

Umělá inteligence srozumitelně

Co umí a co ne, kdo ji ovládá a jak mění práci, peníze i demokracii

Citujte: NIČ, Miloslav. *Umělá inteligence srozumitelně*. 1. vyd. Slaný: Miloslav Nič, 2026. ISBN 978-80-909954-6-8. <https://doi.org/10.5281/zenodo.21937480>



Naskenuj
a stáhni knihu

Celá kniha zdarma (PDF · ePub · HTML):

<https://vydavatelstvi.mila.guide/umela-intelligence/> · Licence CC BY 4.0

OBČANSKÝ KOMPAS · KAPITOLA 1

Co je vlastně umělá inteligence

Co se za zkratkou AI ve skutečnosti skrývá, jak fungují programy typu ChatGPT, kdy si AI vymýšlí a kdy naopak může být užitečným pomocníkem.

Miloslav Nič · Stav: 12. srpna 2026 · obcansky-kompas.mila.guide/umela-intelligence/co-je-ai

Stručně

S **umělou inteligencí** – zkráceně **AI**, z anglického *Artificial Intelligence* – se dnes setkáte skoro všude. Může to být ChatGPT v telefonu, hlasový asistent, doporučování filmů na Netflixu, generování obrázků nebo automatické překlady.

Pod stejnou zkratkou se ale skrývá více různých technologií. Jedno však mají společné: počítač se „naučil“ něco dělat z velkého množství příkladů. Člověk mu tedy nemusel krok za krokem napsat všechna pravidla.

Programy jako ChatGPT, Claude nebo Gemini fungují díky takzvaným velkým jazykovým modelům – zkráceně LLM, z anglického Large Language Models. Jsou to systémy umělé inteligence určené

hlavně k práci s textem: umějí ho vytvářet, shrnovat, překládat, vysvětlovat nebo na něj navazovat.

Tyto modely se učily z obrovského množství textů na internetu, z knih a článků. Neznamená to ale, že přemýšlejí jako člověk. Když jim položíte otázku nebo zadáte úkol, rozloží si text na malé části a postupně počítají, jaká další část má nejspíš následovat.

Odborně se takové části říká **token**. Může to být celé slovo, část slova, znak nebo interpunkční znaménko. Model při tvorbě odpovědi vychází ze vzorců, které si osvojil během tréninku na velkém množství textů. Tyto vzorce mu člověk nepíše ručně jako pravidla v klasickém programu. Vznikají automaticky při tréninku modelu, kdy systém postupně upravuje své vnitřní nastavení podle příkladů, na nichž se učí.

Navenek to proto může působit jako rozhovor, uvnitř jde ale o velmi složitý statistický systém pro práci s jazykem.¹²

1. **AI si umí vymyslet věrohodně znějící nesmysly.** Když odpoví, nezná, může si ji domyslet a podat ji se stejnou sebejistotou jako pravdivou informaci. Tomu se říká **halucinace**.

Nejde o chybu, kterou by bylo možné „opravit jednou provždy“. Souvisí to se způsobem, jak dnešní AI funguje: model skládá odpověď podle vzorců, které se naučil při tréninku, ale sám neověřuje, jestli je výsledek pravdivý.

U důležitých věcí si proto vždy ověřte fakta jinde – například na oficiální stránce úřadu, na ověřeném zpravodajském webu nebo v učebnici.

2. **„AI“ není jedna věc.** Umělá inteligence v autě, která rozpoznává dopravní značky, umělá inteligence v nemocnici, která čte rentgenové snímky, a umělá inteligence v ChatGPT, která s vámi vede rozhovor, jsou tři úplně jiné programy. Liší se tím, k čemu slouží, jaká nesou rizika i podle jakých pravidel se mají řídit. Když někdo říká „AI nás nahradí“ nebo „AI je nebezpečná“, dává smysl se zeptat: *kterou AI máte na mysli?*

3. **AI už nezačíná jen odpovídat, ale také sama jednat.** Od roku 2024 se rychle rozvíjejí takzvaní *AI asistenti* neboli *AI agenti* – systémy, které za vás mohou provést konkrétní úkol. Vyhledají informaci na webu, vyplní formulář nebo naplánují schůzku. V některých oborech, například v programování, rešerších a práci s dokumenty, už jsou prakticky užiteční. Jinde jsou ale stále pomalí, drazí a chybují. Na samostatné plnění běžných úkolů bez dohledu se proto zatím spolehnout nedá. Vývoj je ale rychlý.⁴

4. **AI nerozumí světu stejným způsobem jako člověk.** I když její odpověď zní chytře, AI nemá vědomí, záměry ani city. Nevypočítá si, že vás chce oklamat – prostě „neví“, co říká.
5. **Hranice slova „AI“ se neustále posouvá.** V 60. letech se za AI považovaly i programy jako chatbot **ELIZA** nebo program **DOCTOR**, který napodoboval psychoterapeuta jednoduchým trikem: věty uživatele převáděl zpátky na otázky.

Informatikovi Larrymu Teslerovi se připisuje výrok: *„Intelligence je to, co stroje ještě nedokázaly.“* Jakmile začne nějaká schopnost spolehlivě fungovat, přestane se jí říkat AI a říká se jí prostě software. Hodně z toho, co dnes působí kouzelně, bude za pár let stejně samozřejmé jako kalkulačka.¹¹

Tip

Praktické pravidlo, které vám ušetří hodně potíží: AI je výborný *rychlý pomocník*, ne autorita. Použijte ji k vytvoření konceptu, který si potom sami zkontrolujete – třeba návrhu e-mailu, shrnutí dlouhého článku nebo prvotního nápadu. **Nepoužívejte ji jako konečný zdroj pravdy** – zvlášť ne u zdravotních rad, právních záležitostí nebo důležitých rozhodnutí.

Podrobně

Jak se počítače naučily pracovat s jazykem

Na začátku se umělá inteligence dnešním chatovacím robotům vůbec nepodobala. Od 50. do 80. let minulého století se počítače učily hlavně tak, že jim lidé ručně zapisovali pravidla: „*Když má pacient horečku a kašel, může jít o chřipku.*“ Takových pravidel postupně vznikaly tisíce. Říkalo se jim **expertní systémy** – programy, které dokázaly překvapivě dobře radit v úzkém oboru, například při zdravotní diagnostice nebo chemické analýze.

Expertní systémy ale měly zásadní slabinu. Když narazily na situaci, se kterou žádné pravidlo nepočítalo, nevěděly si rady. V první polovině 80. let přesto zažily komerční rozmach. Brzy se však ukázalo, že jsou drahé na údržbu a jejich hlavní slabinu nelze snadno obejít. Na konci 80. a začátku 90. let proto opadlo nadšení i financování výzkumu. Historici toto období dodnes označují jako „*zimou umělé inteligence*“.⁵

Později se změnil samotný přístup. Člověk už počítači nepředepisoval každé pravidlo předem, ale začal mu dávat příklady. Už ne: „*Takhle přesně poznáš kočku.*“ Spíš: „*Tady máš tisíc obrázků koček – najdi, co mají společné.*“ Tím se otevřela cesta ke **strojovému učení** – metodám, při nichž počítač hledá vzory v datech.

Počítač už neměl jen poslušně vykonávat pravidla napsaná člověkem. Měl se z příkladů sám naučit, které znaky jsou důležité. Pro veřejnost se tato proměna stala viditelnou hlavně díky velkým vítězstvím: v roce 1997 porazil počítač Deep Blue od firmy IBM mistra světa v šachu Garryho Kasparova⁶, v roce 2016 program **AlphaGo** od firmy DeepMind porazil Lee Sedola, jednoho z nejlepších hráčů světa ve hře go, a v roce 2022 přišel ChatGPT. Během dvou měsíců získal 100 milionů uživatelů a spustil „AI horečku“, kterou dnes vidíme téměř všude.⁷

Co se mezitím stalo, že se z počítačů založených na ručně psaných pravidlech staly systémy schopné psát texty, překládat, programovat nebo odpovídat na otázky? Sešly se tři věci najednou. **Internet** přinesl obrovské množství textů, obrázků a dalších dat, z nichž se umělá inteligence mohla učit. **Grafické karty**, původně vyvíjené hlavně pro hry, se ukázaly jako ideální výpočetní motor: umějí rozdělit výpočet na tisíce malých kroků a řešit je současně, přesně tak, jak to tyto systémy potřebují. V roce 2017 pak vědci z Googlu představili **transformer**, nový princip, který umožnil mnohem lépe zachycovat

souvislosti v textu. Právě z něj vyrostly dnešní velké jazykové modely.
3

Tím začala nová éra. Umělá inteligence už nestojí hlavně na ručně psaných pravidlech ani na malých souborech příkladů. Učí se z obrovského množství dat – z miliard ukázek textu, obrázků a dalších vzorů. Neučí se jako člověk a nerozumí světu stejným způsobem jako my. Přesto dokáže rozpoznávat zákonitosti v jazyce tak dobře, že umí skládat věty, vysvětlovat složité pojmy, psát kód nebo napodobovat styl lidského textu. Tady začíná příběh umělé inteligence, která se naučila zacházet s jazykem tak přesvědčivě, že její odpovědi začaly působit jako skutečný rozhovor.

Model nepřemýšlí jako člověk

Když napíšete otázku velkému jazykovému modelu, model nepřemýšlí jako člověk. Nevychází z vlastní zkušenosti, úmyslu ani z porozumění světu v lidském smyslu. Místo toho postupně dopočítává, jaké pokračování textu je v dané chvíli nejpravděpodobnější.

Základní jednotkou takového výpočtu je **token**. Token je malý kousek textu – někdy celé krátké slovo, jindy jen část slova nebo jednotlivý znak. Model vybere pravděpodobné pokračování, přidá ho k rozepsané odpovědi a stejný krok opakuje znovu a znovu. Slovo po slově, přesněji token po tokenu, tak postupně vzniká celá odpověď.

Navenek to působí mnohem chytřeji, než by se z tak jednoduchého principu mohlo zdát. Model při tréninku prošel obrovským množstvím textů, a proto při skládání odpovědi velmi dobře odhaduje, jaké pokračování se hodí k tomu, co už bylo napsáno. Když vysvětluje složitý pojem, staví argument, překládá větu, píše program nebo odpovídá na odbornou otázku, nedělá to podle ručně sepsaných pravidel ani díky lidskému porozumění. Přesto z dlouhé řady drobných voleb vzniká odpověď, která působí jako profesionální lidské dílo. Novější modely navržené pro náročnější otázky – říká se jim **uvažovací modely** – si problém nejprve rozloží, projdou několik mezikroků a teprve potom sestaví odpověď pro čtenáře. Stále nejde o lidské porozumění, ale o mimořádně výkonný způsob, jak generovat smysluplný text.

Parametry: naučená čísla uvnitř modelu

Aby vše fungovalo, musí v sobě model obsahovat obrovské množství „naučených čísel“. Odborně se jim říká **parametry**. Právě v nich je nepřímě uloženo to, co se model při tréninku naučil: styl psaní,

jazykové vzorce, části faktických znalostí i obvyklé podoby vysvětlení, argumentu nebo postupu řešení. Největší dnešní modely mívají stovky miliard parametrů. U některých se mluví i o bilionech, i když firmy u uzavřených modelů přesná čísla obvykle nezveřejňují.

Parametry vznikají při tréninku. Model postupně upravuje vnitřní čísla tak, aby stále lépe odhadoval, jak text obvykle pokračuje. U špičkových modelů může samotný výpočetní trénink stát desítky až stovky milionů dolarů a trvat měsíce na tisících speciálních čipů. Některé novější otevřené modely uvádějí výrazně nižší náklady, které však obvykle nezahrnují celý předchozí výzkum, přípravu dat, vývoj infrastruktury ani neúspěšné pokusy.

Doladění lidmi a jeho skrytá cena

Ve druhé fázi přicházejí na řadu lidé. Model po prvním tréninku umí plynule pokračovat v textu, ale sám od sebe ještě nepozná, která odpověď je užitečná, bezpečná nebo vhodná. Dotaz na podvod, zneužití cizího účtu nebo výrobu jedu by pro něj bez dalšího doladování mohl být jen další text, na který má navázat co nejpravděpodobnější odpovědí – podobně jako u receptu, překladu nebo školní úlohy. Proto lidští hodnotitelé posuzují jeho odpovědi, porovnávají lepší a horší varianty a na příkladech mu pomáhají rozlišovat, kdy má pomoci, kdy má být opatrný a kdy má požadavek odmítnout. Tomu se říká **učení s lidskou zpětnou vazbou** a je to jeden z důvodů, proč jazykové modely odpovídají zdvořile a odmítají škodlivé požadavky.

Tato práce má ale i méně viditelnou stránku. Často se zadává externím firmám v zemích s nízkými mzdami. Známý je například případ pracovníků v Keni, kteří pro dodavatele OpenAI označovali toxický obsah za méně než 2 dolary za hodinu. I proto se vede debata o podmínkách lidí, kteří tuto skrytou práci pro vývoj AI dělají.⁸

Když lidský text dochází: syntetická data a učení odměnou

Celý dosavadní recept – víc textu, větší model – přitom naráží na nečekaný strop: kvalitní lidský text dochází. Výzkumná skupina Epoch AI odhadla zásobu použitelného veřejného textu, který kdy lidé napsali, na 510 bilionů tokenů a spočítala, že při současném tempu ji modely vyčerpají někdy mezi lety 2026 a 2032. Internet se nezmenšuje; nový kvalitní text psaný lidmi ale nepřibývá ani zdaleka tak rychle, jak roste „apetit“ tréninku.¹³

Laboratoře na to reagují dvěma způsoby. První odpověď zní: když lidský text dochází, ať si ho umělá inteligence vyrobí sama. Textům a příkladům, které jeden model vygeneruje pro trénink jiného, se

říká **syntetická data** a dnes s nimi pracují všechny velké laboratoře. Firma NVIDIA například uvedla, že při dolaďování jejího modelu Nemotron-4 pocházelo přes 98 % dat od modelů samotných; lidé dodali jen dvacet tisíc ukázek.¹⁴ Stále ale chybí přesvědčivý důkaz, že by syntetická data zvládla významně nahradit lidský text i v základním tréninku.¹

Má to ovšem háček. Když se modely učí stále dokola z výstupů jiných modelů, hrozí, že se kvalita postupně rozpadne jako u kopie kopie: vzácné a neobvyklé jevy z dat mizí a texty se stávají jednotvárnějšími. Studie v časopise Nature tento jev pojmenovala „kolaps modelu“.¹⁵ Navazující výzkum ale ukázal, že při postupném přidávání syntetických dat k lidským se kvalita rozpadat nemusí.¹⁶ Syntetická budoucnost tréninku tedy není zadarmo – vyžaduje pečlivou práci s daty.

Druhá odpověď se bez nového textu obejde úplně. U úloh, kde lze výsledek automaticky zkontrolovat – matematický příklad buď vyjde, nebo nevyjde, program buď funguje, nebo ne – se model učí pokusy: zkouší řešení a za správná dostává odměnu, podobně jako při dolaďování lidmi, jen roli hodnotitele přebírá automatická kontrola. Právě z tohoto přístupu, odborně **učení posilováním**, vyrostly **uvažovací modely** zmíněné na začátku této části. U modelu DeepSeek-R1 zvedl takový trénink úspěšnost na úlohách náročné americké středoškolské matematické soutěže z 15,6 % na 71,0 % – bez jediného nového lidského textu.¹⁷

Tento posun vysvětluje, proč je dnešní pokrok tak nerovnoměrný. Tam, kde jde výsledek ověřit – v matematice, programování, logických úlohách – se modely zlepšují skokově. Tam, kde by pomohl hlavně další kvalitní lidský text, třeba ve znalostech o méně známých tématech nebo o vašem profesním oboru, jde pokrok pomaleji.

Kontextové okno: co má model „po ruce“

Když je model natrénovaný a začneme ho v praxi používat, neodpovídá jen na poslední větu, kterou mu napíšeme. Při každé odpovědi vychází z textu, který má v danou chvíli „po ruce“. Patří do něj aktuální otázka, často i předchozí části téhož rozhovoru, případně vložený dokument, delší zadání a také instrukce určující, jak se má model chovat. Nejde o další trénink modelu, ale o soubor informací, z něhož model při konkrétní odpovědi vychází.

Odborně se tomuto prostoru říká **kontextové okno**. U některých dnešních modelů může být velmi dlouhé: dokážou v jednom zadání pracovat s textem o rozsahu několika knih, tedy se stovkami tisíc až miliony tokenů. Není to ale paměť v lidském smyslu. Model si sám

od sebe „nevzpomíná“ na dávné rozhovory. Při aktuální odpovědi může využít jen to, co se mu vešlo do kontextového okna.

Díky kontextovému oknu ale mohou moderní modely shrnovat dlouhé právní dokumenty, porovnávat výroční zprávy, hledat souvislosti ve vložených podkladech nebo pomáhat s rozbořením celé knihy. Čím větší část textu má model při odpovědi k dispozici, tím lépe může navazovat na předchozí pasáže, vracet se k nim a držet se zadaného úkolu.

Nově vznikající schopnosti

U velkých modelů se navíc ukazuje, že některé schopnosti se výrazně zlepšují až od určité velikosti a kvality modelu. Někdy se jim říká emergentní schopnosti. Neznačená to ale, že se v modelu z ničeho nic objeví lidské porozumění. Vědci se stále přou, zda jde opravdu o nový skok v chování modelu, nebo spíš o důsledek toho, jak tyto schopnosti měříme.⁹

I proto pokračuje debata, zda se podobné modely mohou jednou přiblížit obecné umělé inteligenci – **AGI**, která by dokázala zvládat široký rozsah úloh podobně pružně jako člověk. Této otázce se podrobněji věnuje kapitola o AGI a existenčních rizicích.

Důležité je pamatovat na obě stránky věci. Velký jazykový model nemá vědomí, vlastní záměr ani lidské porozumění světu. Neví, co chce říct, tak jako to ví člověk. Zároveň však dokáže z obrovského množství naučených jazykových vzorců skládat odpovědi, které působí jako výsledek promyšleného uvažování. V tomto napětí – mezi poměrně jednoduchým principem a překvapivě složitým chováním – spočívá podstata našich rozhovorů s umělou inteligencí.

Proč si AI vymýšlí

Upozornění

AI neumí spolehlivě rozlišit pravdivou odpověď od přesvědčivé nepravdy. Když narazí na otázku, na kterou odpověď nezná, může si ji jednoduše vymyslet a podat ji tak jistě, že zní důvěryhodně. Tomu se říká **halucinace** a patří to k největším slabinám dnešních jazykových modelů.

Představte si, že požádáte ChatGPT nebo Claude, aby vám našel odbornou citaci k určitému tématu. Model obratem vytvoří záznam, který vypadá naprosto přesvědčivě: uvede jméno autora, název článku, odborný časopis, rok vydání i čísla stránek. Na první pohled působí jako poctivě dohledaný zdroj. Jenže žádný takový článek neexistuje – model si celou citaci vymyslel.

Podobně dopadl v roce 2023 americký právník Steven Schwartz. V soudním podání použil několik smyšlených precedentů, které mu „připravil“ ChatGPT. Nešlo o drobnou chybu v citaci, ale o neexistující soudní případy předložené jako skutečná právní opora. Soud proto uložil právníkům i jejich kanceláři Levidow, Levidow & Oberman pokutu 5 000 dolarů. Případ *Mata v. Avianca* se od té doby často uvádí jako varování před bezhlavou důvěrou v odpovědi umělé inteligence.

10

Z toho plynou dvě praktická pravidla, ke kterým se v tomto průvodci budeme vracet.

1. AI je užitečná tam, kde lze výsledek snadno zkontrolovat.

Programový kód buď běží, nebo neběží. Překlad si může přečíst rodilý mluvčí. Shrnutí dlouhého článku porovnáte s původním textem. Návrh e-mailu sami posoudíte, než ho odešlete.

2. AI je riziková tam, kde se správnost ověřuje obtížně.

Faktická tvrzení o méně známých lidech nebo událostech, lékařské diagnózy, právní rady či finanční doporučení mohou znít velmi jistě, i když jsou chybné. V takových případech může AI posloužit jako zdroj nápadů, první orientace nebo seznam otázek, které je třeba dál ověřit. Neměla by ale být konečnou autoritou.

Halucinací postupně ubývá. Novější modely lépe rozpoznávají situace, kdy si nejsou jisté, a některé systémy si navíc před odpovědí dokážou dohledat informace na internetu nebo ve vložených dokumentech. Tomu se říká **RAG** – zjednodušeně odpovídání s pomocí vyhledaných podkladů. Ani tento postup však halucinace úplně neodstraní. Nejde o dětskou nemoc, ze které modely jednoduše vyrostou, ale o důsledek toho, jak jsou uvnitř postavené.

Co AI dnes opravdu umí – a co ne

Úloha	Stav v roce 2026	Kam to směřuje
Napsat smysluplný text v češtině	Velmi dobře (ChatGPT, Claude, Gemini)	Běžně použitelné
Přeložit běžný text	Velmi dobře u běžných textů a velkých jazyků; odborné, právní a literární překlady vyžadují kontrolu člověkem	Z velké části běžná praxe
Napsat funkční počítačový kód	Dobře u běžných úloh, u velkých projektů ale naráží na problémy	Postupně se zlepšuje
Vyřešit středoškolskou matematiku	Dobře	Zlomový pokrok 2024–2025
Vyřešit olympiádní matematiku	Špičkové výsledky 2025–2026	Rychlý postup
Vygenerovat realistický obraz	Velmi dobře (Midjourney, FLUX, Stable Diffusion, generování přímo v ChatGPT)	Běžně použitelné
Vygenerovat realistické video	Klipy dlouhé zhruba 10–25 sekund, u špičkových modelů s nativním zvukem a kvalitou blízkou filmové; stále drahé	Delší vícezáběrové video; plná filmová délka zůstává výzvou
Bavit se o vašem profesním oboru	Dobře po doplnění kontextu, ale AI si občas vymýšlí detaily	Pomalu – limitem jsou data
Pomáhat s lékařskou diagnostikou	Velmi dobré výsledky v některých úzkých úlohách, zejména v obrazové diagnostice (rentgen, CT); klinické nasazení vyžaduje regulaci a dohled lékaře	Postupně, klíčové jsou regulace a odpovědnost
Bezpečně řídit auto	Samořídící taxi (robotaxi) jezdí komerčně v některých amerických městech	Stále roky cesty

Úloha	Stav v roce 2026	Kam to směřuje
Skutečné porozumění a uvažování	V úlohách s ověřitelným výsledkem modely výrazně pokročily; zda jde o porozumění, nebo jen o velmi dobré napodobení, zůstává sporné	Nejasné – možná čeká velký objev
Obecná inteligence (AGI)	Předmět odborného sporu – viz kapitola o AGI a existenčních rizicích	Nikdo neví

K tabulce ještě jedna novinka z poslední doby: nastupují takzvané **světové modely** (world models) — AI, která z textového popisu vytvoří celé interaktivní trojrozměrné prostředí, jímž se dá procházet jako počítačovou hrou. Průkopníkem je systém Genie od Google DeepMind, který se postupně dostává k prvním uživatelům; na vlastních světových modelech dnes pracuje i řada dalších firem, protože se hodí nejen pro hry, ale třeba i pro trénink robotů. Zatím jde o ranou technologii — vygenerované prostředí vydrží souvisle fungovat jen minuty.¹²

A druhá novinka: umělá inteligence už nejen odpovídá, ale začíná také sama jednat. Takzvaní **AI agenti** za vás dokážou vyhledat informaci na webu, vyplnit formulář nebo naplánovat schůzku. V rutinních úlohách už reálně pomáhají, spolehlivě samostatní ale zatím nejsou. Podrobně se jim věnuje hned následující kapitola o AI, která jedná, vidí i slyší.

Tip

Jednoduché pravidlo vám může ušetřit spoustu potíží:

AI je nejlepší pomocník tehdy, když urychluje práci člověka, ale nenahrazuje jeho úsudek. Můžete ji například využít k přípravě návrhu e-mailu, shrnutí článku nebo prvotního nápadu a výsledek si potom sami zkontrolovat. Práci tím výrazně zrychlíte, aniž byste ztratili kontrolu nad kvalitou.

Mnohem větší riziko vzniká ve chvíli, kdy necháte AI samostatně rozhodnout o něčem důležitém a její výstup si neověříte. Týká se to například lékařské diagnózy, právní rady, finančního doporučení nebo výběru uchazečů o práci. V takových případech může chyba způsobit vážnou škodu.

Zdroje

- 1 STANFORD HAI. (2026). *Stanford AI Index Report 2026*
<https://hai.stanford.edu/ai-index/2026-ai-index-report>
- 2 ZHAO ET AL. (2024). *A Survey of Large Language Models*
<https://arxiv.org/abs/2303.18223>
- 3 VASWANI ET AL., GOOGLE. (2017). *Attention Is All You Need*
<https://arxiv.org/abs/1706.03762>
- 4 ANTHROPIC. (2024). *Computer Use (Anthropic developer documentation)*
<https://docs.claude.com/en/docs/agents-and-tools/tool-use/computer-use-tool>
- 5 ENCYCLOPEDIA BRITANNICA. *A Brief History of Artificial Intelligence*
<https://www.britannica.com/technology/artificial-intelligence/Symbolic-vs-connectionist-approaches>
- 6 IBM. *Deep Blue*
<https://www.ibm.com/history/deep-blue>
- 7 REUTERS. (2023). *ChatGPT sets record for fastest-growing user base — analyst note*
<https://www.reuters.com/technology/chatgpt-sets-record-fastest-growing-user-base-analyst-note-2023-02-01/>
- 8 TIME. (2023). *OpenAI Used Kenyan Workers on Less Than \$2 Per Hour to Make ChatGPT Less Toxic*
<https://time.com/6247678/openai-chatgpt-kenya-workers/>
- 9 WEI ET AL., GOOGLE RESEARCH. (2022). *Emergent Abilities of Large Language Models*
<https://arxiv.org/abs/2206.07682>
- 10 US DISTRICT COURT, S.D.N.Y. (2023). *Mata v. Avianca, Inc. — Opinion and Order (US District Court, S.D.N.Y., 22. 6. 2023)*
<https://law.justia.com/cases/federal/district-courts/new-york/nysdce/1:2022cv01461/575368/54/>
- 11 QUOTE INVESTIGATOR. (2024). *As Soon As It Works, No One Calls It AI Anymore — pátrání po Larrym Teslerovi a AI effect*
<https://quoteinvestigator.com/2024/06/20/not-ai/>
- 12 GOOGLE DEEPMIND. (2025). *Genie 3: A new frontier for world models*
<https://deepmind.google/blog/genie-3-a-new-frontier-for-world-models/>
- 13 VILLALOBOS ET AL., EPOCH AI. (2024). *Will we run out of data? Limits of LLM scaling based on human-generated data*
<https://arxiv.org/abs/2211.04325>
- 14 NVIDIA. (2024). *Nemotron-4 340B Technical Report*
<https://arxiv.org/abs/2406.11704>
- 15 SHUMAILOV ET AL., NATURE. (2024). *AI models collapse when trained on recursively generated data*
<https://www.nature.com/articles/s41586-024-07566-y>

- 16 GERSTGRASSER ET AL. (2024). *Is Model Collapse Inevitable? Breaking the Curse of Recursion by Accumulating Real and Synthetic Data*
<https://arxiv.org/abs/2404.01413>
- 17 DEEPSEEK-AI, NATURE. (2025). *DeepSeek-R1 incentivizes reasoning in LLMs through reinforcement learning*
<https://www.nature.com/articles/s41586-025-09422-z>