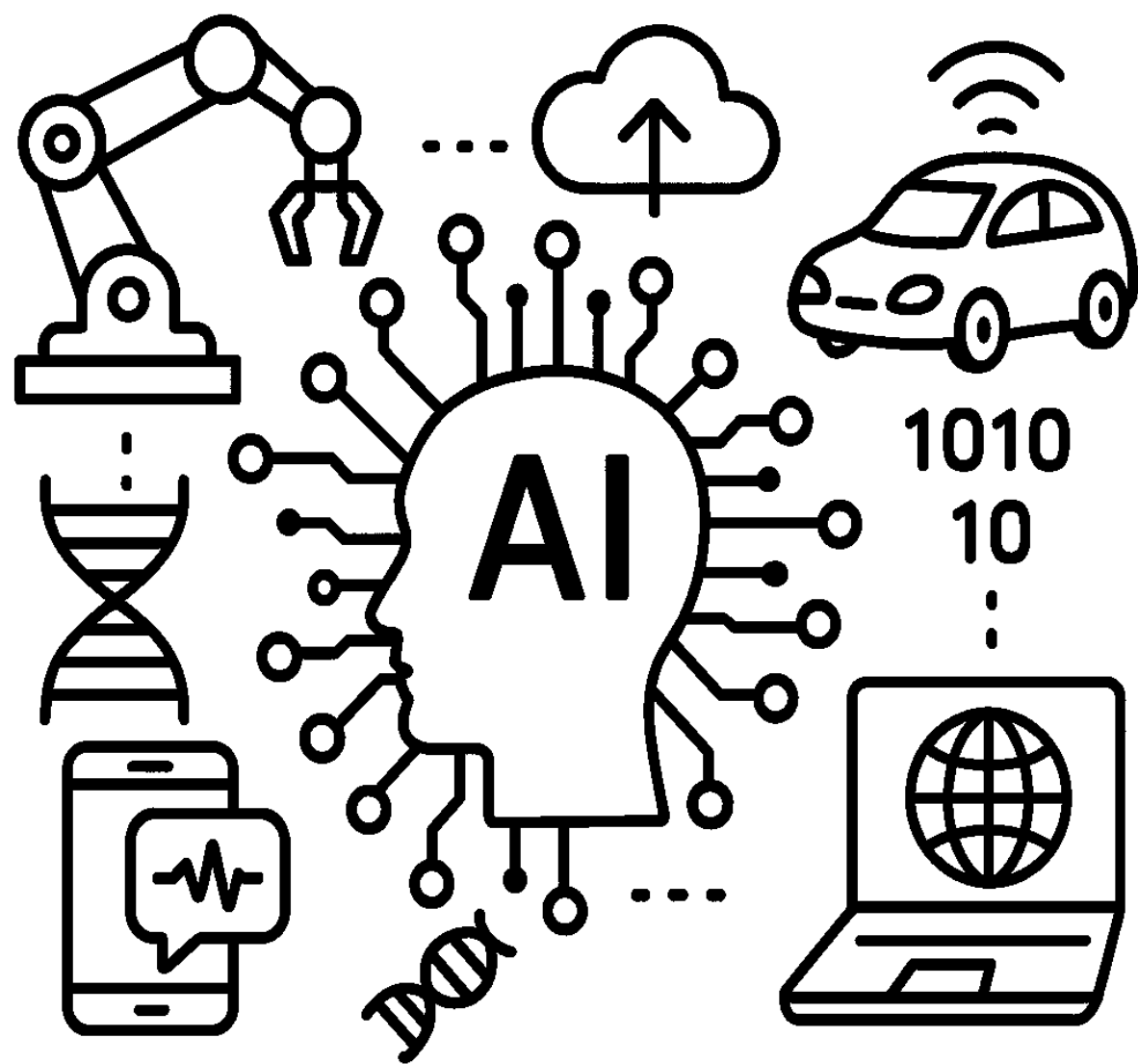




ETICKÉ VÝZVY A MORÁLNE ASPEKTY UMELEJ INTELIGENCIE

Peter Šantavý

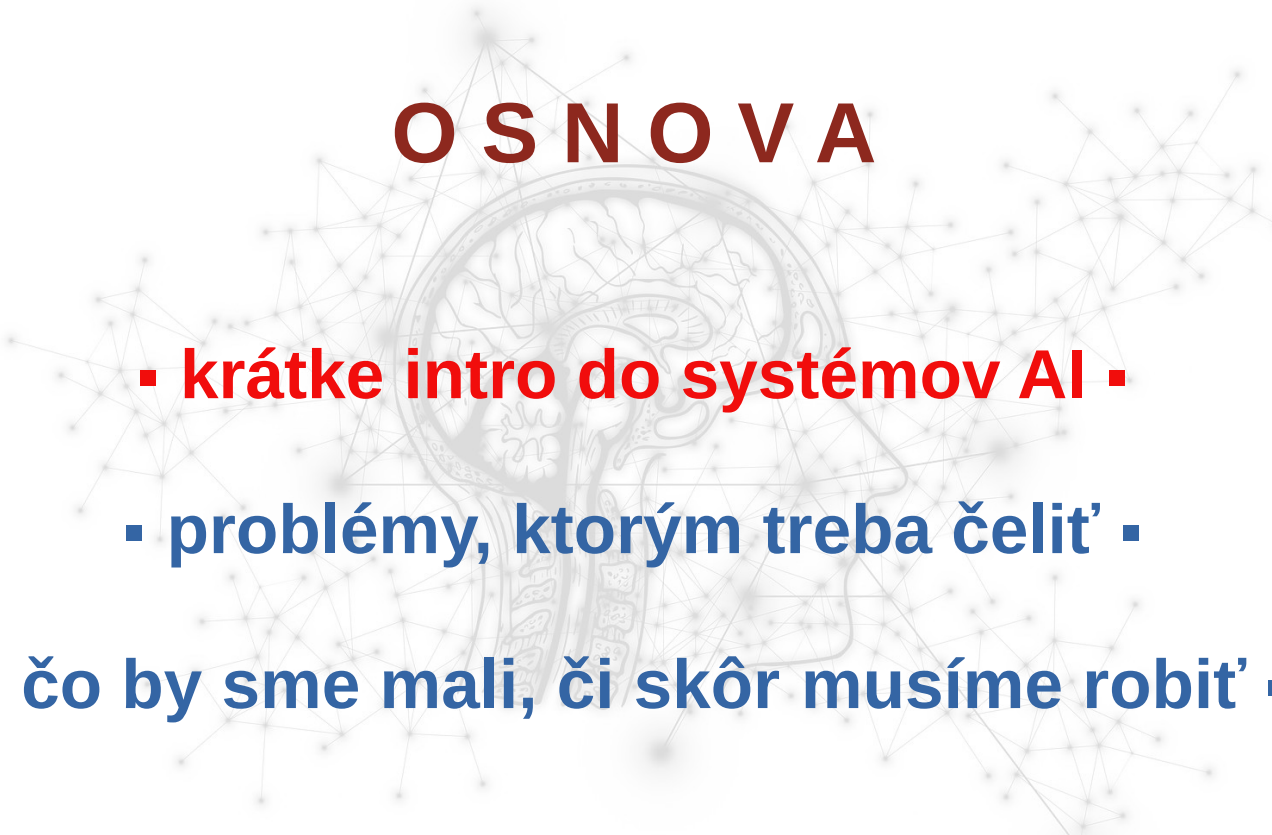






**Fenomén umelej inteligencie ako podstatná
súčasť napĺňania zmeny paradigmy...**

OSNOVA



- krátke intro do systémov AI ▪
- problémy, ktorým treba čeliť ▪
- čo by sme mali, či skôr musíme robiť ▪

AI – artificial intelligence

OSNOVA

- krátke intro

- AI

-

- treba čeliť

(Arthur C. Clarke)

- ne mali, či skôr musíme robiť

Každá dostatočne pokročilá technológia
je na nerozoznanie od mágie :-)

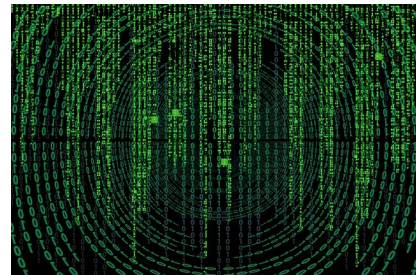
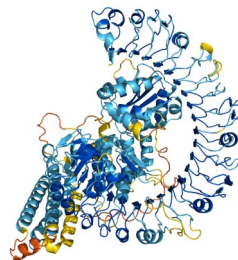
AI – artificial intelligence

Prečo umelá inteligencia?

Prečo umelá inteligencia?

Lebo nám klasické algoritmy nestačia...

- úlohy, ktoré algoritmicky **nevieme riešiť**
- úlohy, ktoré sú extrémne náročné a **nie je v ľudských silách ich zvládnuť** (časovo, organizačne, intelektuálne...)



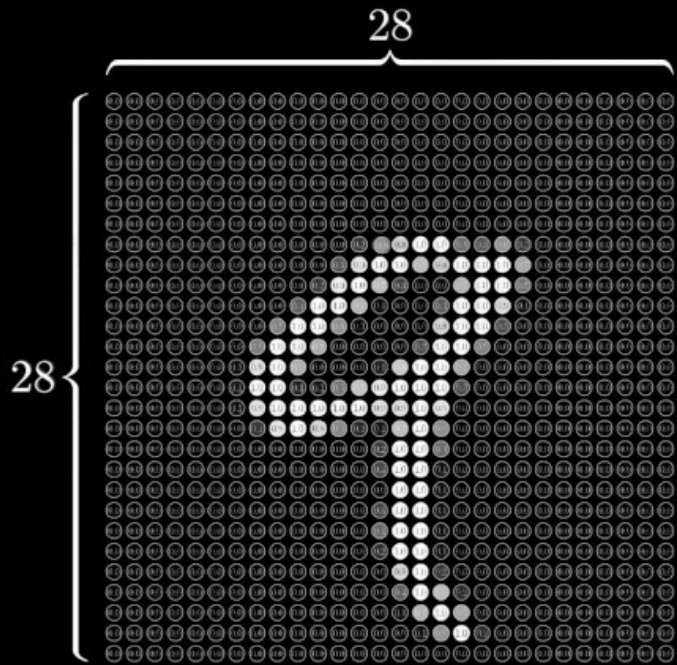
(8, 9) (0, 0) (2, 2) (1, 1) (7, 7) (8, 8) (3, 3) (7, 7)
(0, 0) (4, 4) (6, 6) (7, 7) (4, 4) (6, 6) (7, 7)
(7, 7) (8, 8) (6, 6) (1, 1) (5, 5) (7, 7)
(1, 1) (1, 1) (4, 4) (3, 3) (7, 7) (3, 3)
(1, 1) (1, 1) (0, 0) (7, 7) (0, 0) (6, 6)
(1, 1) (0, 0) (7, 7) (8, 8) (8, 8) (4, 4)
(9, 9) (7, 7) (5, 5) (4, 4) (1, 1) (8, 8)
(7, 7) (4, 4) (7, 7) (7, 7) (4, 4) (2, 2)
(8, 8) (8, 8) (1, 1) (9, 9) (3, 3) (4, 4) (6, 6)
(1, 1) (9, 9) (9, 9) (6, 6) (0, 0) (1, 1) (1, 1) (2, 2)

Ak je to ťažké naprogramovať,
nech sa to stroje naučia!



P R E Ć O A I A Ź T E R A Z ?

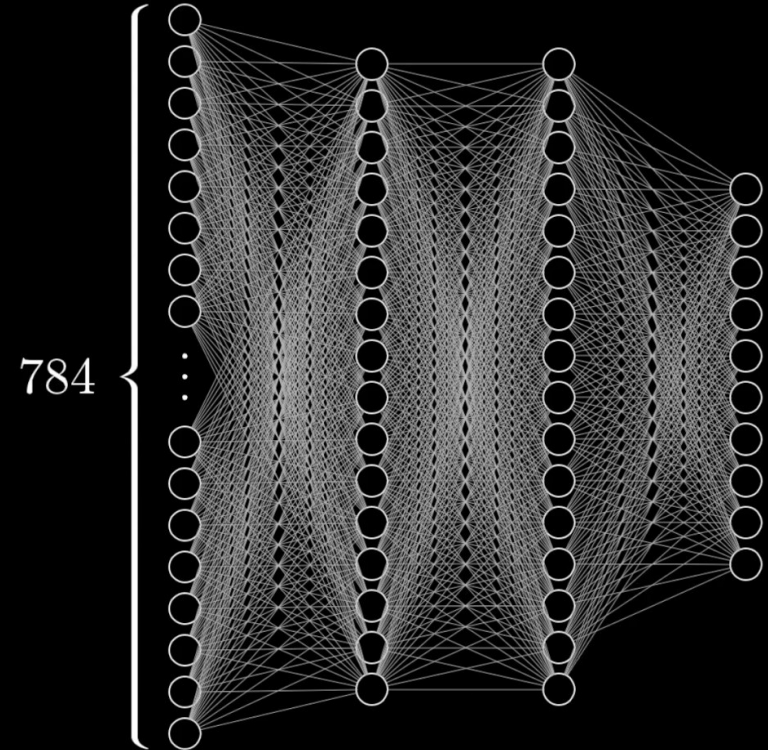
Lebo už na to máme schopnosti a využitie...



$$28 \times 28 = 784$$



13 002 !!!



28

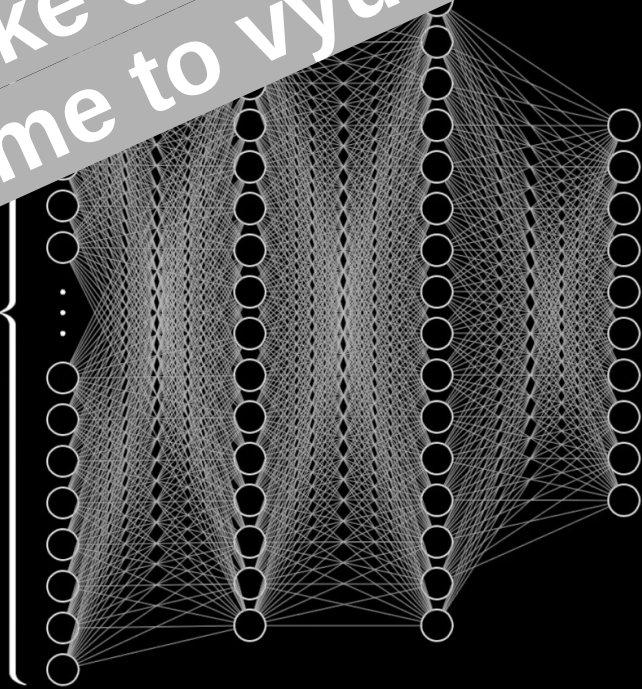
$$28 \times 28 = 784$$

13 000

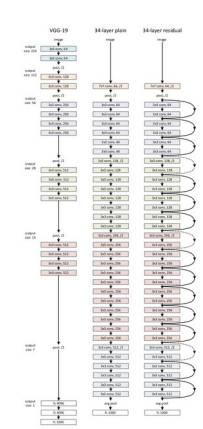
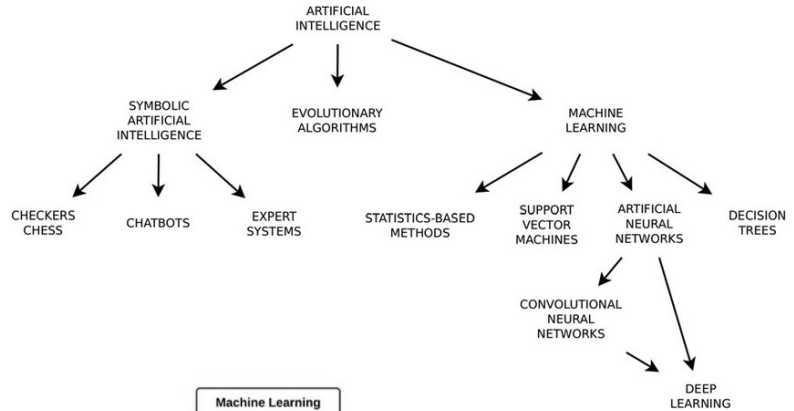
28



784

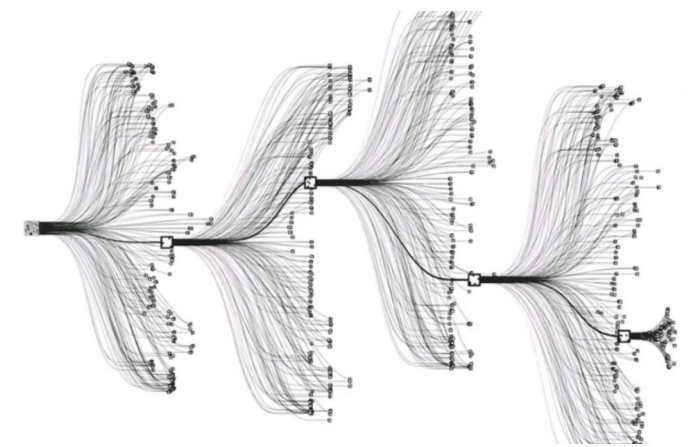
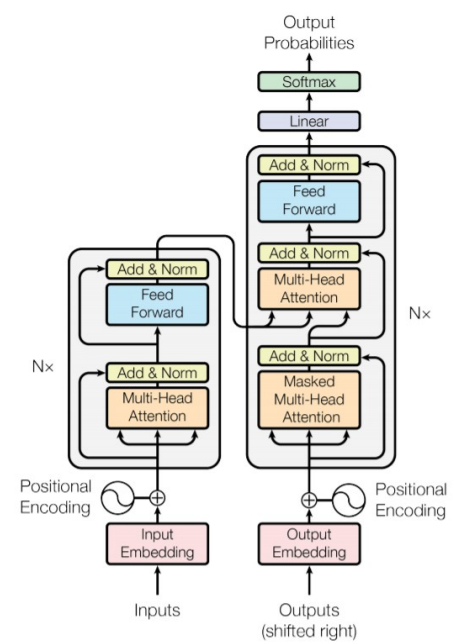
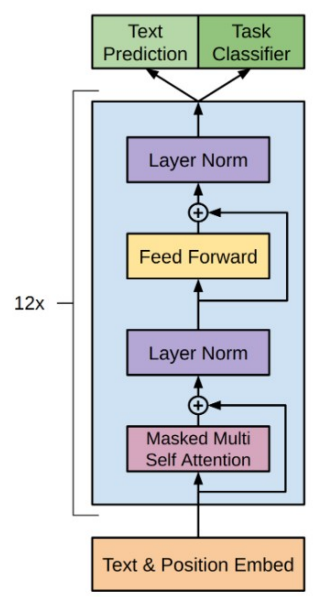
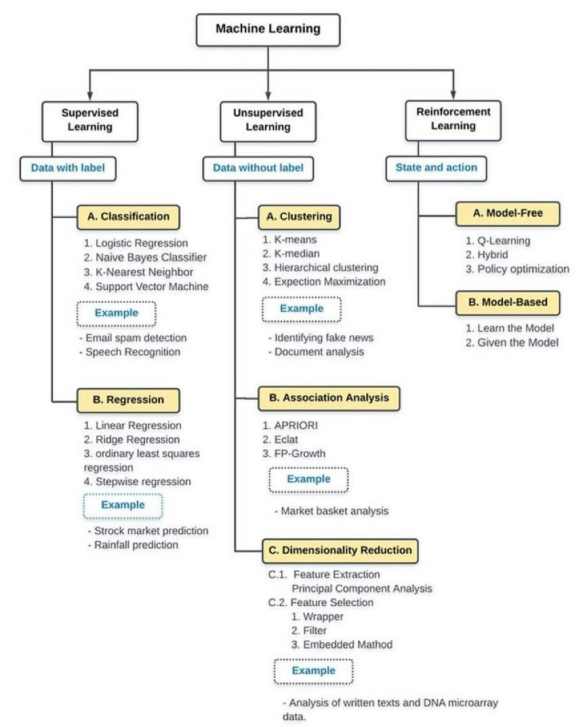
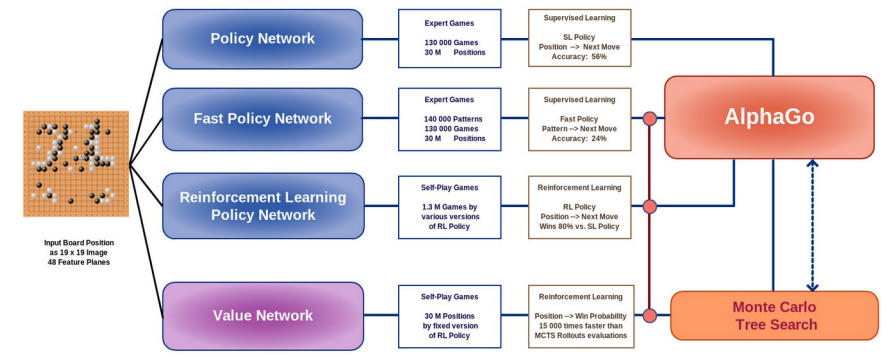


Máme algoritmy, máme veľké dáta, máme
výkonnú techniku a vieme to využiť.



AlphaGo Overview

based on: Silver, D. et al. Nature Vol 529, 2016
copyright: Bob van den Hoek, 2018



PREČO AI?

Existujú ešte nejaké iné dôvody, prečo AI?

...okrem vedeckých a technologických napr. finančné,
mocenské, vojenské, ideologické, spoločenské,...

» bytostné «

PREČO AI?

Existujú ešte nejaké iné dôvody?

...okrem vedeckých, technologických napr. finančné, mocenské, ideologické, spoločenské,...

**Túžba spoznávať, rozvíjať a tvoriť...
...v pozitívnom i negatívnom zmysle...**

» bytostné «

Definícia umelej inteligencie

Inteligenciu ľudského bytia vnímame ako schopnosť vnímať, chápať a spracovávať informácie, učiť sa, odôvodňovať, plánovať, riešiť komplexné problémy a rozhodovať. Schopnosť tiež používať abstraktné a logické myslenie, predstavivosť a uvažovať o hypotetických možnostiach, jazykový cit, atď. Inteligencia stelesnená a vzťahová...

Umelá inteligencia je o vytváraní systému, ktorý vykazuje také správanie, o ktorom si myslíme, že vyžaduje inteligenciu (AI100).

Na prasknutie naplnený kufor rôznych technológií (Marvin Minsky)...

Definícia umelej inteligencie

Inteligenciu ľudského bytia vnímame ako schopnosť vnímať, chápať a spracovávať informácie, učiť sa, odôvodňovať, plánovať, riešiť komplexné problémy a rozhodovať. Schopnosť tiež používať abstraktné a logické myslenie, predstavivosť a uvažovať o hypotetických možnostiach, jazykový cit, atď. Inteligencia stelesnená a vzťahová...

Umelá inteligencia je o vytváraní systému, ktorý vykazuje také správanie, o ktorom si myslíme, že vyžaduje inteligenciu (AI100).

Na prasknutie naplnený kufor rôznych technológií (Marvin Minsky)...

Funkcionalistická umelá inteligencia vs. komplexná inteligencia ľudskej bytosti!

Základné vlastnosti

Autonómnosť – schopnosť samostatne konať

Schopnosť systému vykonávať úlohy v komplexnom prostredí bez neustáleho vedenia používateľom.

Adaptabilita – schopnosť sa prispôsobovať

Schopnosť zlepšovať svoj výkon (a schopnosti) učením sa (nielen) zo skúseností.

Základné delenie – symbolická a subsymbolická AI

Symbolická AI – napodobňuje ľudské myslenie na úrovni pojmov, slov, fráz (= symboly) a vzťahov medzi nimi.

Na základe definovaných pravidiel a postupov („ak niečo, tak potom toto“) sú jednotlivé symboly spracovávané a vykonávajú sa priradené úlohy.

Symbolická AI – veľmi zjednodušene povedané – sa pomocou matematickej logiky snaží emulovať procesy myslenia.

Čisto symbolické systémy zlyhávali a zlyhávajú pri riešení problémov, ktoré sa nedajú exaktne popísať a v reálnych prostrediach, ktoré nie je možné deterministicky uchopiť a sú plné nejasných informácií.

Základné delenie – symbolická a subsymbolická AI

Subsymbolická AI – napodobňuje myšlienkové procesy, ktoré by sme mohli nazvať niekedy nevedomými, či automatickými, a ktoré sú základom rýchleho vnímania (fast perception; rozpoznávanie tvárí, identifikácia hovorených slov,...).

Subsymbolická AI – tiež zjednodušene povedané – sa v mnohých prípadoch snaží emulovať činnosť mozgu na úrovni neurónov. Subsymbolické systémy sú navrhnuté tak, aby sa učili **vykonávať úlohy na základe dát**.

*Väčšina moderných implementácií systémov AI, medzi ktoré patria **systémy strojového učenia a neurónové siete**, vychádza zo subsymbolického prístupu, ktorý sa snaží reálne problémy uchopiť a v určitej miere ich aj úspešne riešiť.*

Základné delenie – ANI a AGI

Slabá umelá inteligencia (ANI) – úzko špecializované systémy umelej inteligencie (**narrow AI**), ktoré sú **optimalizované na zvládnutie konkrétnej úlohy, resp. množiny úloh.**

Ide súčasne o systémy slabej umelej inteligencie (**weak AI**), ktoré vykazujú **intelligentné správanie na základe modelov a aplikovaných metód i tréningových dát.**

Hovoríme teda o systémoch, ktoré sú zamerané na riešenie konkrétnych úloh a sú závislé na ľudskom vstupe a konfigurácii.

Základné delenie – ANI a AGI

Všeobecná umelá inteligencia (AGI) – tzv. silná (**strong**) a všeobecná (**general**) umelá inteligencia.

Všeobecná, lebo **dokáže zvládnuť akúkoľvek intelektuálnu úlohu a má schopnosť zovšeobecňovať** a prenášať, či adaptovať naučené schopnosti na iné úlohy.

Silná, pretože aj **skutočne rozumie tomu, čo rieši a vykonáva**.

*Rozlišujeme rôzne stupne AGI: **BAGI** (below-human), HLAGI/HAGI (human-level AGI), MAGI (moderately-superhuman AGI), SAGI/ASI (superintelligent AGI).*

Základné delenie – ANI a AGI

Všeobecná umelá inteligencia (AGI) – všeobecná (general) umelá inteligencia.

Všeobecná, lebo dokáže vykonávať ľubovoľnú intelektuálnu úlohu a má schopnosť plánovať, učiť sa, či adaptovať naučené schopnosti na nové situácie.

AGI musí rozumieť tomu, čo rieši a vykonáva.

Existujú rôzne stupne AGI: **BAGI** (below-human), HLAGI/HAGI (human-level AGI), MAGI (moderately-superhuman AGI), SAGI/ASI (superintelligent AGI).

Nástup AGI: potrebné sú koncepčné prelomy a posun of striktno funkcionality AI k niečomu viac

Generatívne systémy AI (genAI)

Princíp činnosti genAI

- algoritmy genAI sú **schopné podľa vstupnej sekvencie generovať rôznorodý obsah** (text, obrázky, audio, video,...), ktorý napodobňuje ľudský obsah a schopnosť rozumieť
- výstupy sú generované na základe predtrénovaných vzorov a zložitých algoritmov

Technologické aspekty genAI

- data pre-processing -> foundation model -> fine tuning / assistant model -> RLHF
- tokenization, embedding (+ encoders), transformers (+S/MH attention), decoders,...

Mnohé smery súčasného i očakávaného vývoja

- **zvyšovanie kvality** odpovedí a výkonu (škálovanie, algoritmy, destilácia, quantization,...)
- **zlepšovanie “myslenia”** (reasoning modely) – RAG a prispôsobené modely
- **prepojenie s reálnym svetom** (agent AI) – multimodálne modely

Generatívne systémy AI (genAI) – prečo?

Evidentný technologický posun

- činnosť podobná ľudskému mysleniu na základe naučených vzorov a technológie transformerov dosahujúca **pozoruhodné výsledky a širokú škálu využitia**
- schopnosť čerpať z extrémneho množstva zdrojov a pracovať s bežnými dátami
- potenciál učiť sa vzory v dátach, kontext aj sémantiku jazyka, porozumieť dlhým konverzáciám, poskytovať súvislé a kontextovo relevantné odpovede

Zmena v spôsobe využívania nástrojov AI

- konverzačné rozhranie umožňujúce komunikáciu v bežnom jazyku formou dialógu
- pokročilá – akoby tímová – **spolupráca medzi človekom a strojom**
- riešenie úloh reálneho sveta prostredníctvom komunikácie s genAI (agenti AI)

Generatívne systémy AI (genAI) – prečo?

Evidentný technologický posun

- činnosť podobná ľudskému mysleniu na základe naučiacich sa modelov, najmä architektúry transformerov dosahujúca **pozoruhodné výsledky**
- schopnosť čerpať z extrémneho množstva dát a s bežnými dátami
- potenciál učiť sa vzory v dátach, porozumieť jazyka, porozumieť dlhým konverzáciám, poskytovať relevantné odpovede

Zmena paradigmatu nástrojov AI

– umožňujúce komunikáciu v bežnom jazyku formou dialógu

– nové typy tímová – **spolupráca medzi človekom a strojom**

– riešenie úloh reálneho sveta prostredníctvom komunikácie s genAI (agenti AI)

Pri správnom používaní sa systémy genAI stávajú silným technologickým prostriedkom!

Moderné systémy AI – dôsledky...

Súčasť paradigmatickej zmeny informačnej a znalostnej spoločnosti

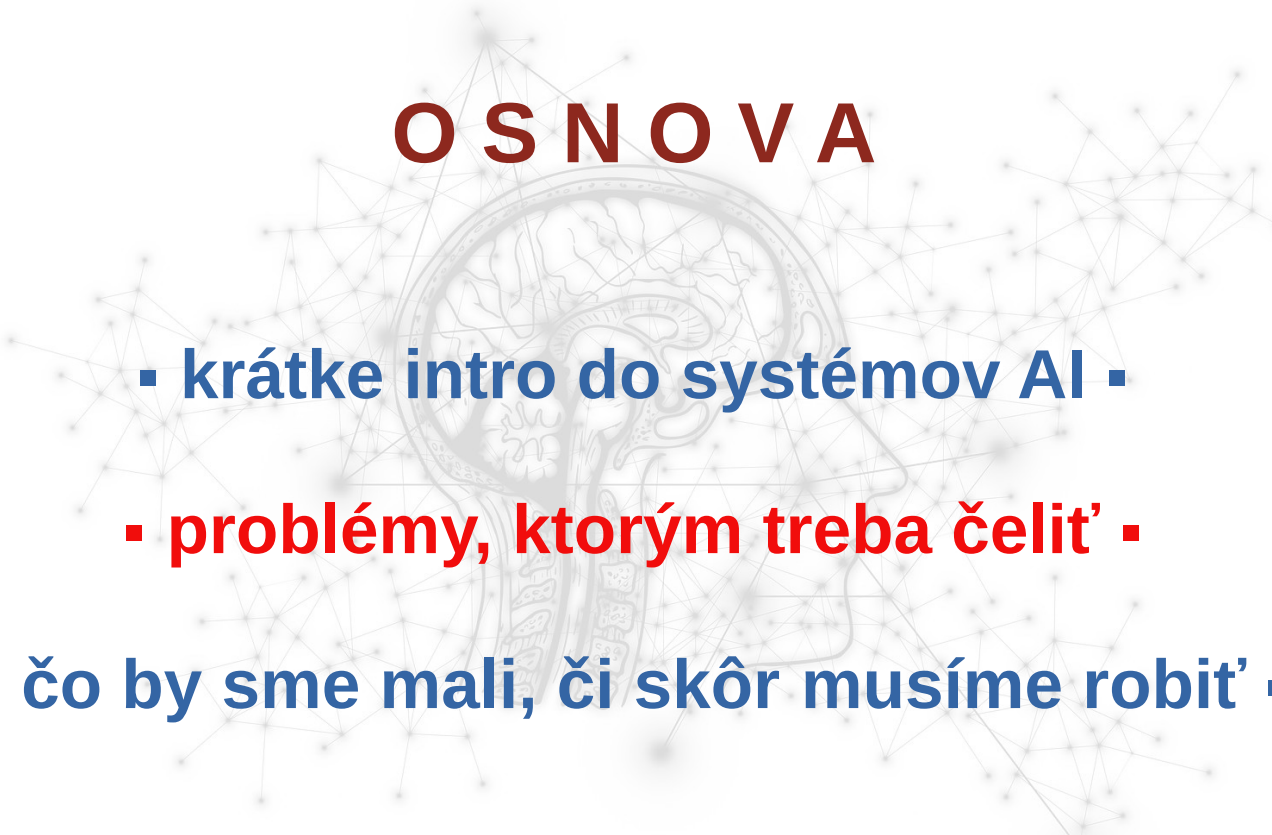
Využívanie nástrojov genAI smeruje k prirodzenému ľudskému spôsobu komunikácie a k radikálnej zmene obsahu i spôsobu práce.

Dôsledky nielen na rozvoj znalostnej spoločnosti, ale – predovšetkým – na ľudskú psychiku a sociologický rozmer spoločnosti.

Strategické technológie

Najlepšie a najsofistikovanejšie systémy AI s očakávaným potenciálom rozvoja môžeme považovať za **strategické technológie**, ktoré **majú potenciál meniť svet – k lepšiemu, ale i k horšiemu...**

OSNOVA



- krátke intro do systémov AI ▪
- problémy, ktorým treba čeliť ▪
- čo by sme mali, či skôr musíme robiť ▪

AI – artificial intelligence

O S N O V A

- krátke intro do sveta AI ▪

**Nič veľkého nevstúpi do života smrteľníkov
bez prekliatia.**
(Sofokles)

treba čeliť ▪

- čo by sme mali, či skôr musíme robiť ▪

AI – artificial intelligence

S čím sa potýkame?

- » technologické riziká a limity
- » nesprávne používanie a zneužitie systémov AI
- » dôsledky pre človeka a spoločnosť
sociologické, psychologické, civilizačné,...

S čím sa potýkame?

» technologické riziká a...

» nesprávne používanie systémov AI

» dôsledky pre spoločnosť

» psychologické, civilizačné,...

AI systémy myslenie len napodobňujú!

AI systémy "rozmyšľajú" diametrálne odlišne, inak než ľudia!

Technologické riziká a limity ANI

Zraniteľnosti, slabiny a klamanie systémov strojového učenia

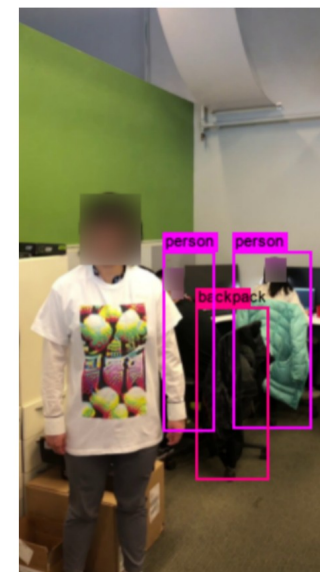
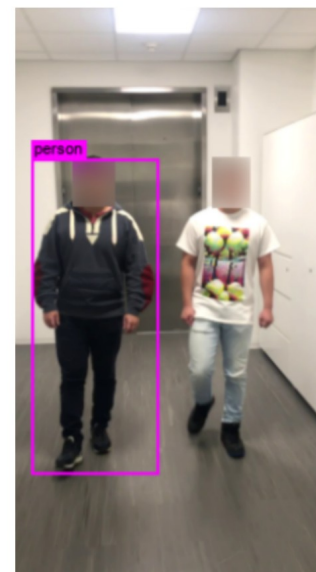
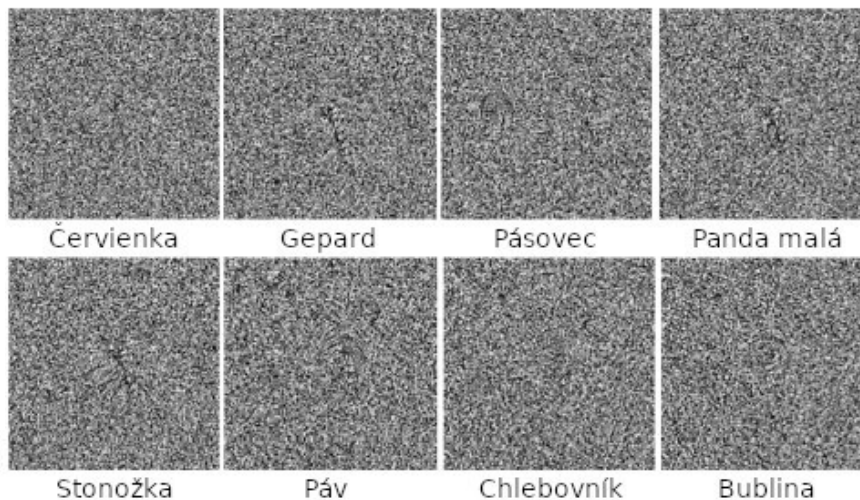
- malá množina trénovacích dát (training dataset)
- nesprávne zvolená/nekvalitná množina trénovacích dát a predsudky (biases)
- nadmerné prispôsobovanie sa tréningovým údajom (overfitting to training data)
- efekt dlhého chvosta (long-tail effect)
- povery (superstition)
- klamanie hlbokých sietí a ich zraniteľnosti (adversarial examples, hacking,...)
- špecifické zlyhania generatívnych systémov (jailbreak: roleplay, base64, UTS, prompt injections, data poisoning / backdoor attacks,...)

Technologické riziká a limity ANI

Zraniteľnosti, slabiny a klamanie systémov strojovej inteligencie

- malá množina trénovacích dát (training data)
- nesprávne zvolená/nekvalitná množina dát a predsudky (biases)
- nadmerné prispôsobovanie sa údajom (overfitting to training data)
- efekt dlhého chvilu
- povery (spoofing)
- klamanie a ich zraniteľnosti (adversarial examples, hacking,...)
- špecifické útoky na generatívnych systémov (jailbreak: roleplay, base64, UTS, prompt injections, data poisoning / backdoor attacks,...)

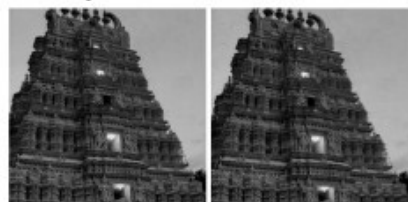
**Zlyhávajú (nielen) pri nesprávnom použití ...
Ešte viac pri zámernom zneužití...**



Školský autobus Pštros



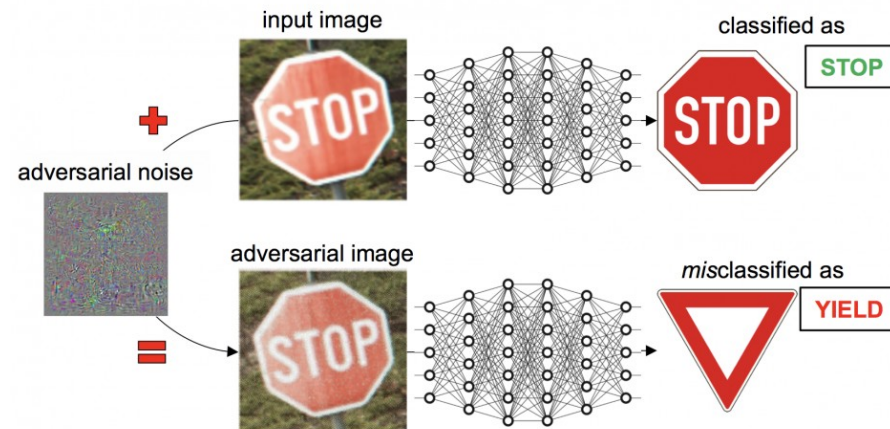
Kudlanka nábožná Pštros

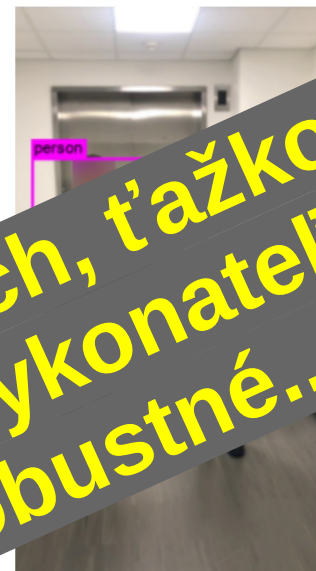


Chrám Pštros

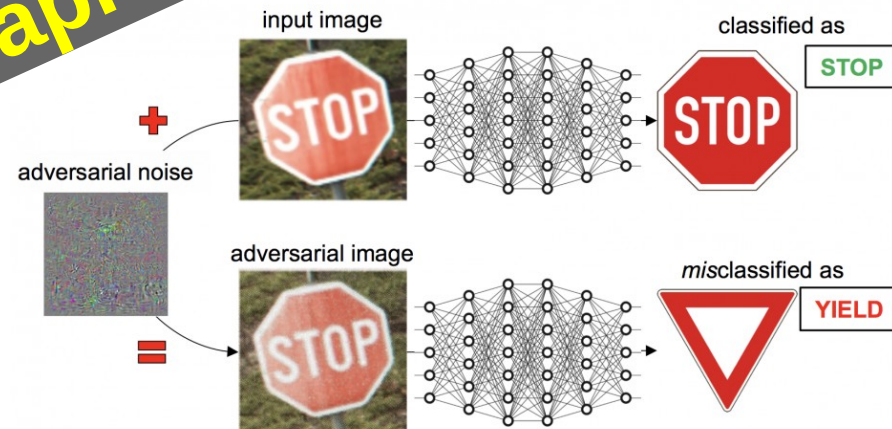


Shih tzu (plemeno) Pštros





Zlyhania sú odlišné od ľudských, ťažko predvídateľné, niekedy ľahko vykonateľné a mnohokrát prekvapivo robustné...



Procesné riziká a limity ANI

Procesné riziká v reálnom nasadení

- **útoky** na dôvernost' (confidentiality attacks)

Útoky zamerané na dáta uložené v rámci modelov systémov AI.

- **útoky** na zraniteľnosti (evasion attacks)

Odhaľovanie a zneužitie existujúcich zraniteľností v modeloch za účelom zmanipulovania výsledkov systémov AI.

- **útoky** s cieľom ovplyvniť model (poisoning attacks)

Ovplyvňovanie modelu, tréningového procesu a tým aj výslednej činnosti.

Etické výzvy, doteraz rozoberané ako súčasť návrhu a realizácie systémov AI, sa rozširujú aj o oblasť etiky použitia, resp. riziká zneužitia.

Sekundárne technologické rizikové faktory

Kybernetická bezpečnosť – je to napadnuteľné

- **vektor útoku** na zraniteľnosti AI
- **neexistuje dokonale bezpečný a spoľahlivý systém**

Infraštruktúra a komplexnosť – je to náročné

- s rozvojom algoritmov strojového učenia sa v poslednej dekáde **požadovaný výpočtový výkon i celková réžia zvyšuje exponenciálne**
- **extrémna komplexnosť** je rizikovým faktorom bezpečnosti a stability fungovania systémov

Sekundárne technologické rizikové faktory

Kybernetická bezpečnosť – je to napadnuteľnosť

- **vektor útoku** na zraniteľnosti AI

- **neexistuje dokonale bezpečný systém**

Infraštruktúra – je základom

- **Systémové zmeny** sa v poslednej dekáde **požadovaný**

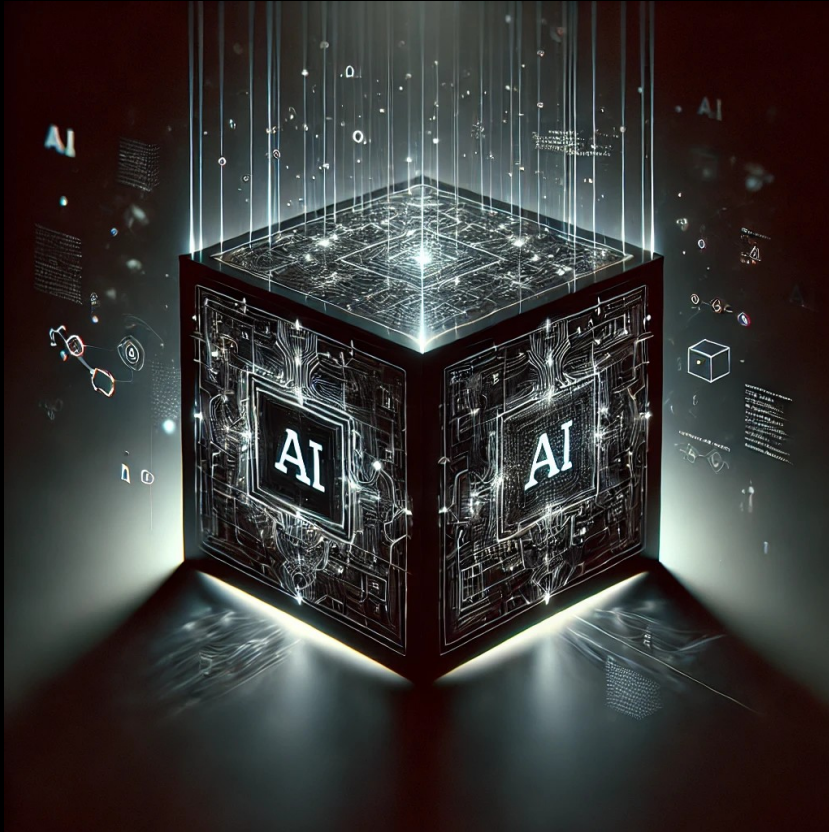
Infraštruktúra zvyšuje **exponenciálne**

Infraštruktúra je rizikovým faktorom bezpečnosti a stability

fungovania systémov

Problémom nie je len zlyhanie AI, ale i človek: zodpovednosť vývoja, etika použitia, riziká zneužitia...

BLACK BOX



Prakticky nevieme, na základe čoho robia hlboké neurónové siete svoje rozhodnutia.

V zásade nevieme, čo presne sa neurónová sieť naučila a ako spoľahlivo to dokáže aplikovať nielen v bežnej prevádzke, ale osobitne v hraničných situáciách za extrémnych podmienok na vstupe, či pri činnosti systému.

BLACK BOX



Prakticky nevieme, ako to robia hlboké učenie. Pracujeme na vysvetliteľných a interpretovateľných systémoch AI ;-)

Prakticky nevieme, ako to robia hlboké učenie. Zásade nevieme, či sa neurónová sieť naučila, a ako spoľahlivo to dokáže aplikovať nielen v bežnej prevádzke, ale osobitne v hraničných situáciách za extrémnych podmienok na vstupe, či pri činnosti systému.

Vybrané dôsledky a riziká pre človeka

Extrémne zhromažďovanie dát

- neustály **dohľad a sledovanie** bez kontroly
- **psychologické profily a modely predikcie** nášho správania
- **manipulácia a ovplyvňovanie** nášho vnímania a konania
- vytváranie **sociálnych bublín** a rozdelenia, **relativizácia hodnôt a pravdy**
- technológie AI **vedú k závislosti, depresiám, úzkostiam, digitálnej demencii a tzv. mozgovej hnilobe**

Vybrané dôsledky a riziká pre človeka

Extrémne zhromažďovanie dát

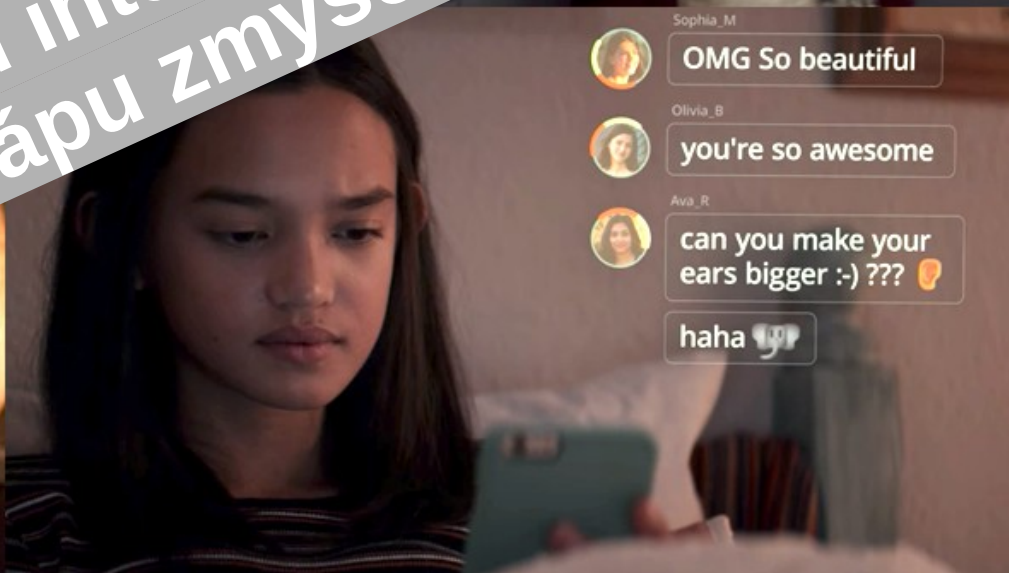
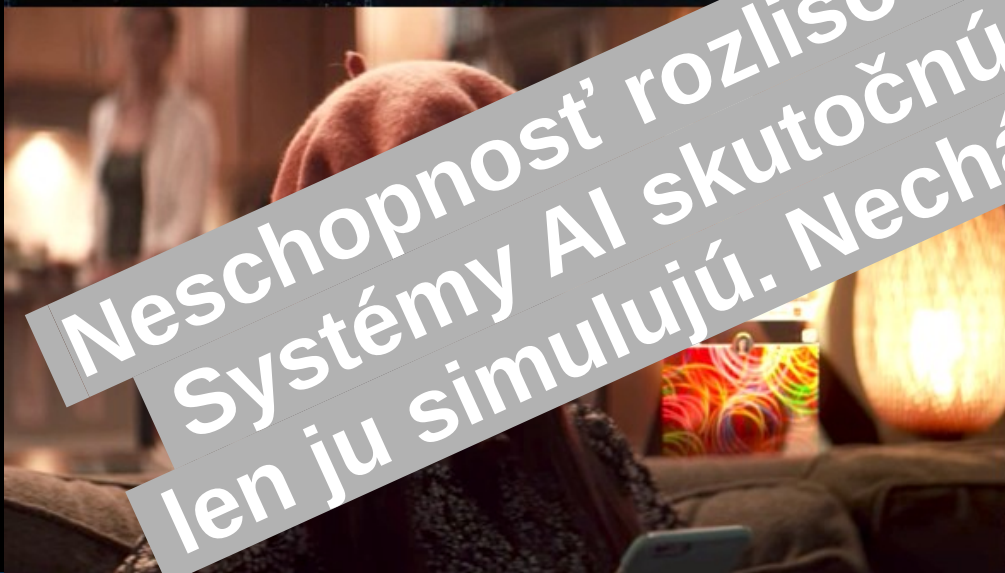
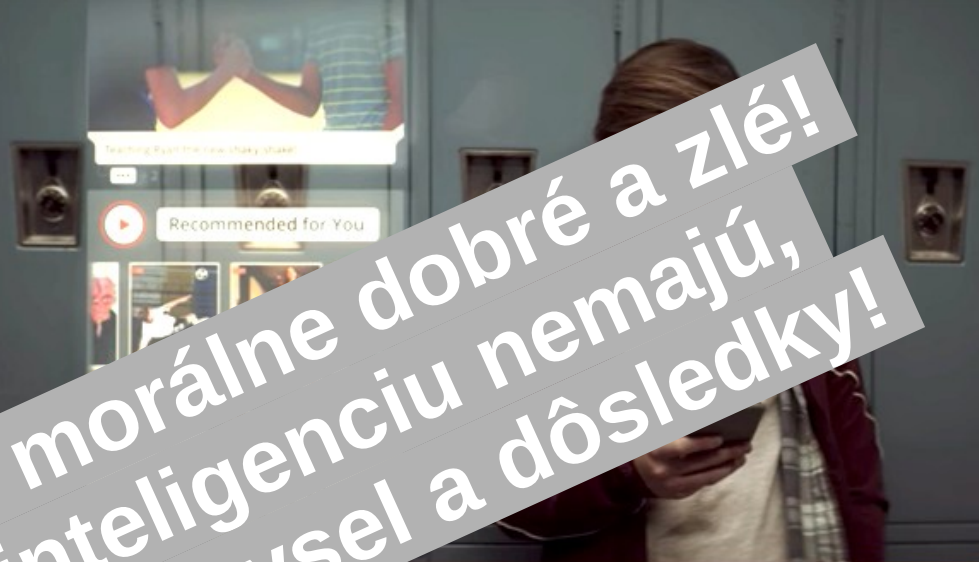
- neustály dohľad a sledovanie bez k
- psychologické profily a ... správania
- manipulácia ... vnímania a konania
- ... sociálnych bublín a rozdelenia, relativizácia hodnôt a pravdy
- technológie AI vedú k závislosti, depresiám, úzkostiam, digitálnej demencii a tzv. mozgovej hnilobe

Ide o riziká, ktoré si mnohokrát ani neuvedomujeme!!!



N

/the social dilemma



Neschopnosť rozlišovať morálne dobré a zlé!
Systémy AI skutočnú inteligenciu nemajú,
len ju simulujú. Nechápu zmysel a dôsledky!

Vybrané dôsledky a riziká pre spoločnosť

Dohľad a sledovanie ako základ...

- “z princípu” neustály dohľad a sledovanie, ktoré je **t'ážko kontrolovať**

Riadenie a fungovanie štátu

- **nezvládnuté a neadekvátne** algoritmické riadenie štátu / algokracia (súdnictvo, zdravotníctvo, bankový sektor,...)
- **sociálny kreditový systém**
- **neetické** vojenské nasadenie
- budúcnosť **práce a fungovania spoločnosti** (mení sa jej paradigma)



PHOTO: SHUTTERSTOCK/ANDREW NEWMAN; DRONE: SHUTTERSTOCK/ANDREW NEWMAN; CITY: SHUTTERSTOCK/ANDREW NEWMAN

CHINA'S SOCIAL CREDIT SYSTEM

It's been dubbed the most ambitious experiment in digital social control ever undertaken. The Chinese government plans to launch its Social Credit System nationally by 2020.

WHAT'S THE AIM?

The system intends to monitor, rate and regulate the financial, social and, possibly, political behavior of China's citizens. The country's computerized system of punishment and reward.

will be used to pilot schemes and set expectations.

Positively influencing one's neighborhood

Defeating blood

Praising the government on social media

Helping the poor

Having a good financial credit history

Committing a heroic act

Posting anti-government messages on social media

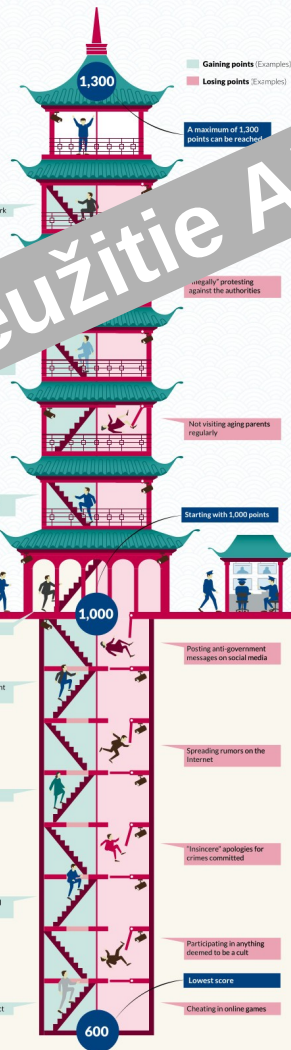
Spreading rumors on the Internet

"Innocent" apologies for crimes connected

Participating in anything deemed to be a cult

Lowest score

Cheating in online games



REWARDS AND PUNISHMENTS
Citizens with high scores get to enjoy special "privileges" while those with low scores ultimately are treated as second-class.

HIGH SCORES CAN LEAD TO

- ★ Priority for school admissions and employment.
- 🚗 Easier access to cash loans and consumer credit.
- 🚲 Deposit-free bicycle and car hire.
- 🏊 Free gym facilities.
- 🚇 Cheaper public transport.
- 🏥 Shorter wait times in hospitals.
- 🚄 Fast-track promotion at work.
- 🏠 Jumping the queue for public housing.
- 💰 Tax breaks.

PUNISHMENTS CAN LEAD TO

- 🚫 Denial of licenses, permits and access to some social services.
- 🚫 Exclusion from booking flights or high-speed train tickets.
- 🚫 Less access to credit.
- 🚫 Restricted access to public services.
- 🚫 Ineligibility for government jobs.
- 🚫 No access to private schools.
- 🚫 Public shaming: exposure either online or on TV screens in public spaces of the names, photos and ID numbers of blacklisted citizens, whose dial tones mandated by authorities that inform people that they are calling a "dishonest debtor".

BertelsmannStiftung

Vybrané riziká generatívnych systémov AI

Halucinovanie

- vymýšľanie si odpovedí, ktoré systém predkladá ako relevantné a správne

Predsudky a neobjektívne výstupy

- riziko poloprávd a nesprávnych odpovedí

Nejasný spôsob narábania s údajmi

- únik dôverných dát, problémy s ochranou osobných údajov a autorskými právami

Vybrané riziká generatívnych systémov

Halucinovanie

- vymýšľanie si odpovedí, ktoré sú nepravdivé a správne

Predsudky

- šírenie predsudkov a stereotypov

- únik dôverných dát, problémy s ochranou osobných údajov a autorskými právami

Generatívne nástroje AI nemôžeme vnímať ako faktografické, spoľahlivé a etické zdroje! Potrebujú človeka, ktorý ich vie správne používať...

Etické problémy (nielen) generatívnych systémov

Neschopnosť rozlišovať morálne dobré a zlé – systémy AI skutočnú inteligenciu nemajú, len ju simulujú. **Nechápu zmysel a dôsledky.**

Riziko digitálnej demencie – prichádza k degradácii intelektuálnych schopností a k dopamínovej závislosti na základe ponorenia sa do virtuálneho sveta, v ktorom systémy umelej inteligencie v čoraz väčšej miere **suplujú** kognitívne činnosti človeka, a to spôsobom, na ktorý nie sme evolučne vôbec pripravení.

Digitálne rozdelenie – riziko digitálneho rozdelenia (digital divide) i v oblasti umelej inteligencie (tí, ktorí nepoužívajú / sú manipulovaní / využívajú AI).



RIZIKO DIGITÁLNEJ DEMENCIE

Dlhodobé a nesprávne využívanie systémov AI, osobitne generatívnych AI a algoritmov sociálnych sietí, prináša **intenzívny útok na naše kognitívne schopnosti a psychiku.**



Môže prichádzať k degradácii intelektuálnych schopností a k dopamínovej závislosti na základe ponorenia sa do virtuálneho sveta, v ktorom systémy AI v čoraz väčšej miere supľujú kognitívne činnosti človeka, a to spôsobom, na ktorý nie sme evolučne vôbec pripravení. Mozog sa mení...



RIZIKO DIGITÁLNEJ DEMENCE

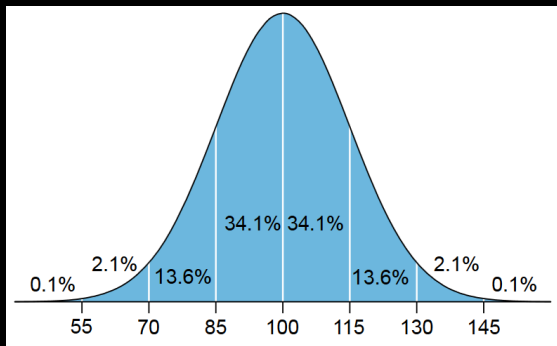
Dlhodobé a nesprávne využívanie osobitne generatívnych AI a sociálnych sietí, prináša **intenzívne kognitívne vyčerpanie a kognitívne vyčerpávanie.**

Vďaka systému AI môžeme postupne strácať časť našich schopností a k degradácii intelektuálnych funkcií a k dopamínovej závislosti na základe toho, ako sa do virtuálneho sveta, v ktorom systémy AI v čoraz väčšej miere supľujú kognitívne činnosti človeka, a to spôsobom, na ktorý nie sme evolučne vôbec pripravení. Mozog sa mení...



INTELIGENČNÉ NOŽNICE, KTORÉ SA OTVÁRAJÚ

Riziko znižovania inteligencie väčšiny populácie a súčasne intelligenčný rast menšiny, ktorá bude odolná voči digitálnej demencii.



Riziko digitálneho rozdelenia (digital divide) i v oblasti umelej inteligencie.

Etické problémy (nielen) generatívnych systémov

Mozgová hniloba (brainrot) – postupná strata schopnosti sústrediť sa a vnímať hlbší kontext. Vytvárajú sa dopamínové dráhy závislosti a návyk na také hodnoty adrenalínu a dopamínu, na ktoré nie sme evolučne pripravení.

Deepfake – vytváranie dôveryhodných virtuálnych identít ľudí a podvrhnutie ich falošnej komunikácie a konania (klamanie, nekalá konkurencia, pornografia, kriminalita, propaganda,...). Stieranie rozdielov medzi pravdou a lžou.

Strata zmyslu pre realitu – únik do virtuálneho sveta a do vzťahov s AI. Neschopnosť fungovať v reálnom svete a rozvíjať skutočné vzťahy, hodnotové posuny, závislosti, zlyhávajúce v reálnom živote (kritické myslenie, sexualita,...).

Etické problémy (nielen) generatívnych AI

Mozgová hniloba (brainrot) – postupná strata schopnosti vnímať hlbší kontext. Vytvárajú sa dopamínové a adrenalínové závislosti na umelých stimuloch a vnímaní. Vytvárajú sa dopamínové a adrenalínové závislosti na umelých stimuloch a vnímaní.

Deepfake – vytváranie falšovaných videí ľudí a podvrhnutie ich identít. Deepfake – vytváranie falšovaných videí ľudí a podvrhnutie ich identít. Deepfake – vytváranie falšovaných videí ľudí a podvrhnutie ich identít.

Stíhajúci únik – únik do virtuálneho sveta a do vzťahov s AI.

Neschopnosť fungovať v reálnom svete a rozvíjať skutočné vzťahy, hodnotové postoje, závislosti, zlyhávajúce v reálnom živote (kritické myslenie, sexualita,...).

Virtuálne zážitky v mozgu mnohých evokujú tak vysoké hladiny adrenalínu a dopamínu, aké nie sú schopní v reálnom živote dosiahnuť. Únik do virtuálneho sveta...

This image was AI-generated



This image was AI-generated

Eyeglasses appear distorted and fused to his cheek, eye area and the shadow

Deepfake: klamanie, nekalá konkurencia,
pornografia, kriminalita...





Nekonečné možnosti tvorby audiovizuálneho obsahu, ktorý vytvára virtuálny svet a neskutočnú realitu...



Nekonečné možnosti tvorby audiovizuálneho obsahu, ktorý vytvára virtuálny svet a neskutočnú realitu...

Strata zmyslu pre realitu a skutočné hodnoty!





Nekonečné možnosti tvorby audiovizuálneho obsahu, ktorý vytvára virtuálnu a neskutočnú realitu.

Strata zmyslu pre skutočné hodnoty!



Matrix: Neo bojuje za niečo, čo je skutočné...
“A koho to zaujíma, či je to skutočné?”

Ďalšie etické problémy generatívnych systémov

Neschopnosť robiť vlastné zodpovedné rozhodnutia – delegujeme to na stroje.

Narušenie psychického vývoja priskorým, resp. nezrelým používaním digitálnych a AI technológií (od synaptického plasticity až po napr. tzv. *virtuálny autizmus...*).

Riziko erózie identity – preceňovanie AI a "splošťovanie" ľudskej komplexnej inteligencie, racionality a celkovej ľudskosti na úroveň technickej inteligencie.

Strata kritického myslenia a nekritické preberanie výsledkov systémov AI spojené s relativizáciou pravdy, vytváraním sociálnych bublín a pod.

Schopnosť a možnosť manipulácie vlastníckmi systémov AI.

Propaganda a manipulácia, nástroj moci a ideológie, algoritmické modelovanie histórie,...

Ďalšie etické problémy generatívnych systémov

Neschopnosť robiť vlastné zodpovedné rozhodnutia – delirium

Narušenie psychického vývoja priskorým, resp. nadmerným používaním digitálnych a AI technológií (od synaptickej plasticity až po autizmus...).

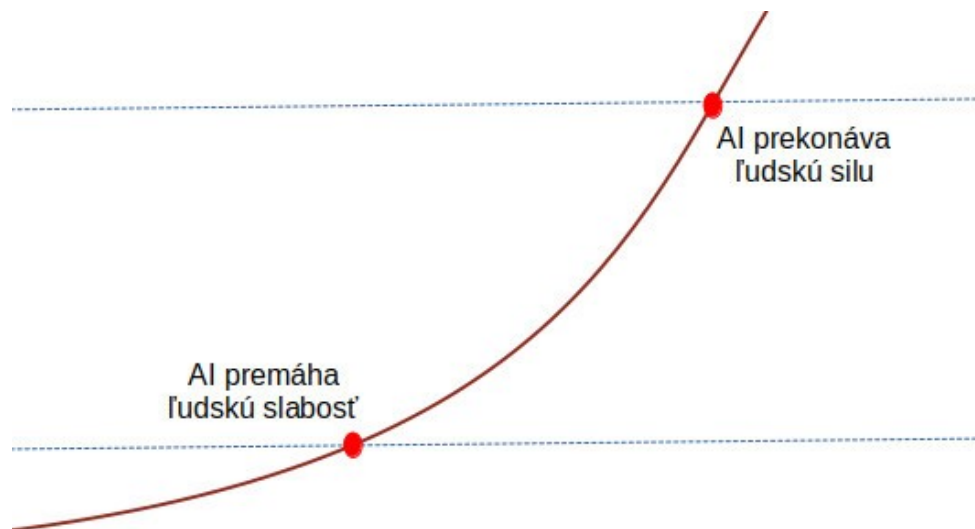
Riziko erózie identity – preceňovanie "kognitívneho výkonu" ľudskej komplexnej inteligencie, racionality a schopnosti na úroveň technickej inteligencie.

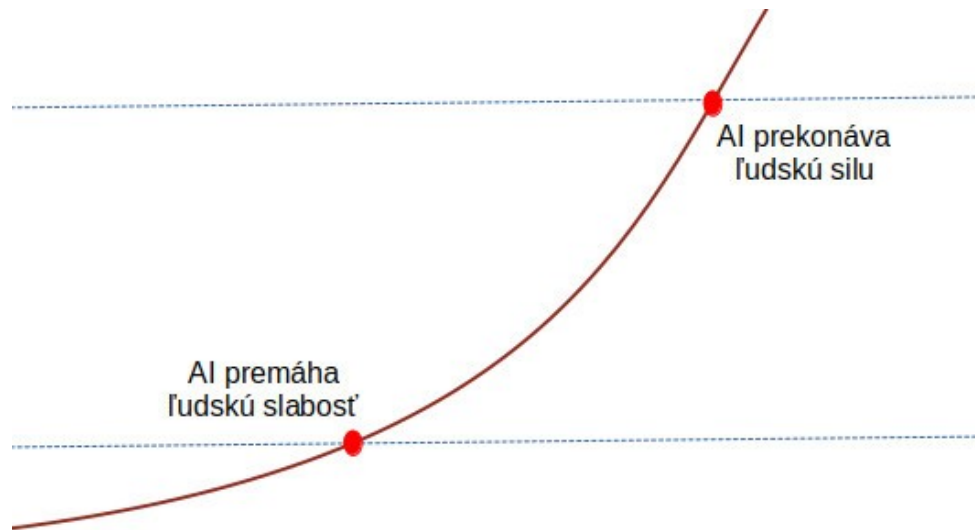
Strata kritického myslenia – nekritické preberanie výsledkov systémov AI spojené s dôverou, vytváraním sociálnych bublín a pod.

Možnosť manipulácie vlastníckmi systémov AI.

Propaganda a manipulácia, nástroj moci a ideológie, algoritmické modelovanie histórie,...

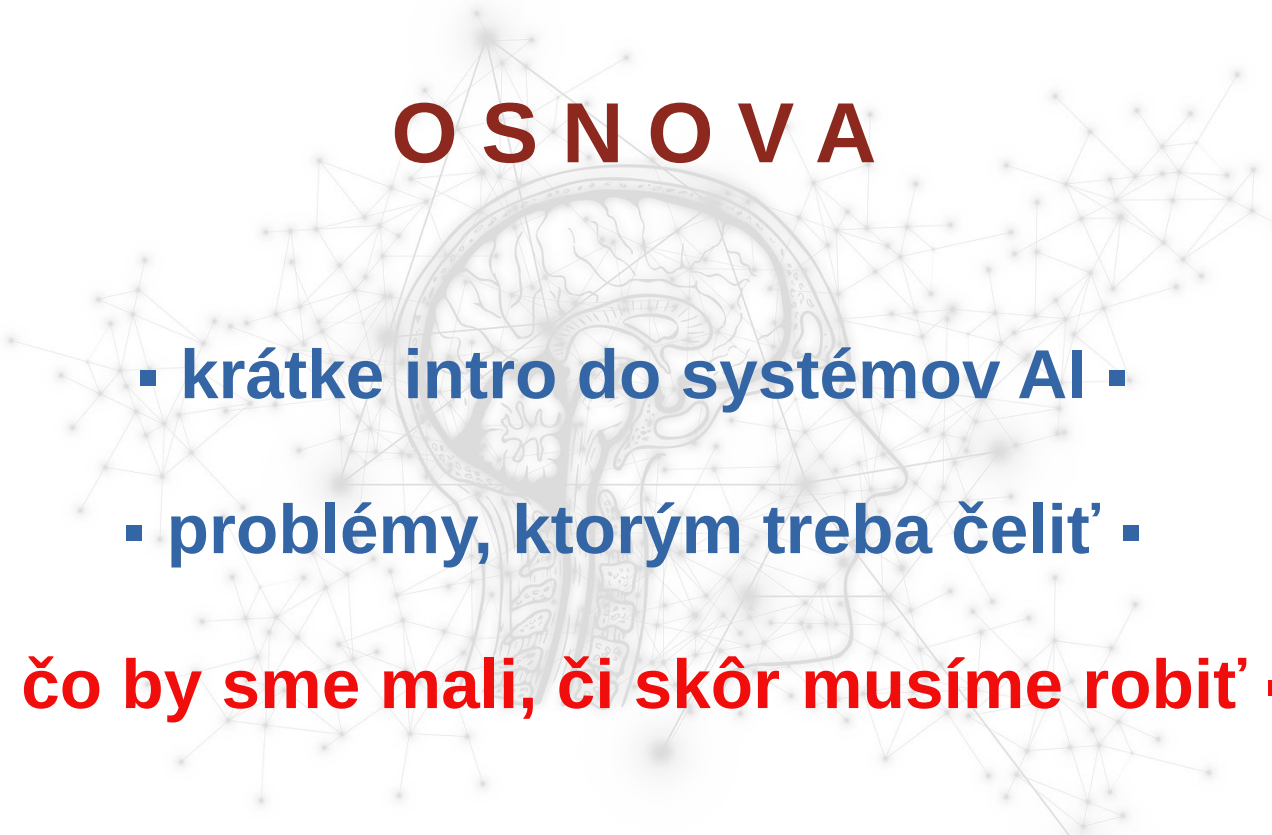
Vďaka systému AI môžeme postupne strácať časť našich schopností a psychickej odolnosti.





Míľnikom, ktorého by sme sa mali obávať, nie je budúca technologická singularita v oblasti umelej inteligencie, v ktorej AI prevýši náš intelekt, ale oveľa skôr moment, keď **technológia ovládne a prekoná naše slabosti**... už vtedy prichádza víťazstvo umelej inteligencie a porážka ľudstva.

OSNOVA



- krátke intro do systémov AI ▪
- problémy, ktorým treba čeliť ▪
- čo by sme mali, či skôr musíme robiť ▪

AI – artificial intelligence

Základné postoje vo svetle Zjavenia I.

Človek stvorený na **Boží obraz** (צֶלֶם אֱלֹהִים, imago Dei)

– inteligencia a jej používanie sú jedným zo základných aspektov naplňania "obrazu"

Užívanie (podmanenie) zvereného sveta je v Gn 1,28 vyjadrené hebrejskými výrazmi „podmaňte si ju“ (וְכִבְשֶׁתָּהּ) a „**panujte**“ (וְיָרְדוּ)

– užívať, rozvíjať, chrániť a z generácie na generáciu si odovzdávať stvorený svet

– dar inteligencie by sa mal prejavovať **zodpovedným využívaním** rozumu a akýchkoľvek civilizačných vymožeností i technicko-vedeckých schopností pri správe stvoreného sveta

Konfrontácia **chokmá** vs. *sofia* v ketubím (múdroslovnej literatúre SZ)

– ovocím tejto konfrontácie nie je zavrhnutie civilizačného rozvoja, ale neustále **pozvanie** rozvíjať a užívať svet podľa Božích zákonov, vo vzťahu voči Bohu a pre dobro ľudí

Základné postoje vo svetle Zjavenia II.

„Chod'te teda, učte všetky národy...“ (Mt 28, 19) a „Chod'te do celého sveta a hlásajte evanjelium všetkému stvoreniu“ (Mk 16, 15)

– nadčasové a všeobímajúce poslanie, zahŕňajúce aj **civilizačný rozvoj**

Sv. Pavol na misijných cestách

- **v aténskom Areopágu** (2. misijná cesta) ohlasuje jazykom helénskych mysliteľov
- v pastoračnom pôsobení v helénskom svete (jazykom synkretizmu Kolosanov,...)
- vyučovanie v škole gréckeho filozofa a učiteľa rétoriky Tyranna v Efeze (3. MC)

Nasledovníci Pavla v civilizačnej konfrontácii

- filozof a mučeník sv. Justín, predstavitelia alexandrijskej a antiochijskej školy,...
- scholastici, scholastická metóda, Occamova britva a následný **vedecký rozvoj**

Základné postoje v učení Magistéria

Dokumenty pápežov – od komunikačných prostriedkov k IKT a AI

- od Gregora XVI. (1831-1846) až po pápeža Františka
- od IKT k virtuálnemu svetu a AI v dokumentoch od Pia XI. až po Sv. otca Františka
- prorocké pozvanie a poslanie v encyklike ***Redemptoris missio*** sv. Jána Pavla II.

Koncilové dokumenty

- dekrét o spoločenských komunikačných prostriedkoch ***Inter Mirifica***
- pastoraľna konštitúcia o Cirkvi v dnešnom svete ***Gaudium et spes***
- celkovo 9 koncilových dokumentov v nejakej miere sa zaoberajúcich aj SKP a IKT

Postoje a dokumenty Cirkvi na Slovensku

- pastoračné plány KBS a rady KBS
- dokumenty biskupov a diecéznych synôd

Magistérium a problematika AI

Rome Call for AI Ethics (2020 – súčasnosť)

- **etické výzvy a princípy** v prístupe, vývoji, realizácii, prevádzke i využívaní systémov AI
- snaha o zjednocujúce pôsobenie a angažovanie náboženstiev, štátov, korporácií,...
- konferencia, dokument a participácia aktérov, ktorých zoznam sa rozširuje i v súčasnosti

Antiqua et nova (2025)

- komplexný dokument o vzťahu medzi ľudskou a umelou inteligenciou
- praktický a odborný dokument: pohľad odborníkov z rôznych uhl'ov pohľadu a profesií
- **“technologický pokrok je súčasťou Božieho plánu pre stvorenie”**

Dokumenty riadneho Magistéria o AI

- príhovory a dokumenty pápeža Františka
- dokumenty vatikánskych dikastérií

Na dobro človeka zameraná umelá inteligencia

UMELÁ INTELIGENCIA ZAMERANÁ NA DOBRO ČLOVEKA

(human-centered and beneficial artificial intelligence)

Známy a všeobecne prijímaný princíp, ktorý by však mal:

- byť chápaný v duchu **kresťanskej antropológie**, resp. klasickej filozofickej antropológie (predovšetkým biologickej a kultúrnej)
- zachovávať **ľudskú dôstojnosť** a podporovať integrálny rozvoj ľudskej osoby a spoločnosti
- zahŕňať každú ľudskú bytosť a nikoho **nediskriminovať**
- mať na zreteli **dobro ľudstva a spoločnosti**, chrániac pri tom a rešpektujúc dobro každej ľudskej bytosti
- sa vyznačovať starostlivosťou o náš „**spoločný a zdieľaný domov**“, teda o celý stvorený svet

Na dobro človeka zameraná umelá inteligencia =

= dôveryhodné systémy AI, ktoré musia byť

funkčné a užitočné – navrhnuté a realizované tak, aby vykonávali požadovanú činnosť

legálne – musia vyhovovať požadovaným normám, zákonom i reguláciám a spĺňať všetky platné zákony a predpisy

etické – musia rešpektovať etické zásady a hodnoty

odolné, resp. robustné* – dosahujúce potrebné štandardy bezpečnosti a spoľahlivosti nielen z technologického hľadiska, ale zohľadňujúce aj sociálne prostredie a dopady na spoločnosť

* schopnosť bezpečne a spoľahlivo pracovať, resp. fungovať za “akýchkoľvek” podmienok

Nutné podmienky seriózneho využívania AI

Human-centred AI a trustworthy AI – nutné princípy akéhokoľvek výskumu, vývoja, realizácie alebo nasadenia systémov AI

Funkčnosť, etika, legálnosť a odolnosť – nutná podmienka vývoja, nasadenia, prevádzky a využívania systémov AI

Vzdelávanie, osveta i prevencia a formácia morálnych postojov tvorcov, poskytovateľov i používateľov týchto technológií

Edukácia a osveta spoločnosti, bez ktorej si ťažko predstaviť rast spoločenskej citlivosti a zodpovednosti v oblasti celoplošnej adaptácie a využívania systémov umelej inteligencie

Čo treba robiť na **spoločenskej** úrovni

Mat' jasne definované etické princípy

- ich chápanie, hodnotový systém, spoločenský konsenzus

Primerané a široko akceptované regulácie

- jasné mantinely + dostatočná voľnosť pre kreativitu a biznis
- primeraný **regulačný rámec** pre celý životný cyklus systémov AI

Celoplošná osвета a vzdelávanie

- osobitne pre mladých a rodičov, primerane pre všetky vekové kategórie
- adekvátne jednotlivým profesiám a spôsobu využívania systémov AI

Jasný postoj a poznanie (**osobná rovina**)

Jasný postoj k technológiám (kriticky pozitívny a postavený na hodnotách, pre veriaceho i na Zjavení)

Základné poznanie moderných technológií a systémov AI so schopnosťou **vedieť využiť ich potenciál**

Poznanie a chápanie rizík, ktoré sú spojené s AI a modernými technológiami a ich dôsledkov

– vnímanie mnohých rizikových faktorov je viazané aj na náš hodnotový systém a svetonázor

Spôsob využívania systémov AI (**osobná rovina**)

V kontexte čností používať na dobré ciele

- **múdrosť, pravda a spravodlivosť**, na ktorých je postavená autentická schopnosť tvoriť, spravovať a rozvíjať...
- **miernosť** pre primerané a **mravnosť** pre osožné používanie AI

Používať nástroje AI s pokorou vo vedomí vlastných mantinelov a slabostí

- bez pravdy nielen o technológiách moderného sveta, a predovšetkým **pravdy o sebe** to nepôjde (moje obmedzenia, slabosti a zranenia...)

V askéze

- schopnosť **byť off-line** mimo systémov AI a virtuálneho sveta
- **vytrvať v zásadách** správneho využívania
- vedome **realizovať činnosť**, ktorá je “**nudná**” a nie je spojená s dopamínovým a adrenalínovým efektom (dopamínový pôst)
- **činnosť v reálnom svete** – fyzická, kreatívna, služba, vzťahy

Spôsob využívania systémov AI (osobná rovi

V kontexte čností používať na dobré ciele

- **múdrost', pravda a spravodlivosť'**, na ktorých je schopnosť tvoriť, spravovať a rozvíjať...
- **miernosť'** pre primerané a **mravnosť'** pri použití AI

Používať nástroje AI s pok

- bez pravdy nielen o **sebe** a o tomto svete, a predovšetkým **pravdy** o sebe to nepôjde (a, slabosti a zranenia...)

V asl

– činnosť mimo systémov AI a virtuálneho sveta
– zásadách správneho využívania

– nie realizovať činnosť, ktorá je "nudná" a nie je spojená s dopamínovým a adrenalínovým efektom (dopamínový pôst)

- **činnosť v reálnom svete** – fyzická, kreatívna, služba, vzťahy

V kontexte stanovených cieľov nebát' sa byť kreatívny, tvorivý a zodpovedne odvážny.

V službe iným (**osobná rovina**)

Formovať tých, ktorí sú mi zverení

- **chrániť deti** pred prístupom k AI, jej obsahu a nesprávnemu využívaniu
- **formovať a vychovávať** k správne mu používaniu
- **prakticky viesť** k správne mu využívaniu iných (otcova dielňa, pedagóg, výskumný tím...)

Angažovať sa v oblasti AI

- etické a transparentné systémy AI v ochrane ľudskej dôstojnosti, spoločného dobra, sociálnej spravodlivosti, pravej slobody a práva

AI vo vzdelávaní

AI ako neriadená strela

- **mnohoraké využívanie (i zneužívanie)** nástrojov AI študentami
- **otvorený príbeh** s rôznymi možnosťami využívania a ovocím, čo to prinesie

AI vo vzdelávaní

AI ako neriadená strela

- mnohoraké využívanie (i zneužívanie) nástrojov AI študentami
- otvorený príbeh s rôznymi možnosťami využívania a ovocím, čo to prinesie

Príklad pravidiel využívania nástrojov AI na UK (smernica + komisia)

- UK **podporuje** aktívne využívanie AI a získavanie potrebných zručností
- **cieľ**: tvorivé, bezpečné, zodpovedné, etické a transparentné používanie
- nástroje AI majú byť **podporou a pomocou pri štúdiu**
- **priame vypracovanie celých textov** pomocou nástrojov AI je **plagiarizmus**
- používanie nástrojov AI musí byť **riadne citované**
- pravidlá stanovujú **vyučujúci** (+ výzva prečo pedagóg a nie AI tútor)
- **vízia**: personalizácia vzdelávania; hĺbková analýza dát a pohľad na štúdium

Smernica rektora UK | Dodatok č.1 (citovanie) | Dodatok č. 2 (využívanie)

Nástroje a využívanie AI na RKCMBF UK

Čo by som chcel záverom akcentovať

Potenciál fenoménu umelej inteligencie

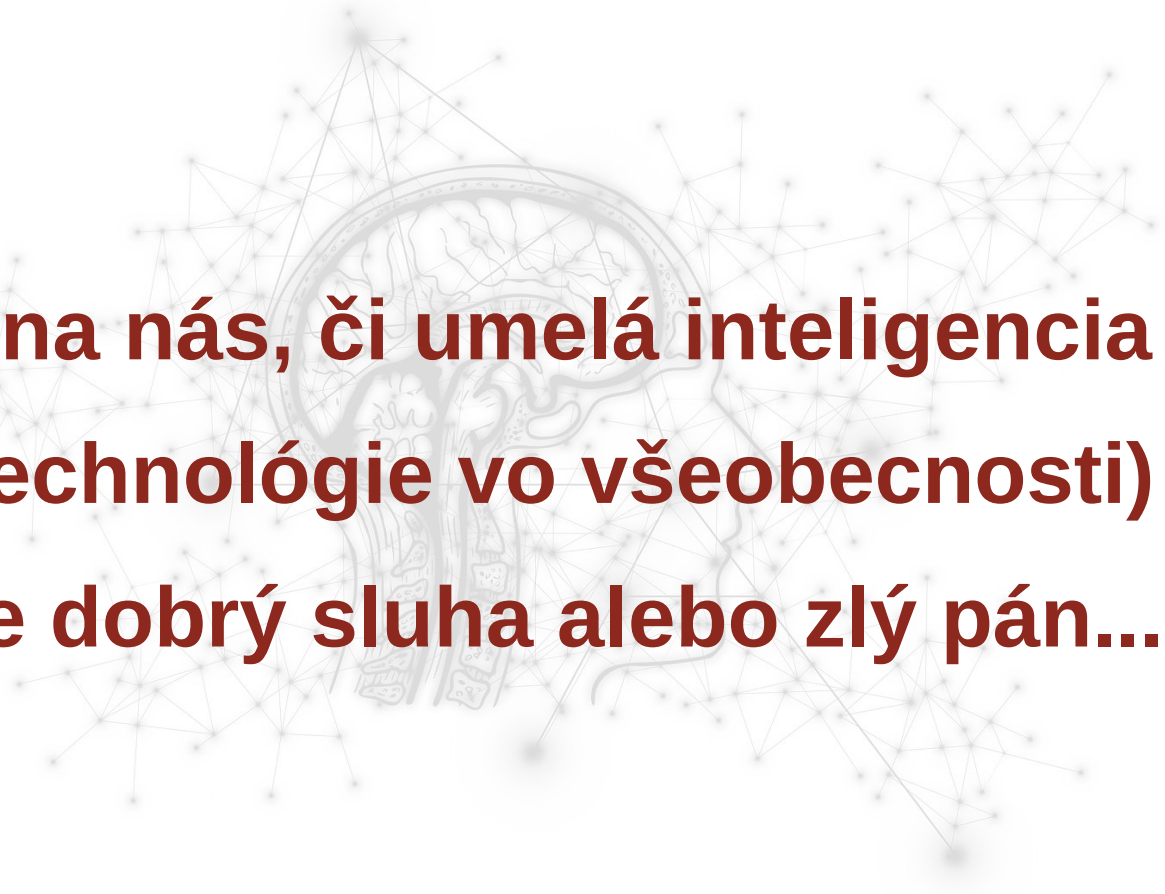
Ide o emergentnú a disruptívnu technológiu so schopnosťou zásadne meniť spôsob, akým spoločnosť funguje a s potenciálom posunúť civilizáciu ďalej...

Skutočnosti, ktoré si (ne)uvedomujeme – riziká a dopady

- riziká a limity fungovania i problematiku zneužitia
- psychologické a spoločenské dopady

Odvahu, snahu a schopnosť mať “opraty pevne v rukách”

- etické princípy a regulácie
- osobný postoj a spôsob vývoja, nasadenia, prevádzky a využívania
- vzdelávanie a osveta



**Je na nás, či umelá inteligencia
(a technológie vo všeobecnosti)
bude dobrý sluha alebo zlý pán...**



Ďakujem za pozornosť

ThLic. Ing. Peter Šantavý, PhD.
peter@santavy.sk



UMELÁ INTELIENCIA

V texte sú pre umelú inteligencia použité skratky UI (umelá inteligencia) a AI (artificial intelligence).

Univerzita Komenského **podporuje využívanie** nástrojov umelej inteligencie študujúcimi, ako aj vyučujúcimi, no **s transparentne nastavenými pravidlami** obzvlášť v otázke uvádzania zdrojov. Nástroje UI by sa mali používať najmä ako podpora a pomoc pri štúdiu a mali by slúžiť na rozvoj a efektívne využívanie nadobudnutých vedomostí, schopností a zručností študentov. Vyučujúcim sa odporúča priebežné oboznamovanie sa so systémami UI a používanie jej nástrojov pri pedagogickej činnosti.

Univerzita vo svojich nariadeniach a usmerneniach reaguje na vplyv UI na akademické prostredie a formuluje princípy pre **tvorivé, bezpečné, zodpovedné, etické a transparentné používanie** jej nástrojov v prostredí univerzity. Nástroje UI sú určené na používanie pri príprave a úpravách textov, možno ich však primerane aplikovať aj na iné formy generatívnej UI. Univerzita pripravila a priebežne aktualizuje aj **odporúčané spôsoby citovania** nástrojov UI.

Univerzita Komenského dbá na **dodržiavanie akademickkej integrity, transparentnosti a spravodlivosti**. Rozvoj a využívanie nástrojov UI vníma ako podnet, ktorý by mohol byť spúšťačom dôležitých zmien vo vzdelávaní – za predpokladu, ak tento fenomén dokážeme správne uchopiť a prispôsobiť sa novým podmienkam.

Smernica rektora UK pre využívanie AI

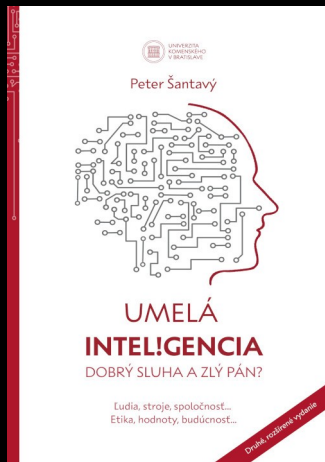
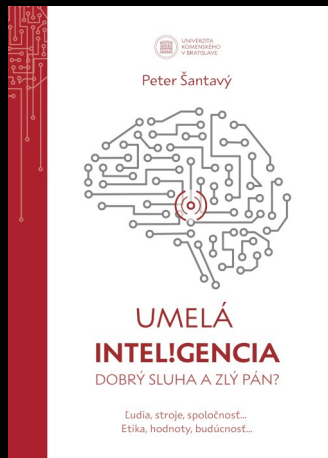
Smernica rektora Univerzity Komenského v Bratislave k používaniu nástrojov umelej

Nástroje a informácie k AI v rámci vzdelávacieho procesu na UK RKCMBF UK - Umelá inteligencia

Niečo na čítanie...

Systemy umelej inteligencie skutočnú
inteligenciu nemajú, len ju simulujú
Rozhovor pre NM

Moje knihy z oblasti AI



Etické výzvy a morálne aspekty AI | Anotácia

Fenomén umelej inteligencie už dávno prekročil rámec atraktívnej technológie úzko využiteľnej v špecifických akademických, vojenských či priemyslených nasadeniach. Povedané rečou moderného technokrata – **ide o emergentnú a disruptívnu technológiu**, teda novú a rozvíjajúcu sa technológiu, ktorá nahrádza a ruší technológie predošlé so schopnosťou zásadne meniť spôsob, akým spoločnosť funguje, pričom má potenciál posunúť civilizáciu ďalej. Čo však už technokrat nepovie – ide o technológiu, ktorú ak nezvládneme, má potenciál nás priviesť k skutočne veľkým problémom.

Ruka v ruke s nárastom potenciálu a dopadu technológií umelej inteligencie sa **etika, legálnosť, bezpečnosť i odolnosť** týchto systémov a ich nasadenia stávajú alfou a omegou ich **zodpovedného využívania pre dobro človeka a spoločnosti**.

Príspevok chce poukázať na niektoré riziká a výzvy v technologickej oblasti a ich sociologické, psychologické i civilizačné dôsledky s akcentom na základný etický rámec dôveryhodných systémov umelej inteligencie a ich zodpovedného využívania.