhne uvěřit myslí, která za ní není. To, co dnešní model doopravdy umí navíc oproti Markovovu řetězu z roku 1913, rozebere kapitola patnáctá. A detektor Whitehouse/Kelvin se vrátí v úplném finále knihy, kde z něj bude jeden z pilířů toho, komu a jak moc věřit.

→ kap. 3: přesvědčivost jako nová veličina — plynulá řeč jako spouštěč víry v mysl za slovy.

→ kap. 15: co přesně transformer umí navíc oproti Markovovu řetězu.

→ kap. 22: detektor Whitehouse/Kelvin jako test důvěryhodnosti zdroje.

„Jak bych poznal, že se pletu?“

Konkrétní test k druhé kapitole: vezmi úplně příští „průlom“ v AI, o kterém uslyšíš — nový model, nové demo, nový titulěk — a než mu uvěříš, polož mu dvě otázky. První: *je to víc voltů, nebo citlivější galvanometr?* Přidal někdo jen sílu a hlasitost (víc parametrů, sebejistější tvrzení, oslnivější ukázka), nebo přidal způsob, jak nezávisle *změřit*, že to doopravdy funguje a kdy selže? Druhá: *viděl jsem produkci, nebo demo?* Byl ten příklad vybraný prodejcem na jeho nejlepším běhu, nebo to běželo ve tmě, na vstupech, které nikdo nevybral? Tvrdím, že když na obě otázky nedokážeš odpovědět něčím konkrétním, nedíváš se na průlom — díváš se na newyorský ohňostroj, který je nádherný přesně do chvíle, než ti zapálí radnici. A když odpovědět dokážeš, právě jsi udělal to, co Thomson: přestal jsi křičet a začal poslouchat.