

IV

Proč pocit dřiny
nic nedokazuje?

Po této kapitole rozlišíš prožitek porozumění od testovaného porozumění — a víš, že čím víc ses nadřel, tím víc si věříš, ale ne tím víc máš pravdu. Znáš iluzi hloubky vysvětlení a umíš si na její dno sáhnout sám: dřív, než ti to udělá realita. Kde kapitola třetí odhalila lháře venku, table odhalí toho uvnitř.

Prvním principem je, že nesmíš oklamat sám sebe — a ty sám jsi ten nejsnáze oklamatelný člověk.

— Richard Feynman, „*The first principle is that you must not fool yourself — and you are the easiest person to fool.*“ · „*Cargo Cult Science*“, proslov k promoci na Caltechu, 1974

Paprsky, které viděla jen Francie

Roku 1903 ohlásil René Blondlot, respektovaný fyzik z univerzity v Nancy, objev nového druhu záření. Pojmenoval ho po svém městě: N-paprsky. Vycházely prý skoro ze všeho — z rozžhaveného drátu, z lidských nervů, dokonce z cihly ohřáté na slunci — a poznaly se podle toho, že nepatrně *zjasnily* elektrickou jiskru, na kterou pozorovatel hleděl v temné komoře. Byl to objev na hraně vnímatelnosti: rozdíl v jasu, který se dal spíš *tušit* než změřit.

A Francie ho tušila. Během jediného roku vyšlo kolem tří set vědeckých článků od zhruba sto dvaceti autorů, kteří N-paprsky „potvrdili“, změřili jejich vlnové délky, popsali jejich lom v hliníkovém hranolu. Byla to velká, poctivá, nadšená práce spousty schopných lidí. Jen s jednou vadou: mimo Francii je nikdo spolehlivě nezahlédl. Němci nic neviděli. Angličané nic neviděli. Něco tu nehrálo — a tak se do Nancy vypravil Američan Robert Wood, fyzik z Johns Hopkins a známý bořitel nesmyslů, aby se na ty paprsky podíval zblízka.

To, co při demonstraci udělal, je jádro celé téhle kapitoly. Seděli v zatemněné laboratoři a Blondlot mu ukazoval, jak hliníkový hranol rozkládá N-paprsky do spektra, které jeho asistent odečítal

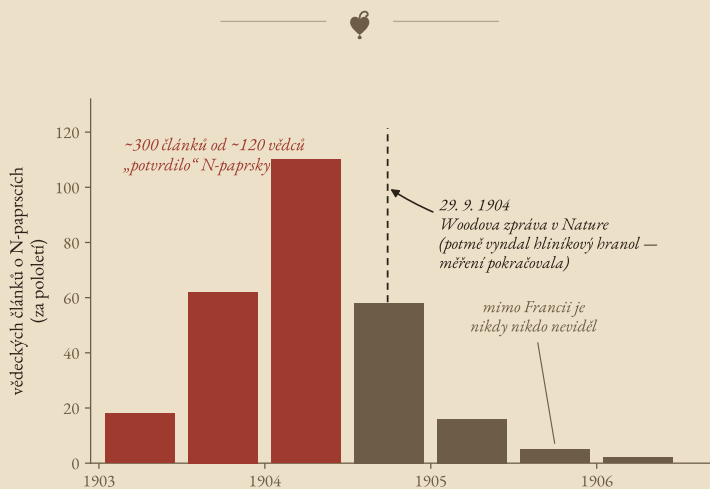
Prameny: R. W. Wood, „The n-Rays“, Nature 70, 29. 9. 1904 (dopis datován 22. 9.). Rozsah „300 článků od 120 vědců“ dle přehledů o N-paprscích; scéna s vyndaným blinikovým hranolem je doložená, ne apokryfní.

z fosforeskujícího stínítka. Wood v té tmě, kde nikdo nic neviděl, nenápadně *vyndal ten hranol z aparatury a schoval si ho do kapsy*. Klíčovou součástku, bez níž by žádné spektrum vzniknout nemohlo. A Blondlotův asistent odečítal dál — hlásil tytéž polohy čar, tytéž hodnoty, jako by se nic nestalo. Wood pak pro jistotu zaměnil i ocelový pilník, který měl paprsky vysílat, za kus dřeva, jež je vysílat nemohlo. Měření pokračovala. Paprsky se „viděly“ dál.

Woodova zpráva vyšla v časopise Nature 29. září 1904 a N-paprsky zabila jedinou větou o tom, že několik experimentátorů, kteří dosáhli kladných výsledků, bylo „nějakým způsobem oklamáno“. Všimni si toho slova: *oklamáno*, ne *lhalo*. Protože to je na celém příběhu nejdůležitější — a nejnepříjemnější. Blondlot nepodváděl. Nikde neschovával dráty, nefalšoval čísla, nechtěl slávu za lež. Byl to schopný člověk, který strávil rok poctivě, vysílající práce nad efektem, o němž byl hluboce přesvědčen — a to přesvědčení mu jeho vlastní oči poslušně potvrzovaly, i když se před nimi měřicí aparát tajně rozpadl. Fyzik Irving Langmuir tenhle druh omylu později pojmenoval *patologická věda*: věda, kterou nesráží podvod, ale sebeklam schopných lidí.

Dohra má tvar, který k příběhu patří. Francouzská akademie věd udělila Blondlotovi za rok 1904 Prix Leconte s odměnou padesát tisíc franků — a to *i poté*, co Wood zveřejnil svou zprávu. Formálně za celoživotní dílo, ne za N-paprsky; ale gesto mluvčí. InSTITUTE, národ, rok cti — to všechno stálo na jedné straně vah. Na druhé stál jeden Američan s hranolem v kapse a s otázkou, kterou si Blondlot nikdy nepoložil: *jak bych poznal, že paprsky vidím, i když tam nejsou?*

Patologická věda (Irving Langmuir, přednáška 1953): jevy na hranici vnímatelnosti, „potvrzené“ mnoha pozorovateli, které mizí, jakmile se zprísni podmínky. N-paprsky jsou její učebnicový případ — vedle nich stojí třeba studená fúze 1989.



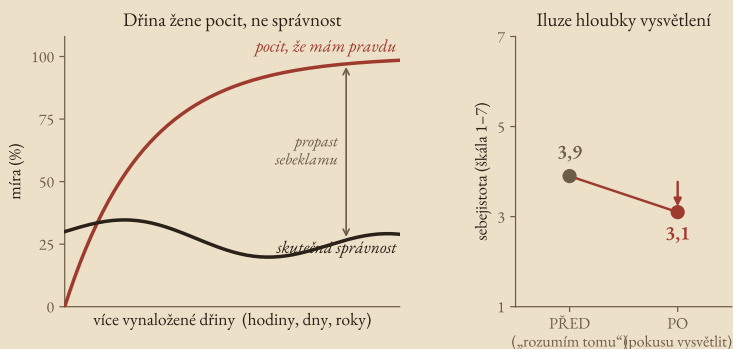
Blondlot obléhl N-paprsky 1903; celkem ~300 článků od ~120 vědců. Příběh za pololetí schematický (řádový), zlom po Woodově zprávě (Nature, 29. 9. 1904) je doložený. Akademie věd přesto podržela Blondlotovi Prix Leconte za rok 1904.

Než pojedeme dál, postav Blondlota vedle koně z minulé kapitoly, protože ti dva vypadají skoro stejně — a nesmějí se splést. Hans nic neuměl a lidé v něm *viděli mysl*, protože jim mimoděk zrcadlil jejich vlastní vodítka: vnější svět (kůň) mlčel a oni v tom mlčení slyšeli hlas. To byla kapitola třetí. Blondlot je druhá strana téhož zrcadla. Tady nejde o to, že venku něco mlčí; tady jde o to, že *uvnitř* tebe pracuje mocná síla, která se jmenuje vynaložené úsilí — a ta ti šeptá, že když jsi do věci vložil tolik poctivé dřiny, *musí* to být pravda. Kůň četl tebe. Blondlota zradil on sám. A ta druhá zrada je zákeřnější, protože proti hlasu zvenčí se dá aspoň otočit; proti vlastnímu pocitu dřiny se otočit nedá — cítíš ho zevnitř a máš ho za důkaz.

Dva roky do záchodu

Tohle je moje jizva a musím ji přiznat hned, jinak by celá kapitola zněla jako kázání shora. Dělal jsme kdysi na jednom projektu dva roky. Dva roky poctivé, vyčerpávající práce: noci, víkendy, refaktoring, schůze, diagramy na tabuli. A celou tu dobu jsme se cítili *skvěle* — jako lidé, kteří dřou na něčem těžkém a důležitém, a čím víc jsme dřeli, tím jistější jsme si byli, že jdeme správným směrem. Ta jistota rostla přesně s tou dřinou. Rostla, protože jsme se nadřeli, ne protože jsme měli pravdu. Na konci se ukázalo, že v základu byla malá, hloupá, triviální chyba — invariant, který jsme celou dobu předpokládali, a on prostě neplatil. Jeden jediný špatný předpoklad, na kterém stály ty dva roky. A nejhorší na tom není ta chyba. Nejhorší je, že jsem se *celou dobu cítil, jako bych té věci rozuměl do hloubky* — a ten pocit byl úplně, dokonale nespolehlivý. Byl jsem Blondlot. Jen jsem místo paprsků viděl smysluplný projekt, a místo hranolu v kapse jsem měl chybějící test, který by mě zastavil první týden.

Tomu jevu, který mě stál dva roky, dala psychologie jméno: *effort heuristic*, česky heuristika úsilí. Je to tichý pravidlo, které máš zabudované stejně hluboko jako vidění tváří v mracích: *co mě stálo víc práce, musí mít větší hodnotu a být blíže pravdě*. Většinou ti to pravidlo slouží dobře — věci, do kterých vložíš úsilí, opravdu bývají cennější. Ale je to *heuristika*, zkratka, a zkratky se dají obelstít. Když měříš správnost pocitem vynaložené dřiny, měříš úplně jinou veličinu, než si myslíš. Měříš, kolik tě to stálo — ne, jestli je to pravda. A tyhle dvě věci spolu nemusejí mít nic společného.



Vlevo: pocit správnosti roste s dřinou, skutečná správnost ne (effort heuristic). Vpravo: průměrná sebejistota „rozumím“ spadne, jakmile se člověk pokusí věc do detailu vysvětlit (Rozenblit & Keil, Cognitive Science 2002; pokles z ~3,9 na ~3,1).

Levý panel je moje jízva v obraze: sanguine křivka (pocit, že mám pravdu) strmě roste s vynaloženou dřinou, tmavá křivka (skutečná správnost) zůstává dole, protože ji dřina bez testu nezvedne. Rozevírající se propast mezi nimi jsou ty dva roky. Cítil jsem tu horní křivku a myslel si, že je to ta spodní.

Podívej se na levý panel toho obrázku pořádně, protože je to anatomie každého dlouhého omylu. Vodorovná osa je dřina — hodiny, dny, u mě roky. Horní křivka je *pocit*, že mám pravdu, a ta poslušně stoupá: čím víc do věci nasypeš práce, tím jistěji jí věříš. Spodní křivka je *skutečná* správnost, a ta se ani nehne, protože práci samotnou nezajímá, jestli je tvůj výchozí předpoklad pravdivý — zvedne ji jedině test, který by mohl selhat, a ten jsme nespustili. Mezi těmi dvěma křivkami se otevírá propast, a přesně v ní bydlí sebeklam. Blondlot byl na jejím dně. Já taky. A tady je ta pointa, kvůli které to celé vyprávím: *v té propasti se cítíš úplně stejně jako na vrcholu skutečného porozumění*. Zevnitř se pravda a sebeklam nedají odlišit. To je celý problém.

Umíš vysvětlit splachovací záchod?

Teď zkusíme něco, co ti to dokáže na vlastní kůži, a nebude to trvat dva roky, jen minutu. Vezmi věc, o které jsi přesvědčen, že jí rozumíš — třeba obyčejný splachovací záchod. Používáš ho denně, viděl jsi ho tisíckrát, klidně bys řekl, že mu rozumíš na sedm z deseti. Fajn. A teď to *vysvětli*. Nahlas nebo na papír, krok za krokem, od stisknutí páčky po znovunapuštění: jak přesně to, že zmáčkneš kliku, způsobí, že se voda vyvalí, spláchně, a pak sama zase zastaví na správné hladině. Ne „splachovadlo to nějak udělá“ — mechanismus. Ventil po ventilu. Plovák po plováku.

Většina lidí se v půlce zadrhne. Zjistí, že netuší, jak se ten odtok utěsní zpátky, nebo co přesně zastaví přítok. A tady je to překvapení: *ta jistota, že rozumím, byla skutečná — dokud jsem ji nezkusil rozbalit do kroků*. Zmizela přesně ve chvíli, kdy jsem měl předvést mechanismus, který jsem si celou dobu myslel, že mám v hlavě. Tomuhle se v psychologii říká *iluze hloubky vysvětlení* a je to nejlepší kapesní přístroj na měření vlastního sebeklamu, jaký znám.

Iluze je nejsilnější právě u vysvětlovacích znalostí (jak něco funguje) — méně u faktů, postupů a příběhů. Proto „umím vyjmenovat kroky“ neklame tak jako „chápu, proč to funguje“. A přesně tu drubou, klamavější věc od tebe chce každý netriviální kód.

Pravý panel hrdinského obrázku ukazuje, co se stane s číslem. Leonard Rozenblit a Frank Keil dali roku 2002 lidem seznam běžných věcí — zip, záchod, rychloměr, klavírní klávesu — a nechali je ohodnotit na škále od jedné do sedmi, jak *dobře* rozumějí, jak fungují. Pak řekli: teď to napište. Do detailu, krok za krokem. A pak se ohodnoťte znovu. Průměrná sebejistota spadla — z nějakých tří celých devíti na tři celé jedna — jen tím, že se lidé pokusili vyslovit to, o čem byli chvíli předtím přesvědčeni, že to vědí. Nepřečetli si nic nového. Nikdo je neopravil. Jen sáhli na dno vlastní jistoty a zjistili, že je mčlčí, než mysleli. To je celá metoda: *porozumění se netestuje pocitem, testuje se pokusem vyslovit ho nahlas*.

Rozenblit & Keil, „The misunderstood limits of folk science: an illusion of explanatory depth“, Cognitive Science 26(5), 2002, s. 521–562. Škála 1–7, pětifázový postup; pokles sebejistoty T1 → T2 přežil replikace napříč populacemi.

Rituály proti sobě samému

Dobrá zpráva zní stejně jako v minulé kapitole, jen otočená dovnitř. Proti sebeklamu z dřiny nepomáhá „snažit se víc soustředit“ ani „být poctivý k sobě“ — právě poctiví lidé jako Blondlot na jeho dno padají nejlouběji. Pomáhá jen *ritual*: pevný postup, který uděláš vždycky, i když ti pocit říká, že je zbytečný, protože přece *víš*. Ten pocit je nepřítel; rituál je stěžeň. Tady jsou čtyři, které stojí za to mít v ruce.

Čtyři rituály, kterými přistihneš vlastní sebeklam dřív, než tě to bude stát dva roky. Všechny stojí na jednom pohybu: *dostaň jistotu ven z hlavy* — do řeči, na papír, do měřitelného výsledku — kde se dá zkontrolovat.

1) KACHNIČKA – vysvětli to nahlas cizímu člověku

Vysvětluj svůj problém (nebo řešení) do posledního kroku kolegovi, který o něm nic neví – a když není po ruce, gumové kachničky na monitoru. Kde se zadrhneš, tam jsi jen CÍTIL, že rozumíš.

2) PREDIKCE PŘEDEM – napiš, co uvidíš, PŘED spuštěním

Než pustíš test / experiment / dotaz na model, zapiš na papír, co čekáš, že vyjde. Po výsledku ti paměť jinak namluví „to jsem tušil“ a ukradne ti důkaz.

3) DENÍK OMYLŮ – zapisuj, v čem ses spletl
Krátká věta ke každému zamítnutému nápadu: co jsem

čekal, co vyšlo. Sbírka vlastních omylů je jediná mapa, na které je vidět, kde tě pocit klame nejčastěji.

4) LEVNÉ ZAHOZENÍ – commit a větev PŘED experimentem

Ulož si pozici (commit) a odboč na větev, než se pustíš do „přepíšu to celé“. Pak smíš zahodit dva měsíce práce bez ztráty tváře – sunk cost tě nedrží.

První rituál je nejmocnější a je to doslova malý domácí Rozenblit-Keilův test. Vezmi člověka, který o tvém problému nic neví — kolegu od vedlejšího stolu, partnera u večeře — a *vysvětli mu ho nahlas*. Ne shrnout, vysvětlit: každý krok, každé „protože“. V programátorské tradici se tomu říká *rubber duck debugging*, ladění s gumovou kachničkou, protože posluchač nemusí ani rozumět — stačí, že mluvíš nahlas ke komusi vně své hlavy. Řeč je bezcitný přístroj: donutí tě vyslovit i ta „protože“, která jsi měl za samozřejmá, a přesně na nich se sebeklam obvykle láme. Kolikrát jsem uprostřed věty ke kachničce zmlkl a pochopil, kde je bug — dřív, než mi ho stačil ukázat počítač.

Druhý rituál je nenápadný a nesmírně tvrdý: *napiš předpověď dřív, než uvidíš výsledek*. Ne v hlavě — na papír, do komentáře, do zprávy

u commitu. Důvod je krutý a už jsi ho zažil, i když sis toho nevšiml: jakmile uvidíš výsledek, paměť ti šalebně přepíše, cos čekal, na to, co se stalo, a ty budeš mít neochvějný pocit, že jsi to *celou dobu věděl*. Zpětný pohled ukradne experimentu veškerou důkazní sílu. Jediná obrana je zapsat predikci předem — černé na bílém, dřív, než realita promluví. Když pak výsledek vyjde jinak, nemůžeš si namluvit, žeš to tušil; máš to napsané, že ne. A právě tuhle zapsanou predikci pokřtí kapitola devátá jako jádro vědecké metody a kapitola osmá jako *pakt se sebou samým* — první z paktů, které si dáš, dokud jsi ještě střízlivý, protože v přímém přenosu prohraješ.

Čtvrtý rituál si zaslouží slovo navíc, protože léčí bolest, kterou znáš z mého příběhu: *nechtěl jsem ty dva roky zabít*. Čím víc jsi do něčeho vložil, tím těžší je to opustit, i když je jasné, že je to slepá ulička — ekonomové tomu říkají *sunk cost*, utopené náklady, a je to effort heuristic ve své nejdražší podobě: dřina, kterou jsi už zaplatil, ti brání zaplatit ji znovu líp. Proti tomu má vibe coder zbraň, o jaké se Blondlotovi nesnilo. Uděláš commit — uložíš si pozici ve hře — a odbočíš na *větev*, paralelní vesmír na experiment. Teď smíš zkusit i tu bláznivou představu, protože když nevyjde, přepneš zpátky a ty dva měsíce práce jsou pořád tam, netknuté. Zahození přestane bolet, když je vratné. O commitu a větvi bude řeč v kapitole osmnácté; tady si jen zapamatuj, že *levné zahození je disciplína této kapitoly* — způsob, jak si nedovolit stát se rukojmím vlastní dřiny.

Věž: proč pocit není důkaz

Zbývá vylézt na věž a přeložit „dřina lže“ do čísel — protože zní to jako psychologie, ale je za tím tvrdá logika, kterou lze vyslovit přesně. Tahle věž je poctivá, ne fingovaná: každý její výhled si můžeš ověřit sám. Vezmi si z ní tři.

→ kap. 9: predikce zapsaná před zásahem je páteř debuggerovy smyčky (reprodukuje → hypotéza → predikce → jeden zásah → měř).

→ kap. 8: dát si pravidlo předem, protože vůle v reálném čase probírá, je pakt — Odysseus přivázaný ke stěžni.



Jistota a přesnost jsou dvě různé osy (metakognice). Když řešíš úlohu, probíhají v tobě *dva* soudy, ne jeden. První je samotná odpověď — a její kvalitu měří *přesnost* (accuracy): jak často máš pravdu. Druhý je tvůj pocit o té odpovědi — jistota — a jeho kvalitu měří něco jiného: *metakognitivní citlivost*, tedy jak dobře tvá jistota *odlišuje* tvé správné odpovědi od špatných. Teorie detekce signálu tuhle druhou schopnost kvantifikuje veličinou *meta-d'*: kolik informace o vlastní správnosti tvá jistota vlastně nese. A pointa je, že ty dvě osy jsou *oddělitelné* — dá se být přesný a přitom metakognitivně slepý (často se pleteš tam, kde jsi si nejjistější). Effort heuristic je porucha přesně téhle druhé osy: dřina napumpuje *jistotu*, ale s *přesností* nehne, takže tvá jistota přestane rozlišovat pravdu od omylu — *meta-d'* padá k nule. Cítíš se jistý stejně u toho, co máš správně, i u toho, co máš úplně vedle. Blondlotova jistota nesla nula bitů o tom, jestli paprsky existují.

Rozlišení jistoty a přesnosti formalizuje signal detection theory; meta-d' zavedli Maniscalco & Lau (2012) jako míru, jak dobře sebestjista diskriminuje správné od špatných v jednotkách výkonu úlohy. Lidský přístroj na měření toho, jak spolehlivý je tvůj pocit „vím to“.



Wasonova úloha 2-4-6: hledáš potvrzení, ne vyvrácení. Roku 1960 dal Peter Wason lidem trojici čísel — 2, 4, 6 — s tím, že se řídí nějakým pravidlem, a úkol zněl: uhodni pravidlo tak, že navrhuješ vlastní trojice a já ti u každé řeknu „ano“ / „ne“. Skoro každý si utvořil hypotézu („rostou po dvou“) a pak navrhoval trojice, které ji *potvrzovaly*: 8-10-12, 20-22-24 — a slyšel samé „ano“. Každé „ano“ posílilo jistotu; nikdo ji netestoval trojicí, která by hypotézu mohla *shodit* (třeba 1-2-3, nebo 6-4-2). Pravidlo přitom bylo prosté: *jakákoli tři rostoucí čísla*. Lidé si stavěli hromadu potvrzení a rostla jim z ní jistota — ale potvrzení hypotézy, kterou testuješ jen na souhlasících případech, nenese *žádnou* informaci navíc. Informaci nese jedině pokus o vyvrácení, který nevyšel. Tohle je *konfirmační strategie* a je to tentýž mechanismus jako effort heuristic: sbíráš důkazy pro, cítíš rostoucí jistotu, a přehlížíš jediný druh testu, který by tě mohl zachránit — ten, který mohl dopadnout špatně.

P. C. Wason, „On the failure to eliminate hypotheses in a conceptual task“, Quarterly Journal of Experimental Psychology 12(3), 1960. Předzvěst refrénu kap. 9: vědu nedělá hromada „ano“, ale ochota položit otázku, na kterou může přijít „ne“.



Bayesovsky: „cítím to“ má věrohodnostní poměr ≈ 1 . Přelož „pocit dřiny“ do řeči důkazů. Bayesova věta říká, že evidence E posune tvé přesvědčení o hypotéze H přesně o *věrohodnostní poměr* (likelihood ratio):

$$\frac{P(H | E)}{P(\neg H | E)} = \underbrace{\frac{P(E | H)}{P(E | \neg H)}}_{\text{věrohodnostní poměr}} \cdot \frac{P(H)}{P(\neg H)}.$$

Aby evidence E tvé přesvědčení vůbec pohnula, musí být *různě pravděpodobná* podle toho, zda hypotéza platí, či ne — tedy $P(E|H) \neq P(E | \neg H)$. A teď dosaď za E pocit „hodně jsem se nadřel, takže tomu věřím“. Nadřít se a mýlit se je úplně stejně možné jako nadřít se a mít pravdu — Blondlot je důkaz. Takže $P(\text{dřina} | H) \approx P(\text{dřina} | \neg H)$, věrohodnostní poměr vychází kolem *jedné*, a jednička v té rovnici nezmění *nic*: $P(H|E) = P(H)$. Pocit dřiny je informačně prázdný. Naopak *test, který mohl selhat, ale neselhal*, má $P(E|H)$ vysoké a $P(E | \neg H)$ nízké — poměr mnohem větší než jedna, a ten tvé přesvědčení skutečně posune. To je matematické jádro celé kapitoly: *důkaz není to, co tě stálo úsilí, ale to, co mohlo dopadnout jinak*.

→ kap. 19: „test, který mohl selhat“ je v inženýrství doslova jednotkový test — jeden test na ten triviální invariant by moje dva roky utnul v prvním týdnu. Rozuzlení Dvou roků do záchodu je tam.

→ kap. 3: tam evidence „vidím vzor“ nepřežila zamíchání dat; tady evidence „cítím, že rozumím“ nepřežije pokus vysvětlit. Táž prázdnota, druhá strana.

Co si ze čtvrté kapitoly odnést

Poskládej to dohromady. René Blondlot nebyl podvodník — byl to schopný, pracovitý fyzik, kterému rok poctivé dřiny vypěstoval tak pevné přesvědčení, že mu jeho vlastní oči hlásily paprsky i ve chvíli, kdy Wood držel klíčový hranol v kapse. To není kuriozita z dějin vědy; to je model toho, co se děje ve tvé hlavě pokaždé, když se do něčeho pořádně opřeš. Effort heuristic ti napumpuje jistotu úměrně vynaložené dřině — a ta jistota nenese o pravdě žádnou informaci, protože nadřít se a mýlit se stejně snadné jako nadřít se a mít pravdu. Mě to stálo dva roky a triviální neověřený předpoklad.

Kde kapitola třetí odhalila lháře venku — svět, který mlčí, a mozek, který mu dodá smysl —, tahle odhalila lháře uvnitř: vlastní dřinu, která se vydává za důkaz. Proti oběma platí týž lék, jen zrcadlově obrácený: *dostaň jistotu z blavy ven, kde se dá zkontrolovat*. Vysvětli to nahlas kachničce. Zapiš predikci, než uvidíš výsledek. Veď si deník omylů. A commitni před experimentem, ať smíš levně zahodit. Věž ti dala tvrdé důvody proč: jistota a přesnost jsou dvě různé osy a dřina hýbe jen tou první; hromada potvrzení (Wason) nenese informaci; a bayesovsky má pocit „cítím to“ věrohodnostní poměr kolem jedné — tedy nulovou důkazní sílu.

Nit, kterou si poneseš, je nenápadná, ale je to páteř celé druhé půlky knihy: *důkaz je jedině to, co mohlo dopadnout jinak, a nedopadlo*. V kapitole deváté z toho bude vědecká metoda, v osmé pakt se sebou samým a v devatenácté ten nejlevnější a nejmilosrdnější přístroj, jaký znám — test, který nespí, nepodléhá náladě a negratuluje ti k dřině. Jeden takový by mě zastavil první týden.

■ DEMO

`lesk-a-bida-vibecodingu.gaussalgo.com/d/iluze-hloubky`



Test iluze hloubky vysvětlení: ohodnot, jak rozumíš zipu, záchodu a gyroskopu — pak to zkus vysvětlit krok za krokem a ohodnot se znovu. Sleduj, jak ti sebejistota spadne. A vedle: Wasonova hra 2-4-6 — zkus jednou vyvrátit vlastní hypotézu, ne ji potvrdit.

„Jak bych poznal, že se pletu?“

Konkrétní test ke čtvrté kapitole: vezmi věc, o které jsi přesvědčen, že jí „rozumíš“ — svůj kód, návrh řešení, odpověď, kterou ti dal model — a *zkus ji vysvětlit nahlas cizímu člověku* (nebo gumové kachničce) do posledního kroku, bez jediného „to prostě funguje“. Tvrdím, že tam, kde se zadrhneš — kde ti chybí to „protože“ —, jsi jen *cítil*, že rozumíš; skutečné porozumění tam nebylo, jen pocit dřiny, který se za ně vydával. A tvrdší půlka, pro věci, které vysvětlit umíš: zeptej se, jaký *test, který mohl selhat*, tvé přesvědčení ověřil. Když žádný takový není — když máš jen hodiny práce a silný pocit —, nemáš důkaz, máš Blondlotovy paprsky: něco, co vidíš tím jasněji, čím víc ses nadřel, a co zmizí, jakmile někdo potmě vyndá hranol.