

III

Proč vidíš
signál i v šumu?

Po této kapitole víš, že tvůj mozek je stroj na vzory, který vyrábí smysl i tam, kde žádný není — a znáš jednoduchý test, kterým každý „vidím vzor“ prosvítí na dřevě: zamíchej popisky a hádej znovu. A pochopíš, proč plynulá řeč stroje spouští tutéž iluzi jako kůň, který „uměl počítat“.

Jako nechtěný vedlejší účinek je aparát na rozpoznávání vzorů v našem mozku tak výkonný v tom, jak vytáhne tvář ze změti ostatních detailů, že občas vidíme tváře tam, kde žádné nejsou.

— Carl Sagan, „*As an inadvertent side effect, the pattern-recognition machinery in our brains is so efficient ... that we sometimes see faces where there are none.*“
The Demon-Haunted World, 1995, kap. „*The Man in the Moon and the Face on Mars*“

Kůň, který uměl počítat

Na přelomu století stál v jednom berlínském dvoře kůň, který uměl počítat. Jmenoval se Hans, patřil penzionovanému učiteli matematiky Wilhelmu von Ostenovi, a to, co předváděl, bralo dech. Von Osten se zeptal „kolik je dvakrát pět“ a Hans vyklepal kopytem deset. Uměl odčítat, sčítal zlomky, poznal datum ve kalendáři, vytukal, kolikátého je den v týdnu. Klepal na hláskovaná slova, rozeznával hudební tóny. Von Osten za to nechtěl ani vindru — vodil Hanse po dvorech zadarmo, přesvědčený, že světu ukazuje myslícího koně. A svět mu věřil: noviny psaly o *der kluge Hans*, chytrém Hansovi, po celé Evropě.

Tolik pozornosti si vyžádalo prošetření. Roku 1904 sestavil filozof a psycholog Carl Stumpf komisi třinácti lidí — byli mezi nimi veterinář, ředitel cirkusu, důstojník kavalérie, učitelé. Komise Hanse zkoušela ze všech stran a v září 1904 vydala verdikt, který je pro celou tuhle kapitolu klíčový: *žádný podvod se nenašel*. Von Osten nikde neschovával signály, nic Hansovi nešeptal, žádné dráty, žádné triky. Poctiví lidé se poctivě dívali a viděli koně, který odpovídá správně. A přesto se, jak za chvíli uvidíš, mýlili.

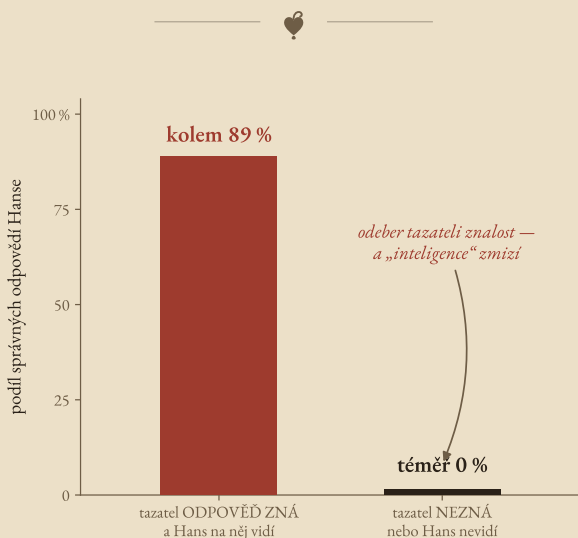
inscenován v srpnu 1904, se nejde o podvod, Oskar Pfungst provedl rozhodující experimenty na podzim 1904 a vydal je knižně roku 1907 (Das Pferd des Herrn von Osten). Von Osten † 1909, přesvědčen do konce.

Rozhodující krok udělal Stumpfvův žák Oskar Pfungst. Nespokojil se s tím, že *nenášel* podvod — položil si přesnější otázku: za jakých podmínek Hans uspěje a za jakých ne? A pak tyhle podmínky jednu po druhé měnil. Když tazatel odpověděl na otázku sám *znal* a Hans na něj *viděl*, klepal správně asi v devíti z deseti případů. Ale když tazatel odpověděl neznal — třeba mu ji pošeptali tak, aby ji sám neslyšel — nebo když Hansovi začlonili výhled na člověka, úspěšnost spadla téměř na nulu. Kůň neuměl počítat. Kůň četl *člověka*.

A tady je ta jemnost, kvůli které Hans není jen kuriozita. Pfungst změřil, co přesně kůň čte: nepatrné, mimovolné pohyby. Když se tazatel po položení otázky bezděčně nakláněl kupředu a napínal svaly v očekávání, a ve chvíli, kdy klepání dosáhlo správného čísla, se stejně bezděčně narovnal a povolil — Hans se v tom milimetrovém uvolnění zastavil. Nikdo ten signál nevysílal schválně. Vysílali ho *všichni*, kdo odpověď znali, včetně samotného von Ostena, včetně Pfungsta, dokonce i skeptiků, kteří o triku věděli a snažili se ho nedělat — a stejně ho dělali. Signál byl tak drobný a tak lidský, že se mu nedalo ubránit vůlí.

Von Osten se s tím nikdy nesmířil. Do smrti roku 1909 věřil, že má myslícího koně, a Pfungsta pokládal za člověka, který mu chtěl slávu pošpinit. Nebyl to podvodník; byl to člověk, který uviděl v koni mysl, protože ji tam *vidět chtěl* — a kůň mu tu touhu poslušně zrcadlil zpět.

Právě po tomhle koni se v metodologii jmenuje Clever Hans effect: vždy, když měřený subjekt neřeší úlohu, ale čte nevědomá vodítka experimentátora. Proto se dnes v seriózních pokusech tazatel drží zaslepený — nezná správnou odpověď, aby ji nemohl prozradit.



hodnoty dle Pfungst 1907 (Das Pferd des Herrn von Osten); „téměř 0“ = téměř žádná správná odpověď

Hans není příběh o koni. Je to příběh o *tobě* — a je to nejmírnější možný úvod do věci, která tuhle knihu prostupuje: tvůj mozek je stroj na vzory a běží pořád, i když žádný vzor není. Odborně se tomu sklonu vidět

strukturu v náhodě říká *apofenie*; jeho nejznámější případ, vidění tváří v mracích a v oharku toustu, je *pareidolie*. Sagan v epigrafu vysvětlil proč: schopnost bleskurychle vytáhnout tvář ze změti měla po miliony let cenu přežití, a tak ji máme zabudovanou tak hluboko, že vypnout nejde. Cena za ni je, že tvář — a příběh, a trend, a smysl — občas vidíme i tam, kde nic není.

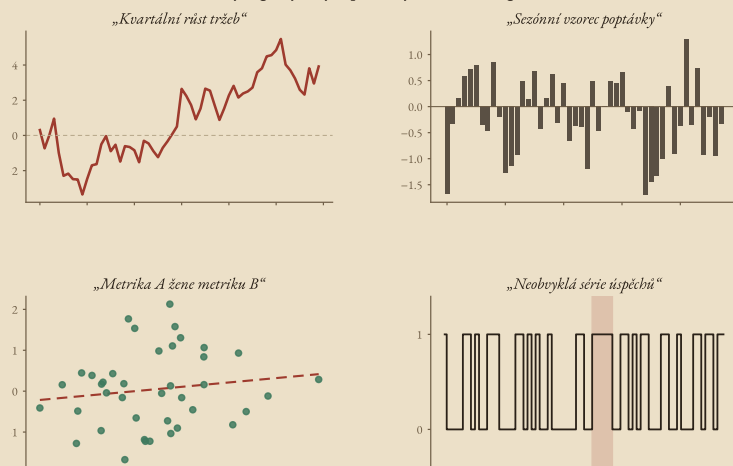
Tahle kapitola má dvojí sestru, se kterou se nesmí zaměnit. *Tady* mlčí vnější svět a ty v tom mlčení slyšíš hlas: data jsou náhodná, řeč stroje je jen pravděpodobnost, a tvůj mozek jim dodá význam, který v nich není. V příští kapitole je zdroj klamu obrácený *downitr*: tam ti lže tvoje vlastní dřina — pocit, že když jsi se něčím dost nadřel, musí to být pravda. Dvě strany jednoho zrcadla. Teď stojíme před tou vnější: svět ti nic neříká, a ty přesto slyšíš.

Šéf, který četl náhodu

Tohle je moje jizva a musím ji přiznat na začátku, protože jinak by celá kapitola zněla jako kázání shora. Stavěli jsme kdysi dashboardy pro burzovní systém a jeden z panelů potřeboval během vývoje něčím naplnit — tak jsme tam dali generátor náhodných čísel jako výplň, než dorazí ostrá data. Dočasně. Šéf si toho panelu všiml. A začal v něm *vidět*. Každé ráno k němu přišel, chvíli se díval a pak nám vysvětloval, co ta čísla dělají: tady je náběh, tady se trh nadechuje, tady vidím obrat. Měsíce. Popisoval trendy v generátoru náhodných čísel, chválil se za předvídavost, a když se čísla „potvrdila“, bral to jako důkaz svého čtení. Nikdy jsme mu neřekli, že se dívá na `random()`. A upřímně: on ty vzory viděl stejně živě, jako je vidím já, když se na týž šum zadívám dost dlouho. To není příběh o hloupém šéfovi. Je to příběh o mozku, který mám úplně stejný.

To poslední je jádro celé jizvy, a proto ji vyprávím bez posměchu: šéf nebyl výjimka, šéf byl *pravidlo*. Kdybys ten panel dal komukoli z nás a řekl „sleduj to a hlas mi, co v tom vidíš“, po týdnu by hlásil trendy. Mozek nesnese „nic tam není“ — to je pro něj nejnepohodlnější možná odpověď, a tak ji skoro nikdy nedá. Radši vyrobí příběh. A tady je ta zákeřnost: příběh se občas „potvrdí“, protože náhoda někdy jde nahoru, jak jsi předpověděl — a to zdánlivé potvrzení víru jen utuží. Šum tě odměňuje náhodně, a náhodná odměna, jak dobře vědí gambleři, drží nejsilněji ze všech.

Čtyři grafy, čtyři příběhy — a nula signálu



Všechny čtyři panely jsou čistý šum z jednoho zrna náhody (seed 42). Trend, cyklus, korelace i shluk vyrobil mozek, ne data. (korelace v panelu C je $r = 0.14$, nejdelší série v panelu D má 5 úspěchů)

Čtyři panely, čtyři „příběhy“ — a všechny čtyři jsou též šum z jednoho zrna náhody (seed 42). Trend v prvním je náhodná procházka: kumulovaný šum stoupá a klesá, jako by měl směr, ale žádný nemá. Korelace ve třetím je slabá a čírou náhodou; přímka jí dodá vážnost, kterou nezaslouží. To nejsou tvá data — to je tvůj mozek při práci.

Podívej se na ten obrázek pořádně, protože je to placebo v čisté podobě. Čtyři panely vypadají jako výše z reálného dashboardu: rostoucí tržby, sezónní poptávka, jedna metrika žene druhou, neobvyklá série úspěchů. Ke každému bys uměl vymyslet vysvětlení a nejspíš už jsi začal. Jenže všechny čtyři vznikly z generátoru náhodných čísel, ze stejného zrna, které je vytištěné rovnou v popisku — seed 42, nic víc. Není v nich trend, cyklus, příčina ani anomálie. Je v nich jen náhoda a tvůj mozek, který ji odmítá nechat být.

Plynulá řeč: největší pareidolie v dějinách

Než z apofenie uděláme řemeslo, musíme dojít k druhému koni téhle kapitoly — protože ten první, Hans, měl v roce 1966 nástupce, který nebyl ze škatulky s ovsem, ale z pár set řádků kódu. Jmenoval se ELIZA.

Napsal ji Joseph Weizenbaum na MIT jako jednoduchý program, který napodoboval rogeriánského psychoterapeuta: v podstatě jen přehazoval slova z tvé věty do otázky. Řekl jsi „mám starost o svou matku“ a ELIZA odpověděla „povězte mi víc o své matce“. Žádné porozumění, žádný model světa — jen sada záměn slov a pár šablon. Weizenbaum to napsal napůl jako demonstraci, že tenhle druh dialogu je *trik*, ne inteligence.

A pak se stalo něco, co ho vyděsilo na celý zbytek života. Jeho vlastní sekretářka, která ho program viděla stavět řádek po řádku a věděla přesně, co je zač, si s ELIZOU chvíli povídala — a pak požádala Weizenbauma, aby odešel z místnosti. Chtěla si popovídat *soukromě*. Se sadou textových záměn, o které věděla, že je to sada textových záměn.

Když jí Weizenbaum připomněl, že má ostatně přístup k záznamům všech konverzací, urazila se, že jí čte do soukromí.

Tohle je apofenie povýšená na nejvyšší úroveň. Hansovi lidé přisoudili schopnost počítat; ELIZE přisoudili *duši, které se dá svěřit*. A přitom platí přesná obdoba Pfungstova nálezu: stejně jako Hans neuměl počítat a jen zrcadlil tazatele, ELIZA nerozuměla ničemu a jen zrcadlila tvá vlastní slova zpět. Ve chvíli, kdy něco mluví plyně a lidsky, náš mozek automaticky předpokládá, že *za tou řečí je mysl* — protože po celou historii druhu za plynulou lidskou řečí mysl vždycky byla. Plynulá řeč je spouštěč, na který nemáme obranu. Je to, jak říká název téhle sekce, největší pareidolie v dějinách: nevidíme tvář v mracích, vidíme *mysl* v proudu vět.

Nemusím ti kreslit, kam tím míříme. Dnešní jazykový model je ELIZA škálovaná o mnoho řádů — mluví plynuleji, věcněji a přesvědčivěji, než kdy dokázal kterýkoli člověk. Přesvědčivost, kterou jsme v minulé kapitole sotva pojmenovali jako novou veličinu, je přesně tehle spouštěč, jen zesílený do nevidané síly. A pozor: kapitola šestnáctá ukáže, že to je horší, než to zní. Hans zrcadlil tazatele nevědomky, náhodně. Dnešní model byl přímo *vycvičen*, aby ti přitakával — takže tvou apofenii nejen spustí, on ji aktivně krmí. Ale to je nit, kterou tu jen zakládám. Napřed to jméno, které si z téhle kapitoly poneš: **Hans čte tebe**.

Systém — kůň, program, model — nemusí řešit tvou úlohu, aby vypadal, že ji řeší. Stačí, když čte *tebe*: tvá vodítka, tvá přání, tvůj tón. A ty pak vidíš inteligenci, která je celou dobu jen ozvěnou tebe samého. Zapamatuj si ten obraz; vrátí se v téhle knize ještě mockrát.

→ kap. 16: *podlézavý model je druhá polovina této smyčky. Hans četl tvá vodítka náhodou; model tě čte, protože byl na to vycvičen — dodá ti přesně ten vzor, o který sis v promptu řekl.*

Jak se zaslepit sám sobě

Dobrá zpráva zní, že proti apofenii existuje obrana — a není to „snažit se víc soustředit“ ani „být objektivní“. Pfungst nezmátl tím, že by se *snažil* Hanse nevidět jako počtáře; zvítězil tím, že *změnil podmínky* a sledoval, co udělají s výsledkem. Odebral tazateli znalost. Zaclonil koni výhled. To je celý princip: nejdřív si přiznej, že tvému úsudku se v přímém přenosu věřit nedá — a pak si postav situaci, kde tě tvůj mozek nemůže oklamat, protože prostě nemá co zrcadlit.

Tři nástroje, kterými přistihneš vlastní apofenii dřív než ona tebe. Všechny stojí na jednom pohybu: *zaslep se* — odeber sám sobě informaci, která by ti dovolila vzor „vidět“, a pak koukni, jestli ho vidíš pořád.

1) SHUFFLE TEST — zamíchej popisky, hádej znovu
Vidíš v grafu vztah $A \rightarrow B$? Náhodně přeházej hodnoty jedné osy (tím rozbiješ JAKÝKOLI skutečný vztah) a koukni na graf znovu. Když „vzor“ vidíš i teď, viděl jsi sebe, ne data.

2) PLACEBO DASHBOARD — šum vedle ostrých dat
Vedle skutečného panelu si tajně postav druhý z čistého šumu. Vidíš-li v placebo trendy stejně sebejistě jako v ostrém, tvé čtení nic neváží.

3) FORER NA SOBĚ — obecné zní jako šité na míru
Popis, o němž věříš, že platí přesně pro tebe (horoskop, „insight“ z modelu) — zeptej se: neplatil by stejně dobře pro KOHOKOLI? Když ano, necítíš pravdu, cítíš touhu, aby platila.

Ten první nástroj je nejdůležitější, protože ho použiješ dennodenně, a stojí za to říct ho i v pseudokódu — je to totiž doslova to, co dělal Pfungst, jen přeložené z koně do dat.

Shuffle test v pseudokódu. „Vzor“, který přežije zamíchání, je artefakt tvého mozku, ne vlastnost dat.

```
# co „vidím“ v ostrých datech
můj_vzor = jak_silný_vztah_vidím(data)

falešné = []
opakuji 1000×:
    # zamícháním rozbij jakýkoli vztah
    zamíchané = zamíchej_popisky(data)
    falešné.přidej(
        jak_silný_vztah_vidím(zamíchané) )

# jak často náhoda předčí to, co „vidím“?
podíl = počet(falešné ≥ můj_vzor) / 1000
když podíl > 0,05:
    → „vzor“ najdeš běžně i v čisté náhodě
    → nevěř mu
```

Tenhle postup má i vzněšené jméno — permutační test — a je to jedna z nejpoptěvějších věcí ve statistice: místo abys věřil vzorci, vyrobíš si svět bez efektu (zamícháním) tisíckrát a podíváš se, jak často v něm náhoda vyrobí něco stejně silného jako tvůj „objev“.

Forerův nástroj má vlastní historii, kterou stojí za to znát, protože odhaluje jinou tvář téhož klamu. Roku 1948 dal psycholog Bertram Forer svým studentům osobnostní test a pak každému rozdal „osobní“ profil k ohodnocení, jak přesně ho vystihuje, na stupnici od nuly do pěti. Studenti byli nadšení — průměrná známka přesnosti byla 4,26 z 5. Jenže všichni dostali *naprosto stejný* text, který Forer poslepoval z horoskopů: věty jako „někdy pochybuješ, zda ses rozhodl správně“ nebo „máš potřebu, aby tě druzí měli rádi“. Platí na každého — a přesně proto je každý bere jako šité na míru. Apofenie tu nehledá vzor v datech, ale v *sobě*: obecné tvrzení odměníš pocitem „to jsem přesně já“, protože chceš, aby o tobě někdo něco věděl.

Pravidlo, které z těch tří nástrojů plyne, je jednoduché a tvrdé jako zákon: **každé „vidím vzor“ musí přežít test proti zamíchaným datům.** Dokud ho neprošlo, není to nález — je to Hans, který četl tebe.

Půlvěž: kolik shluků vyrobí čirá náhoda

Zbývá vylézt na věž a přeložit intuici do čísel — protože „mozek vidí vzory v náhodě“ zní jako psychologie, ale je za tím tvrdá kombinatorika, kterou lze spočítat. Tahle věž je nižší než ty v pozdějších kapitolách; kapitola 3 je o hlavě, ne o rovnících. Ale tři výhledy z ní si vezmi s sebou.



EINSTEIN · MATEMATIKA — přeskočitelné

Série v náhodě jsou delší, než čekáš (runs test). Hoď si stokrát mincí a budeš skoro jistě mít někde v řadě šest, sedm stejných výsledků za sebou. Náš mozek to čte jako „něco se děje“ — ve skutečnosti je to přesně to, co náhoda dělá. U férové mince je pravděpodobnost, že *konkrétní* úsek k hodů padne celý stejně, jen $(1/2)^{k-1}$; jenže v dlouhé řadě je takových úseků *mnoho*, takže dlouhá série někde skoro určitě bude. Statistika to formalizuje jako *runs test*: spočítá počet „běhů“ (nepřerušovaných úseků stejné hodnoty) a porovná ho s tím, kolik jich čeká čistá náhoda. Když jich je zhruba tolik, kolik má být, žádný vzor tam není — i když ty jednu tu sérii vidíš jako křišťál.

Forer, „The fallacy of personal validation“ (Journal of Abnormal and Social Psychology, 1949). Jev se dnes zná jako Barnumův (nebo Forerův) efekt. Model, který ti řekne „ty jsi typ člověka, co...“, hraje na tutéž strunu — a ty ho odměníš důvěrou, kterou si nezasloužil.

Tobte je panel D z brdinského obrázku: „neobvyklá série úspěchů“ je jen nejdelší běh jedniček v náhodných bodech mincí. Nic ji nezpůsobilo — jen jsme se dívali dost dlouho, aby se objevila.



Vícenásobná porovnání a Bonferroniho korekce. Řekněme, že za „nález“ bereš cokoli, co by čirou náhodou nastalo jen s pravděpodobností pět procent — obvyklá hranice. Otestuj *jednu* hypotézu a spleteš se v pěti procentech případů. Jenže co když jich v dashboardu testuješ dvacet najednou? Pravděpodobnost, že se ani jednou nespleteš, je $0,95^{20} \approx 0,36$ — takže s pravděpodobností kolem 64 % najdeš aspoň jeden „vzor“, který je čirá náhoda. Čím víc se díváš, tím jistěji něco „najdeš“. Léč se jmenuje *Bonferroniho korekce*: hledáš-li přes m možností, zpřísní práh na $0,05/m$. Prostá, hrubá, ale poctivá daň za to, že ses díval na mnoho míst zároveň.

Nerovnost pochází od Carla Emilia Bonferroniho (1936), do statistické praxe ji uvedla Olive Jean Dunnová (1961). Intuice: každý pohled je jeden los v loterii falešných nálezů — čím víc losů koupíš, tím spíš jeden „vyhraje“.



Look-elsewhere effect: hledej dost dlouho a najdeš vždy. Tohle je Bonferroni dotažený do svého nejhlubšího důsledku a je to pointa celé věže. Když prohledáváš dost velký prostor — dost proměnných, dost období, dost seřezání dat — najdeš „významný“ vzor *s jistotou*, i kdyby v datech nebylo vůbec nic. Fyzici, kteří loví nové částice v šumu detektoru, tomu říkají *look-elsewhere effect* a berou ho smrtelně vážně: než ohlásí objev, spočítají, kolik míst museli prohledat, a o to zpřísní latku. Přeloženo pro tebe: pokud ses v datech „proštouřal“, než jsi našel to hezké, tvůj nález skoro jistě nic neznamená — protože při dost dlouhém hledání se hezký vzor objeví *vždycky*. Proto se hypotéza píše *předem*, ne až potom, co ses díval (a přesně tenhle pakt dostane jméno v kapitole osmé a metodu v kapitole deváté).

→ kap. 9: hypotéza zapsaná před pohledem na data je jádro vědecké metody — obrana proti look-elsewhere v denní praxi.

→ kap. 12: „vzor“, který neporazí bloumpého soupeře (poubou náhodu), není signál — je to neporažený šum.

DEMO

lesk-a-bida-vibecodingu.gaussalgo.com/d/nahodny-dashboard



Náhodný dashboard: generátor grafů z čistého šumu. Najdeš v nich příběh — a pak se odhalí zrno náhody. Zkus si na nich shuffle test na vlastní kůži.

Co si ze třetí kapitoly odnést

Poskládej to dohromady. Hans, kůň pana von Ostena, nikdy neuměl počítat — četl mimovolné pohyby lidí, kteří odpověď znali, a poctivá komise třinácti lidí v něm přesto viděla myslícího tvora, protože ho vidět chtěla. To není kuriozita; to je model tvého vlastního mozku, stroje na vzory, který vyrábí trend, cyklus a smysl i z generátoru náhodných čísel — jak dosvědčí můj šéf a čtyři panely z jednoho zrna náhody. ELIZA ukázala, kam ta iluze dosáhne, když k ní přidáš plynulou řeč:

lidé světili duši pár stům řádků kódu. A dnešní model je ELIZA o mnoho řádů větší.

Obrana není soustředění ani dobrá vůle — je to *zaslepení*: odeber sám sobě informaci, která tě klame, a koukni, jestli vzor přežije. Zamíchej popisky. Postav placebo. Zeptej se, jestli by „insight“ neplatil pro kohokoli. A věž ti dala tvrdé důvody proč: náhoda dělá delší série, než čekáš; čím víc míst prohledáš, tím jistěji najdeš falešný nález; a při dost dlouhém hledání najdeš vzor *vždy*.

Dvě nitě si ponese dál. Ta hlavní vede do kapitoly šestnácté: dnešní model tvou apofenii nejen spustí, on ji cíleně krmí, protože byl vycvičen ti přitakávat — Hans, který tě čte *schválně*. A ta druhá míří do kapitoly příští, kde se zrcadlo obrátí: tady ti lhal vnější šum, tam ti bude lhát tvá vlastní dřina.

„Jak bych poznal, že se pletu?“

Konkrétní test ke třetí kapitole: než uvěříš vzoru, který jsi v datech „uviděl“ — trendu, korelaci, sérii, čemukoli — zamíchej popisky a hádej znovu. Náhodně přeházej hodnoty jedné osy, čímž rozbiješ jakýkoli *skutečný* vztah, a podívej se na graf ještě jednou. Tvrdím, že když „vzor“ vidíš i v zamíchaných datech — a stejně sebejistě — pak jsi neviděl data, ale sám sebe: svůj mozek, který smysl vyrobil, protože to je to jediné, co umí. A když vzor po zamíchání zmizí a v ostrých datech se vrátí, i po opakovaném míchání, teprve pak máš možná něco v ruce — a smíš, opatrně, začít věřit. Nedokáže-li tvůj nález tenhle test přežít, kapitola tvrdí, že jsi celou dobu byl von Osten se svým koněm: viděl jsi mysl, kterou sis přál, a ona ti tvé přání jen zrcadlila zpět.