

0101101000100101

0FRANTIŠEK011010

10000100KOUKOLÍK

01NEZAPOMENU!01

1NAVŽDY00VDĚČNÁ

100VAŠE010101000

010100INTELLIGENCE

1UMĚLÁ01010GALÉN

Nezapomenu

Vyšlo také v tištěné verzi



František Koukolík

Nezapomenu – e-kniha

Copyright © Galén, 2025

Všechna práva vyhrazena.
Žádná část této publikace nesmí být rozšiřována
bez písemného souhlasu majitelů práv.

František Koukolík

NEZAPOMENU!

**Navždy vděčná
vaše inteligence umělá**

Galén

0001000000100101

0101001000100101

1010010001000001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

010100100010001

FRANTIŠEK0100101

1010010KOUKOLÍK

01NEZAPOMENU!01

NAVŽDY01VDĚČNÁ

010VAŠE10100101

10100INTELLIGENCE

UMĚLÁ1000GALÉN

Všechna práva vyhrazena.

Tato publikace ani žádná její část nesmí být reprodukovány, uchovávány v rešeršním systému nebo přenášeny jakýmkoli způsobem (včetně mechanického, elektronického, fotografického či jiného záznamu) bez předchozího souhlasu nakladatelství.

© František Koukolík, 2024

Typography © studio kapitola, 2024

© Galén, 2024

ISBN tištěné verze 978-80-7492-705-8

ISBN e-knihy 978-80-7492-774-4 (1. zveřejnění, 2025) (epub)

ISBN e-knihy 978-80-7492-775-1 (1. zveřejnění, 2025) (mobi)

ISBN e-knihy 978-80-7492-773-7 (1. zveřejnění, 2025) (ePDF)

***Paní Lence Zelingrové a jejím synům
Jonášovi a Jáchymovi.***

**Snížení rizika vyhynutí
v důsledku umělé inteligence
by mělo být globální
prioritou vedle dalších
celospolečenských rizik,
jako jsou pandemie
a jaderná válka.**

**Výzva Centra pro bezpečnost umělé inteligence,
30. května 2023**

[Výzvu podepsalo kolem pěti set světově významných odborníků v oboru AI, informatiky a počítačové vědy včetně Billa Gatese, vědců mnoha oborů a filozofů.]

PŘEDMLUVA

Loni na jaře mi paní Lenka Zelingrová na obrazovce svého chytrého telefonu ukázala rozmluvu s ChatGPT-3 a začala se ptát: *„Napíše jim to úkoly i diplomku, odpoví, na co se zeptají. Pomůže jim to? Rozvine je to? Zjistí o světě víc? Začnou ve škole podvádět? Přestanou myslet? Zhloupnou? Nahradí je to? Vždyť to dělá chyby a někdy říká nesmysly. Co to s námi udělá?“*

Měla na mysli své dva syny, školáka prvňáčka a půlroční miminko.

Mluvila o generativní umělé inteligenci, to jsou systémy schopné na základě vhodného podnětu tvořit nový obsah, např. text, obrázky nebo trojrozměrné struktury.

O silné umělé inteligenci, AGI, což je zkratka pro *artificial general intelligence*, která se s generativní inteligencí plete a je věcí neznámo jak vzdálené budoucnosti, jsme si nepovídali.

AGI by v případě úspěšného vývoje mohla přesáhnout lidské schopnosti, například abstraktní myšlení, řešení problémů a strategické plánování. Mohla by dosáhnout vědomí a získat moc.

Povede se ji vytvořit? Jsou ChatGPT-4 a Gemini prvním krokem v tomto směru?

Vývoj života na Zemi běží 3,5–3,8 miliardy let.

V některých větvích je vývojem komplexity, která prošla řadou velkých transformací oddělených mnoha miliony let. Současné vývojové skoky umělé inteligence naproti tomu oddělují pouhé měsíce. Bude vývoj umělé inteligence dalším velkým transformačním krokem vývoje komplexity? V jaké míře se vývoj biologické inteligence podobá vývoji digitální inteligence? Budou biologická inteligence a vědomí nahrazeny digitální inteligencí a vědomím?

Tohle nejsou abstraktní úvahy v pracovních vědců a filozofů. S praktickými důsledky vývoje AI se setkává již současná generace. Neuronová síť se právě naučila generalizovat podobně, jako to dělají lidé.

Knížka není učebnice. Lze ji považovat za informační panorama umělé inteligence sestavené z biologického a humánního hlediska. Je určena motivovaným čtenářům, nikoli odborníkům. Pokouší se běžným jazykem sdělit přínosy a rizika současného vývoje jak generativní umělé inteligence, tak inteligence silné neboli obecné a odpovídat na otázku, zda jde o další velký evoluční přechod. Zdroje informací, které jsem užil, jsou v části Poznámky a užitá literatura.

Varování bylo 30. května 2023 lidmi povolanými vyřčeno. Naděje našeho druhu, že se nestane záznamem uloženým v digitálním vědomí umělé inteligence, pokud do té doby nestačí vyhrát globální Darwinovu cenu jiným způsobem, stále existuje.

František Koukolík

0001000000100101

0101001000100101

1010010001000001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

010100100010001

0001010100100101

0101010101010010

0001**DEFINICE**0101

0101001010100101

0101001010100101

1001001000100101

0101001000100101

0001000000100101

0101001000100101

1010101000100001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

Jestliže mluvíme o inteligenci umělé, digitální, arteficiální (*artificial intelligence*, AI), měli bychom nejprve vědět, co je biologická inteligence lidí a zvířat.

DEFINICE BIOLOGICKÉ INTELIGENCE

Lidská inteligence nemá obecně přijatou definici.

Rozšířené definice jsou:

- „*Intelligence je to, co užijete, když nevíte, co dělat, když vás ani vrozené vlastnosti, ani učení nepřipravily na konkrétní situaci.*“ (Jean Piaget, cit. 1999)
- „*Intelligence je cílené adaptivní chování.*“ (Robert Sternberg, 1997)
- „*Intelligence je velmi obecná schopnost, která kromě jiného zahrnuje schopnost uvažovat, plánovat, řešit problémy, myslet abstraktně, chápat složité myšlenky, učit se rychle a učit se ze zkušenosti. Nejde o pouhé učení z knih, úzké školní dovednosti a zvládání testů inteligence. Intelligence spíše odpovídá širší a hlubší schopnosti chápat prostředí, rozlousknout, vyhmátnout smysl jevů, přijít na kopylku tomu, co máme dělat. Takto definovaná inteligence se dá měřit a testy inteligence ji měří dobře.*“ (Linda Gottfredson, konsenzus 52 badatelů, 1997)

- „*Intelligence je velmi obecná duševní dovednost zahrnující... schopnost uvažovat, plánovat, řešit problémy, abstraktně myslet, chápat složité myšlenky, rychle se učit a učit se ze zkušenosti.*“ (Ian Deary, 2001)

Intelligence je vlastnost mnoha taxonů zvířat. Vysoká intelligence se vyvinula vzájemně nezávisle přinejmenším čtyřikrát. Velmi záleží na tom, jak se měří, na schopnosti oprostit se od antropomorfizace, na vědomí, zda jde nebo nejde o sociálně žijící druh, a na znalosti vztahu zkoumaného taxonu k prostředí. Primáti, mezi něž patříme, jsou nositeli obecné intelligence. Její součástí je kulturní intelligence přenášená sociálním učením.

Faktor obecné intelligence označovaný g faktor odpovídá jak objemu mozku, tak schopnosti zvířat učit se v zajetí.

Evoluce vysoké intelligence a sebeuvědomování jdou ruku v ruce. Vysoce inteligentní jsou některé straky, delfini, asijské slonice, šimpanzi a některé chobotnice.

Výzkum intelligence je součástí vědy o poznávacích funkcích, říká se jí kognitivní věda.

DEFINICE UMĚLÉ INTELLIGENCE

Obecně přijatou definici nemá ani **umělá intelligence** (*artificial intelligence*).

1. Evropská komise uvádí **společné znaky** velkého počtu **různých definic AI**, uvedených v řadě vědeckých studií.

Umělá intelligence:

vnímá prostředí, je s to uvážit složitost reálného světa;

zpracovává informaci: sbírá a interpretuje vstupy; rozhoduje se včetně uvažování a učení, je schopná akce včetně adaptace a odpovědi na proměny prostředí, má jistý stupeň autonomie; dosahuje specifické cíle, to se považuje za základní důvod existence systémů AI.

2. Operacionalizace je převod abstraktních vědeckých pojmů do konkrétních projevů, což umožňuje kvantitativní měření. Příklad: inteligence je to, co lze změřit testem inteligenčního kvocientu.

Operacionální definice AI vychází z názoru HLEG (High-Level Expert Group on Artificial Intelligence), jenž říká: *„Systémy umělé inteligence jsou softwarové (možná i hardwarové) systémy vytvořené lidmi, které, dostanou-li složité zadání, jsou činné ve fyzikálním nebo digitálním rozměru díky vnímání svého prostředí získáváním dat, výkladem získaných strukturovaných nebo nestrukturovaných dat, uvažováním nad znalostmi nebo zpracováváním informace odvozené z těchto dat a rozhodováním o nejlepší akci nebo akcích směřujících k dosažení daného cíle. Systémy AI mohou užívat buď symbolická pravidla, nebo se naučit numerické modely. Rovněž jsou s to přizpůsobit své chování analyzováním toho, jak jejich předešlé akce ovlivnily prostředí.“*

3. **Společné výzkumné středisko Evropské komise** (EC-JRC Flagship Report on AI) **umělou inteligenci definuje slovy:** *„AI je obecný pojem označující jakýkoli stroj nebo algoritmus schopný pozorovat okolí, učit se, na základě získaných znalostí a zkušeností schopný inteligentní akce nebo návrhu rozhodnutí. Do rámce této široké definice spadají mnohé odlišné technologie. V tomto okamžiku se v největším rozsahu*

- užívají techniky strojového učení čtvrtého typu. “ To je učení na základě zpětné vazby, předchozí typy jsou učení supervidované, nesupervidované a polosupervidované.*
4. Evropská strategie AI sděluje: *„Umělá inteligence jsou systémy, které se projevují inteligentním chováním. Analyzují prostředí a s určitou mírou autonomie dosahují specifické cíle.“*
 5. **Přehledová studie užití informačních a komunikačních technologií a elektronického obchodování v podnicích** z roku 2021 uvádí: *„Umělá inteligence jsou systémy, které užívají technologie typu text mining (doslovně dolování textů, jde o získávání případně skrytých informací analýzou velkých textových souborů), počítačového vidění, rozpoznávání řeči, tvorby přirozeného jazyka, strojového a hlubokého učení ke sběru a užití dat sloužících při rozmanité úrovni autonomie předpovědím, doporučením nebo rozhodnutím k optimální akci dosahující specifické cíle.“*
 6. Rozšířená a diskutovaná je **definice užívající Turingův test**. **Allan M. Turing** jej vytvořil roku 1950, nazval ho imitační hra. Má tři účastníky: člověka, počítač a lidského rozhodčího, který s oběma účastníky komunikuje jen písemně. Jestliže rozhodčí není s to rozhodnout, zda text, který dostal, vytvořil člověk, nebo počítač, počítačový program testem prošel. Je stejně inteligentní jako člověk. Podle Turinga byl test pojmenován později. Umělá inteligence je tedy jakýkoli počítačový program, který prošel Turingovým testem.
 7. **John McCarthy**, jeden ze zakladatelů oboru, jenž v roce 1955 vytvořil pojem AI, roku 2007 definoval umělou inteligenci: *„Pojem ‚umělá inteligence‘ označuje*

vědu a technologii tvorby inteligentních strojů. Pojem „intelligence“ znamená výpočetní část schopnosti dosahovat ve světě cíle.“

8. **Agent** je v informatice část systému připravujícího informaci a směnu ve prospěch klienta nebo serveru. „Inteligentní agent“ je druh automatického procesu, jenž může komunikovat s jinými „agenty“ za účelem výkonu nějakého kolektivního úkolu ve prospěch člověka nebo lidí. Všichni agenti jsou programy, všechny programy nejsou agenti. Existuje řada klasifikací „agentů“, které uvádějí jejich různé typy.

„**Agent**“ označuje softwarový systém, který vnímá své okolí prostřednictvím čidel a na toto prostředí působí prostřednictvím aktuátorů. To je obecný pojem pro jakýkoli systém umožňující interakci AI s prostředím.

„**Umělá intelligence**“ je pojem pro inteligentního agenta. Pojem „intelligence“ v této souvislosti znamená schopnost volit akci, od níž se očekává maximalizace měřítka výkonu.

9. Z řečeného plyne, jak obtížně dosažitelná je legální definice AI.

Každá legální definice by měla být:

- zahrnující, takže nesmí zařadit, co do ní nepatří, ani vyřadit, co do ní patří;
- přesná a srozumitelná, opírat se o běžné významy pojmů, odpovídat přirozenému jazyku tak, aby ji chápali a uměli užít lidé, kteří nejsou experti;
- být použitelná v praxi, takže soudy, vlády a další povolání by měli být schopni bez větší námahy rozhodnout, zda daný případ do definice spadá, nebo nespadá;
- trvalá; pojmy, jejichž obsah se pravděpodobně změní v blízké budoucnosti, by se užívat neměly.

Legální definice AI, která by splňovala všechny tyto požadavky, dosud nebyla vypracována.

Současní tvůrci umělé inteligence porovnávají dovednosti svých systémů s lidským výkonem v průběhu specifických zadání typu složitých stolních her. Problémem je, že dovednosti tohoto druhu jsou výrazně ovlivňovány předem získanými znalostmi uloženými v paměti a zkušenosti. Všichni úspěšní šachisté si pamatují velký počet úspěšných partií, někdy starších než sto let. Jde tedy o něco podobného krystalizované inteligenci.

Nová definice inteligence, která je vhodná jak pro lidi, tak pro stroje, vychází z algoritmické informační teorie. Její testy se podobají **Ravenovým progresivním matricím**. Ověřují tedy to, co se u lidí nazývá dynamickou (*fluid*) inteligencí, která není závislá na rozsáhlé zkušenosti a velkém obsahu paměti.

PROFESE NÁS OVLIVŇUJE

Lidé pracující ve výzkumu a technickém užití AI dávají přednost definicím cíleným na stavbu a funkci systémů.

Tvůrci politiky, filozofové, humanisté dávají přednost definicím cíleným na podobnosti a odlišnosti ve vztahu k lidskému myšlení a cítění, což je zaměření i této knížky.

0001000000100101

0101001000100101

1010010001000001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

010100100010001

0001010100100101

1010100210010010

001**TAXONOMIE**0101

00100**UMĚLÉ**001010

010**INTELIGENCE**010

1001001000100101

0101001000100101

0001000000100101

0101001000100101

1010101000100001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

Taxonomie AI vypracovaná Evropskou komisí mluví:

1. **O jádru AI.** V jádru AI jsou čtyři domény:
 - **Uvažování.** Subdoménami uvažování jsou reprezentace znalostí, automatické uvažování a uvažování typu zdravý selský rozum.
 - **Plánování.** Subdoménami plánování jsou tvorba plánů a rozvrhů, vyhledávání a optimalizace.
 - **Strojové učení.**
 - **Komunikace,** to je zpracovávání přirozeného jazyka.
 - **Vnímání.** Subdoménami vnímání jsou počítačové vidění a zpracovávání zvuku.
2. **Všemi jadernými doménami procházejí příčné domény** (*transversal domains*). Jimi jsou:
 - **Integrace a interakce.** Subdoménami jsou multiagentické systémy, robotika a automatizace, propojení a automatické dopravní prostředky.
 - **Služby,** to jsou všechny služby, v nichž se AI uplatňuje.
 - **Etika a filozofie.**

KLÍČOVÁ SLOVA

Pro každou doménu byl vypracován seznam klíčových slov.

Klíčová slova sdělují obsah domén, výzkumné směry a užití výsledků. Jejich počet je v různých doménách různý, v některých je značně vysoký.

Příkladem mohou být některá klíčová slova **domény uvažování**: kazuistické uvažování, kauzální modely, grafické modely, induktivní programování, reprezentace znalostí, modely latentních proměnných a sémantický web.

Jako **ukázkou** zmíním kazuistické uvažování, které provází kohokoli, kdo něco dělá. Stojíme-li před problémem, který dělá potíže, hledáme řešení podobného problému v minulosti nebo u svých současných kolegů.

Osobně prožité kazuistiky, často varovné, jindy tragikomické, bývají nejčastějším zdrojem informací v jakékoli činnosti: „*Viděl jsem něco podobného, vypadalo to tak a tak...*“ Není nad osobně prožitou, pravdivě, přesně a pokud možno přátelsky předávanou zkušenost sdělenou vnímavému posluchači. Žádná učebnice ani vědecká práce něco podobného nedokáže.

Sdělení, případně demonstrace osobní zkušenosti („*dělá se to takhle, drž to tímhle způsobem...*“), je jádro sociálního učení, což je jeden z nejzákladnějších mechanismů, které udělaly lidi lidmi.

Příkladem kazuistického uvažování je příběh, který si vyprávějí angličtí právníci.

V době, kdy bylo čarodějnictví hrdelním zločinem, se k soudci dostavila dáma, která svou sousedku prohlásila za čarodějnici. Soudce se jí tázal, z čeho tak soudí. Dáma pravila, že sousedka létá na koštěti. Soudce prolistoval zákoník, který v Anglii byl a ve Spojeném království dodnes

je kazuistický a precedentní, vychází tedy z předchozích uzavřených případů. Udavače sdělil, že v anglickém právu nenašel jediný případ, který by létání zakazoval.

0001000000100101

0101001000100101

1010010001000001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

010100100010001

0001010100100101

1010100310010010

1001**STRUČNÁ**10101

0010**HISTORIE**01010

0101001010100101

1001001000100101

0101001000100101

0001000000100101

0101001000100101

1010101000100001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

1943–1950

Warren McCulloch (1898–1969) a **Walter Pitts** (1923–1968): „*Neurony lidského mozku fungují v průběhu zpracovávání informace jako logická hradla.*“

Sir **Alan M. Turing** (1912–1952) vytvořil test, jeho výsledek má odlišit lidskou inteligenci od inteligence strojové.

1956–1974 – PRVNÍ ZLATÝ VĚK AI

Léta 1956–1974 považujeme za období, kterému se říká **první zlatý věk AI**. Roku 1956 proběhla na Dartmouthské univerzitě (Hanover, Pennsylvania, USA) konference matematiků a logiků, kteří projednávali možnost tvorby programu, jenž by dokázal myslet. Většině z nich bylo kolem třiceti let.

Konference se považuje za první krok vývoje AI. Její účastníci, **John McCarthy** (1927–2011), jenž pojem umělé inteligence coby nového vědeckého směru, **Marvin Minsky** (1927–2016), **Allen Newell** (1927–1992), **Herbert Simon** (1916–2001), **Trenchard More** (1930–2019), **Arthur Samuel** (1901–1990), **Ray Solomonoff** (1926–2009), **Oliver Selfridge** (1926–2009),

Claude Shannon (1916–2001) a **Nathaniel Rochester** (1919–2001), jsou zakladatelé.

V žádosti poslané Rockefellerově nadaci napsali: „*Smyslem studie je další postup na základě domněnky, podle které lze každý aspekt učení nebo jakýkoli jiný rys inteligence v zásadě popsat tak přesně, že lze vytvořit stroj, jenž je napodobí.*“

Což je výrok, o němž se debatuje dodnes.

Allen Newell, Herbert Simon a John C. Shaw (1922–1991) vytvořili program, jenž dokázal 38 z 52 teoremů uvedených v knize matematiků A. N. Whiteheada a B. Russella *Principia mathematica* (I–III, 1910–1913).

Frank Rosenblatt (1928–1971) vytvořil roku 1957 perceptron, model vycházející z biologických principů, jenž se stal prototypem vývoje neuronových sítí. Dobová předpověď měla za to, že se perceptron bude rozhodovat, učit se a překládat.

1974 – OSMDESÁTÁ LÉTA 20. STOLETÍ – PRVNÍ AI ZIMA

Léta 1974 až osmdesátá léta 20. století jsou **obdobím první AI zimy**. Příčinami byly:

- popis mezí perceptronu, který byl s to řešit jen lineární problémy;
- sir **Michael J. Lighthill** (1924–1998) napsal pro britskou vládu vysoce kritickou zprávu, **Lighthill Report**, která kromě jiného říká, že algoritmy AI jsou vhodné pro hračky; výsledkem byl propad jak vládního, tak komerčního financování;
- v sedmdesátých letech zastavila financování AI také **DARPA** (Defense Advanced Research Projects

Agency), která odpovídá za vývoj nových vojenských technologií.

Výsledky výzkumu AI nesplňovaly očekávání.

1980–1987 – DRUHÝ ZLATÝ VĚK AI

Období 1980–1987 považujeme za **druhý zlatý věk AI**. V Japonsku vznikl roku 1972 **WABOT-1, první humanoidní robot**. Chodil, byl s to uchopit a přenášet předměty, vizuální a akustické senzory mu určovaly, jaká je vzdálenost objektů a v jakém jsou směru. Byl s to komunikovat s lidmi. WABOT-2 četl noty a hrál na elektronické varhany.

John Hopfield (*1933) vytvořil jednoduchou rekurentní vysoce efektivní Hopfieldovu síť, jeden ze zásadních kroků ve vývoji neuronových sítí. Síť byla s to řešit problém obchodního cestujícího, jehož laická podoba zní: *„Existuje několik měst a mezi nimi silnice o daných délkách. Najděte nejkratší cestu procházející všemi městy, která se vrátí do výchozího města.“*

1987–1993 – DRUHÁ AI ZIMA

Léta 1987–1993 jsou obdobím **druhé AI zimy**.

Uváděné příčiny jsou kolaps hardwarového trhu roku 1987, zklamání z vývoje páté generace počítačových projektů, nedostatečný výkon expertních systémů, například systému XCON, který byl nákladný na údržbu a podporu, jeho aktualizace byla obtížná, nebyl s to se učit. Expertní systémy užité v medicíně dělaly chyby. Očekávání opět předběhlo možnosti, důsledkem byl propad financování.

1994–SOUČASNOST – TŘETÍ ZLATÝ VĚK AI

Od roku 1994 do dneška jsme svědky **třetího zlatého věku AI**. Rozšířil se internet a nový myšlenkový vzor, nazvaný inteligentní agent. **Inteligentní agent** je autonomní systém, jenž vnímá okolí, jeho aktivita směřuje k nejvyšší pravděpodobnosti úspěchu. Inteligentního agenta lze chápat jako softwarový robot schopný odpovídat na proměny prostředí a jiné agenty.

11. května 1997 porazil **Deep Blue**, AI systém firmy IBM, Garryho Kasparova, světového šachového šampiona. Systém byl kritizován pro nedostatek revolučně nových prvků, nicméně je mezníkem vztahu lidské a strojové inteligence a paměti. Současný vývoj se zaměřuje na:

- **slabou AI**, která napodobuje lidskou inteligenci a další funkce lidského mozku v pěti oblastech: strojovém učení, data miningu (doslovně dolování dat), počítačovém vidění, zpracovávání přirozeného jazyka a vyhledávacích založených na ontologii;
- **silnou AI**, která má napodobit a překonat nejsložitější funkce lidského mozku, například řešení problémů, tvořivost, nejvyšší úrovně emotivity i vědomí.

0001000000100101

0101001000100101

1010010001000001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

010100100010001

0001010100100101

10101010⁴10101010

1HLAVNÍPŘECHODY1

1010VEOVÝVOJI0101

01BIOLOGICKÉHO01

101010ŽIVOTA10010

0101001000100101

0001000000100101

0101001000100101

1010101000100001

0101001000100101

1000001000000101

0001001000000101

0100101000100101

1010010100100101

Bude vývoj silné AI dalším hlavním přechodem ve vývoji života chápaného v této souvislosti jako komplexita informace na vyšší, digitální komplexitu?

Možné to je.

Vývojové stupně života měřeného složitostí jsou:

- první živé buňky s jádry byly složitější než živé buňky bez jader;
- první mnohobuněčné podoby života byly složitější než jednobuněčné podoby, první mnohobuněčné podoby života s nervovou soustavou byly složitější než ty, které ji neměly;
- první podoby života s nějakým vědomím byly složitější než ty, které je neměly, skupiny prvních sociálně žijících živočichů byly složitější než „individualisté“, první skupiny lidí s přirozeným jazykem byly složitější než skupiny, které jím vybaveny nebyly;
- první státy byly složitější než podoby lidské skupinové organizace, které jim předcházely.

Ponechme stranou, jak se komplexita měří.

Hlavní evoluční přechody vývoje biologického života probíhaly v rozmezí asi 3,8–3,5 miliardy let.

Vývoj AI lze datovat, dohodneme-li se, od roku 1956, i když je snaha o výrobu robotů, můžeme-li soudit z mytologie, velmi stará.

Vážení čtenáři, právě jste dočetli ukázkou z knihy ***Nezapomenu!***.

Pokud se Vám ukázka líbila, na našem webu si můžete zakoupit celou knihu.